
ECONtribute

Discussion Paper No. 369

Model Uncertainty

Robin Musolff

Florian Zimmermann

July 2025

www.econtribute.de



**UNIVERSITÄT
ZU KÖLN**

Model Uncertainty

Robin Musolff and Florian Zimmermann*

This version: July 30, 2025

Abstract

Mental models help people navigate complex environments. This paper studies how people deal with model uncertainty. In an experiment, participants estimate a company's value, facing uncertainty about which one of two models correctly determines its true value. Using a between-subjects design, we vary the degree of model complexity. Results show that in high-complexity conditions people fully neglect model uncertainty in their actions. However, their beliefs continue to reflect model uncertainty. This disconnect between beliefs and actions suggests that complexity leads to biased decision-making, while beliefs remain more nuanced. Furthermore, we show that complexity, via full uncertainty neglect, leads to higher confidence in the optimality of own actions.

1 Introduction

People rely on mental models to interpret and navigate the complexities of external reality. These models are mental frameworks that shape how individuals process information, form expectations, and make decisions. In contexts of economic decision-making, mental models play a critical role, as individuals attempt to simplify and make sense of complex environments. However, in many situations, economic agents face not only uncertainty about future outcomes, but also model uncertainty—uncertainty regarding which mental model best captures the true dynamics of the environment.

*Musolff: Bonn Graduate School of Economics and IZA (robin.musolff@uni-bonn.de); Florian Zimmermann: University of Bonn, IZA and Max Planck Institute for Research on Collective Goods (florian.zimmermann@iza.org). We are grateful to Chiara Aina, Kai Barron, John Conlon, Benjamin Enke, Tilman Fries, Nicola Gennaioli, Alex Imas, Botond Koszegi, Alexander Laubel and Chris Roth for very helpful comments and suggestions. Funding by the Deutsche Forschungsgemeinschaft (DFG) through CRC TR 224 (Project A01) and under Germany's Excellence Strategy - EXC 2126/1 - 390838866 is gratefully acknowledged. This project has received funding from the European Research Council (ERC) under the European Union's Horizon 2020 research and innovation programme (ERC Starting Grant: 948424 - MEMEB).

Model uncertainty is prevalent in many areas of economics. For instance, when forming expectations about future inflation or returns on investments, individuals may be uncertain about which underlying model of the economy is most appropriate to use. Similarly, when assessing the effectiveness of government policies, people may lack clarity on which mental model best explains how government actions translate into economic outcomes. In these cases, individuals face a dual problem: not only must they process new information, but they must also decide how to map this information into a decision, given uncertainty about competing models of the world. This problem arguably is especially severe in environments where working with models is complex. The question hence arises how people function under model uncertainty, in particular when model complexity is high. The uncertainty about which model is correct can have profound implications for decision-making, beliefs, and ultimately, economic outcomes.

Existing research on mental models has largely focused on understanding the determinants of model selection, how mental models can be used to persuade others, and the consequences of relying on misspecified models. For example, Schwartzstein and Sunderam (2021) and Barron and Fries (2024) examine how individuals might be persuaded to adopt specific mental models under different contexts, while Heidhues, Kőszegi, and Strack (2018) explore how biased models can lead to systematic mistakes in decision-making. Similarly, recent work by Frick, Iijima, and Ishii (2022) investigates the strategic use of mental models in shaping economic beliefs. However, less is known about how individuals cope with model uncertainty in the first place.

In this paper, we aim to fill this gap by studying how individuals operate under model uncertainty. We hypothesize that the complexity of mental models determines whether people neglect model uncertainty. Our experimental results indeed reveal that, when complexity is high, people simplify the world by operating as if one model of the world is correct, hence fully neglecting model uncertainty. This is echoed by hovering data which suggest that when complexity is high, people have a tendency to predominantly attend to one specific model of the world. These results are robust to variations in the signal space as well as variations in the specific task people need to perform. Turning to implications of this complexity-induced simplification, we document two key results: (i) neglect of model uncertainty does not translate into a distorted view of the world. When we directly elicit beliefs about which model of the world is correct, these beliefs do reflect

model uncertainty despite model neglect in actions, creating a wedge between actions and beliefs; (ii) model neglect creates an illusion of certainty. The complexity-induced simplification leads to higher levels of confidence in the optimality of own actions such that, perhaps counter to intuitions, complexity increases rather than decreases decision confidence in our setting.

We design an experiment that allows us to infer from their actions how people deal with model uncertainty. We implement our experiment in a financial decision-making context where participants task is to provide valuations for fictitious companies. There is a set of variables that are potentially relevant to determine company values. Participants face uncertainty about which of two models of the world makes correct use of these variables to estimate the value of a company. Participants are endowed with a 50-50 prior and then receive a noisy signal about which model is correct. In a between-subjects design, we manipulate the complexity of these models, keeping all other aspects of the decision environment constant. In the low-complexity condition, each model provides a direct estimate of the company's value using the variables as input, requiring no computations from participants to get at each model's estimate. In the high-complexity condition, we add a layer of computational complexity, as participants need to compute the value estimates of the models themselves using the variables as inputs. Participants provide value estimates for eight different companies. Afterwards, we also measure beliefs about which model participants think is correct. This allows us to examine the extent to which complexity influences both actions (i.e., the valuation) and beliefs (i.e., the perception of which model is correct).

Our findings suggest that complexity significantly alters how people deal with model uncertainty. Compared to the low-complexity condition, participants in the high-complexity condition have a more pronounced tendency to simplify the world by neglecting model uncertainty in their actions, behaving as if the more likely model is definitely correct. This tendency to act as if there is no model uncertainty implies an overreaction to the signal about the correct model. In other words, when complexity is high, participants place too much weight on the value estimate provided by the more likely model. Data on hovering times indicate that complexity affects decision-making through an attention channel (Bordalo, Gennaioli, and Shleifer (2012), Bordalo, Conlon, Gennaioli, Kwon, and Shleifer (2023b), Bordalo, Conlon, Gennaioli, Kwon, and Shleifer (2023a), Ba, Bohren,

and Imas (2022)). In the high-complexity condition, participants spend significantly more time focusing on the signal-congruent model, whereas in the low-complexity condition, their attention is more evenly distributed between the two models.

In another experiment, we verify that simplification does not require one model to be more likely than the other. When both models of the world are equally likely, complexity nonetheless leads to a neglect of model uncertainty. Additional experiments further underscore the robustness of our results. When we replace the value estimation task with an investment task, we continue to see that complexity leads to a neglect of model uncertainty.

Next, we turn to implications of this complexity-induced simplification. We first document that neglect of model uncertainty in actions does not translate into beliefs. Participants' stated beliefs continue to reflect model uncertainty. This disconnect between beliefs and actions creates a systematic wedge between what individuals believe and how they behave under model uncertainty. This echoes earlier work by Giglio, Maggiori, Stroebel, and Utkus (2021), Ameriks, Kézdi, Lee, and Shapiro (2020), Beutel and M. Weber (2023), Laudenbach, A. Weber, R. Weber, and Wohlfart (forthcoming) who have identified a gap between subjective beliefs and economic behavior in different contexts.¹ Notably, this gap arises directly after participants stated their valuations, and persists even after a 1 day delay, as we verify in a separate experiment that spans over 2 days.

Second, using additional experiments where we measure participants' confidence in the optimality of their actions, we document that confidence in the optimality of own actions is higher in the high-complexity conditions compared to the low-complexity conditions. This holds for an unincentivized confidence measure, where participants state the probability that their guesses were optimal, as well as an incentivized measure, where participants place a bet on the optimality of their actions. Prior literature in contexts different from ours has shown that complexity tends to increase cognitive uncertainty (Enke and Graeber (2023), Enke, Graeber, and Oprea (forthcoming), Enke, Graeber, and Oprea (2023)). In contrast, it seems that in the context of model uncertainty, the possibility to respond to complexity by simplifying the world through full neglect of model uncertainty leads to an illusion of certainty and hence increased confidence in action optimality. In the final part of the paper we present a simple model that formalizes

¹Yang (2023) provides an explanation for the attenuated relation between beliefs and actions based on a specific form of cognitive uncertainty.

this intuition and can generate the key results from our experiments. The model is an augmented and simplified version of Bordalo, Gennaioli, Lanzani, and Shleifer (2025).

Our work directly relates to a growing literature on (misspecified) mental models (Schwartzstein and Sunderam (2021), Montiel Olea, Ortoleva, Pai, and Prat (2022), Gagnon-Bartsch, Rabin, and Schwartzstein (2023), Mailath and Samuelson (2019), Heidhues, Kőszegi, and Strack (2018) Heidhues, Kőszegi, and Strack (2023) Frick, Iijima, and Ishii (2022), Aina (2024), Barron and Fries (2024)).² Until now, this literature has largely focused on an analysis of the determinants of model selection, how mental models can be used to persuade others, and the consequences of relying on misspecified models. In contrast, our focus is on how people deal with model uncertainty. A common assumption in the literature is that people work with a single (possibly misspecified) model when making decisions, rather than entertaining multiple weighted models simultaneously (e.g. Schwartzstein (2014), Schwartzstein and Sunderam (2021), Montiel Olea, Ortoleva, Pai, and Prat (2022)).³ We empirically document this kind of simplification, show that it increases with decision complexity and study its implications.

Our work also ties to a literature that studies how limited attention shapes how people react to information.⁴ Bordalo, Conlon, Gennaioli, Kwon, and Shleifer (2023b) and Bordalo, Conlon, Gennaioli, Kwon, and Shleifer (2023a) provide formal frameworks as well as experimental evidence of how attention (and memory) patterns shape belief formation and information processing. Ba, Bohren, and Imas (2022) show that attention processes can lead to overreaction to information. Esponda, Oprea, and Yuksel (2023) show that a form of representativeness heuristic has important implications in contexts of statistical discrimination. Enke and Zimmermann (2019), Enke (2020), and Graeber (2023) provide evidence that people systematically fail to attend to key aspects of

²Relatedly a growing theoretical and empirical literature studies the role of stories and narratives in economics (e.g., Shiller (2017), Eliaz and Spiegler (2022), Andre, Haaland, Roth, and Wohlfahrt (2024), Graeber, Roth, and Zimmermann (2024), Graeber, Roth, and Schesch (2024)).

³Aina and Schneider (2025) study how people update their beliefs in the presence of competing models that could explain observed data. They provide evidence that the majority of people select the model that explains the observed data best. While our focus is on the role of complexity and the implications of model uncertainty neglect, consistent with our finding of full neglect of model uncertainty, they also find that many participants in their experiments build their belief formation process on only one model of the world.

⁴More broadly, Bordalo, Gennaioli, and Shleifer (2012), Bushong, Rabin, and Schwartzstein (2021), Kőszegi and Szeidl (2013) formalize how contextual features steer attention and focus and hence influence behavior. (See Dertwinkel-Kalt, Gerhardt, Riener, Schwerter, and Strang (2022) as well as Somerville (2022) for experimental tests.) Gabaix (2014), Kőszegi and Matějka (2020), Caplin, Dean, and Leahy (2019) formalize attention as a “top-down” process where decision-makers decide to limit attention to reduce the complexity of a problem.

the information environment when processing new information.⁵ We document a related phenomenon in the context of model uncertainty where people selectively attend to only one model of the world.

Our work also relates to a literature that studies the effect of complexity on attention. Recent research on complexity has made substantial progress in defining and quantifying complexity and studying implications in different decision contexts (e.g., Oprea (2020), Kendall and Oprea (2024), Enke and Shubatt (2024), Shubatt and Yang (2024)), Arieta and Nielsen (2024). A common finding is that limited attention on a subset of relevant decision parameters is a simplification response to complexity (Enke (2024), Ba, Bohren, and Imas (2022), Enke (2020), Enke and Zimmermann (2019), Graeber (2023)). We show that in the presence of model uncertainty, complexity induces people to fully neglect model uncertainty. Furthermore, we document that due to the simplification response of uncertainty neglect, confidence in the optimality of own actions is higher in the high-complexity conditions compared to the low complexity conditions.

The rest of the paper is structured as follows. In Section 2, we describe the baseline experimental design. Section 3 presents the results on how model complexity leads to the simplification of model uncertainty in actions. Section 4 then delves into implications of this simplification, studying beliefs and cognitive uncertainty. In Section 5 we present a short model that can generate the observed pattern of results, before concluding in Section 6.

2 Baseline Experimental Design

We designed our experiment with the following goals in mind: (i) implement an economically meaningful and somewhat natural decision environment that allows us to study how people deal with model uncertainty in the face of model complexity; (ii) achieve a well-defined notion of model uncertainty that allows us to exogenously vary model complexity in a straightforward way and (iii) being able to infer the weighting of competing models through the measurement of actions.

⁵Hartzmark, Hirshman, and Imas (2021) show that ownership leads to overreaction to information, an effect that is driven by channeled attention on information. Augenblick, Lazarus, and Thaler (forthcoming) and Fan, Liang, and Cameorn Peng (2024) instead focus on the role of cognitive noise and similarity patterns, respectively, in explaining over- and underreaction to information.

The Task and Model Uncertainty. We chose a financial decision-making task. In the experiment, respondents had to estimate the value of 8 fictitious companies. The correct company value was determined by one of two models, "The CEO is key" or "Products are crucial". The two models used different variables as inputs. One of the models was correct, meaning that it provided the correct company values for all 8 companies, while the other model produced uninformative values.

Each of the two models consists of a formula to calculate the proposed values. "The CEO is key" had the variables CEO competence C and Supporting Staff S as inputs, which could be used to calculate the company value as $C \times S - C - S + 10$. "Products are crucial" had the inputs number of products P and research cost R , and the company value was given by $P \times (10 - R) + R - P$.

To implement model uncertainty, the correct rule was determined in secret by the computer with a simulated coin flip. Therefore, without any additional information, the probability that either rule produced the correct company values was 50%. Respondents then received a noisy but informative signal about the correct model. This signal corresponded to the truth with a probability of 65%.⁶ This was visualized using a ball drawn from an urn containing 65 balls with the correct model and 35 balls with the incorrect model. Afterwards, to measure actions, respondents were asked to provide value estimates for the 8 companies, each featuring different sets of variable realizations.⁷

Company value estimates were incentivized through a binarized scoring rule.⁸ In this way, the chance of respondents to win a bonus was maximized by stating their best guess. Danz, Vesterlund, and Wilson (2022) document empirically that the binarized scoring rule can lead to systematic bias. Notice that such bias (if present in our setting) would not compromise our identification which relies on the comparison of value estimates between conditions of high and low complexity. Furthermore, our results are robust to using investment behavior rather than value estimates as an outcome, which features a different incentive scheme (see Section 3.3).

⁶We did not provide any feedback between rounds. Hence, the only information respondents receive about which model is correct is the signal described above.

⁷The variable realizations were integers between 0 and 10 and all implemented configurations were selected to yield company values between 10 and 100.

⁸Every tenth participant was eligible to receive a bonus payment, in which case a random decision from the survey was incentivized. If a company value guess was incentivized, respondents received \$10 with a probability (in percent) of $100 - 100 \times (\text{Truth}/100 - \text{Guess}/100)^2$, where Truth is the true company value and Guess the company value guess.

Complexity Manipulation. We varied the implementation complexity of the two models between-subjects. In treatment condition *HighComplexity*, working with the models was complex. Specifically, respondents had to calculate the company values under both models themselves. To obtain the Bayesian company value guess, they then needed to weight both values by the respective probabilities of 65% and 35%. In treatment condition *LowComplexity*, they were provided with the calculated company values for both models. Hence, respondents only needed to combine and weight the two values to make their guess, which significantly reduced the complexity of working with the models.⁹

Cognitive Uncertainty and Beliefs about Model Uncertainty. We measured cognitive uncertainty, i.e., people’s confidence in the optimality of their value estimates similar to Enke and Graeber (2023). Specifically, after the series of 8 value estimation tasks, we asked people: *‘How certain are you that, on average, your guesses were no more than 10 points away from the best possible guess given the information you received?’* Respondents indicated their answer on a scale from 0 to 100 percent. The cognitive uncertainty measure was only implemented in a subset of experiments (see Table A.1 and Section 4.2) and was not incentivized. We conducted an additional experiment that features an incentivized confidence measure where participants place a bet on the optimality of their actions (see Section 4.2)).

In the last part of the experiment, respondents had to guess the probabilities that either of the two models generated the correct company values. Since respondents received a noisy signal during the company valuation task, this corresponds to the standard bookbag and chips belief updating task with a signal precision of 65% (cf. Benjamin (2019)). The stated belief was incentivized using a binarized scoring rule.¹⁰ Participants were also asked an un-incentivized direct recall question where they were asked to state which rule was indicated to be more likely by the signal.

Design Details and Procedures. Respondents in both treatment conditions initially received the same set of instructions, explaining both treatments. Afterwards, they went through compulsory comprehension checks and two test runs, where they could practice

⁹Indeed, average guessing times in the Baseline experiment were significantly shorter at 17.09 seconds in *LowComplexity*, compared to 49.11 seconds in *HighComplexity* ($p < 0.001$).

¹⁰If a probability guess was incentivized, respondents received \$10 with a probability of $100 - 100 \times (\text{Truth}/100 - \text{Guess}/100)^2$ %, where Truth is either 100 or 0, and Guess is the stated probability between 0 and 100.

applying each of the model formulas. After they completed the test runs, they were randomly assigned their treatment, received the noisy signal, and subsequently completed the 8 rounds of the task.

On the decision screens, respondents had to hover their mouse over the name of the respective model to uncover the variables and formula or calculated company value (cf. Ba, Bohren, and Imas (2022)). This allows us to study the attention paid to both of the models and signal realizations.

The experiments in this paper were pre-registered on the platform AsPredicted. The pre-registrations include the experimental design, hypotheses, analyses, sample sizes, and exclusion criteria. Table A.1 provides links to each pre-registration.

We conducted the experiments online using the survey provider Prolific. Respondents were recruited from the United States and restricted to be fluent in English. To qualify for the survey, participants had to pass comprehension checks after reading the instructions. The experimental instructions can be found in Appendix D. The median completion time in the Baseline Experiment was 22 minutes. Respondents received a fixed payment of \$4 for the initial study. Every tenth respondent had the chance of winning an additional bonus of up to \$10.

As preregistered, we focus our analysis on two different samples. After reading the instructions, but before treatment assignment, we presented respondents with two test runs for how to calculate the estimates of the models under high complexity (as described above). There, they could familiarize themselves with how to calculate the estimates under both decision models. In order to ensure that we have a respondent pool that is in principle able to solve the formulas in the high-complexity condition, we pre-registered to restrict our sample to respondents who correctly answered both of the two example decisions. The restricted sample of the Baseline Experiment includes 230 respondents. All figures and tables in the main text and Appendix A are based on the restricted sample. Appendix B reproduces all exhibits using the also pre-registered more lenient sample that only requires one of the two questions to be answered correctly, featuring 319 respondents for the Baseline Experiment.

3 Complexity and Simplification of Model Uncertainty

3.1 Framework and Hypothesis

Take a decision-maker whose task is to state a value estimate g_{actual} . Recall that in each decision scenario, the correct company value corresponds to one of the two values given by the models "The CEO is key" and "Products are crucial". The noisy signal received by the participant then indicates the model that is more likely to deliver the correct value. We call the signal-consistent value $v_{consistent}$, and the signal-inconsistent value proposed by the other model $v_{inconsistent}$. Since the signal reveals the correct rule with a probability of 65%, the rational guess for the company value is given by

$$g_{rational} = 0.65 \times v_{consistent} + 0.35 \times v_{inconsistent}.$$

Now take a decision-maker who seeks to simplify the world by neglecting model uncertainty. Such a decision-maker will base their value estimates exclusively on the more likely model. The naive benchmark that fully neglects model uncertainty and takes the signal at face value is hence given by

$$g_{naive} = v_{consistent}.$$

As pre-registered, the main statistical measure we employ is the naive weight λ implicitly defined by

$$g_{actual} = \lambda \times g_{naive} + (1 - \lambda) \times g_{rational}, \quad (1)$$

which can be rearranged to obtain

$$\lambda = \frac{g_{actual} - g_{rational}}{g_{naive} - g_{rational}}. \quad (2)$$

The naive weight λ equals 1 if a participant states the naive guess (full neglect of model uncertainty) and 0 if they state the rational guess.

We expect that decision-makers will be more prone to simplify the world if complexity is high. Therefore we state the following, pre-registered, hypothesis:

Hypothesis. *Decision-makers are more likely to neglect model uncertainty when complexity*

is high, compared to when complexity is low.

3.2 Results

Figure 1a plots histograms of decision-level naive weights in our main sample for both treatments. The distribution in *HighComplexity* features more mass at 1 than the one in *LowComplexity*, indicating a larger amount of fully naive guesses and overreaction to information in the former condition. Conversely, the distribution for *LowComplexity* has more mass around 0 than for *HighComplexity*, implying more guesses in close proximity to the rational benchmark. The prevalence of fully naive guesses is 80% in *HighComplexity*, compared to 59% in *LowComplexity* ($p < 0.01$). Hence, the vast majority of guesses is fully naive in *HighComplexity*.¹¹ Figure 1b confirms this pattern by comparing the average naive weight across treatments, finding a significantly larger average naive weight under complexity ($p < 0.01$).

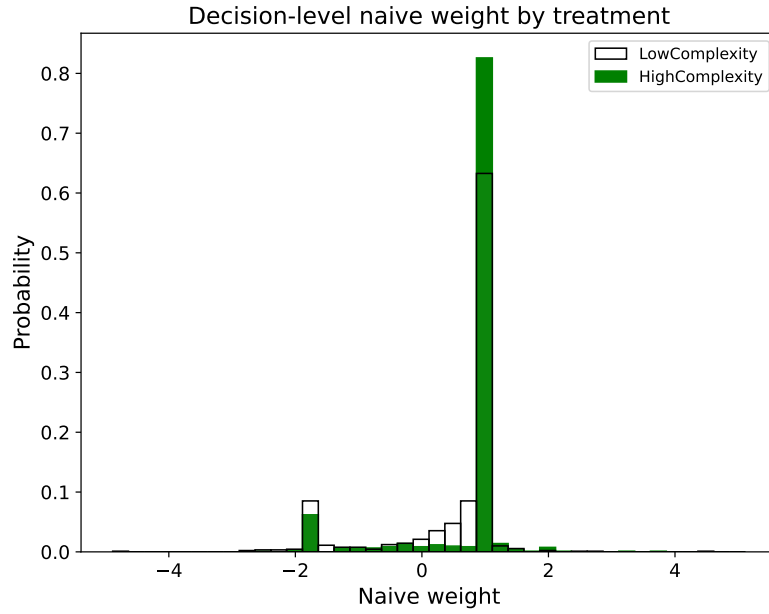
Table 1 complements this analysis by regressing the company value guesses on the rational and naive benchmarks. This can be interpreted as estimating Equation 1 without the restriction that the weights on the rational and naive benchmarks add up to 1. We can see that the fully naive benchmark is relatively more predictive of participants' guesses in treatment *HighComplexity* compared to *LowComplexity*.

Result 1. *There is substantially more neglect of model uncertainty when complexity is high, compared to when it is low.*

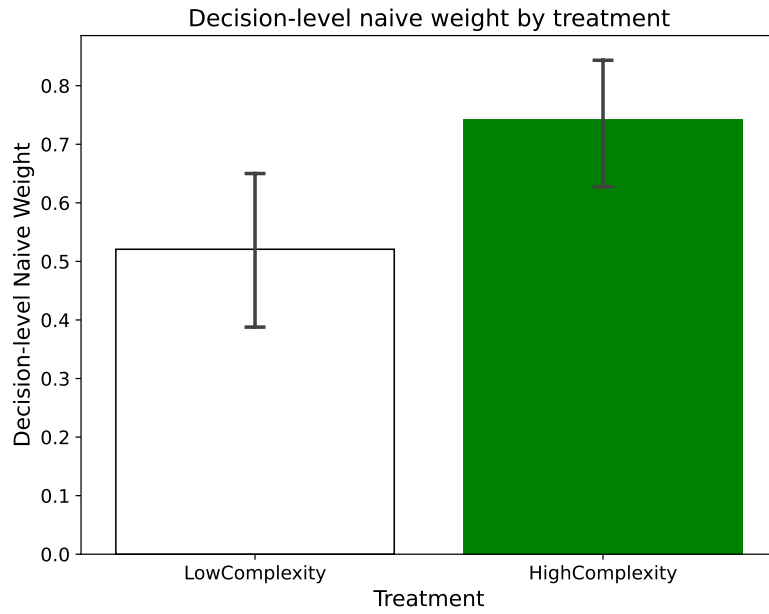
We note that the naive weight defined in Equation (2) is quite sensitive to outliers. To ensure that our results are not driven by outliers, we look at the 8 decision scenarios for each subject and compute the median naive weight. In an analysis that was not preregistered, we plot the histogram of median naive weights for each subject in Figure A.1. We again observe a higher mass around naive weights of 1 in *HighComplexity*, and more mass between 0 and 1 in *LowComplexity*, confirming the earlier results that full neglect of model uncertainty is more prevalent under complexity.

Appendix B.1 replicates all exhibits using the also pre-registered more lenient sample that only requires one of the two questions to be answered correctly. In Appendix A.7

¹¹A large fraction of guesses in *LowComplexity* also reveal full naivete, which may be caused by the residual complexity of the general experimental set-up or more specifically of the need to combine the signal-consistent and signal-inconsistent values into a value estimate.



(a) Distribution of decision-level naive weights



(b) Mean decision-level naive weights

Figure 1: Decision-level naive weights in LowComplexity and HighComplexity conditions. Panel (a) plots the distribution of naive weights λ calculated as specified in Equation 2, using the restricted sample of the Baseline Experiment with 230 participants. Panel (b) plots average naive weights.

Table 1: Company Value Guesses

<i>Dependent variable:</i>	Company Value Guess		
<i>Sample:</i>	LowComplexity (1)	HighComplexity (2)	Pooled (3)
Rational Benchmark	0.593*** (0.079)	0.343*** (0.065)	0.593*** (0.079)
Naive Benchmark	0.481*** (0.070)	0.703*** (0.057)	0.481*** (0.070)
Rational B. $\times$ HighComplexity			-0.250** (0.102)
Naive B. $\times$ HighComplexity			0.222** (0.090)
R^2	0.881	0.918	0.900
Observations	904	936	1840

The table presents OLS regressions of respondents' company value guesses on the rational and naive benchmarks as detailed in Section 3.1. The table is based on the restricted sample of the Baseline Experiment. Column (1) uses observations from the LowComplexity treatment, column (2) from the HighComplexity treatment, and column (3) from both. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

we show that we replicate all the findings for the Baseline Confidence Experiment (see Section 4.2 and Table A.1).

Role of Attention. The above results indicate that people simplify complexity by neglecting model uncertainty. To further corroborate this finding, we study the role of attention. Recall that respondents had to hover their mouse over either "The CEO is key" or "Products are crucial" to observe the parameters needed to make company value guesses. In *LowComplexity*, the company value was displayed only when hovering over the respective rule. In *HighComplexity*, the respective variable realizations needed to calculate the company value were displayed. The resulting data on hover times allows us to study how much attention participants paid to either rule.

Figure 2 shows the distribution of the median consistent hovering shares per subject separately for both treatments. In treatment *HighComplexity*, an overwhelming majority of participants only considers the variables of the signal-consistent model. In treatment *LowComplexity*, most participants consider the proposed values by both models, with the mode being at equal hovering times for both the signal-consistent and inconsistent model.

Table 2 analyzes more formally how attention differs between the two treatments. The first column shows that the average hover time for the signal-consistent rule more than triples with higher complexity. The second column shows that the hover time for the signal-inconsistent rule also increases, but by much less. The third column confirms that the share of consistent hover time increases in *HighComplexity*, meaning that participants in this condition pay relatively more attention to the signal-consistent rule.

3.3 Robustness

Investment Behavior. In the baseline experiment, our main outcome measure is respondents' value estimates for the hypothetical companies. We ran two additional pre-registered versions of the Baseline Experiment that replace the company value estimate with an investment decision. The design of the experiments was exactly as described in Section 2, with only one change. Instead of providing a guess for the value of the company, participants were asked to submit an investment bid for each company. For each decision, they received a budget of 100 cents and could subsequently decide how much

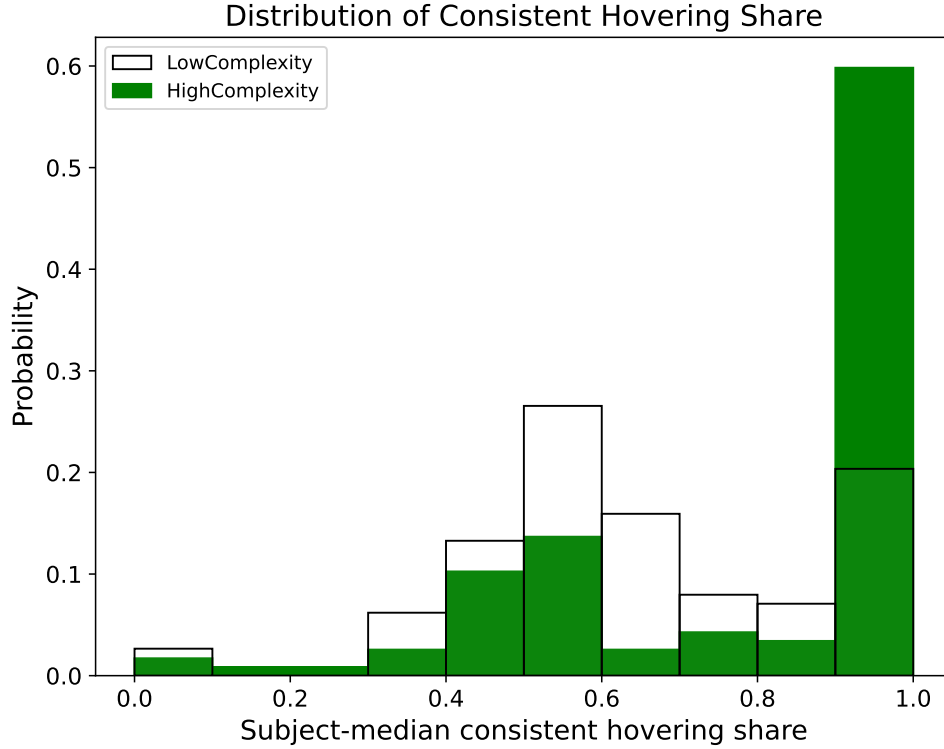


Figure 2: Distribution of subject-medians of the consistent hovering shares. *The figure plots the distribution of the share of time that respondents spent looking at the values of the signal-consistent rule, using the restricted sample of the Baseline Experiment with 230 participants. Only the median consistent share for each participant is plotted.*

Table 2: Hover Times

<i>Dependent variable:</i>	Consistent Hover Time	Inconsistent Hover Time	Consistent Share
<i>Sample:</i>	Pooled (1)	Pooled (2)	Pooled (3)
Constant	2.641*** (0.232)	1.533*** (0.153)	0.635*** (0.018)
HighComplexity	11.611*** (0.954)	1.588*** (0.392)	0.153*** (0.027)
R^2	0.262	0.033	0.079
Observations	1840	1840	1840

The table presents OLS regressions using the restricted sample of the Baseline Experiment. Hover times were winsorized at the top at the 97.5% quantile. All columns use observations from both the HighComplexity and LowComplexity conditions. In column (1), the time that respondents spent looking at the values for the signal-consistent rule is regressed on a constant and a treatment dummy for the HighComplexity condition. In column (2), the dependent variable is the time spent looking at the signal-inconsistent rule. In column (3), it is the share of time that was spent looking at the signal-consistent rule. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

to bid for the company. The bids were incentivized through a random auction mechanism: a random price between 10 and 100 cents was drawn for the company. If the bid was greater than or equal to the price, the company was bought by the participant, paying the drawn price, and receiving the true company value. If the bid was lower than the price, there was no transaction.

Both experiments were pre-registered (see Table A.1). To ensure comparability with the other experiments, for both experiments we focus on the results from the restricted sample. For experiment 1 this yields 193 participants, for experiment 2 this sample features 323 participants.¹²

Taken together, the results go in the same direction as for the Baseline Experiment. The overall prevalence of fully naive bids in experiment 1 is 73% in *HighComplexity* and 47% in *LowComplexity* ($p < 0.001$). Figure A.5b shows that the mean naive weight is significantly higher under complexity. Hover times show the same patterns as in the Baseline Experiment, with there being significantly more consistent hovering in *HighComplexity*. In experiment 2, the share of fully naive guesses is 70% in *HighComplexity* and 41% in *LowComplexity*, and the mean naive weight is also significantly higher under complexity as can be seen in Figure A.9b, confirming the results from the first experiment. Again, hover times also show the same pattern as before (see Appendix A.5).¹³

Result 2. *Our results are robust to using a different outcome, namely investment behavior. People’s investment bids reflect substantially more full neglect of model uncertainty when complexity is high, compared to when it is low.*

Equally Likely Models. In the baseline experiment, we endow respondents with a natural candidate for simplification, namely the objectively more likely model. Hence, the question arises whether simplification also occurs if both models are equally likely. To address this, we conducted an additional preregistered study, where participants completed a version of the Baseline Experiment that excluded the informative signal. Instead, they were only endowed with a 50-50 prior when estimating company values. The link

¹²Notice that for experiment 1, we deviate from the specification in the pre-registration where we pre-registered the use of the lenient sample and the full sample. Appendix B.3 and Appendix B.4 produce the corresponding results. Also notice that for experiment 1 we pre-registered a smaller sample size than for the other experiments.

¹³Appendices B.3 and B.5 produce results for the more lenient sample. Overall, results are similar to the restricted sample, although the treatment difference in mean naive weights in experiment 1, while directionally present, fails to be significant in the pre-registered lenient sample ($p = 0.207$).

to the preregistration can be found in Table A.1.

The experiment followed the design described in Section 2 but omitted both the signal and the belief elicitation regarding the rule determining company values. To maintain the experiment's length and incentive structure, we replaced the belief elicitation with an unrelated belief-updating task.

In the absence of a noisy signal, there is no clear reason to expect participants to simplify model uncertainty in a specific direction when faced with complexity. Therefore, our hypothesis for this study was that participants in the *HighComplexity* condition would more frequently state guesses equal to the values proposed by either rule, whereas more intermediate guesses would be observed in the *LowComplexity* condition.

To test this, we examine the implicit decision weight, γ , assigned to the company value proposed by the "The CEO is key" rule, defined as

$$g_{actual} = \gamma \times v_{CEO} + (1 - \gamma) \times v_{Products}. \quad (3)$$

Additionally, we define weight extremity as

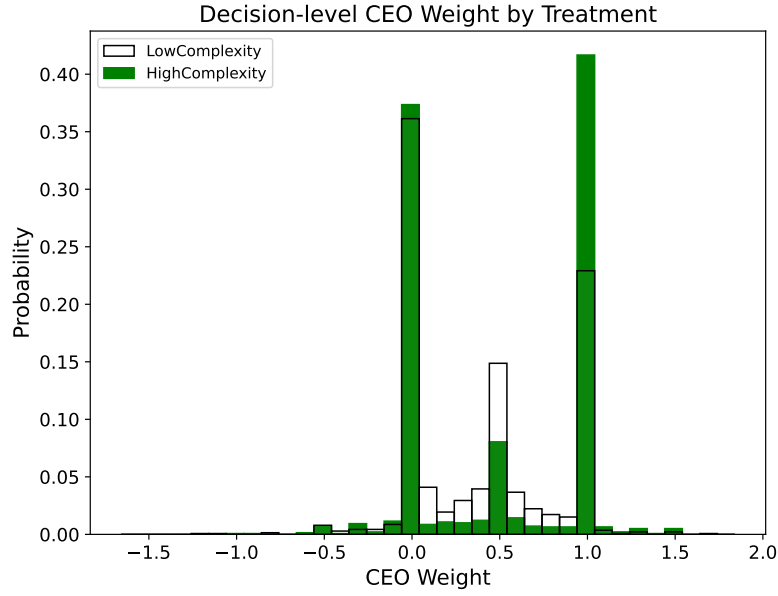
$$\text{Weight extremity} = \left| \gamma - \frac{1}{2} \right|. \quad (4)$$

Our pre-registered hypothesis then is that weight extremity is higher in the *HighComplexity* condition compared to the *LowComplexity* condition.

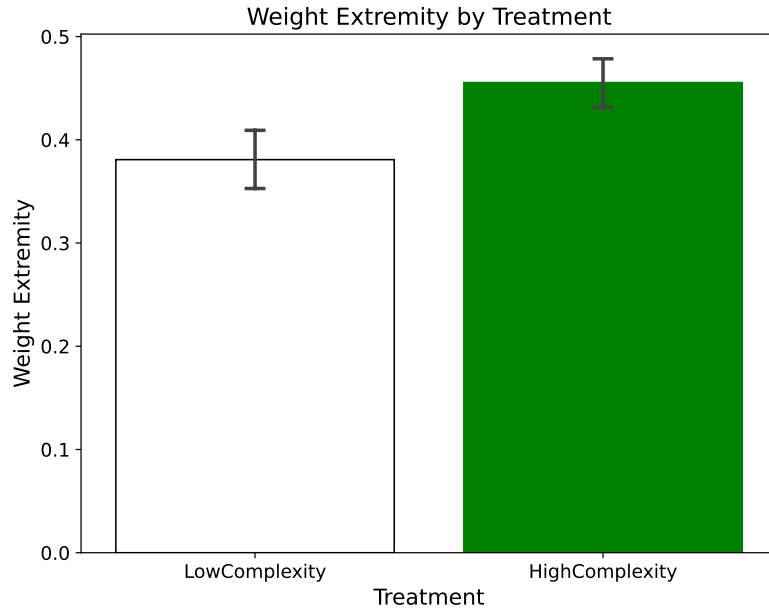
The pre-registered restricted sample consists of 348 participants who correctly solved both *HighComplexity* example decision screens.

Figure 3a displays the distribution of decision-level CEO weights for both treatments. In the *LowComplexity* condition, the distribution has greater mass at intermediate values, particularly near the rational CEO weight of 0.5. In contrast, the *HighComplexity* distribution shows more mass at weights corresponding to simplified guesses, especially at a CEO weight of 1. Figure 3b further confirms that the average weight extremity in *HighComplexity* is significantly higher than in *LowComplexity*.

Figure A.2 displays the distribution of the share of time participants spent hovering over the "The CEO is key" model in both treatments. Participants in *HighComplexity* are more likely to focus on a single model, whereas those in *LowComplexity* typically attend to both models.



(a) Distribution of decision-level CEO weights



(b) Mean decision-level weight extremity

Figure 3: Decision-level CEO weights and weight extremity in LowComplexity and High-Complexity conditions of the Equally Likely Models Experiment. *Panel (a) plots the distribution of CEO weights γ calculated as specified in Equation 3, using the restricted sample of the Equally Likely Models study with 348 participants. Panel (b) plots the average weight extremity $|\gamma - \frac{1}{2}|$ calculated as specified in Equation 4.*

In Appendix B.2, we replicate the analysis using the more lenient sample of 452 participants who answered at least one of the two HighComplexity example screens correctly. All results continue to hold in this sample.¹⁴

Result 3. *People also neglect model uncertainty as a response to complexity when both models are equally likely.*

4 Results: Implications of Simplification

4.1 Beliefs about Models

We have shown that model complexity leads to a simplification response: people tend to fully neglect model uncertainty in their actions. As pre-registered, we now ask whether this simplification also causes a distorted view of reality, namely that people misperceive model uncertainty when directly asked.¹⁵

For this purpose, beliefs about which model is correct were elicited immediately after all company value guesses had been submitted.

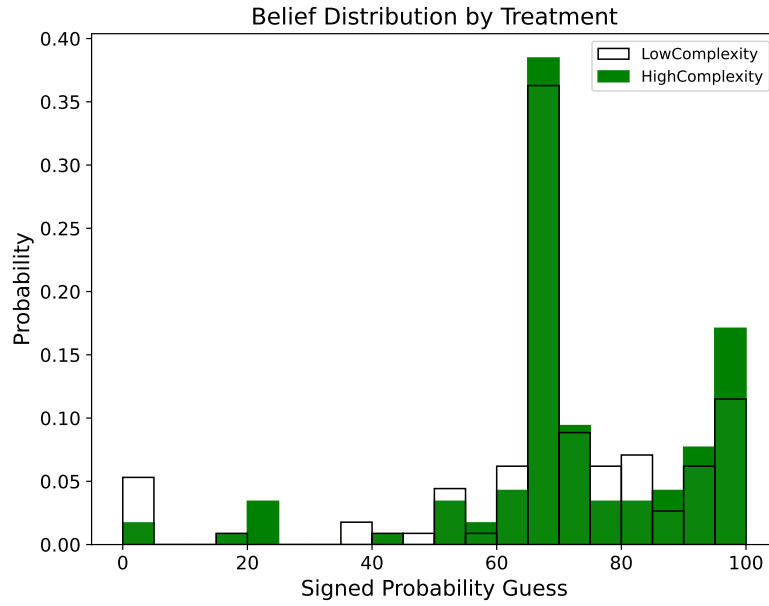
Figure 4a shows histograms of signed probability guesses (signed in the direction of the received signal), separately for treatments *LowComplexity* and *HighComplexity* in the Baseline experiment. There are no visible differences in beliefs. This is confirmed by Figure 4b which plots average guesses for the probability that "The CEO is key" is the correct model, stratified by the signal participants received and their assigned treatment. There is no significant treatment difference in average beliefs for either signal.

Table 3 presents regressions for signed probability guesses and correct recall of the signal. The regressions confirm that there are no significant treatment differences in beliefs. Similarly, there are no differences in the accuracy of recall of the received signal, as column (2) reveals.

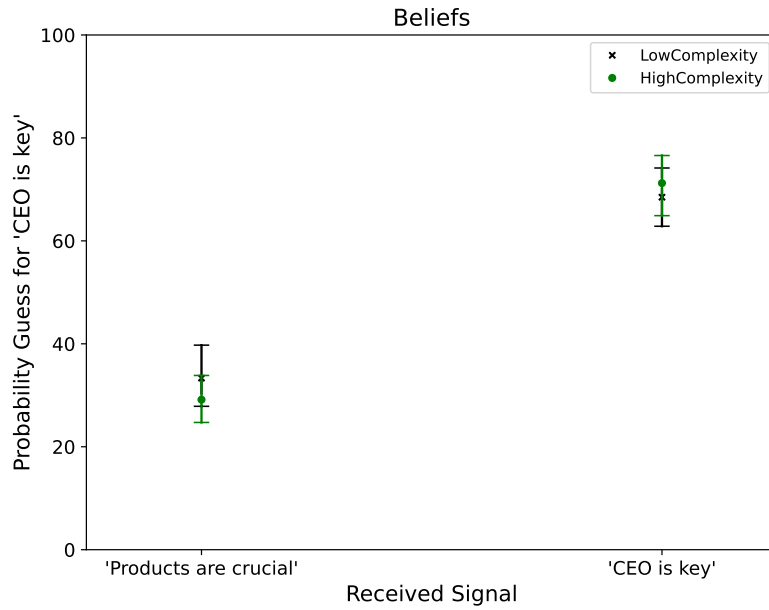
We also investigate belief patterns in the Investment 1, Investment 2 and Baseline

¹⁴A perhaps interesting follow-up question is whether people that simplify model uncertainty as a response to complexity do so in a consistent fashion, always focusing on the same model. To shed some light on this, in an analysis that was not preregistered, Figure A.3 restricts the sample to participants in the *HighComplexity* condition who made only fully naive guesses—that is, who consistently guessed a value corresponding to either the "The CEO is key" or "Products are crucial" rule for every decision. This applies to 91 out of all 174 participants in *HighComplexity*. The figure suggests a large fraction of respondents (about half) consistently simplify in the same direction, always choosing values from the same model, while a sizeable fraction alternates between models.

¹⁵Notice that, while we pre-registered this analysis, we did not pre-register a specific hypothesis.



(a) Histogram of signed probability guesses.



(b) Mean probability guesses for "The CEO is key".

Figure 4: Beliefs in the Baseline Experiment. *In Panel (a) beliefs are converted into the direction of the more likely model, so that 65 corresponds to the Bayesian probability. Panel (b) plots mean probability guesses for the model "The CEO is key". The figure is based on the restricted sample of the Baseline Experiment with 230 participants.*

Table 3: Beliefs and Recall in the Baseline Experiment

<i>Dependent variable:</i>	Probability Guess	Correct Recall
<i>Sample:</i>	Pooled (1)	Pooled (2)
Constant	67.679*** (2.078)	0.965*** (0.018)
HighComplexity	3.351 (2.819)	-0.007 (0.026)
R^2	0.006	0.000
Observations	230	230

The table presents OLS regressions using the restricted sample of the Baseline Experiment. In column (1), the dependent variable is the probability guess for the likely state. Column (2) uses a dummy for whether respondents correctly recall the more likely model. All columns use observations from both the HighComplexity and Low-Complexity conditions. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

Confidence experiments. In the restricted samples of all of these studies, we find the same result of no significant treatment differences in beliefs and recall, as can be seen in the respective sections of Appendix A.¹⁶ Taken together, these results suggest that the complexity-induced neglect of model uncertainty in actions does not translate into corresponding belief patterns.

Result 4. *The neglect of model uncertainty in actions does not translate to beliefs about model uncertainty. These beliefs appear to reflect model uncertainty and are largely unaffected by complexity.*

Robustness: Delayed Belief Elicitation. In an additional pilot study that was not pre-registered (see Table A.1), we investigated whether adding a delay between actions (value guesses) and belief elicitation about model uncertainty could potentially induce treatment differences in beliefs. The rationale was that it might take time for simplification-induced misperceptions to form. The design was exactly as described in Section 2, with an additional second part that took place one day after the initial survey. Respondents in

¹⁶In Appendix B we present the results for the lenient samples. Here, we find a (marginally) stronger belief response under complexity in the Baseline and Baseline Confidence studies, and no significant treatment differences in Investment Experiment 1 and Investment Experiment 2.

the *Immediate* condition stated their beliefs immediately after the company value guesses as in the Baseline Survey, and completed unrelated tasks during the second survey. Respondents in the *Delay* condition completed the company value guesses and unrelated tasks during the first part of the survey, and the belief elicitation during the second survey one day later.

We first note that also in this experiment, we see that complexity induces people to neglect model uncertainty in actions. However, as Appendix A.6 reveals, both with and without delay, simplification does not induce biased beliefs about model uncertainty.

To summarize, adding a one day delay does not induce treatment differences both in stated beliefs and correct recall of the received signal.

Results on Belief-Action Link. Notice that our results imply a complexity-induced wedge between actions and beliefs. A recent literature has investigated the link between beliefs and actions (Giglio, Maggiori, Stroebel, and Utkus (2021), Ameriks, Kézdi, Lee, and Shapiro (2020), Beutel and M. Weber (2023), Laudenbach, A. Weber, R. Weber, and Wohlfart (forthcoming)), and the determinants of how strongly beliefs translate into actions. Charles, Frydman, and Kilic (2024) find that increased complexity of forming a belief weakens the transmission of beliefs into actions. Similarly, Yang (2023) as well as Enke, Graeber, Oprea, and Yang (2024) highlight a link between information processing constraints and a weak elasticity of decisions with respect to economic fundamentals.

To investigate this wedge in our experiment more formally, Figure 5 presents a binned scatterplot with beliefs on the horizontal axis, and corresponding actions on the vertical axis, using data from the Baseline Experiment. Beliefs are given by the probability guess for the "The CEO is key" model. To ensure comparability, actions are represented by the implicit decision weight γ on the company value proposed by the "The CEO is key" model, as defined in Equation 3.

The plot illustrates the overreaction in actions caused by the complexity of calculating the optimal guess: when moving away from a probability guess of 50%, the decision weights in *HighComplexity* quickly move to extreme values, while the response is more muted in *LowComplexity*. This implies that the effect of complexity on the belief-action link is not straightforward in our setting. In a neighborhood around 50%, complexity increases responsiveness of actions to beliefs. However, when moving to more extreme beliefs, complexity renders the action response rather flat, since naive guessing is already

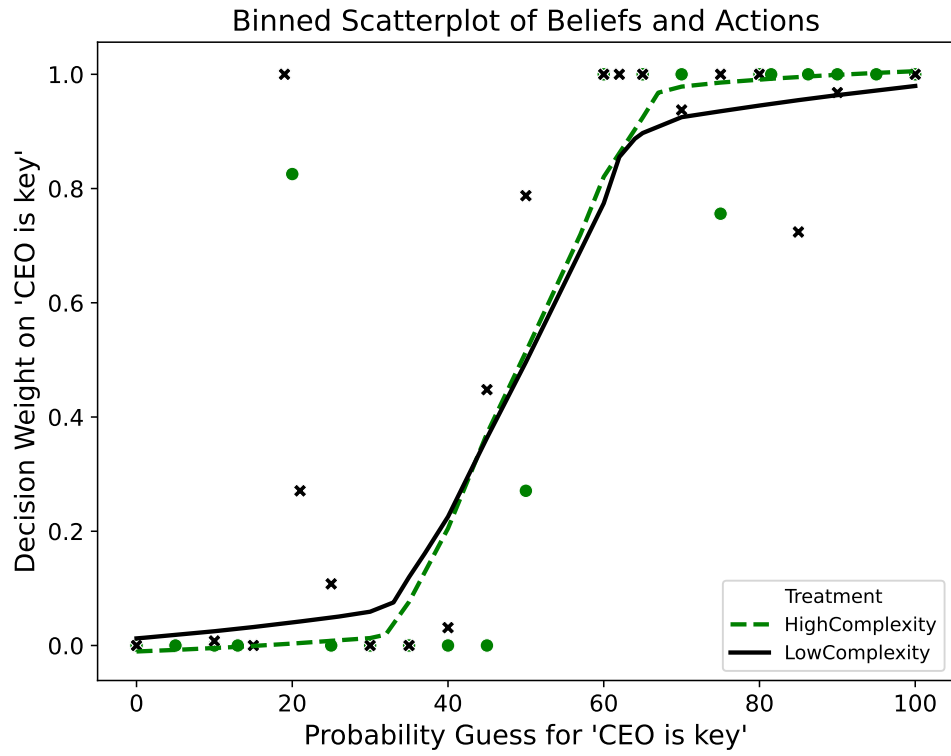


Figure 5: The relationship between decisions and beliefs. *The figure shows a binned scatterplot using the restricted sample of the Baseline Experiment with 230 participants. It has the probability guess that "The CEO is key" is the more likely model on the horizontal axis, and the decision weight γ as defined in Equation 3 on the vertical axis. The lines show LOWESS regressions based on all datapoints.*

triggered starting from fairly moderate beliefs.

We show analogous plots for the Investment 1, Investment 2 and Baseline Confidence studies in Appendix A, as well as results for the lenient samples in Appendix B, each showing qualitatively identical results.

Result 5. *Complexity induces a wedge between beliefs and actions. When complexity is high, beliefs continue to reflect model uncertainty, actions tend to be based on neglect of model uncertainty.*

4.2 Confidence

Here we investigate whether the complexity-induced simplification of neglect of model uncertainty in actions affects how people view the optimality of their actions. This type of confidence has been shown to explain a broad range of behavioral anomalies and also mediates to what extent individual biases matter for aggregate outcomes (Enke and Graeber (2023), Enke, Graeber, and Oprea (forthcoming), Enke, Graeber, and Oprea (2023)). A typical and highly intuitive finding in the literature is that complexity reduces confidence in action optimality.

Baseline Confidence. In a replication of the Baseline Experiment, we elicited cognitive uncertainty, i.e. people’s confidence in the optimality of their value estimates. The experiment was pre-registered, including the analysis of the cognitive uncertainty measure (see Table A.1). As pre-registered, we use the same sample restriction as for the Baseline study, resulting in a restricted sample with 336 participants.¹⁷

Figure 6 shows the average confidence levels by treatment condition. We find that participants in the *HighComplexity* condition are significantly *more* confident in their guesses than those in the *LowComplexity* condition—despite exhibiting greater simplification and a lower rate of rational guesses under complexity, a seemingly contradictory pattern.

We also elicited the same measure in the Equally Likely Models study and featured its analysis in the preregistration. Figure A.4 confirms that the finding replicates in this study. Additional results presented in Appendix B show that the results in both studies also hold true when using the more lenient sample restrictions.

¹⁷Similar to the analysis of beliefs, while we pre-registered this analysis, we did not pre-register a specific hypothesis.

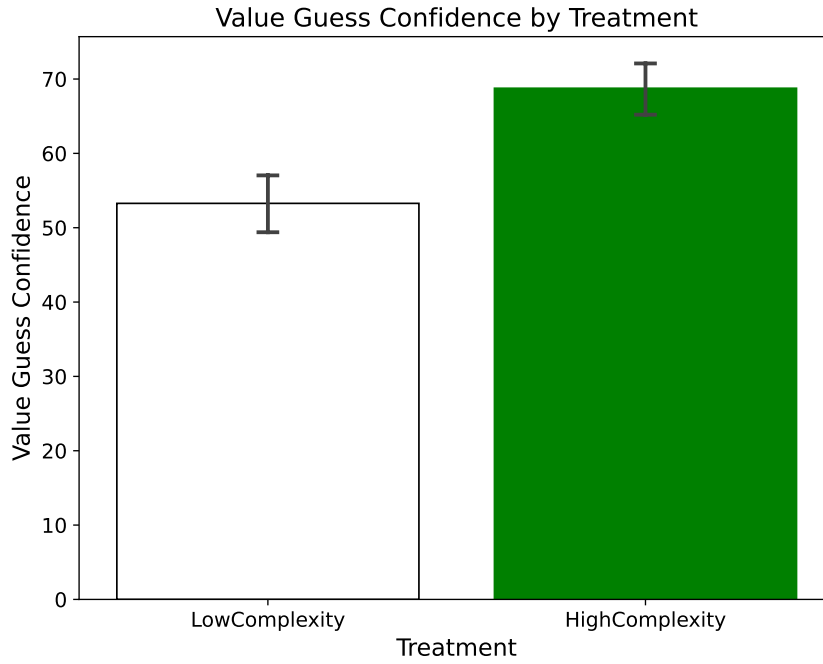


Figure 6: Average value guess confidence in the Baseline Confidence Experiment. *The figure plots the average confidence that respondents had in their company value guesses, using the restricted sample of the Baseline Confidence study with 336 participants.*

Incentivized Confidence. In another replication of the Baseline Confidence experiment, we further added an incentivized version of the confidence elicitation. Here, participants could bet on the optimality of their guesses. They received an endowment of \$10 and could choose how much to bet (between \$0 and \$10) on the event that their guesses had been, on average, no more than 10 points away from the best possible guess. The bet was multiplied by 3 if their guesses had indeed been accurate and was lost otherwise. After completing all company value guesses, participants answered the non-incentivized (probability) and incentivized (bet) versions on the same screen. The experiment was pre-registered, including the analysis of the incentivized and non-incentivized cognitive uncertainty measures (see Table A.1). As pre-registered, we use the same sample restriction as for the Baseline study, resulting in a restricted sample with 203 participants. We report the results on confidence for this experiment in the main text and refer the reader to Appendix A.8 for the remaining analyses.

Figure 7 plots the average confidence measures. The lower panel replicates the previous result that participants in the high-complexity condition stated a higher confidence level in the optimality of their company value guesses than their counterparts in the low-complexity condition. The upper panel presents results from the incentivized bet-

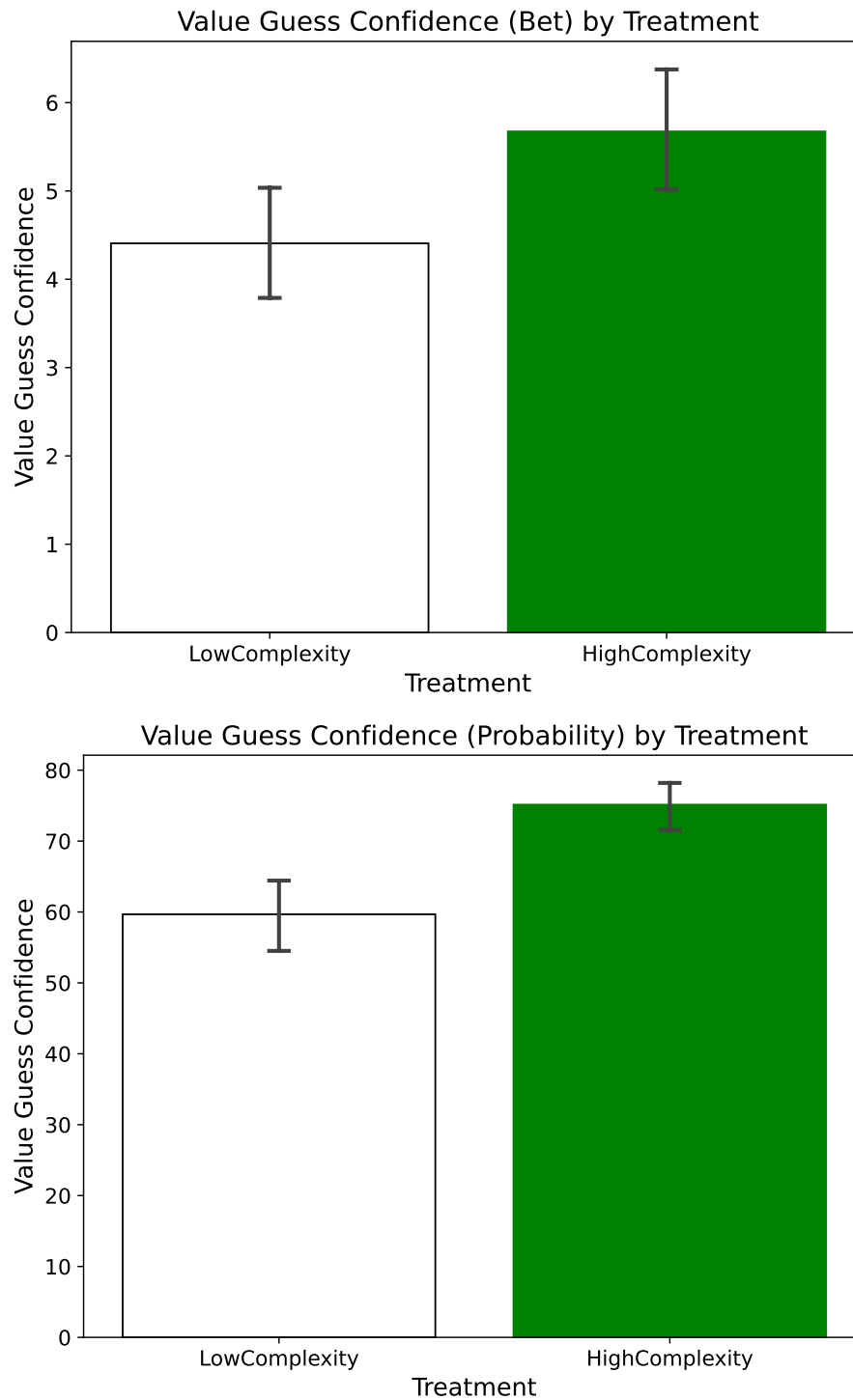


Figure 7: Average value guess confidence in the Incentivized Confidence Experiment. *The top figure plots the average incentivized confidence measure, while the bottom figure plots the non-incentivized measure, both using the restricted sample of the Incentivized Confidence study with 203 participants.*

ting task. While these data are naturally more noisy due to the betting context, the Figure again confirms our result that higher complexity yields higher confidence. Appendix B.7 produces these results under the more lenient sample restrictions, yielding qualitatively identical results.

Taken together, our results indicate that confidence in the optimality of own actions is higher in the high-complexity conditions compared to the low complexity conditions. It seems that in contrast to prior results from different contexts, in the context of model uncertainty, the possibility to respond to complexity by greatly simplifying the world through full neglect of model uncertainty leads to an illusion of certainty and hence increased confidence in action optimality.

Result 6. *Complexity leads to higher confidence in the optimality of own actions.*

5 A Simple Model of Representations

We have shown that respondents simplify model uncertainty when computational complexity is high. This, however, does not carry over to beliefs about model uncertainty. Finally and perhaps most surprisingly, higher computational complexity increases confidence in the optimality of own actions rather than decreasing it. In the final part of the paper we present a simple model that can generate this pattern of results. The model is an augmented and simplified version of Bordalo, Gennaioli, Lanzani, and Shleifer (2025). In the model, when faced with a decision problem, an agent first forms a mental representation of the problem. This process is shaped both by bottom-up and top-down attention. Upon being presented a decision problem, a bottom-up process of cue-dependent memory determines which of the currently stored mental representations is top of mind. Then, in a top-down process, the agent decides whether they want to further simplify this representation. This section contains the basic intuitions, while Appendix C presents the formal model.

This model formalizes economic decision making as a cognitively constrained process operating over a structured internal database of mental representations. Initially, this database comprises only a broad representation, which encodes a probabilistic assessment of the decision environment (e.g., a 0.65 likelihood for the more probable state) along with abstract contextual features (e.g., task framing, informational structure, and

the need to estimate a company's value).

Each decision - whether it concerns an action, confidence judgment, or belief report - invokes a two-step cognitive process to form a mental representation of the decision problem:

1. **Bottom-Up Retrieval:** Agents hold a database of mental representations. When faced with a decision, similarity-based recall (Bordalo, Gennaioli, and Shleifer (2020), Bordalo, Conlon, Gennaioli, Kwon, and Shleifer (2023b), Enke, Schwelter, and Zimmermann (2024), Graeber, Roth, and Zimmermann (2024), Jiang, Liu, Cameron Peng, and Yan (forthcoming)) determines which representation is top of mind. In other words, the representation most similar to the current decision cue is retrieved from memory.
2. **Top-Down Simplification:** Once a representation is top of mind, the agent, in a top-down process, evaluates whether the cognitive costs of computing a response using the retrieved representation exceed the benefits. If so, the agent further simplifies the representation to reduce processing demands.

The key insight from our model is that people make decisions within the mental representation they formed for this specific decision problem. Hence, specific decisions (actions, confidence judgments, or belief reports) in a given underlying environment may be based on very different mental representations of that environment, allowing us to explain our seemingly contradictory pattern of results.

Estimation of Company Values. Upon observing a signal, the agent must compute an optimal guess. Since this is the first decision agents take, the database of representations only contains the broad representation of the problem. Under low complexity, cognitive costs are manageable, so participants compute the optimal guess using the full broad representation of the problem, i.e., fully taking model uncertainty into account. Under high complexity, cognitive costs may exceed the benefits, prompting simplification to a mental representation in which model uncertainty is fully neglected. Once a decision has been taken, the corresponding representation is added to the database, including their contextual features.

Confidence Elicitation. In the low complexity treatment, when faced with the confidence assessment, agents will retrieve the broad representation (their database only consists of broad representations). As confidence judgments are cognitively light, no further simplification occurs. In the high complexity treatment, similarity-based retrieval favors the recall of the simplified representation. This is because the confidence elicitation cues the estimation tasks (it explicitly asks for confidence about the estimation task). As described above, confidence in own action optimality is then assessed by agents within this mental representation. Consequently, confidence can be elevated by complexity if confidence within the simplified high complexity environment is higher than in the full representation of the low complexity environment.¹⁸

Belief Elicitation. In the low complexity treatment, the broad representation is again retrieved and used without simplification, yielding probabilistically grounded belief reports about model uncertainty. In the high complexity treatment, retrieval of mental representations favors the broad representation over the simple one, since the belief task directly asks about model uncertainty and is hence more similar to the broad representation. Since the cognitive costs of belief elicitation are negligible, no further simplification occurs. People then state the beliefs of their mental representation and therefore beliefs in both treatments will tend to reflect model uncertainty.

6 Conclusion

This paper explores how individuals navigate model uncertainty in economic decision-making and demonstrates that complexity significantly influences the way people deal with model uncertainty. Through a controlled experimental framework, we investigate whether individuals account for model uncertainty when making decisions or instead simplify the world by implicitly assuming one model is correct. Our findings reveal that when model complexity is high, individuals tend to neglect model uncertainty in their actions, behaving as if one model is definitively correct. However, this neglect of uncertainty does not translate into distorted beliefs, as participants' stated beliefs continue to

¹⁸While this seems plausible, our model does not formalize why confidence within the simplified high complexity environment may be higher than in the full representation of the low complexity environment. The key insight from our model is that confidence is assessed within the mental representation of the problem.

reflect model uncertainty. This creates a systematic wedge between actions and beliefs. Furthermore, our results show that complexity-induced simplification leads to increased confidence in decision optimality, contradicting prior findings that suggest complexity typically raises cognitive uncertainty.

Our results provide a direct test of the widely held assumption in the theoretical literature on misspecified models that people attend to a single model when making decisions, rather than entertaining multiple weighted models simultaneously.

By systematically manipulating the complexity of models, we provide robust evidence that individuals simplify complex problems by focusing on a single mental model. Data on attention allocation further support this conclusion, showing that participants in high-complexity conditions spend more time attending to one model, rather than considering both models equally. This attention-based mechanism helps explain why complexity amplifies the tendency to neglect model uncertainty in actions. The robustness of these findings is confirmed across different tasks, including an investment decision context, demonstrating that the observed pattern extends beyond the specific valuation task used in our primary experiment.

Our findings contribute to a deeper understanding of how complexity shapes economic cognition. Existing research on mental models has largely focused on selection and persuasion, whereas our study highlights a novel dimension: the impact of complexity on how people handle competing models. Our findings reveal that, when complexity is high, in order to be able to operate and make decisions, people need to simplify the world by acting as if only one model exists. Interestingly, this complexity-induced simplifications can lead to an illusion of certainty, where people are confident about the optimality of their actions.

While we do not study persuasion directly, an intuitive implication of our results may be that, when complexity is high, the presence of model uncertainty might make decision-makers more susceptible to persuasive narratives that present a single, seemingly definitive interpretation of economic realities.

References

- Aina, Chiara (2024). “Tailored Stories”. In: *Working Paper*.
- Aina, Chiara and Florian Schneider (2025). “Weighting Competing Models”. In: *Working Paper*.
- Ameriks, John, Gábor Kézdi, Minjoon Lee, and Matthew D Shapiro (2020). “Heterogeneity in expectations, risk tolerance, and household stock shares: The attenuation puzzle”. In: *Journal of Business & Economic Statistics* 38.3, pp. 633–646.
- Andre, Peter, Ingar Haaland, Christopher Roth, and Johannes Wohlfahrt (2024). “Narratives about the Macroeconomy”. In: *Working Paper*.
- Arrieta, Gonzalo and Kirby Nielsen (2024). “Procedural Decision-Making In The Face Of Complexity”. In: *Working Paper*.
- Augenblick, Ned, Eben Lazarus, and Michael Thaler (forthcoming). “Overinference from Weak Signals and Underinference from Strong Signals”. In: *Quarterly Journal of Economics*.
- Ba, Cuimin, J Aislinn Bohren, and Alex Imas (2022). “Over-and underreaction to information”. In: *Available at SSRN* 4274617.
- Barron, Kai and Tilman Fries (2024). “Narrative Persuasion”. In: *Working Paper*.
- Benjamin, Daniel J (2019). “Errors in probabilistic reasoning and judgment biases”. In: *Handbook of Behavioral Economics: Applications and Foundations* 1 2, pp. 69–186.
- Beutel, Johannes and Michael Weber (2023). “Beliefs and portfolios:Causal evidence”. In: *Working Paper*.
- Bordalo, Pedro, John Conlon, Nicola Gennaioli, Spencer Kwon, and Andrei Shleifer (2023a). “How People Use Statistics”. In: *Working Paper*.
- (2023b). “Memory and Probability”. In: *Quarterly Journal of Economics* 138.1, pp. 265–311.
- Bordalo, Pedro, Nicola Gennaioli, Giacomo Lanzani, and Andrei Shleifer (2025). “A Cognitive Theory of Reasoning and Choice.” In: *Working Paper*.
- Bordalo, Pedro, Nicola Gennaioli, and Andrei Shleifer (2012). “Salience theory of choice under risk”. In: *The Quarterly journal of economics* 127.3, pp. 1243–1285.
- (2020). “Memory, Attention, and Choice”. In: *The Quarterly Journal of Economics* 135.3, pp. 1399–1442. doi: 10.1093/qje/qjaa012.

- Bushong, Benjamin, Matthew Rabin, and Joshua Schwartzstein (2021). “A model of relative thinking”. In: *The Review of Economic Studies* 88.1, pp. 162–191.
- Caplin, Andrew, Mark Dean, and John Leahy (2019). “Rational inattention, optimal consideration sets, and stochastic choice”. In: *The Review of Economic Studies* 86.3, pp. 1061–1094.
- Charles, Constantin, Cary Frydman, and Mete Kilic (2024). “Insensitive investors”. In: *The Journal of Finance* 79.4, pp. 2473–2503.
- Danz, David, Lise Vesterlund, and Alistair Wilson (2022). “Belief Elicitation and Behavioral Incentive Compatibility”. In: *American Economic Review* 112.9, pp. 2851–2883.
- Dertwinkel-Kalt, Markus, Holger Gerhardt, Gerhard Riener, Frederik Schwerter, and Louis Strang (2022). “Concentration bias in intertemporal choice”. In: *The Review of Economic Studies* 89.3, pp. 1314–1334.
- Eliaz, Kfir and Rani Spiegler (2022). “A Model of Competing Narratives”. In: *American Economic Review* 110.12, pp. 3786–3816.
- Enke, Benjamin (2020). “What you see is all there is”. In: *The Quarterly Journal of Economics* 135.3, pp. 1363–1398.
- (2024). “The Cognitive Turn in Behavioral Economics”. In: *Working Paper*.
- Enke, Benjamin and Thomas Graeber (2023). “Cognitive Uncertainty”. In: *Quarterly Journal of Economics* 138.4, pp. 2021–2067.
- Enke, Benjamin, Thomas Graeber, and Ryan Oprea (2023). “Confidence, Self-selection and Bias in the Aggregate”. In: *American Economic Review* 113.7, pp. 1933–1966.
- (forthcoming). “Complexity and Time”. In: *Journal of the European Economic Association*.
- Enke, Benjamin, Thomas Graeber, Ryan Oprea, and Jeffrey Yang (2024). *Behavioral Attenuation*. Tech. rep. National Bureau of Economic Research.
- Enke, Benjamin, Frederik Schwerter, and Florian Zimmermann (2024). “Associative Memory, Beliefs and Market Interactions”. In: *Journal of Financial Economics* 157, p. 103853. DOI: 10.1016/j.jfineco.2024.103853.
- Enke, Benjamin and Cassidy Shubatt (2024). “Quantifying Lottery Choice Complexity”. In: *Working Paper*.
- Enke, Benjamin and Florian Zimmermann (2019). “Correlation neglect in belief formation”. In: *The review of economic studies* 86.1, pp. 313–332.

- Esponda, Ignacio, Ryan Oprea, and Sevgi Yuksel (2023). “Seeing What is Representative”. In: *Quarterly Journal of Economics* 138.4, pp. 2607–2657.
- Fan, Tony, Yucheng Liang, and Cameorn Peng (2024). “The Inference-Forecast Gap in Belief Updating”. In: *Working Paper*.
- Frick, Mira, Ryota Iijima, and Yuhta Ishii (2022). “Dispersed behavior and perceptions in assortative societies”. In: *American Economic Review* 112.9, pp. 3063–3105.
- Gabaix, Xavier (2014). “A sparsity-based model of bounded rationality”. In: *The Quarterly Journal of Economics* 129.4, pp. 1661–1710.
- Gagnon-Bartsch, Tristan, Matthew Rabin, and Joshua Schwartzstein (2023). “Channeled attention and stable errors”. In: *Working Paper*.
- Giglio, Stefano, Matteo Maggiori, Johannes Stroebe, and Stephen Utkus (2021). “Five facts about beliefs and portfolios”. In: *American Economic Review* 111.5, pp. 1481–1522.
- Graeber, Thomas (2023). “Inattentive inference”. In: *Journal of the European Economic Association* 21.2, pp. 560–592.
- Graeber, Thomas, Christopher Roth, and Constantin Schesch (2024). “Explanations”. In: *Working Paper*.
- Graeber, Thomas, Christopher Roth, and Florian Zimmermann (2024). “Stories, statistics, and memory”. In: *Quarterly Journal of Economics*.
- Hartzmark, Samuel, Samuel Hirshman, and Alex Imas (2021). “Ownership, Learning, and Beliefs”. In: *Quarterly Journal of Economics* 136.3, pp. 1665–1717.
- Heidhues, Paul, Botond Köszegi, and Philipp Strack (2018). “Unrealistic Expectations and Misguided Learning”. In: *Econometrica* 86.4, pp. 1159–1214.
- (2023). “Misinterpreting Yourself”. In: *Working Paper*.
- Jiang, Zhengyang, Hongqi Liu, Cameron Peng, and Hongjun Yan (forthcoming). “Investor Memory and Biased Beliefs: Evidence from the Field”. In: *The Quarterly Journal of Economics*.
- Kendall, Chad and Ryan Oprea (2024). “On the Complexity of Forming Mental Models”. In: *Quantitative Economics* 15, pp. 175–211.
- Köszegi, Botond and Filip Matějka (2020). “Choice simplification: A theory of mental budgeting and naive diversification”. In: *The Quarterly Journal of Economics* 135.2, pp. 1153–1207.

- Kőszegi, Botond and Adam Szeidl (2013). “A model of focusing in economic choice”. In: *The Quarterly journal of economics* 128.1, pp. 53–104.
- Laudenbach, Christina, Annika Weber, Rüdiger Weber, and Johannes Wohlfart (forthcoming). “Beliefs About the Stock Market and Investment Choices: Evidence from a Field Experiment”. In: *Review of Financial Studies*.
- Mailath, George J and Larry Samuelson (2019). “The Wisdom of a Confused Crowd: Model-Based Inference”. In: *Working Paper*.
- Montiel Olea, José Luis, Pietro Ortoleva, Mallesh M Pai, and Andrea Prat (2022). “Competing models”. In: *The Quarterly Journal of Economics* 137.4, pp. 2419–2457.
- Oprea, Ryan (2020). “What Makes a Rule Complex?” In: *American Economic Review* 110.12, pp. 3913–3951.
- Schwartzstein, Joshua (2014). “Selective attention and learning”. In: *Journal of the European Economic Association* 12.6, pp. 1423–1452.
- Schwartzstein, Joshua and Adi Sunderam (2021). “Using models to persuade”. In: *American Economic Review* 111.1, pp. 276–323.
- Shiller, Robert (2017). “Narrative Economics”. In: *American Economic Review* 107.4, pp. 967–1004.
- Shubatt, Cassidy and Jeffrey Yang (2024). *Tradeoffs and Comparison Complexity*. Tech. rep. Working Paper.
- Somerville, Jason (2022). “Range-Dependent Attribute Weighting in Consumer Choice: An Experimental Test”. In: *Econometrica* 90.2, pp. 799–830.
- Yang, Jeffrey (2023). “On the Decision-Relevance of Subjective Beliefs”. In: *Available at SSRN 4425080*.

A Additional Results for the Restricted Samples

A.1 Overview of Data Collections

Table A.1: Overview of Data Collections

Collection	Participants	Description	Link to pre-analysis plan
Baseline	600	As described in Section 2. Treatments: HighComplexity and LowComplexity. Outcomes: Company value guesses, hover times, beliefs.	https://aspredicted.org/t79h-2mkj.pdf
Equally Likely Models	600	As Baseline, but without noisy indication, hence 50-50 belief about more likely model, and with additional guess confidence measure. Treatments: HighComplexity and LowComplexity. Outcomes: Company value guesses, guess confidence, hover times.	https://aspredicted.org/92cr-9kwd.pdf
Investment 1	400	As Baseline, but company value guesses replaced by investment decisions. Treatments: HighComplexity and LowComplexity. Outcomes: Investment decisions, hover times, beliefs.	https://aspredicted.org/4c38-7psb.pdf
Investment 2	600	As Investment, but larger sample size. Treatments: HighComplexity and LowComplexity. Outcomes: Investment decisions, hover times, beliefs.	https://aspredicted.org/7vcv-gq8z.pdf
Delayed Belief Elicitation	588	As Baseline, but additional variation the timing of belief elicitation: For half the respondents, beliefs are elicited immediately as in Baseline, for the other half they are elicited with a one-day delay. Treatments: (HighComplexity, LowComplexity) $\times$ (Immediate, Recall). Outcomes: Company value guesses, hover times, beliefs.	Not preregistered
Baseline Confidence	600	As Baseline, but with additional guess confidence measure. Treatments: HighComplexity and LowComplexity. Outcomes: Company value guesses, guess confidence, hover times, beliefs.	https://aspredicted.org/v6cv-y9yh.pdf
Incentivized Confidence	600	As Baseline Confidence, but with additional incentivized guess confidence measure on the same page as the non-incentivized measure. Treatments: HighComplexity and LowComplexity. Outcomes: Company value guesses, guess confidence, hover times, beliefs.	https://aspredicted.org/txsk-3hky.pdf

This Table provides an overview of the different data collections. The sample sizes refer to the size of the original data collection, prior to applying exclusion restrictions.

A.2 Baseline Experiment: Additional Results in Restricted Sample

Here, we present additional results for the restricted sample of the Baseline Experiment, featuring 230 participants who solved both of the example screens.

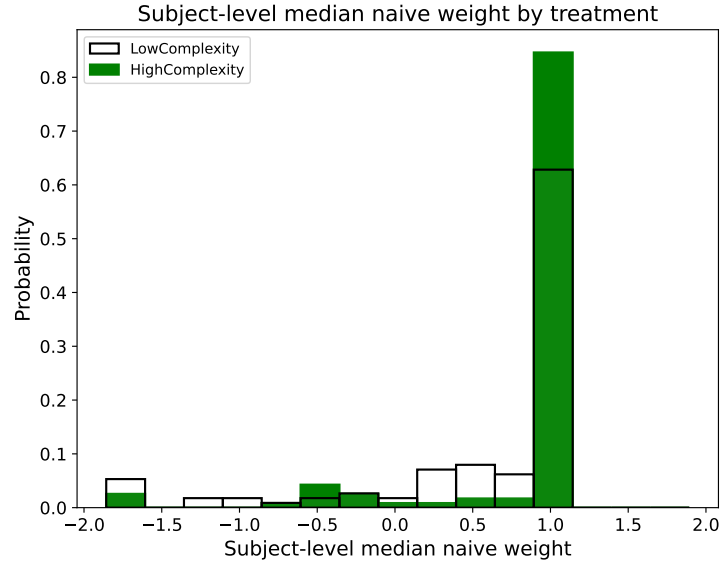


Figure A.1: Distribution of median naive weights, computed for each subject. *The figure plots the distribution of naive weights λ calculated as specified in Equation (2), using the restricted sample of the Baseline Experiment with 230 participants. Only the median naive weight for each participant is plotted.*

A.3 Equally Likely Models: Additional Results in Restricted Sample

Here, we present additional results for the restricted sample of the Equally Likely Models Experiment, featuring 348 participants who solved both of the example screens.

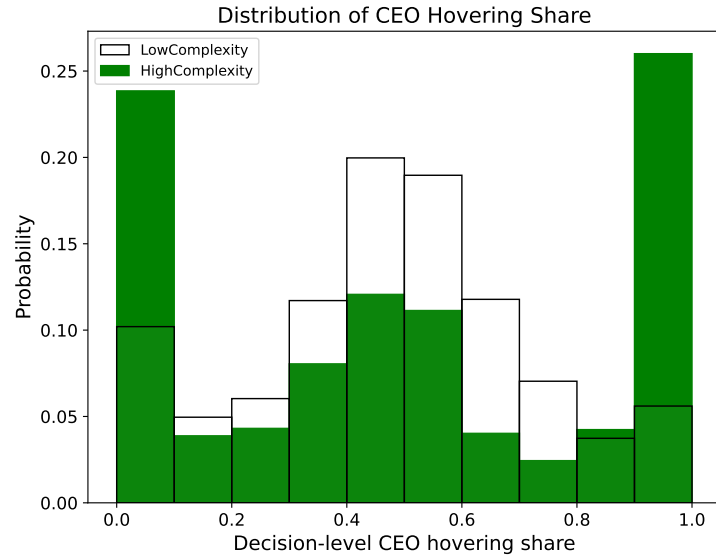


Figure A.2: Distribution of decision-level CEO hovering share in the Equally Likely Models Experiment. *The figure plots the distribution of the share of time that respondents spent looking at the values of the "The CEO is key" model, using the restricted sample of the Equally Likely Models study with 348 participants. Hovering shares are plotted separately for each of the eight guesses made by respondents*

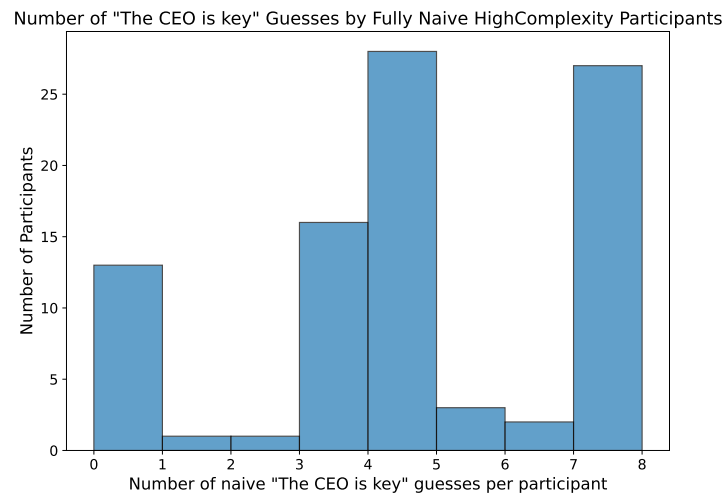


Figure A.3: Distribution of naive decision-making in the Equally Likely Models Experiment. *The figure plots the distribution of the number of times participants selected the value corresponding to the "The CEO is key" model, using the restricted sample of the HighComplexity treatment of the Equally Likely Models study, limited to the 91 participants in the HighComplexity condition who made only fully naive guesses, i.e. who always selected a value corresponding to either the "The CEO is key" or "Products are crucial" model.*

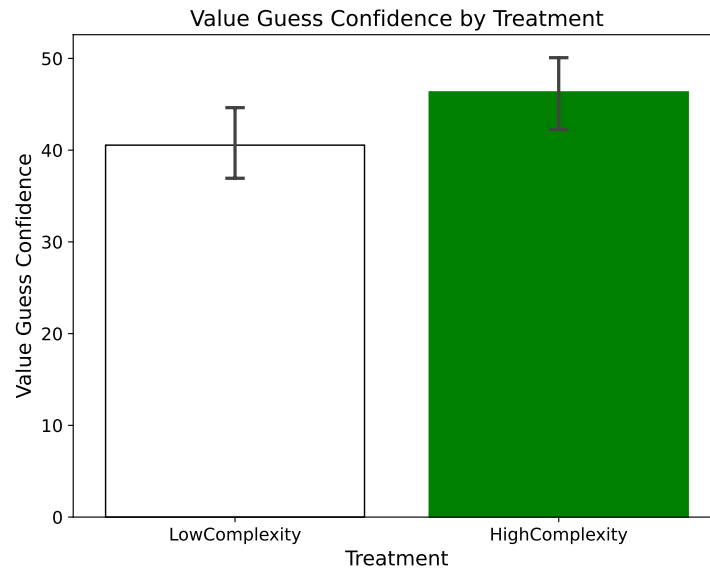
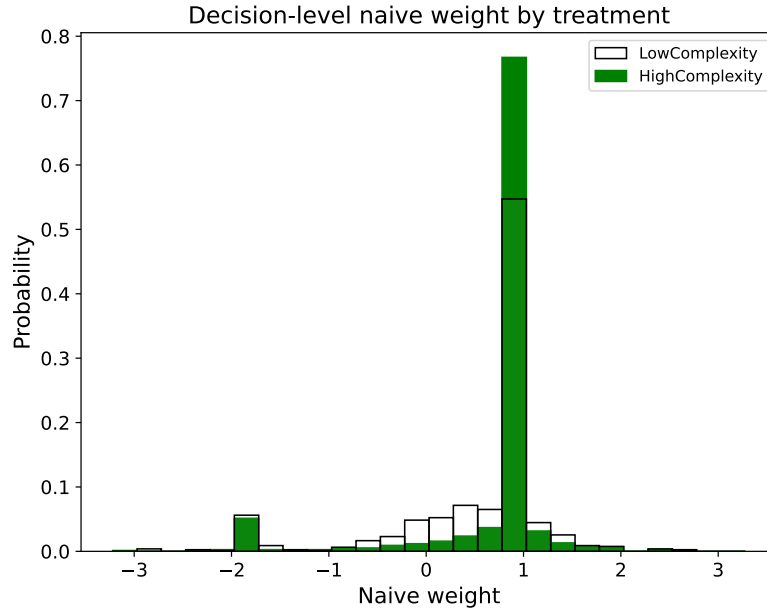


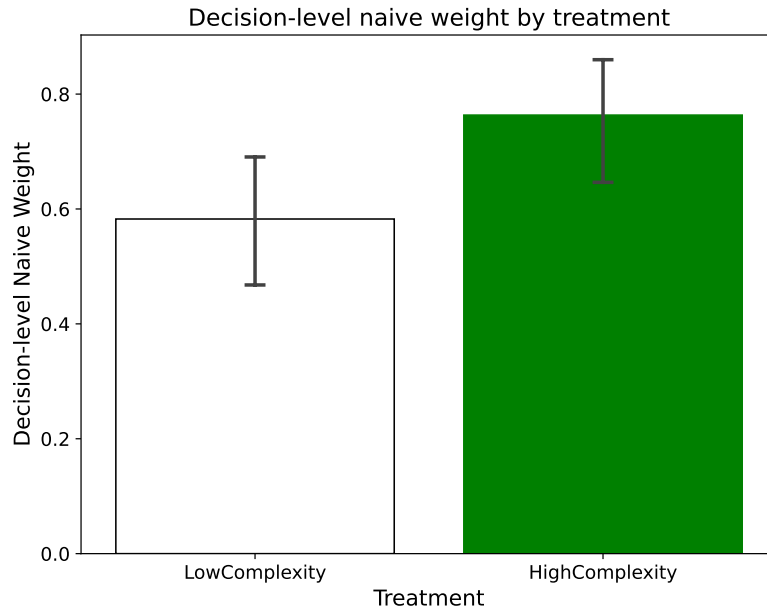
Figure A.4: Average value guess confidence in the Equally Likely Models Experiment. *The figure plots the average confidence that respondents had in their company value guesses, using the restricted sample of the Equally Likely Models study with 348 participants.*

A.4 Investment Experiment 1: Results in Restricted Sample

Here, we present the results for the restricted sample of the Investment Experiment 1, featuring 193 participants who solved both of the example screens.



(a) Distribution of decision-level naive weights



(b) Mean decision-level naive weights

Figure A.5: Decision-level naive weights in the restricted sample of the Investment Experiment 1. Panel (a) plots the distribution of naive weights λ calculated as specified in Equation 2, using the restricted sample of the Investment Experiment 1 with 193 participants. Panel (b) plots average naive weights.

Table A.2: Company Bids in Restricted Sample of Investment Experiment 1

<i>Dependent variable:</i>	Company Bids		
	LowComplexity (1)	HighComplexity (2)	Pooled (3)
Rational Benchmark	0.478*** (0.065)	0.311*** (0.065)	0.478*** (0.064)
Naive Benchmark	0.558*** (0.053)	0.725*** (0.055)	0.558*** (0.052)
Rational B. $\times$ HighComplexity			-0.167* (0.092)
Naive B. $\times$ HighComplexity			0.167** (0.076)
R^2	0.906	0.920	0.913
Observations	784	760	1544

The table presents OLS regressions of respondents' company bids on the rational and naive benchmarks as detailed in Section 3.1. The table is based on the restricted sample of the Investment Experiment 1. Column (1) uses observations from the LowComplexity treatment, column (2) from the HighComplexity treatment, and column (3) from both. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

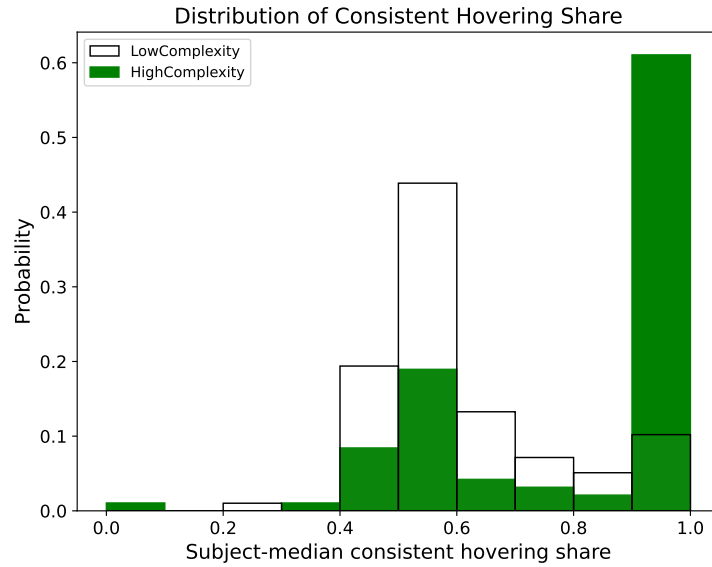
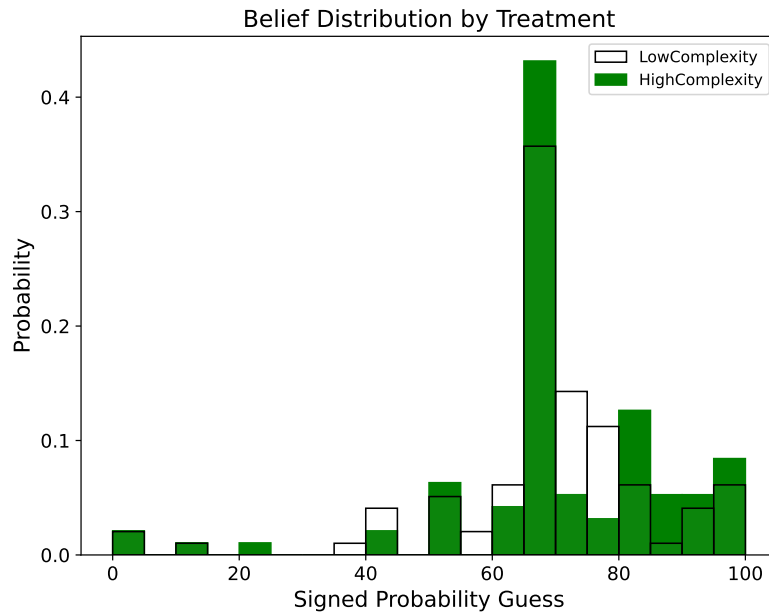
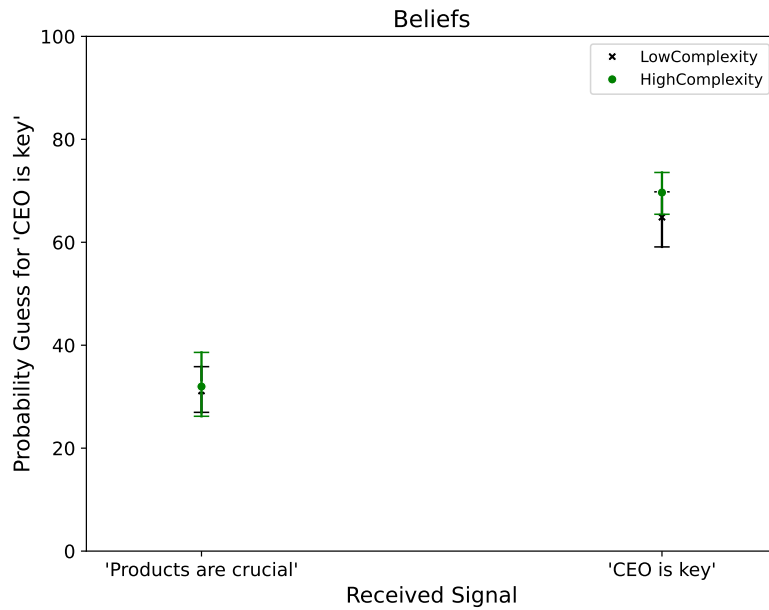


Figure A.6: Distribution of subject-medians of the consistent hovering shares in the restricted sample of the Investment Experiment 1. The figure plots the distribution of the share of time that respondents spent looking at the values of the signal-consistent model, using the restricted sample of the Investment Experiment 1 with 193 participants. Only the median consistent share for each participant is plotted.



(a) Histogram of signed probability guesses.



(b) Mean probability guesses for "The CEO is key".

Figure A.7: Beliefs in the restricted sample of the Investment Experiment 1. In Panel (a) beliefs are converted into the direction of the more likely model, so that 65 corresponds to the Bayesian probability. Panel (b) plots mean probability guesses for the model "The CEO is key". The figure is based on the restricted sample of the Investment Experiment 1 with 193 participants.

Table A.3: Beliefs and Recall in Restricted Sample of the Investment Experiment 1

<i>Dependent variable:</i>	Probability Guess	Correct Recall
<i>Sample:</i>	Pooled (1)	Pooled (2)
Constant	66.844 ^{***} (1.758)	0.959 ^{***} (0.020)
HighComplexity	2.044 (2.598)	0.030 (0.023)
R^2	0.003	0.009
Observations	193	193

The table presents OLS regressions using the restricted sample of the Investment Experiment 1. In column (1), the dependent variable is the probability guess for the likely state. Column (2) uses a dummy for whether respondents correctly recall the more likely model. All columns use observations from both the HighComplexity and LowComplexity conditions. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

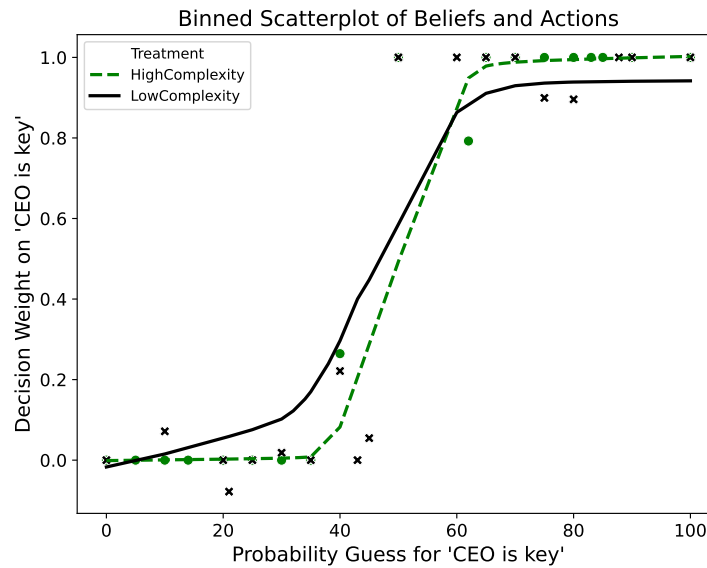
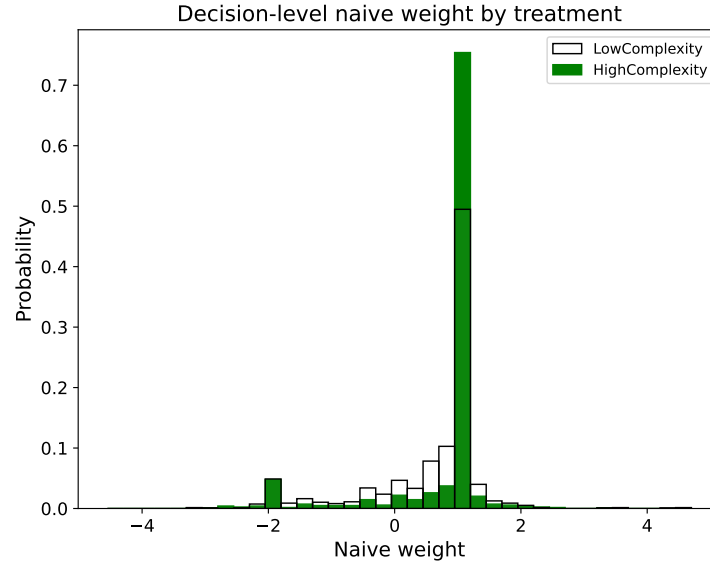


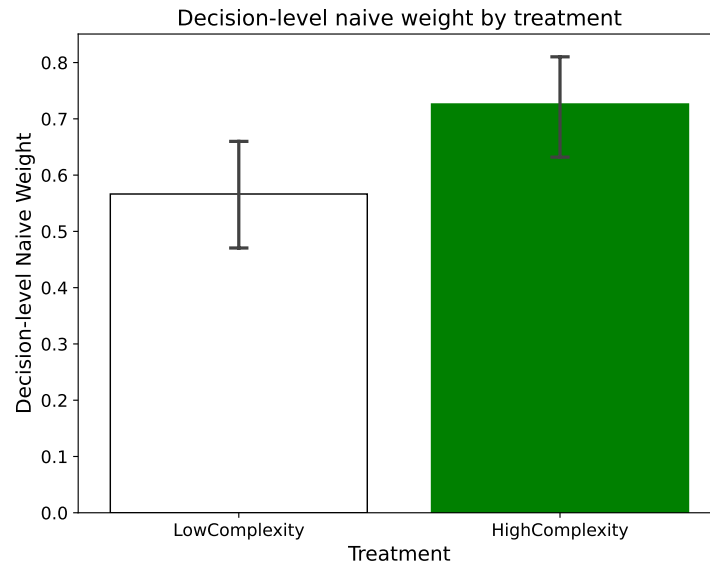
Figure A.8: The relationship between decisions and beliefs in the restricted sample of the Investment Experiment 1. The figure shows a binned scatterplot using the restricted sample of the Investment Experiment 1 with 193 participants. It has the probability guess that "The CEO is key" is the more likely model on the horizontal axis, and the decision weight γ as defined in Equation 3 on the vertical axis. The lines show LOWESS regressions based on all datapoints.

A.5 Investment Experiment 2: Results in Restricted Sample

Here, we present the results for the restricted sample of the Investment Experiment 2, featuring 323 participants who solved both of the example screens.



(a) Distribution of decision-level naive weights



(b) Mean decision-level naive weights

Figure A.9: Decision-level naive weights in the restricted sample of the Investment Experiment 2. Panel (a) plots the distribution of naive weights λ calculated as specified in Equation 2, using the restricted sample of the Investment Experiment 2 with 323 participants. Panel (b) plots average naive weights.

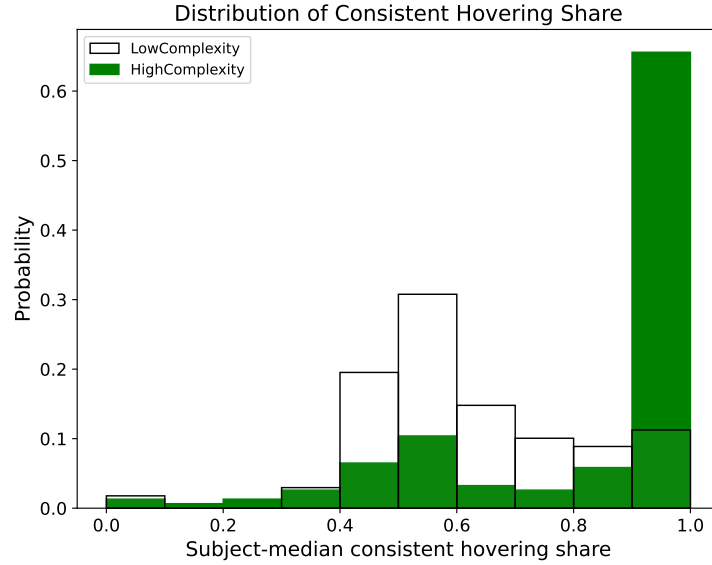
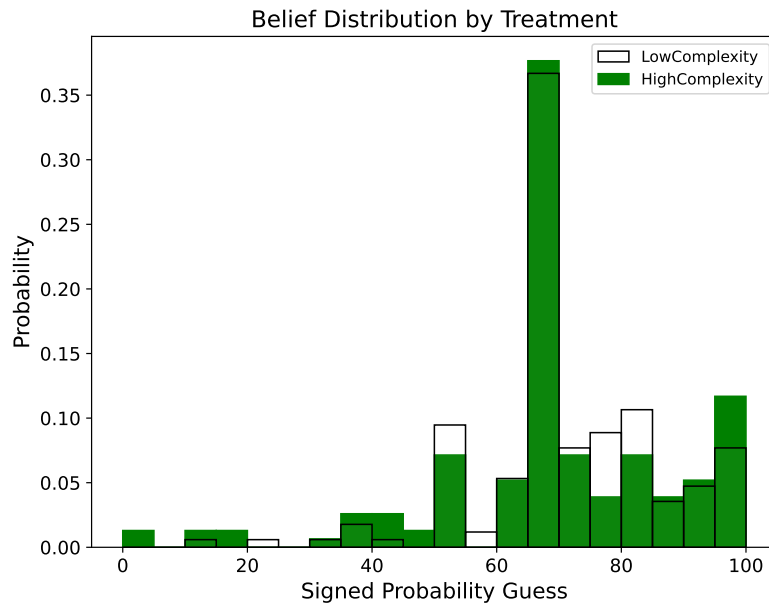


Figure A.10: Distribution of subject-medians of the consistent hovering shares in the restricted sample of the Investment Experiment 2. *The figure plots the distribution of the share of time that respondents spent looking at the values of the signal-consistent model, using the restricted sample of the Investment Experiment 2 with 323 participants. Only the median consistent share for each participant is plotted.*

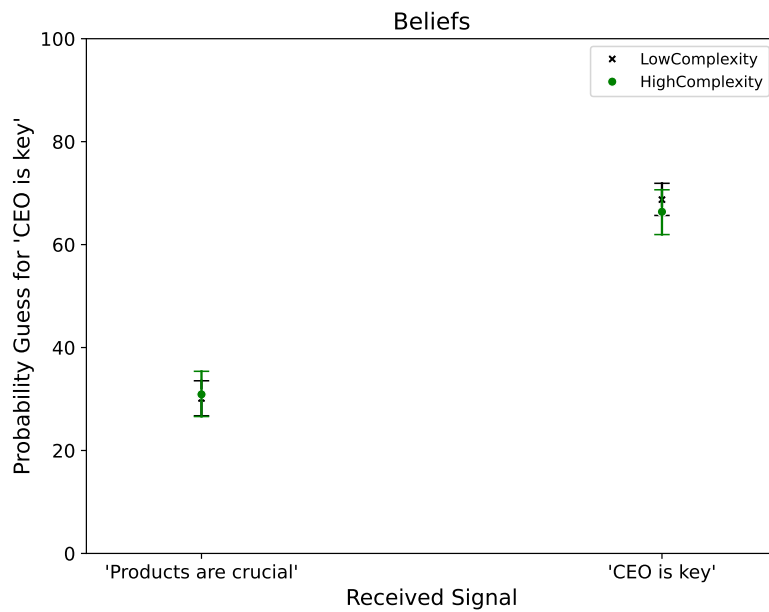
Table A.4: Company Bids in Restricted Sample of Investment Experiment 2

<i>Dependent variable:</i>	Company Bids		
	LowComplexity (1)	HighComplexity (2)	Pooled (3)
<i>Sample:</i>			
Rational Benchmark	0.508*** (0.052)	0.281*** (0.050)	0.508*** (0.052)
Naive Benchmark	0.534*** (0.043)	0.721*** (0.045)	0.534*** (0.043)
Rational B. $\times$ HighComplexity			-0.227*** (0.072)
Naive B. $\times$ HighComplexity			0.187*** (0.062)
R^2	0.892	0.908	0.900
Observations	1352	1232	2584

The table presents OLS regressions of respondents' company bids on the rational and naive benchmarks as detailed in Section 3.1. The table is based on the restricted sample of the Investment Experiment 2. Column (1) uses observations from the LowComplexity treatment, column (2) from the HighComplexity treatment, and column (3) from both. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.



(a) Histogram of signed probability guesses.



(b) Mean probability guesses for "The CEO is key".

Figure A.11: Beliefs in the restricted sample of the Investment Experiment 2. In Panel (a) beliefs are converted into the direction of the more likely model, so that 65 corresponds to the Bayesian probability. Panel (b) plots mean probability guesses for the model "The CEO is key". The figure is based on the restricted sample of the Investment Experiment 2 with 323 participants.

Table A.5: Beliefs and Recall in Restricted Sample of the Investment Experiment 2

<i>Dependent variable:</i>	Probability Guess	Correct Recall
<i>Sample:</i>	Pooled (1)	Pooled (2)
Constant	69.302 ^{***} (1.176)	0.982 ^{***} (0.010)
HighComplexity	-1.563 (1.999)	-0.021 (0.019)
R^2	0.002	0.004
Observations	323	323

The table presents OLS regressions using the restricted sample of the Investment Experiment 2. In column (1), the dependent variable is the probability guess for the likely state. Column (2) uses a dummy for whether respondents correctly recall the more likely model. All columns use observations from both the HighComplexity and Low-Complexity conditions. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

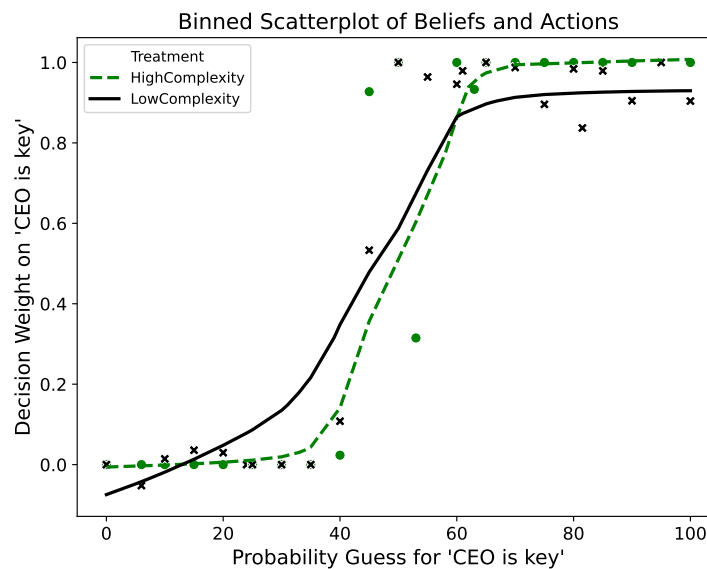


Figure A.12: The relationship between decisions and beliefs in the restricted sample of the Investment Experiment 2. The figure shows a binned scatterplot using the restricted sample of the Investment Experiment 2 with 323 participants. It has the probability guess that "The CEO is key" is the more likely model on the horizontal axis, and the decision weight γ as defined in Equation 3 on the vertical axis. The lines show LOWESS regressions based on all datapoints.

A.6 Delayed Belief Elicitation: Results in Restricted Sample

Here, we present the results for the restricted sample of the Delayed Belief Elicitation Experiment, featuring 342 participants who solved both of the example screens.

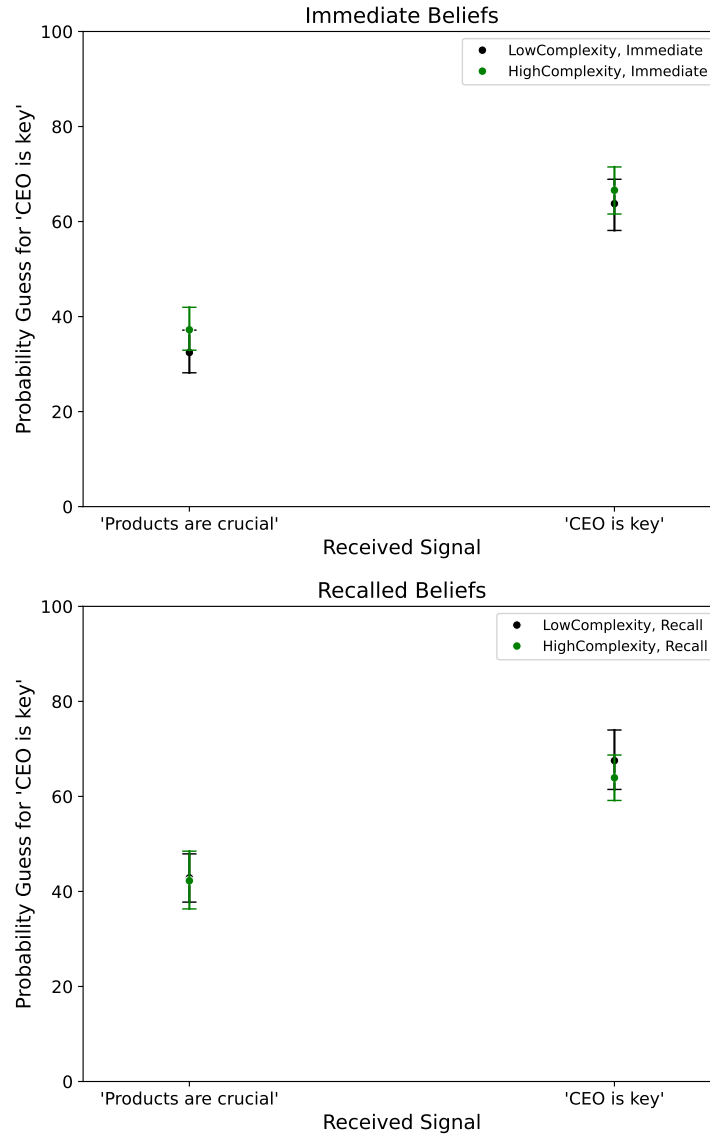


Figure A.13: Immediate and recalled beliefs in the Delayed Belief Elicitation Experiment. The top panel plots the mean probability guess for the "The CEO is key" model in the Immediate condition by received signal separately for the LowComplexity and HighComplexity condition. The bottom panel does the same for the Recall condition, where beliefs were elicited with a one day delay. This figure is based on the restricted sample of the Delayed Belief Elicitation Experiment with 342 participants.

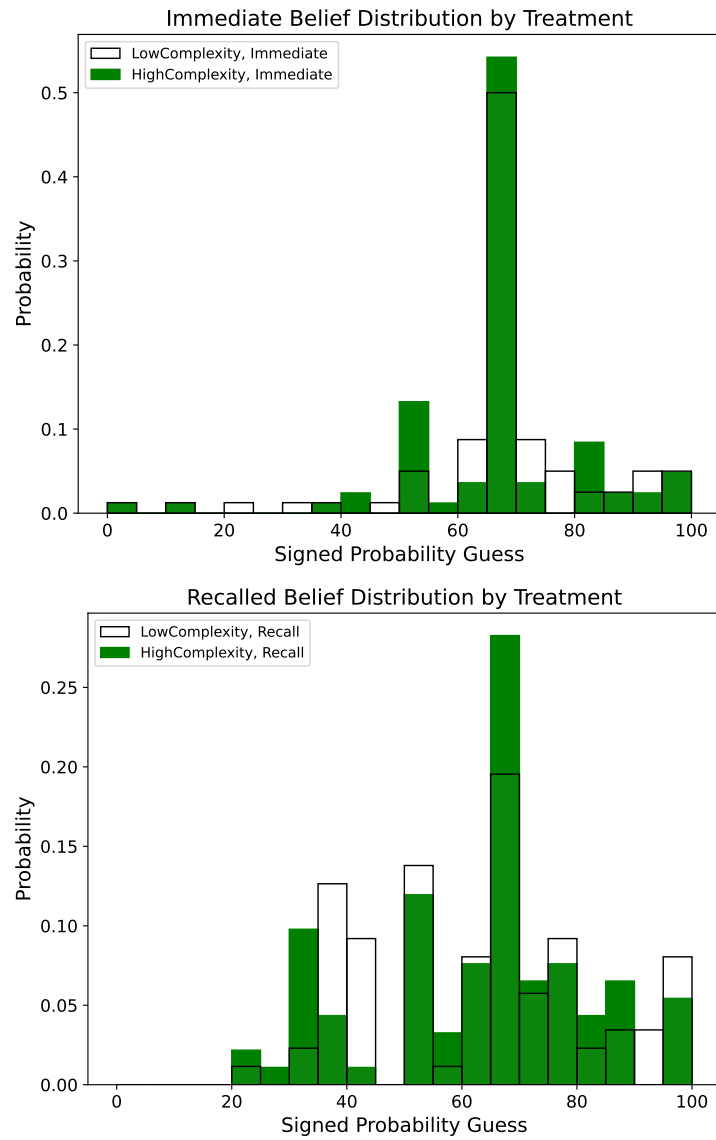


Figure A.14: Distribution of immediate and recalled beliefs in the restricted sample of the Delayed Belief Elicitation Experiment. *The top panel plots the distribution of immediate beliefs and the bottom panel of recalled beliefs. Beliefs are converted into the direction of the more likely model, so that 65 corresponds to the Bayesian probability. This figure is based on the restricted sample of the Delayed Belief Elicitation Experiment with 342 participants.*

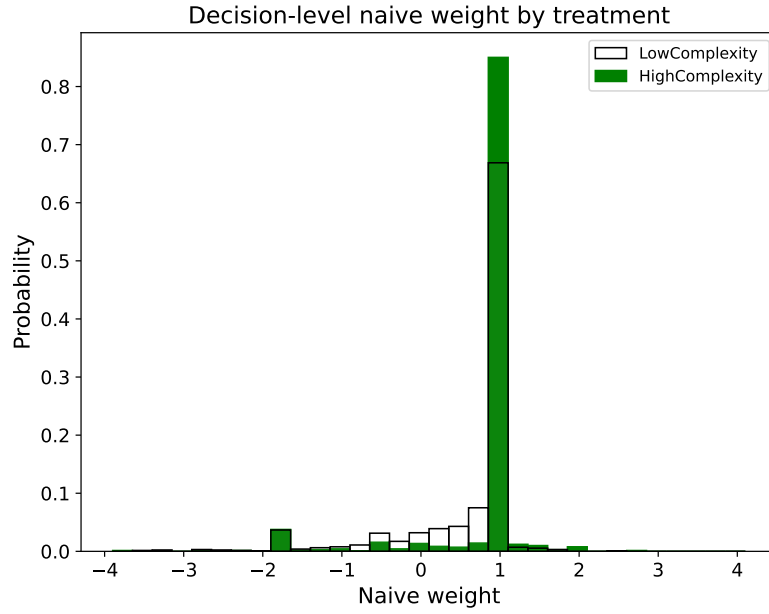
Table A.6: Beliefs and Recall in the Restricted Sample of the Delayed Belief Elicitation Experiment

<i>Dependent variable:</i>	Probability Guess		Correct Recall	
<i>Sample:</i>	Immediate (1)	Recall (2)	Immediate (3)	Recall (4)
Constant	65.713*** (1.865)	61.195*** (2.106)	0.950*** (0.025)	0.736*** (0.048)
HighComplexity	-1.008 (2.547)	0.044 (2.843)	0.026 (0.030)	0.014 (0.066)
R^2	0.001	0.000	0.005	0.000
Observations	163	179	163	179

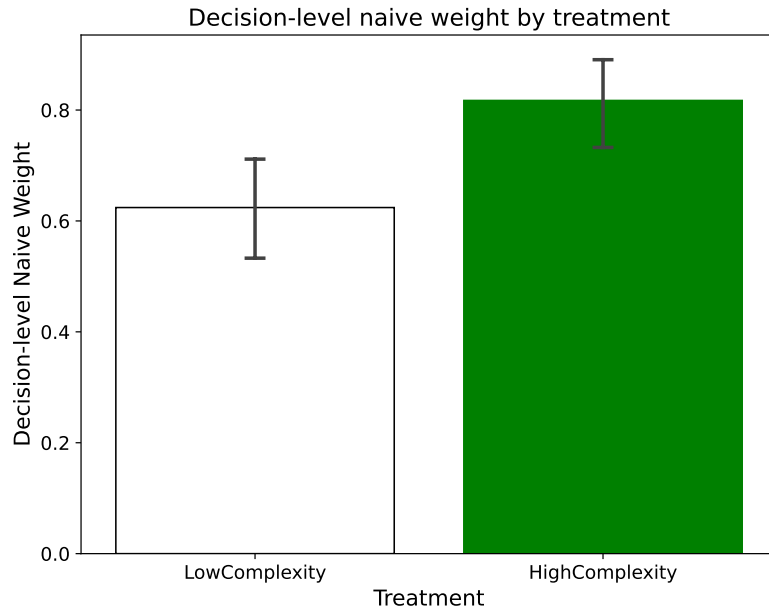
The table presents OLS regressions using the restricted sample of the Delayed Belief Elicitation Experiment. In columns (1) and (2), we regress respondents' probability guesses on a constant and a treatment dummy for the HighComplexity condition. Probability guesses are converted in the direction of the more likely model, so that a guess of 65 corresponds to the Bayesian probability. In columns (3) and (4), the dependent variable is a dummy for whether respondents correctly recall the more likely model. Columns (1) and (3) use observations from the Immediate condition, while columns (2) and (4) use observations from the Recall condition. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

A.7 Baseline Confidence: Additional Results in Restricted Sample

Here, we present the results for the restricted sample of the Baseline Confidence Experiment, featuring 336 participants who solved both of the example screens.



(a) Distribution of decision-level naive weights



(b) Mean decision-level naive weights

Figure A.15: Decision-level naive weights in the restricted sample of the Baseline Confidence Experiment. *Panel (a) plots the distribution of naive weights λ calculated as specified in Equation 2, using the restricted sample of the Baseline Confidence Experiment with 336 participants. Panel (b) plots average naive weights.*

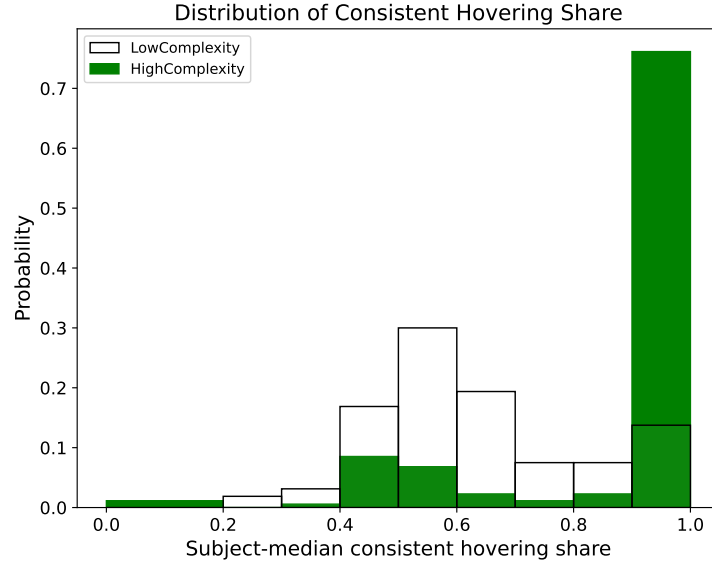
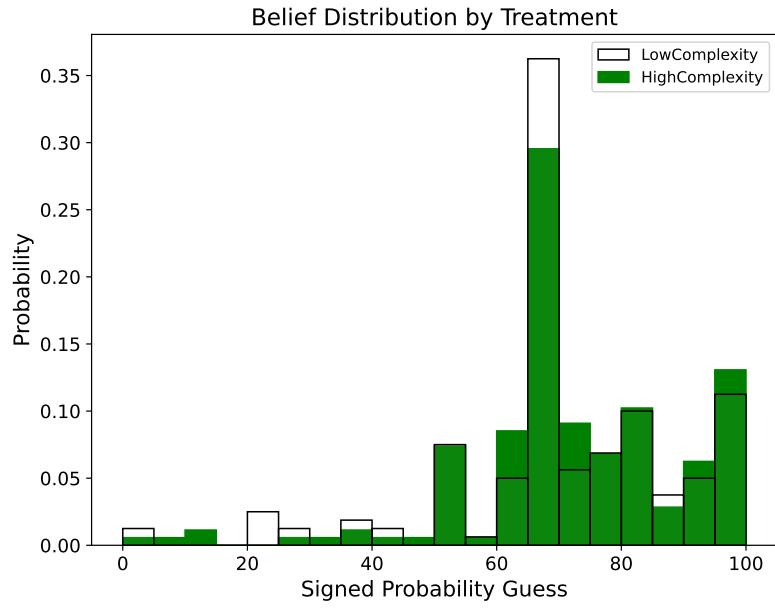


Figure A.16: Distribution of subject-medians of the consistent hovering shares in the restricted sample of the Baseline Confidence Experiment. *The figure plots the distribution of the share of time that respondents spent looking at the values of the signal-consistent model, using the restricted sample of the Baseline Confidence Experiment with 336 participants. Only the median consistent share for each participant is plotted.*

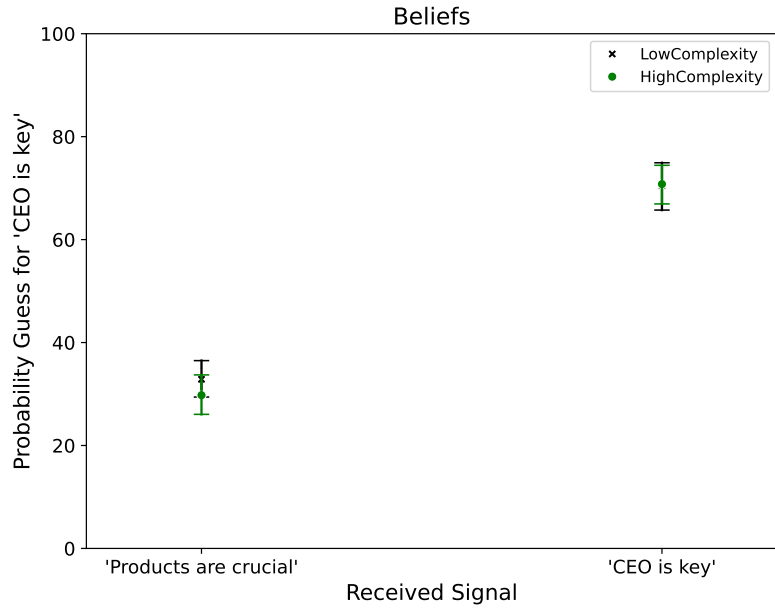
Table A.7: Company Value Guesses in the Restricted Sample of the Baseline Confidence Experiment

Dependent variable:	Company Value Guess		
	LowComplexity (1)	HighComplexity (2)	Pooled (3)
Rational Benchmark	0.428*** (0.052)	0.226*** (0.042)	0.428*** (0.052)
Naive Benchmark	0.613*** (0.044)	0.797*** (0.039)	0.613*** (0.044)
Rational B. $\times$ HighComplexity			-0.202*** (0.067)
Naive B. $\times$ HighComplexity			0.184*** (0.059)
R^2	0.920	0.936	0.928
Observations	1280	1408	2688

The table presents OLS regressions of respondents' company value guesses on the rational and naive benchmarks as detailed in Section 3.1. The table is based on the restricted sample of the Baseline Confidence Experiment. Column (1) uses observations from the LowComplexity treatment, column (2) from the HighComplexity treatment, and column (3) from both. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.



(a) Histogram of signed probability guesses.



(b) Mean probability guesses for "The CEO is key".

Figure A.17: Beliefs in the restricted sample of the Baseline Confidence Experiment. In Panel (a) beliefs are converted into the direction of the more likely model, so that 65 corresponds to the Bayesian probability. Panel (b) plots mean probability guesses for the model "The CEO is key". The figure is based on the restricted sample of the Baseline Confidence Experiment with 336 participants.

Table A.8: Beliefs and Recall in the Restricted Sample of the Baseline Confidence Experiment

<i>Dependent variable:</i>	Probability Guess	Correct Recall
<i>Sample:</i>	Pooled (1)	Pooled (2)
Constant	68.806 ^{***} (1.518)	0.944 ^{***} (0.018)
HighComplexity	1.716 (2.056)	0.022 (0.023)
R^2	0.002	0.003
Observations	336	336

The table presents OLS regressions using the restricted sample of the Baseline Confidence Experiment. In column (1), the dependent variable is the probability guess for the likely state. Column (2) uses a dummy for whether respondents correctly recall the more likely model. All columns use observations from both the HighComplexity and LowComplexity conditions. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

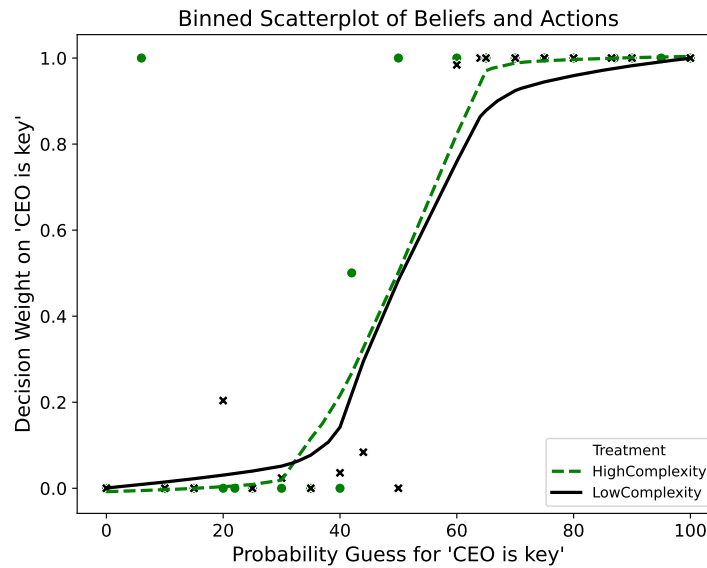
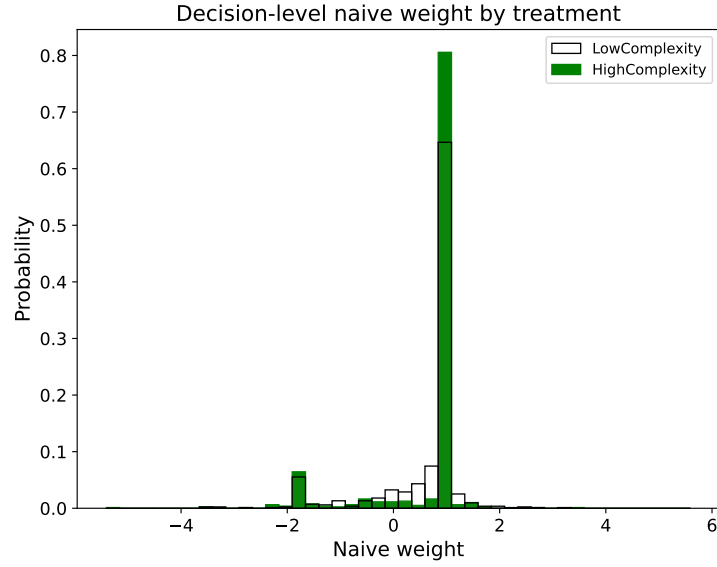


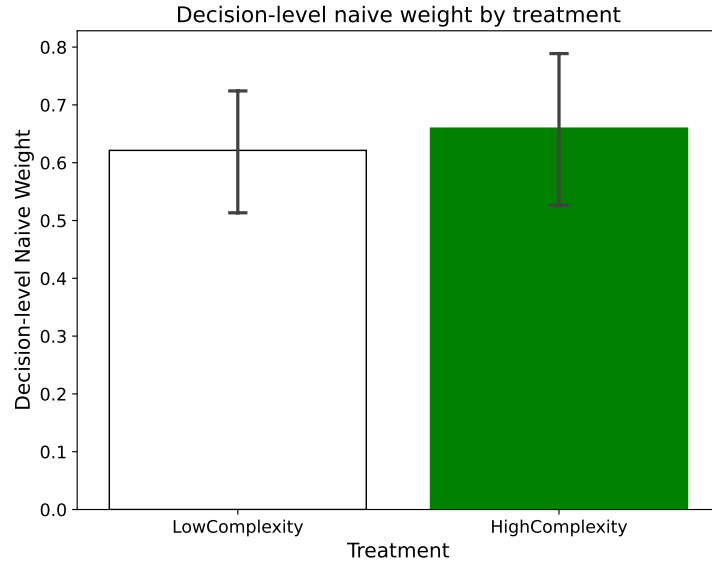
Figure A.18: The relationship between decisions and beliefs in the restricted sample of the Baseline Confidence Experiment. *The figure shows a binned scatterplot using the restricted sample of the Baseline Confidence Experiment with 336 participants. It has the probability guess that "The CEO is key" is the more likely model on the horizontal axis, and the decision weight γ as defined in Equation 3 on the vertical axis. The lines show LOWESS regressions based on all datapoints.*

A.8 Incentivized Confidence: Additional Results in Restricted Sample

Here, we present the results for the restricted sample of the Incentivized Confidence Experiment, featuring 203 participants who solved both of the example screens.



(a) Distribution of decision-level naive weights



(b) Mean decision-level naive weights

Figure A.19: Decision-level naive weights in the restricted sample of the Incentivized Confidence Experiment. Panel (a) plots the distribution of naive weights λ calculated as specified in Equation 2, using the restricted sample of the Incentivized Confidence Experiment with 203 participants. Panel (b) plots average naive weights.

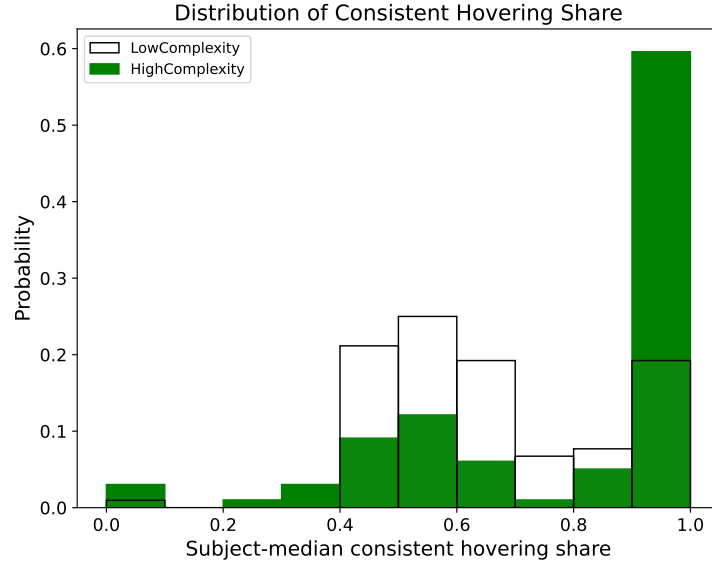
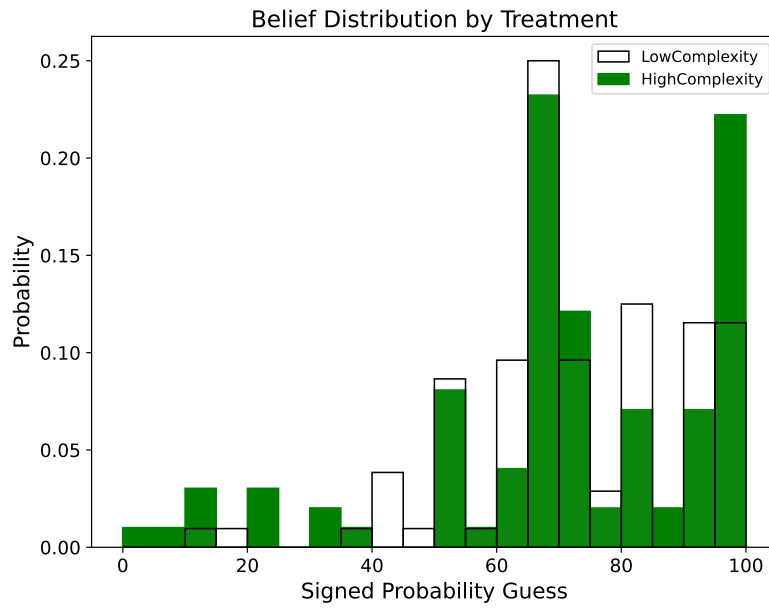


Figure A.20: Distribution of subject-medians of the consistent hovering shares in the restricted sample of the Incentivized Confidence Experiment. *The figure plots the distribution of the share of time that respondents spent looking at the values of the signal-consistent model, using the restricted sample of the Incentivized Confidence Experiment with 203 participants. Only the median consistent share for each participant is plotted.*

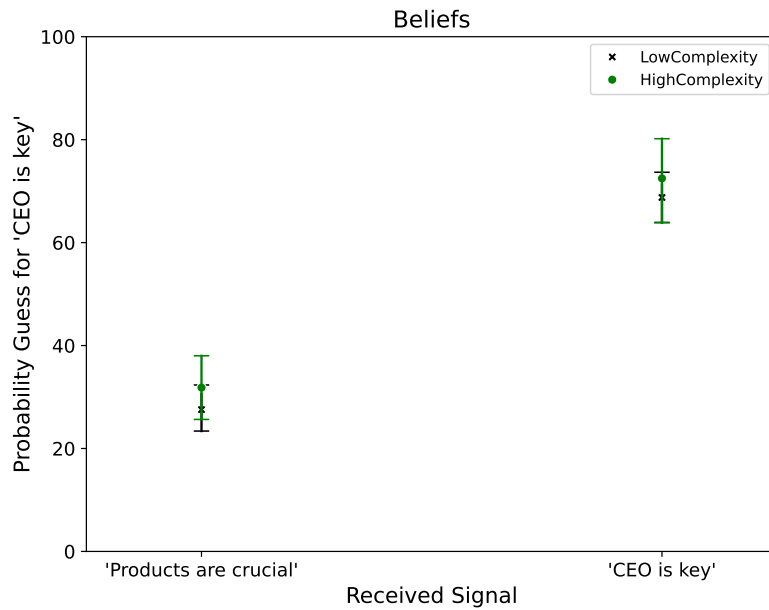
Table A.9: Company Value Guesses in the Restricted Sample of the Incentivized Confidence Experiment

Dependent variable:	Company Value Guess		
	LowComplexity (1)	HighComplexity (2)	Pooled (3)
Rational Benchmark	0.488*** (0.069)	0.404*** (0.080)	0.488*** (0.069)
Naive Benchmark	0.593*** (0.054)	0.638*** (0.069)	0.593*** (0.054)
Rational B. $\times$ HighComplexity			-0.084 (0.106)
Naive B. $\times$ HighComplexity			0.045 (0.087)
R^2	0.912	0.898	0.905
Observations	832	792	1624

The table presents OLS regressions of respondents' company value guesses on the rational and naive benchmarks as detailed in Section 3.1. The table is based on the restricted sample of the Incentivized Confidence Experiment. Column (1) uses observations from the LowComplexity treatment, column (2) from the HighComplexity treatment, and column (3) from both. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.



(a) Histogram of signed probability guesses.



(b) Mean probability guesses for "The CEO is key".

Figure A.21: Beliefs in the restricted sample of the Incentivized Confidence Experiment. In Panel (a) beliefs are converted into the direction of the more likely model, so that 65 corresponds to the Bayesian probability. Panel (b) plots mean probability guesses for the model "The CEO is key". The figure is based on the restricted sample of the Incentivized Confidence Experiment with 203 participants.

Table A.10: Beliefs and Recall in the Restricted Sample of the Incentivized Confidence Experiment

<i>Dependent variable:</i>	Probability Guess	Correct Recall
<i>Sample:</i>	Pooled (1)	Pooled (2)
Constant	70.327*** (1.766)	0.962*** (0.019)
HighComplexity	-0.238 (3.036)	0.018 (0.024)
R^2	0.000	0.003
Observations	203	203

The table presents OLS regressions using the restricted sample of the Incentivized Confidence Experiment. In column (1), the dependent variable is the probability guess for the likely state. Column (2) uses a dummy for whether respondents correctly recall the more likely model. All columns use observations from both the HighComplexity and LowComplexity conditions. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

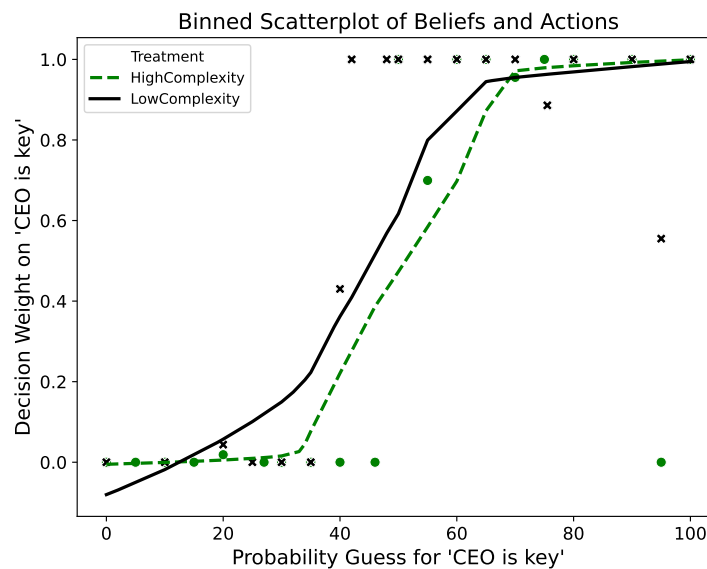
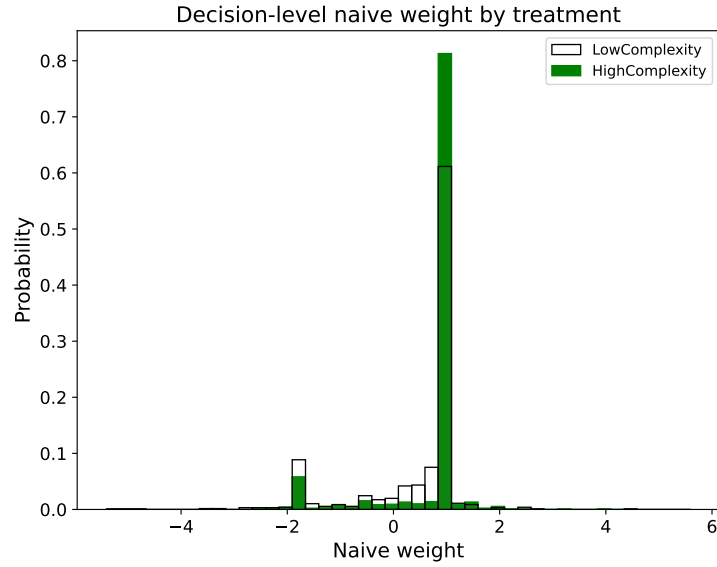


Figure A.22: The relationship between decisions and beliefs in the restricted sample of the Incentivized Confidence Experiment. The figure shows a binned scatterplot using the restricted sample of the Incentivized Confidence Experiment with 203 participants. It has the probability guess that "The CEO is key" is the more likely model on the horizontal axis, and the decision weight γ as defined in Equation 3 on the vertical axis. The lines show LOWESS regressions based on all datapoints.

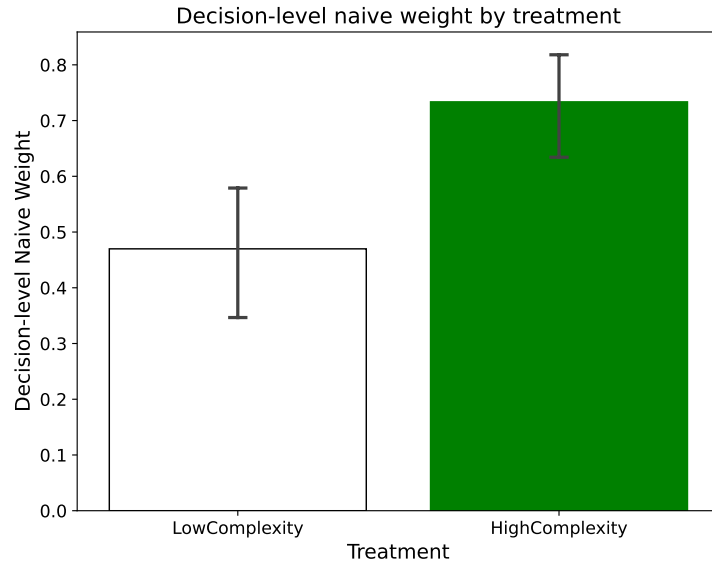
B Additional Results for the Robustness Samples

B.1 Baseline Experiment: Results in Lenient Sample

Here, we present the results for the lenient sample of the Baseline Experiment, featuring 319 participants who solved at least one of the example screens.



(a) Distribution of decision-level naive weights



(b) Mean decision-level naive weights

Figure B.1: Decision-level naive weights in the lenient sample of the Baseline Experiment. Panel (a) plots the distribution of naive weights λ calculated as specified in Equation 2, using the lenient sample of the Baseline Experiment with 319 participants. Panel (b) plots average naive weights.

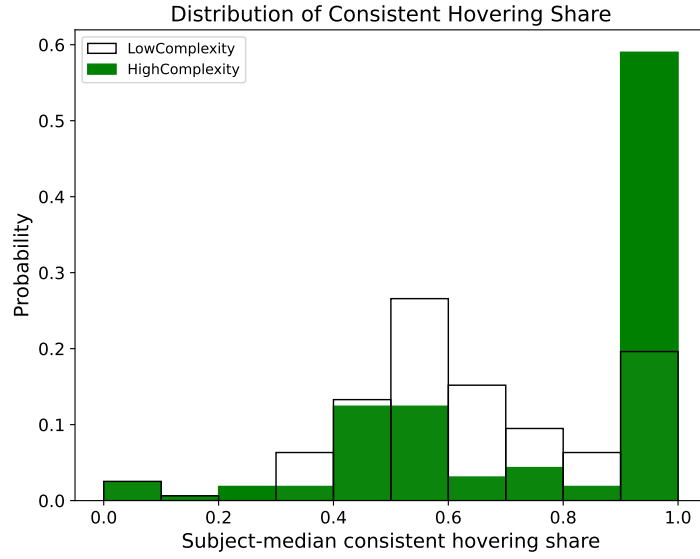


Figure B.2: Distribution of subject-medians of the consistent hovering shares in the lenient sample of the Baseline Experiment. *The figure plots the distribution of the share of time that respondents spent looking at the values of the signal-consistent model, using the lenient sample of the Baseline Experiment with 319 participants. Only the median consistent share for each participant is plotted.*

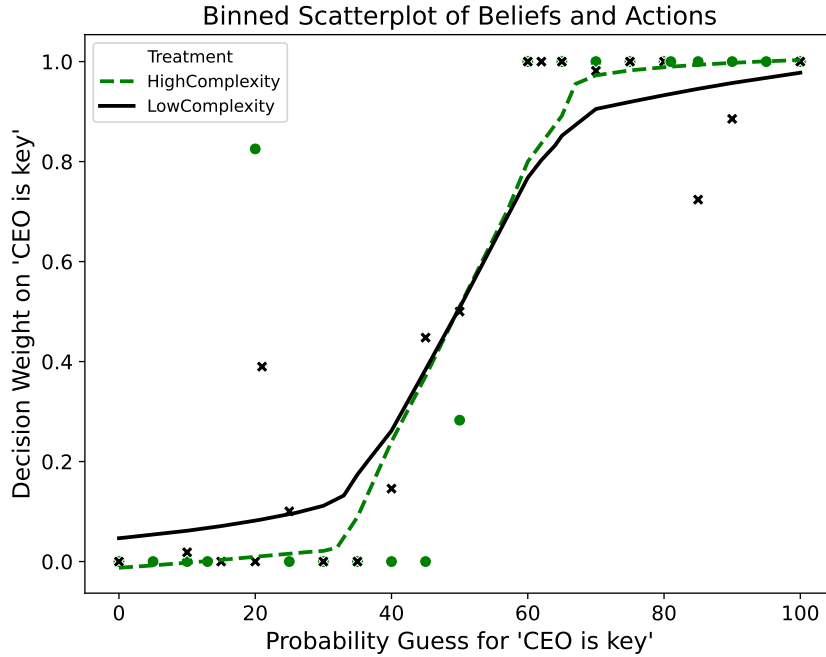
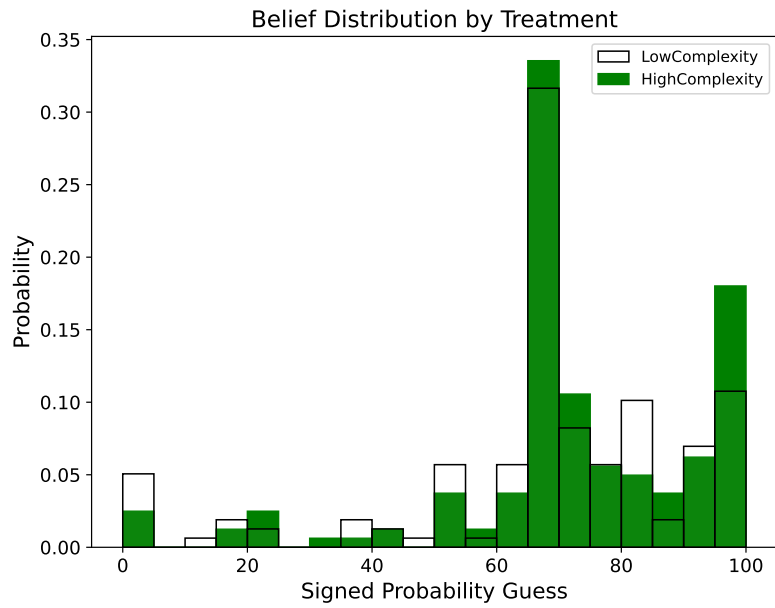


Figure B.4: The relationship between decisions and immediate beliefs in the lenient sample of the Baseline Experiment. *The figure shows a binned scatterplot using the lenient sample of the Baseline Experiment with 319 participants. It has the probability guess that "The CEO is key" is the more likely model on the horizontal axis, and the decision weight γ as defined in Equation 3 on the vertical axis. The lines show LOWESS regressions based on all datapoints.*

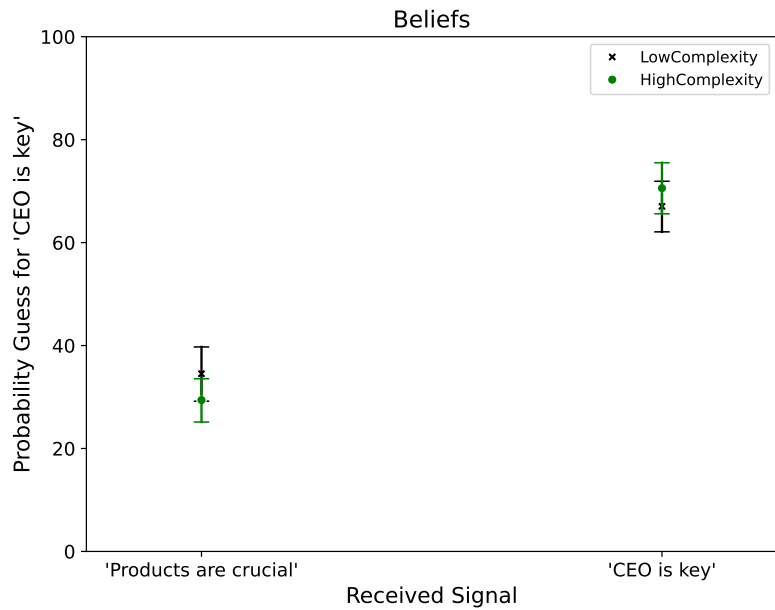
Table B.1: Company Value Guesses in the Lenient Sample of the Baseline Experiment

<i>Dependent variable:</i>	Company Value Guess		
<i>Sample:</i>	LowComplexity (1)	HighComplexity (2)	Pooled (3)
Rational Benchmark	0.658*** (0.070)	0.334*** (0.053)	0.658*** (0.070)
Naive Benchmark	0.432*** (0.060)	0.702*** (0.047)	0.432*** (0.060)
Rational B. $\times$ HighComplexity			-0.323*** (0.088)
Naive B. $\times$ HighComplexity			0.270*** (0.076)
R^2	0.875	0.915	0.894
Observations	1264	1288	2552

The table presents OLS regressions of respondents' company value guesses on the rational and naive benchmarks as detailed in Section 3.1. The table is based on the lenient sample of the Baseline Experiment. Column (1) uses observations from the LowComplexity treatment, column (2) from the HighComplexity treatment, and column (3) from both. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.



(a) Histogram of signed probability guesses.



(b) Mean probability guesses for "The CEO is key".

Figure B.3: Beliefs in the Lenient Sample of the Baseline Experiment. *In Panel (a) beliefs are converted into the direction of the more likely model, so that 65 corresponds to the Bayesian probability. Panel (b) plots mean probability guesses for the model "The CEO is key". The figure is based on the lenient sample of the Baseline Experiment with 319 participants.*

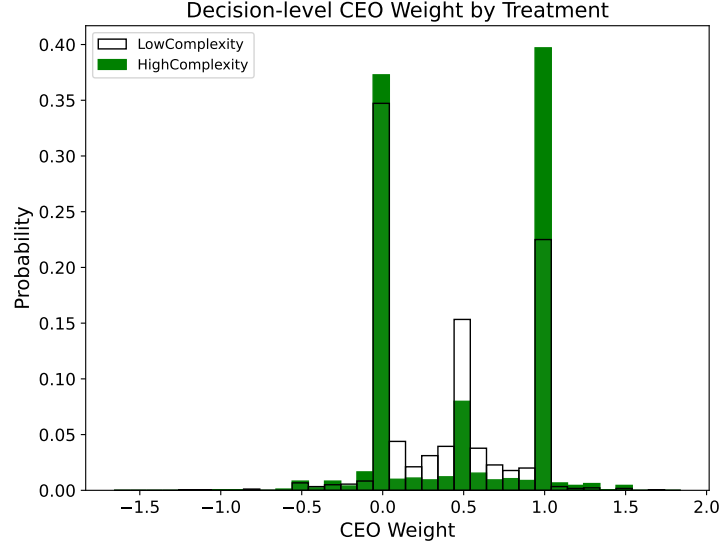
Table B.2: Beliefs and Recall in the Lenient Sample of the Baseline Experiment

<i>Dependent variable:</i>	Probability Guess	Correct Recall
<i>Sample:</i>	Pooled (1)	Pooled (2)
Constant	66.321 ^{***} (1.871)	0.943 ^{***} (0.019)
HighComplexity	4.266 [*] (2.536)	0.020 (0.024)
R^2	0.009	0.002
Observations	319	319

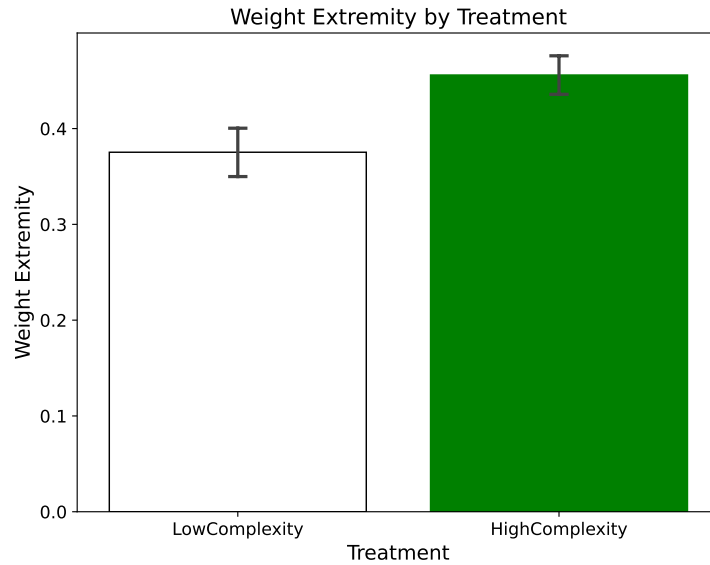
The table presents OLS regressions using the lenient sample of the Baseline Experiment. In column (1), the dependent variable is the probability guess for the likely state. Column (2) uses a dummy for whether respondents correctly recall the more likely model. All columns use observations from both the HighComplexity and Low-Complexity conditions. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

B.2 Equally Likely Models: Results in Lenient Sample

Here, we present the results for the lenient sample of the Equally Likely Models Experiment, featuring 452 participants who solved at least one of the example screens.



(a) Distribution of decision-level CEO weights



(b) Mean decision-level CEO weight extremity

Figure B.5: Decision-level CEO weights and CEO weight extremity in the lenient sample of the Equally Likely Models Experiment. *Panel (a) plots the distribution of CEO weights γ calculated as specified in Equation 3, using the lenient sample of the Equally Likely Models study with 452 participants. Panel (b) plots the average CEO weight extremity $|\gamma - \frac{1}{2}|$ calculated as specified in Equation 4.*

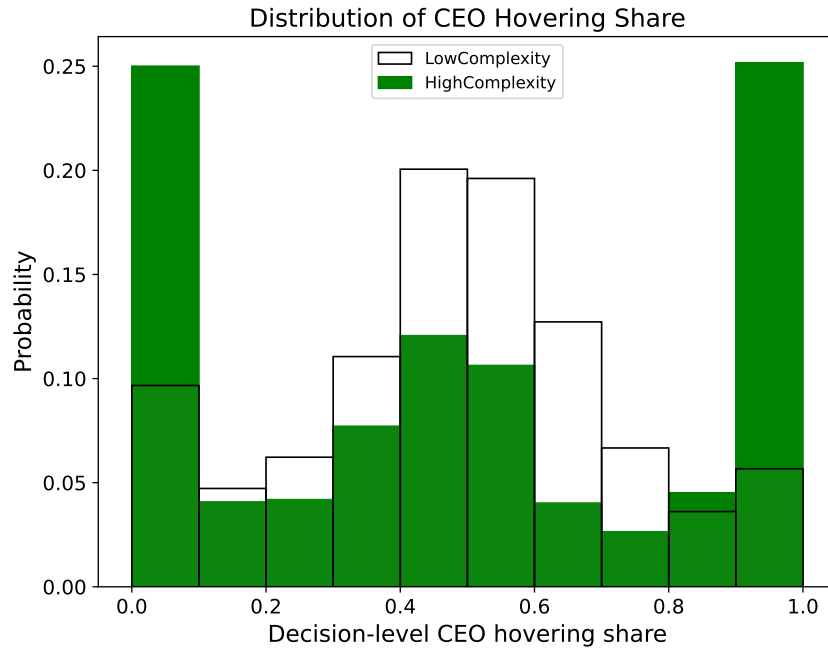


Figure B.6: Distribution of decision-level CEO hovering share in the Equally Likely Models Experiment. *The figure plots the distribution of the share of time that respondents spent looking at the values of the "The CEO is key" model, using the lenient sample of the Equally Likely Models study with 452 participants.*

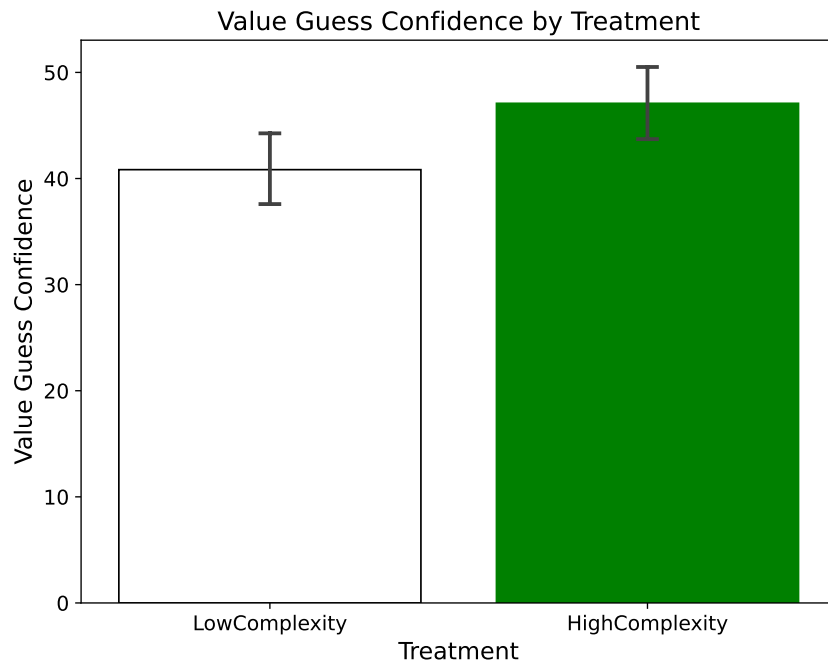


Figure B.7: Average value guess confidence in the lenient sample of the Equally Likely Models Experiment. *The figure plots the average confidence that respondents had in their company value guesses, using the lenient sample of the Equally Likely Models study with 452 participants.*

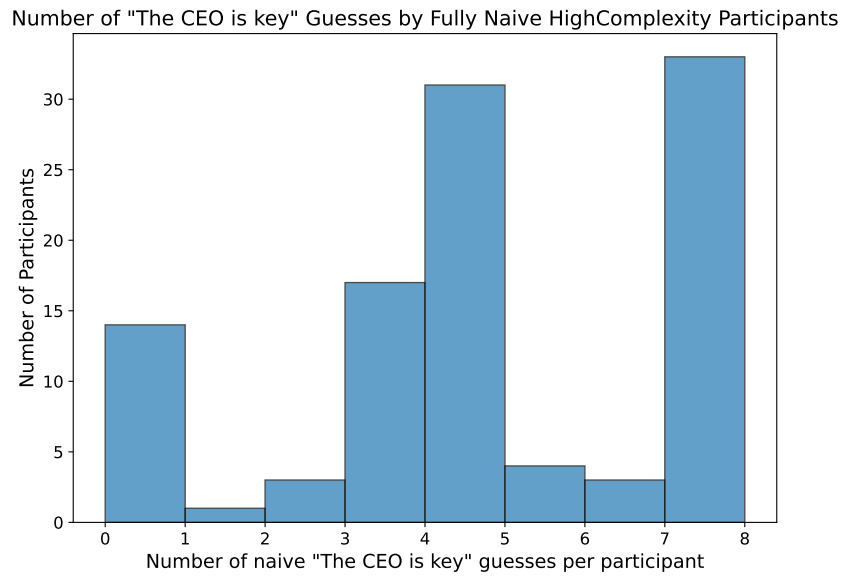
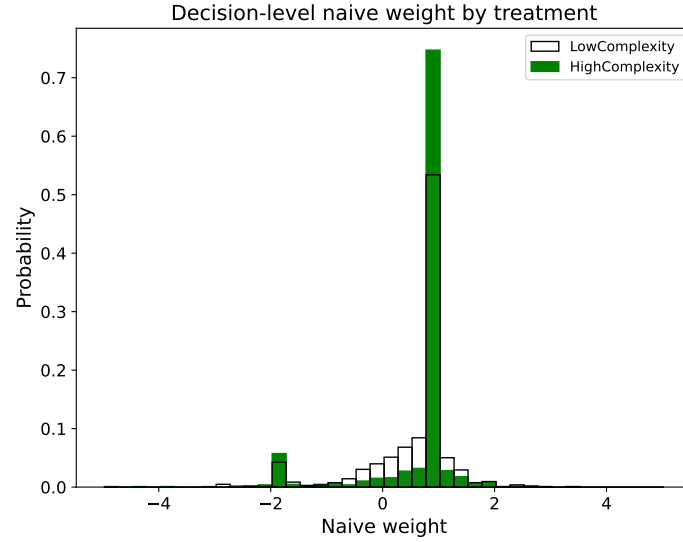


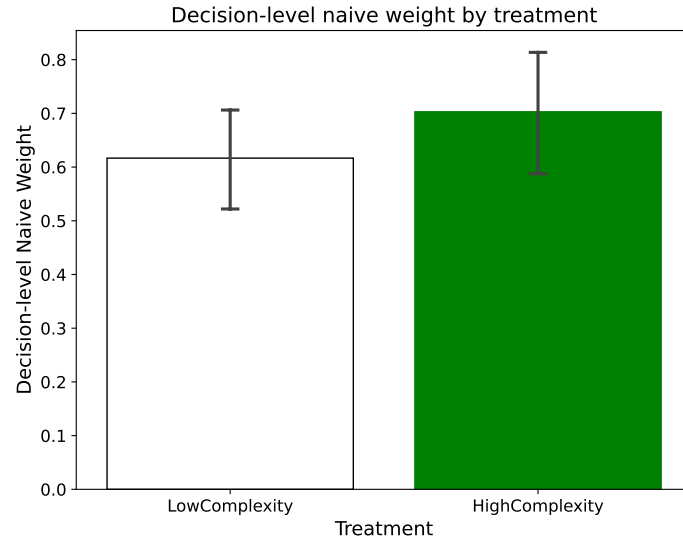
Figure B.8: Distribution of naive decision-making in the lenient sample of the Equally Likely Models Experiment. *The figure plots the distribution of the number of times participants selected the value corresponding to the "The CEO is key" model, using the lenient sample of the HighComplexity treatment of the Equally Likely Models study. The sample is limited to the 106 participants in the HighComplexity condition who made only fully naive guesses, meaning they always selected a value corresponding to either the "The CEO is key" or "Products are crucial" model.*

B.3 Investment Experiment 1: Results in Lenient Sample

Here, we present the results for the lenient sample of the Investment Experiment 1, featuring 265 participants who solved at least one of the example screens.

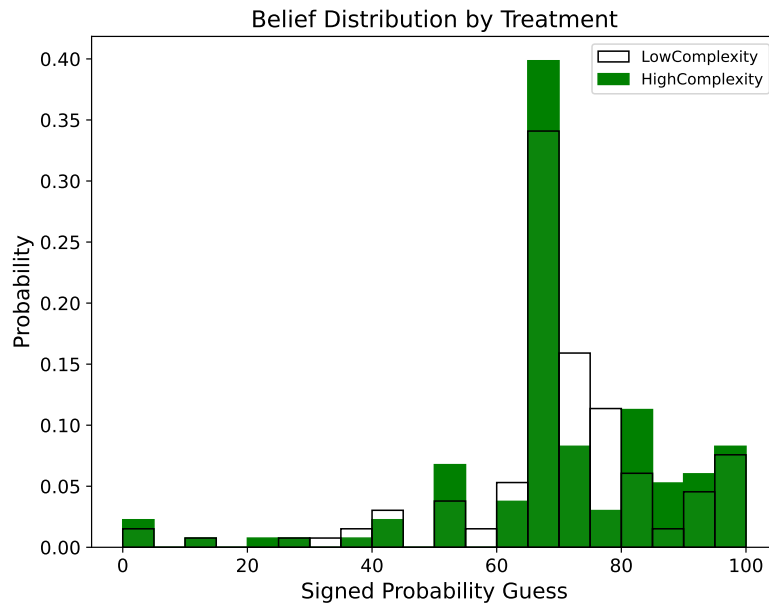


(a) Distribution of decision-level naive weights

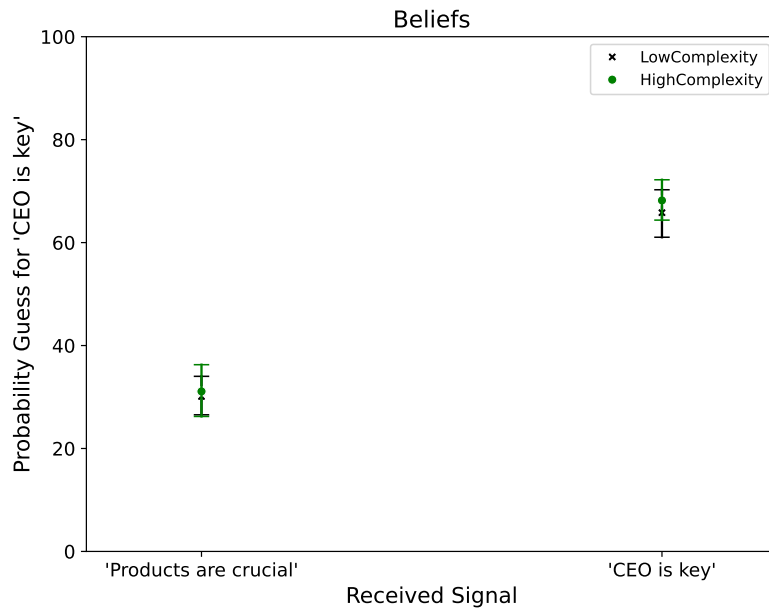


(b) Mean decision-level naive weights

Figure B.9: Decision-level naive weights in the lenient sample of the Investment Experiment 1. Panel (a) plots the distribution of naive weights λ calculated as specified in Equation 2, using the lenient sample of the Investment Experiment 1 with 265 participants. Panel (b) plots average naive weights.



(a) Histogram of signed probability guesses.



(b) Mean probability guesses for "The CEO is key".

Figure B.10: Beliefs in the Lenient Sample of the Investment Experiment 1. In Panel (a) beliefs are converted into the direction of the more likely model, so that 65 corresponds to the Bayesian probability. Panel (b) plots mean probability guesses for the model "The CEO is key". The figure is based on the lenient sample of the Investment Experiment 1 with 265 participants.

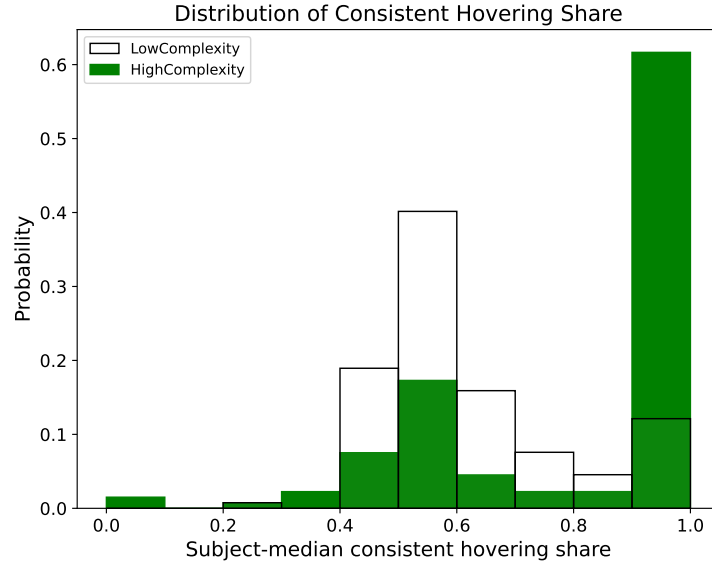


Figure B.11: Distribution of subject-medians of the consistent hovering shares in the lenient sample of the Investment Experiment 1. *The figure plots the distribution of the share of time that respondents spent looking at the values of the signal-consistent model, using the lenient sample of the Investment Experiment 1 with 265 participants. Only the median consistent share for each participant is plotted.*

Table B.3: Company Bids in the Lenient Sample of the Investment Experiment 1

<i>Dependent variable:</i>	Company Bids		
	LowComplexity (1)	HighComplexity (2)	Pooled (3)
<i>Sample:</i>			
Rational Benchmark	0.432*** (0.052)	0.363*** (0.065)	0.432*** (0.052)
Naive Benchmark	0.596*** (0.043)	0.674*** (0.056)	0.596*** (0.043)
Rational B. $\times$ HighComplexity			-0.070 (0.083)
Naive B. $\times$ HighComplexity			0.078 (0.071)
R^2	0.915	0.905	0.910
Observations	1056	1064	2120

The table presents OLS regressions of respondents' company bids on the rational and naive benchmarks as detailed in Section 3.1. The table is based on the lenient sample of the Investment Experiment 1. Column (1) uses observations from the LowComplexity treatment, column (2) from the HighComplexity treatment, and column (3) from both. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

Table B.4: Beliefs and Recall in the Lenient Sample of the Investment Experiment 1

<i>Dependent variable:</i>	Probability Guess	Correct Recall
<i>Sample:</i>	Pooled (1)	Pooled (2)
Constant	67.915*** (1.515)	0.955*** (0.018)
HighComplexity	0.629 (2.237)	0.023 (0.022)
R^2	0.000	0.004
Observations	265	265

The table presents OLS regressions using the lenient sample of the Investment Experiment 1. In column (1), the dependent variable is the probability guess for the likely state. Column (2) uses a dummy for whether respondents correctly recall the more likely model. All columns use observations from both the HighComplexity and LowComplexity conditions. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

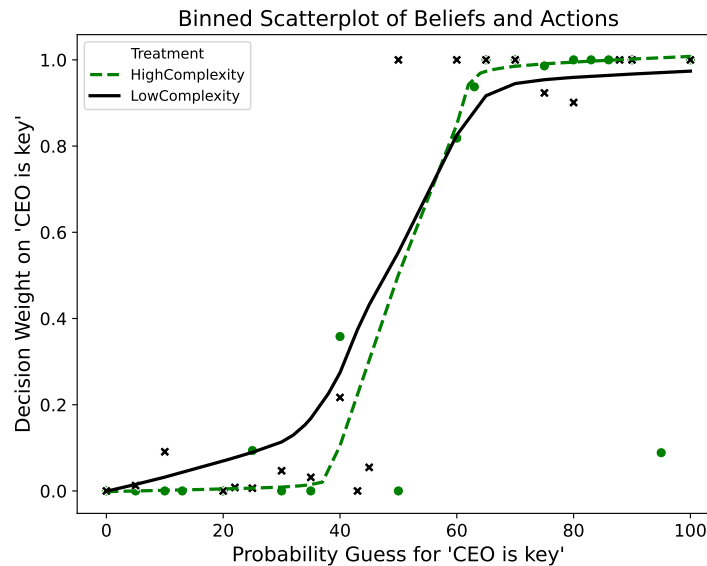
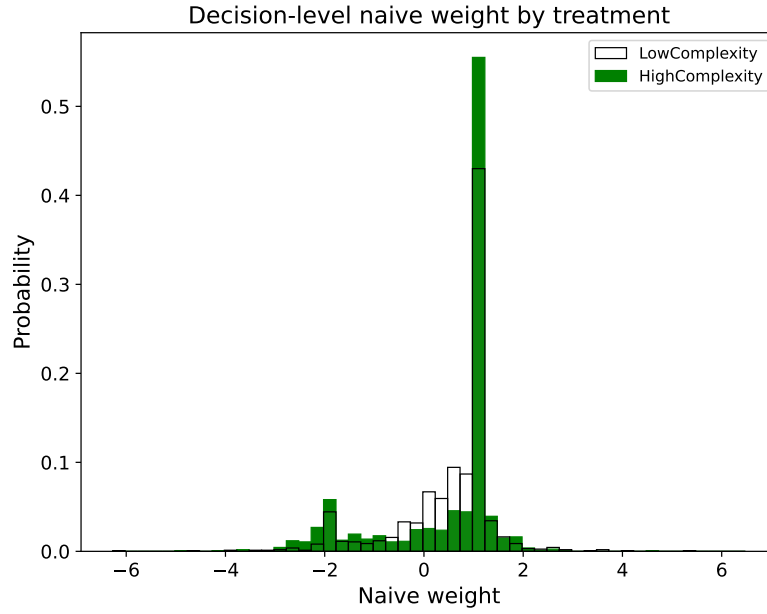


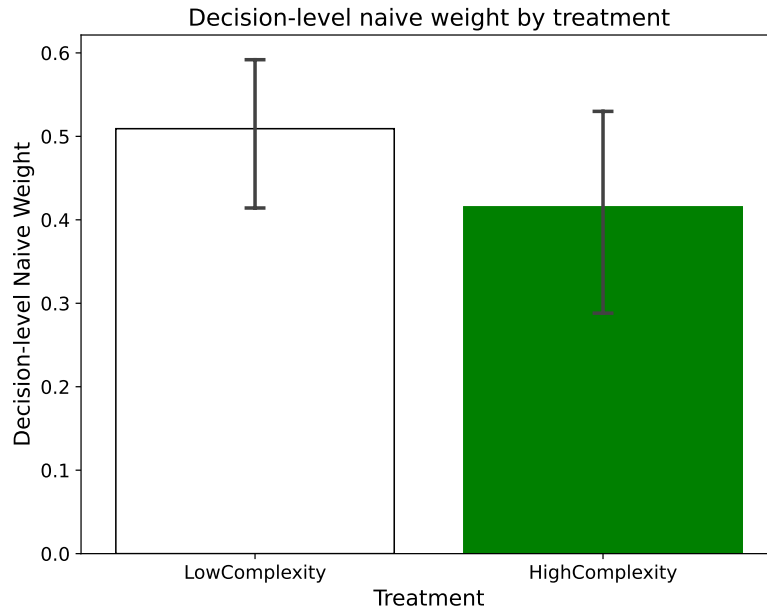
Figure B.12: The relationship between decisions and immediate beliefs in the lenient sample of the Investment Experiment 1. The figure shows a binned scatterplot using the lenient sample of the Investment Experiment 1 with 265 participants. It has the probability guess that "The CEO is key" is the more likely model on the horizontal axis, and the decision weight γ as defined in Equation 3 on the vertical axis. The lines show LOWESS regressions based on all datapoints.

B.4 Investment Experiment 1: Results in Full Sample

Here, we present the results for the full sample of the Investment Experiment 1, featuring all 400 participants.



(a) Distribution of decision-level naive weights



(b) Mean decision-level naive weights

Figure B.13: Decision-level naive weights in the full sample of the Investment Experiment 1. Panel (a) plots the distribution of naive weights λ calculated as specified in Equation 2, using the full sample of the Investment Experiment 1 with 400 participants. Panel (b) plots average naive weights.

Table B.5: Company Bids in the Full Sample of the Investment Experiment 1

<i>Dependent variable:</i>	Company Bids		
	LowComplexity (1)	HighComplexity (2)	Pooled (3)
Rational Benchmark	0.598*** (0.053)	0.551*** (0.057)	0.598*** (0.053)
Naive Benchmark	0.463*** (0.043)	0.423*** (0.055)	0.463*** (0.043)
Rational B. $\times$ HighComplexity			-0.047 (0.078)
Naive B. $\times$ HighComplexity			-0.040 (0.069)
R^2	0.888	0.816	0.854
Observations	1600	1600	3200

The table presents OLS regressions of respondents' company bids on the rational and naive benchmarks as detailed in Section 3.1. The table is based on the full sample of the Investment Experiment 1. Column (1) uses observations from the LowComplexity treatment, column (2) from the HighComplexity treatment, and column (3) from both. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

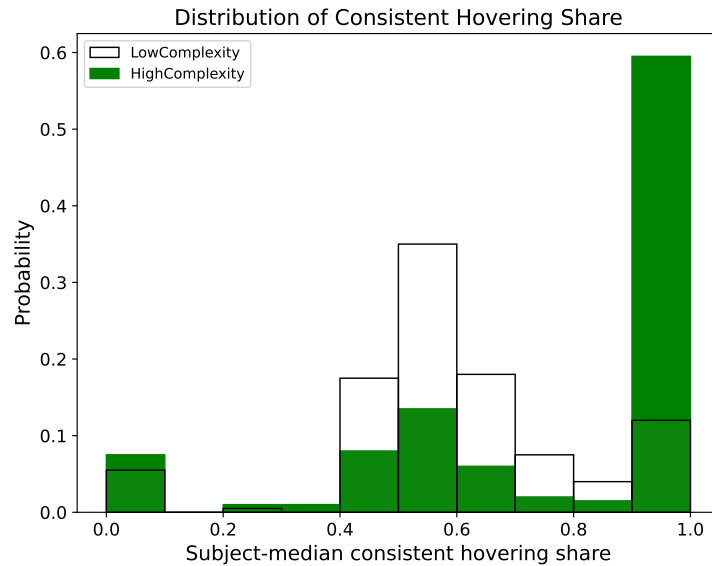
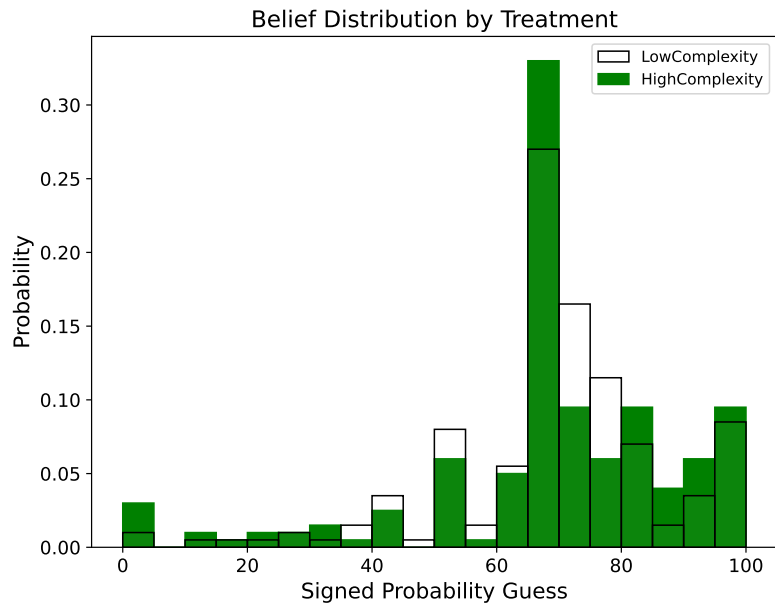
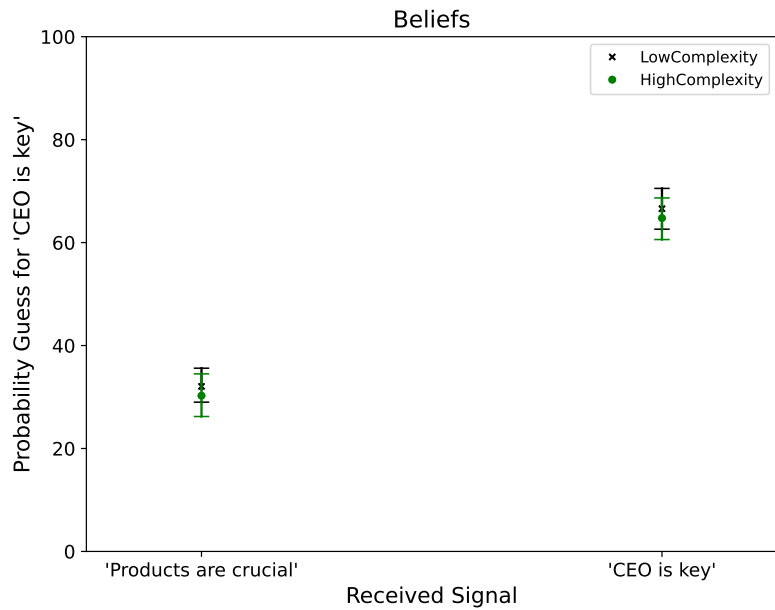


Figure B.14: Distribution of subject-medians of the consistent hovering shares in the full sample of the Investment Experiment 1. The figure plots the distribution of the share of time that respondents spent looking at the values of the signal-consistent model, using the full sample of the Investment Experiment 1 with 400 participants. Only the median consistent share for each participant is plotted.



(a) Histogram of signed probability guesses.



(b) Mean probability guesses for "The CEO is key".

Figure B.15: Beliefs in the full sample of the Investment Experiment 1. In Panel (a) beliefs are converted into the direction of the more likely model, so that 65 corresponds to the Bayesian probability. Panel (b) plots mean probability guesses for the model "The CEO is key". The figure is based on the full sample of the Investment Experiment 1 with 400 participants.

Table B.6: Beliefs and Recall in the Full Sample of the Investment Experiment 1

<i>Dependent variable:</i>	Probability Guess	Correct Recall
<i>Sample:</i>	Pooled (1)	Pooled (2)
Constant	67.319*** (1.259)	0.945*** (0.016)
HighComplexity	-0.092 (1.958)	0.030 (0.020)
R^2	0.000	0.006
Observations	400	400

The table presents OLS regressions using the full sample of the Investment Experiment 1. In column (1), the dependent variable is the probability guess for the likely state. Column (2) uses a dummy for whether respondents correctly recall the more likely model. All columns use observations from both the HighComplexity and LowComplexity conditions. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

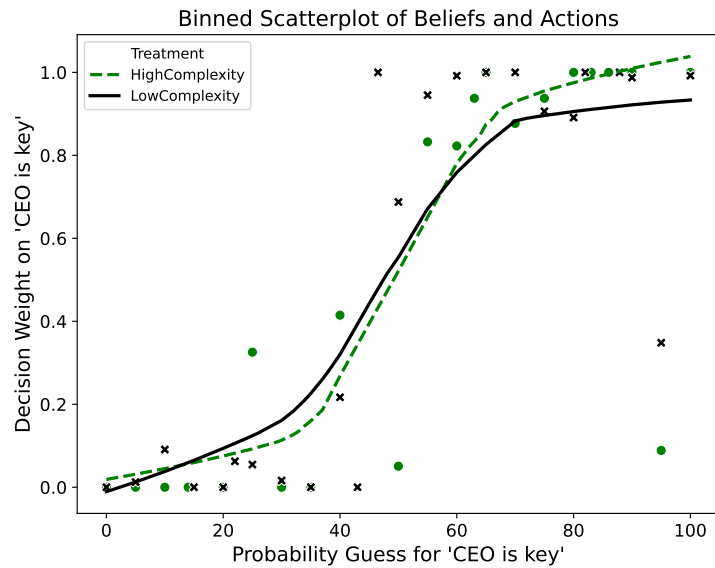
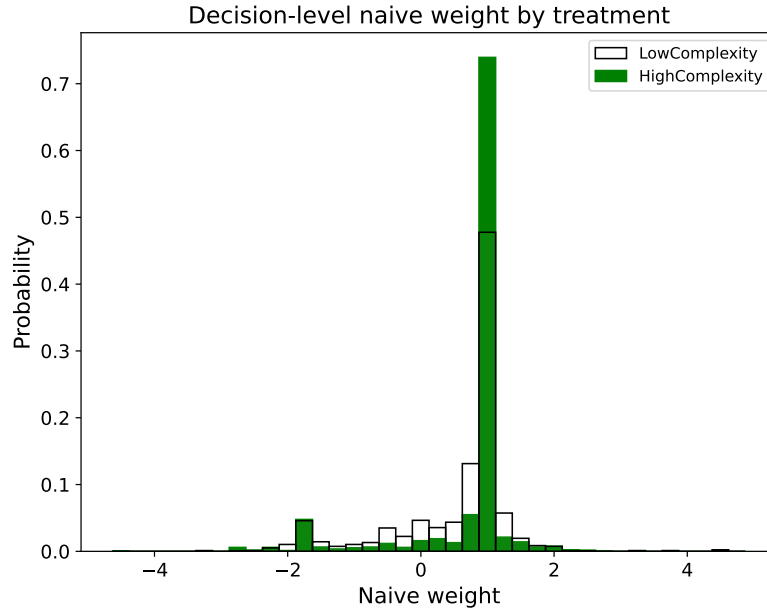


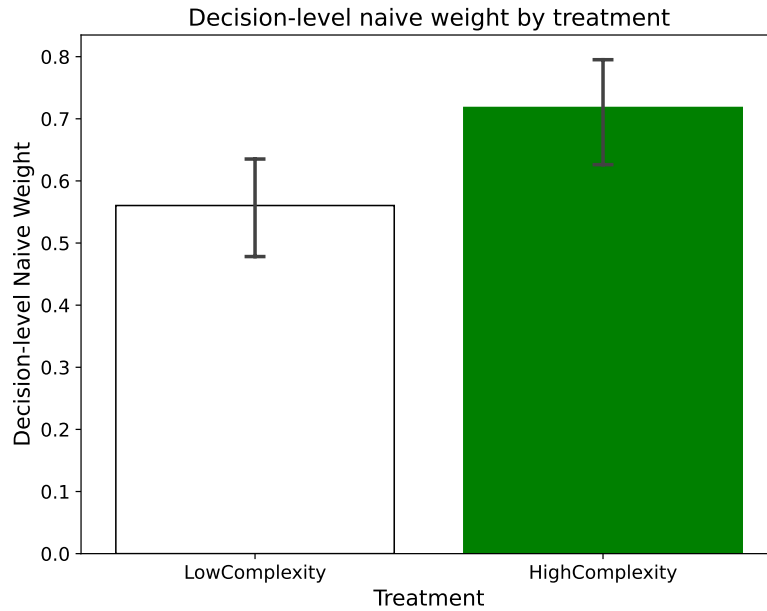
Figure B.16: The relationship between decisions and beliefs in the full sample of the Investment Experiment 1. The figure shows a binned scatterplot using the full sample of the Investment Experiment 1 with 400 participants. It has the probability guess that "The CEO is key" is the more likely model on the horizontal axis, and the decision weight γ as defined in Equation 3 on the vertical axis. The lines show LOWESS regressions based on all datapoints.

B.5 Investment Experiment 2: Results in Lenient Sample

Here, we present the results for the lenient sample of the Investment Experiment 2, featuring 430 participants who solved at least one of the example screens.



(a) Distribution of decision-level naive weights



(b) Mean decision-level naive weights

Figure B.17: Decision-level naive weights in the lenient sample of the Investment Experiment 2. Panel (a) plots the distribution of naive weights λ calculated as specified in Equation 2, using the lenient sample of the Investment Experiment 2 with 430 participants. Panel (b) plots average naive weights.

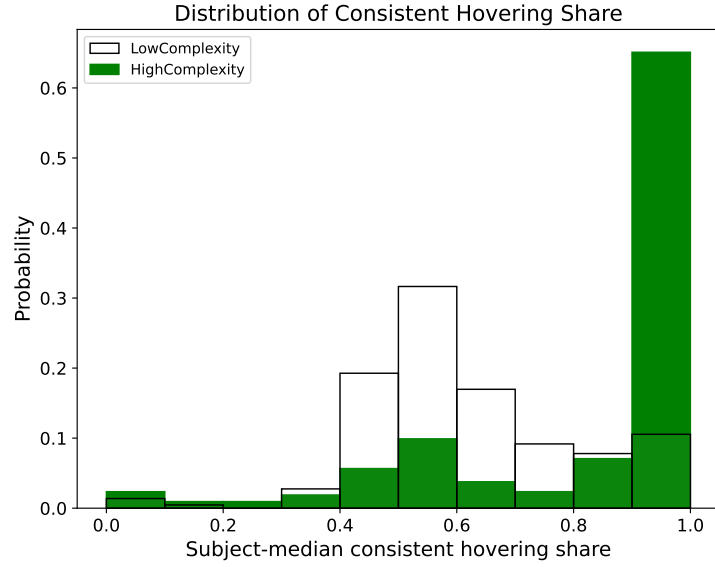
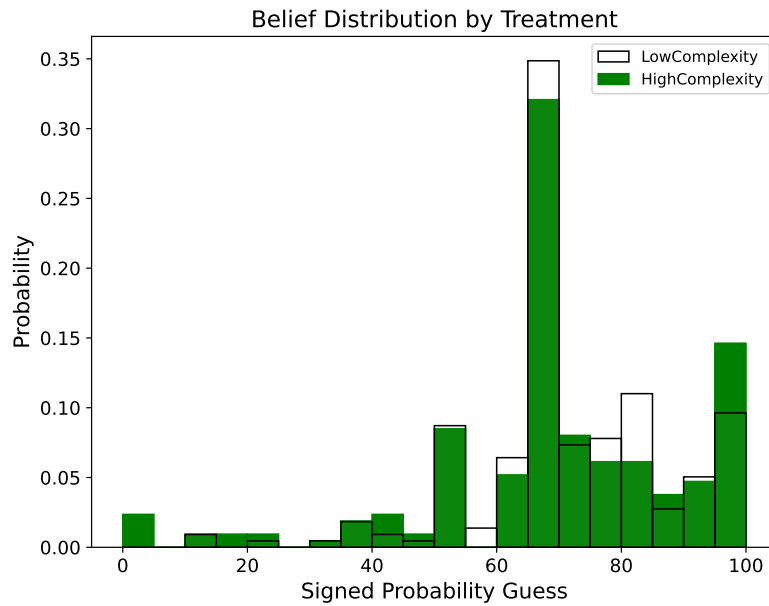


Figure B.18: Distribution of subject-medians of the consistent hovering shares in the lenient sample of the Investment Experiment 2. *The figure plots the distribution of the share of time that respondents spent looking at the values of the signal-consistent model, using the lenient sample of the Investment Experiment 2 with 430 participants. Only the median consistent share for each participant is plotted.*

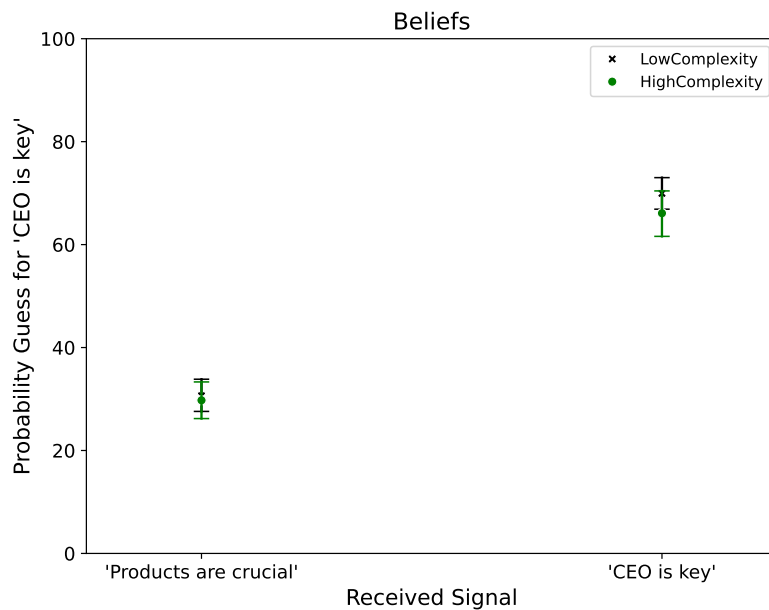
Table B.7: Company Bids in the Lenient Sample of the Investment Experiment 2

<i>Dependent variable:</i>	Company Bids		
	LowComplexity (1)	HighComplexity (2)	Pooled (3)
Rational Benchmark	0.500*** (0.045)	0.283*** (0.043)	0.500*** (0.045)
Naive Benchmark	0.532*** (0.038)	0.715*** (0.039)	0.532*** (0.038)
Rational B. $\times$ HighComplexity			-0.218*** (0.062)
Naive B. $\times$ HighComplexity			0.183*** (0.054)
R^2	0.891	0.904	0.898
Observations	1744	1696	3440

The table presents OLS regressions of respondents' company bids on the rational and naive benchmarks as detailed in Section 3.1. The table is based on the lenient sample of the Investment Experiment 2. Column (1) uses observations from the LowComplexity treatment, column (2) from the HighComplexity treatment, and column (3) from both. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.



(a) Histogram of signed probability guesses.



(b) Mean probability guesses for "The CEO is key".

Figure B.19: Beliefs in the lenient sample of the Investment Experiment 2. In Panel (a) beliefs are converted into the direction of the more likely model, so that 65 corresponds to the Bayesian probability. Panel (b) plots mean probability guesses for the model "The CEO is key". The figure is based on the lenient sample of the Investment Experiment 2 with 430 participants.

Table B.8: Beliefs and Recall in the Lenient Sample of the Investment Experiment 2

<i>Dependent variable:</i>	Probability Guess	Correct Recall
<i>Sample:</i>	Pooled (1)	Pooled (2)
Constant	69.656 ^{***} (1.096)	0.973 ^{***} (0.011)
HighComplexity	-1.464 (1.843)	-0.020 (0.018)
R^2	0.002	0.003
Observations	430	430

The table presents OLS regressions using the lenient sample of the Investment Experiment 2. In column (1), the dependent variable is the probability guess for the likely state. Column (2) uses a dummy for whether respondents correctly recall the more likely model. All columns use observations from both the HighComplexity and LowComplexity conditions. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

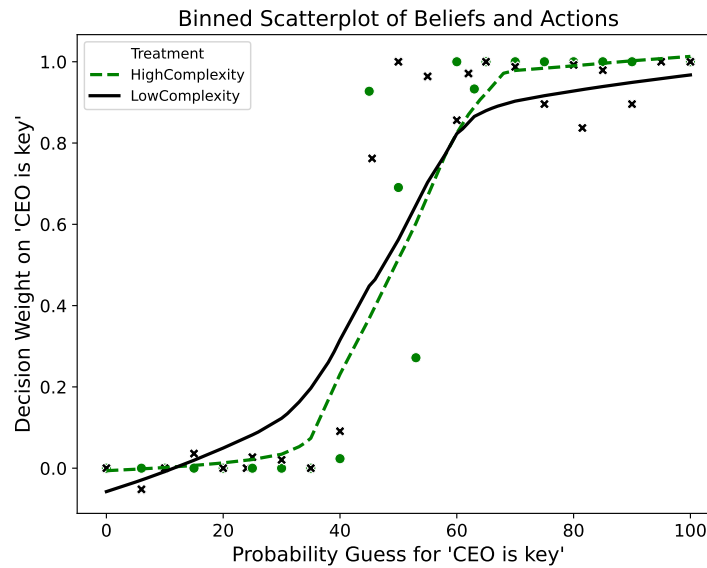
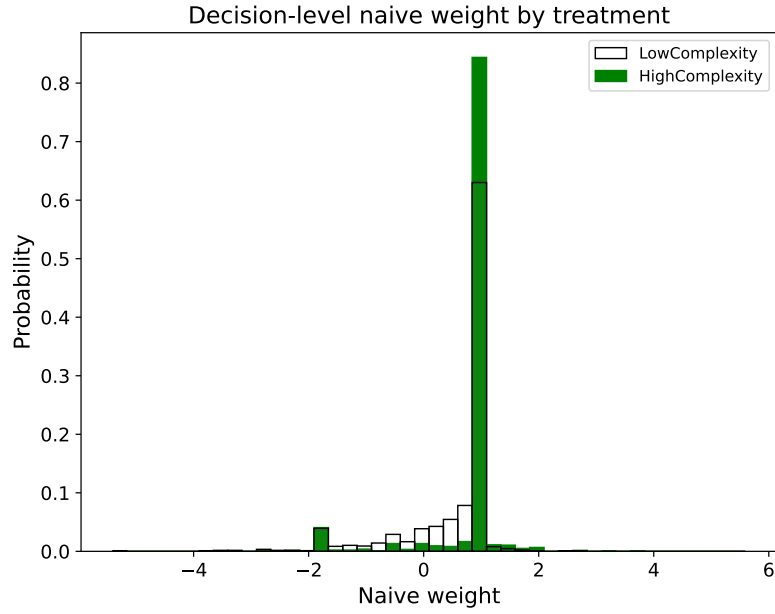


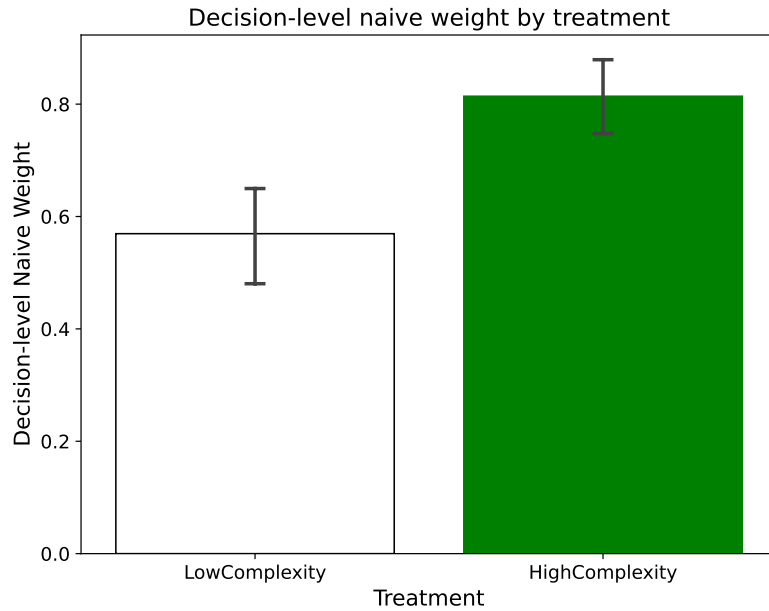
Figure B.20: The relationship between decisions and beliefs in the lenient sample of the Investment Experiment 2. The figure shows a binned scatterplot using the lenient sample of the Investment Experiment 2 with 430 participants. It has the probability guess that "The CEO is key" is the more likely model on the horizontal axis, and the decision weight γ as defined in Equation 3 on the vertical axis. The lines show LOWESS regressions based on all datapoints.

B.6 Baseline Confidence: Results in Lenient Sample

Here, we present the results for the lenient sample of the Baseline Confidence Experiment, featuring 445 participants who solved at least one of the example screens.



(a) Distribution of decision-level naive weights



(b) Mean decision-level naive weights

Figure B.21: Decision-level naive weights in the lenient sample of the Baseline Confidence Experiment. Panel (a) plots the distribution of naive weights λ calculated as specified in Equation 2, using the lenient sample of the Baseline Confidence Experiment with 445 participants. Panel (b) plots average naive weights.

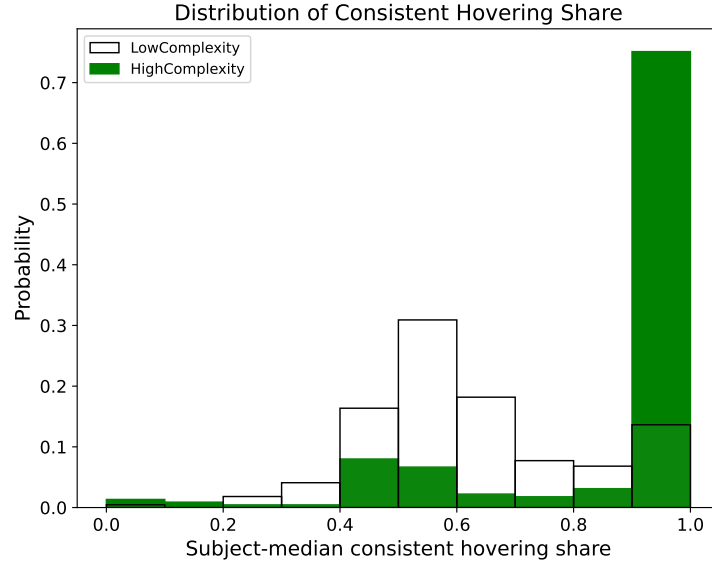
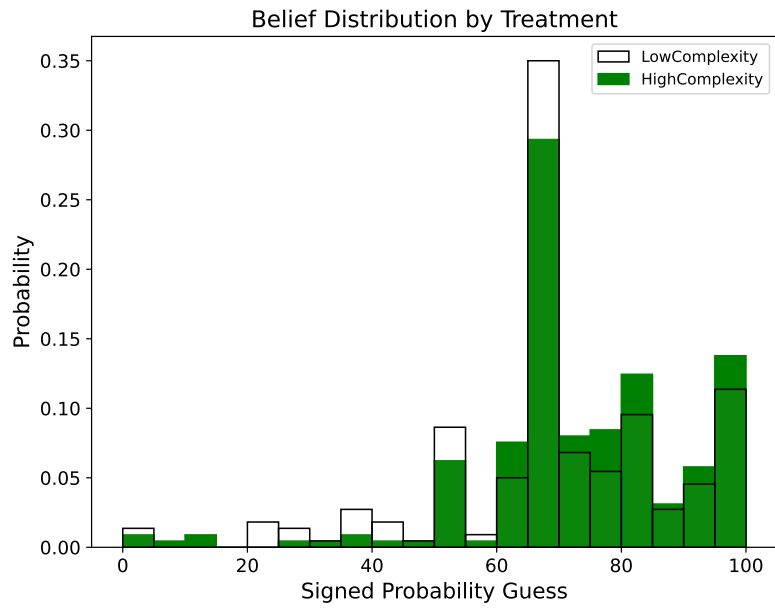


Figure B.22: Distribution of subject-medians of the consistent hovering shares in the lenient sample of the Baseline Confidence Experiment. *The figure plots the distribution of the share of time that respondents spent looking at the values of the signal-consistent model, using the lenient sample of the Baseline Confidence Experiment with 445 participants. Only the median consistent share for each participant is plotted.*

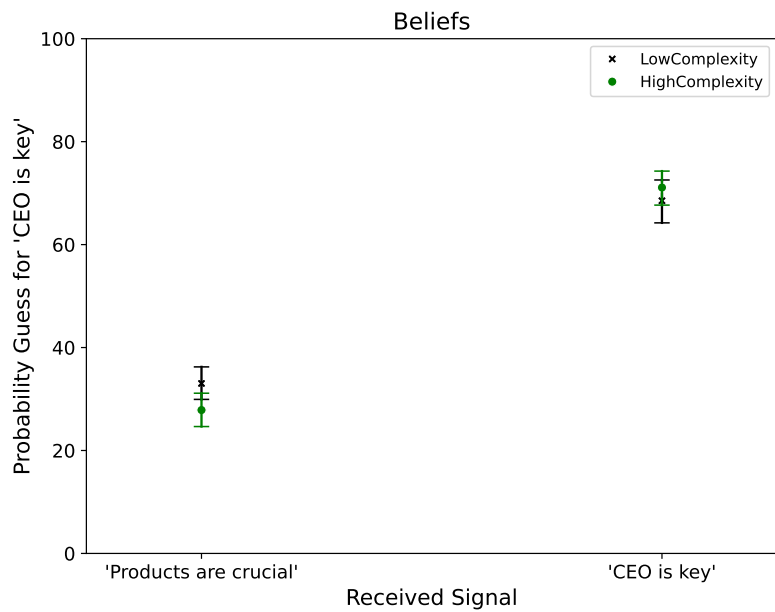
Table B.9: Company Value Guesses in the Lenient Sample of Baseline Confidence Experiment

Dependent variable:	Company Value Guess		
	LowComplexity (1)	HighComplexity (2)	Pooled (3)
Rational Benchmark	0.490*** (0.048)	0.227*** (0.037)	0.490*** (0.048)
Naive Benchmark	0.560*** (0.040)	0.793*** (0.033)	0.560*** (0.040)
Rational B. $\times$ HighComplexity			-0.263*** (0.060)
Naive B. $\times$ HighComplexity			0.233*** (0.052)
R^2	0.912	0.932	0.922
Observations	1760	1800	3560

The table presents OLS regressions of respondents' company value guesses on the rational and naive benchmarks as detailed in Section 3.1. The table is based on the lenient sample of the Baseline Confidence Experiment. Column (1) uses observations from the LowComplexity treatment, column (2) from the HighComplexity treatment, and column (3) from both. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.



(a) Histogram of signed probability guesses.



(b) Mean probability guesses for "The CEO is key".

Figure B.23: Beliefs in the lenient sample of the Baseline Confidence Experiment. In Panel (a) beliefs are converted into the direction of the more likely model, so that 65 corresponds to the Bayesian probability. Panel (b) plots mean probability guesses for the model "The CEO is key". The figure is based on the lenient sample of the Baseline Confidence Experiment with 445 participants.

Table B.10: Beliefs and Recall in the Lenient Sample of the Baseline Confidence Experiment

<i>Dependent variable:</i>	Probability Guess	Correct Recall
<i>Sample:</i>	Pooled (1)	Pooled (2)
Constant	67.759 ^{***} (1.315)	0.936 ^{***} (0.017)
HighComplexity	3.814 ^{**} (1.787)	0.033 (0.020)
R^2	0.010	0.006
Observations	445	445

The table presents OLS regressions using the lenient sample of the Baseline Confidence Experiment. In column (1), the dependent variable is the probability guess for the likely state. Column (2) uses a dummy for whether respondents correctly recall the more likely model. All columns use observations from both the HighComplexity and LowComplexity conditions. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

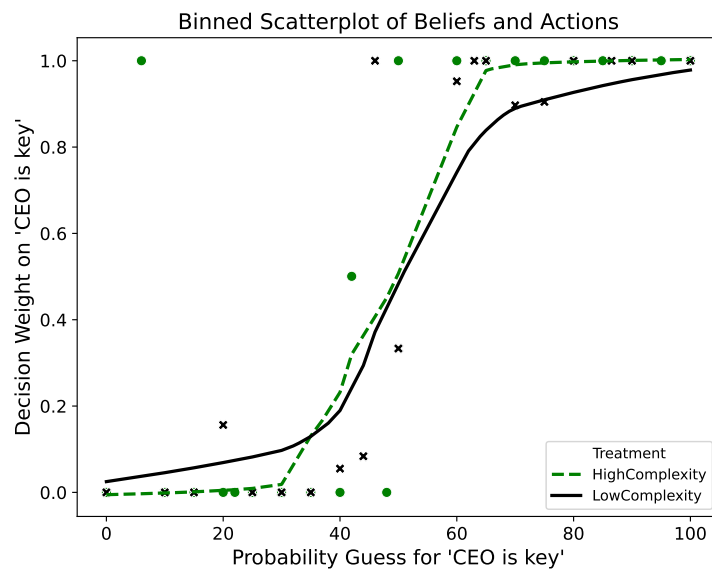


Figure B.24: The relationship between decisions and beliefs in the lenient sample of the Baseline Confidence Experiment. The figure shows a binned scatterplot using the lenient sample of the Baseline Confidence Experiment with 336 participants. It has the probability guess that "The CEO is key" is the more likely model on the horizontal axis, and the decision weight γ as defined in Equation 3 on the vertical axis. The lines show LOWESS regressions based on all datapoints.

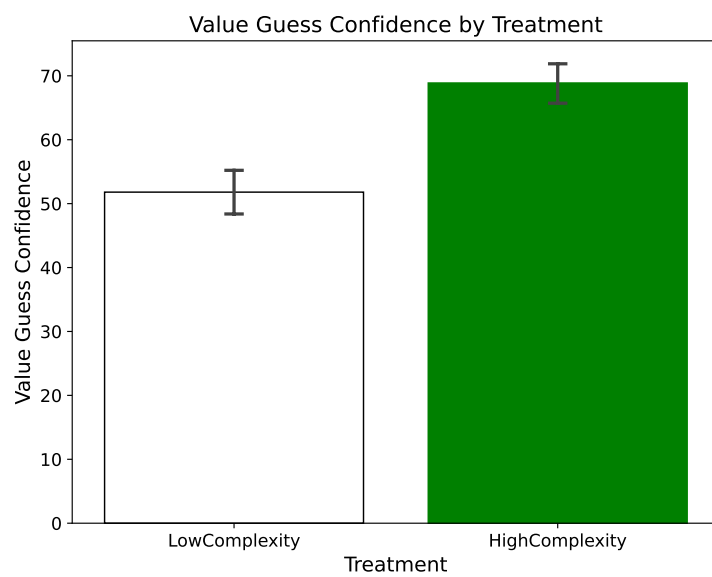


Figure B.25: Average value guess confidence in the lenient sample of the Baseline Confidence Experiment. *The figure plots the average confidence that respondents had in their company value guesses, using the lenient sample of the Baseline Confidence study with 445 participants.*

B.7 Incentivized Confidence: Results in Lenient Sample

Here, we present the results for the lenient sample of the Incentivized Confidence Experiment, featuring 281 participants who solved at least one of the example screens.

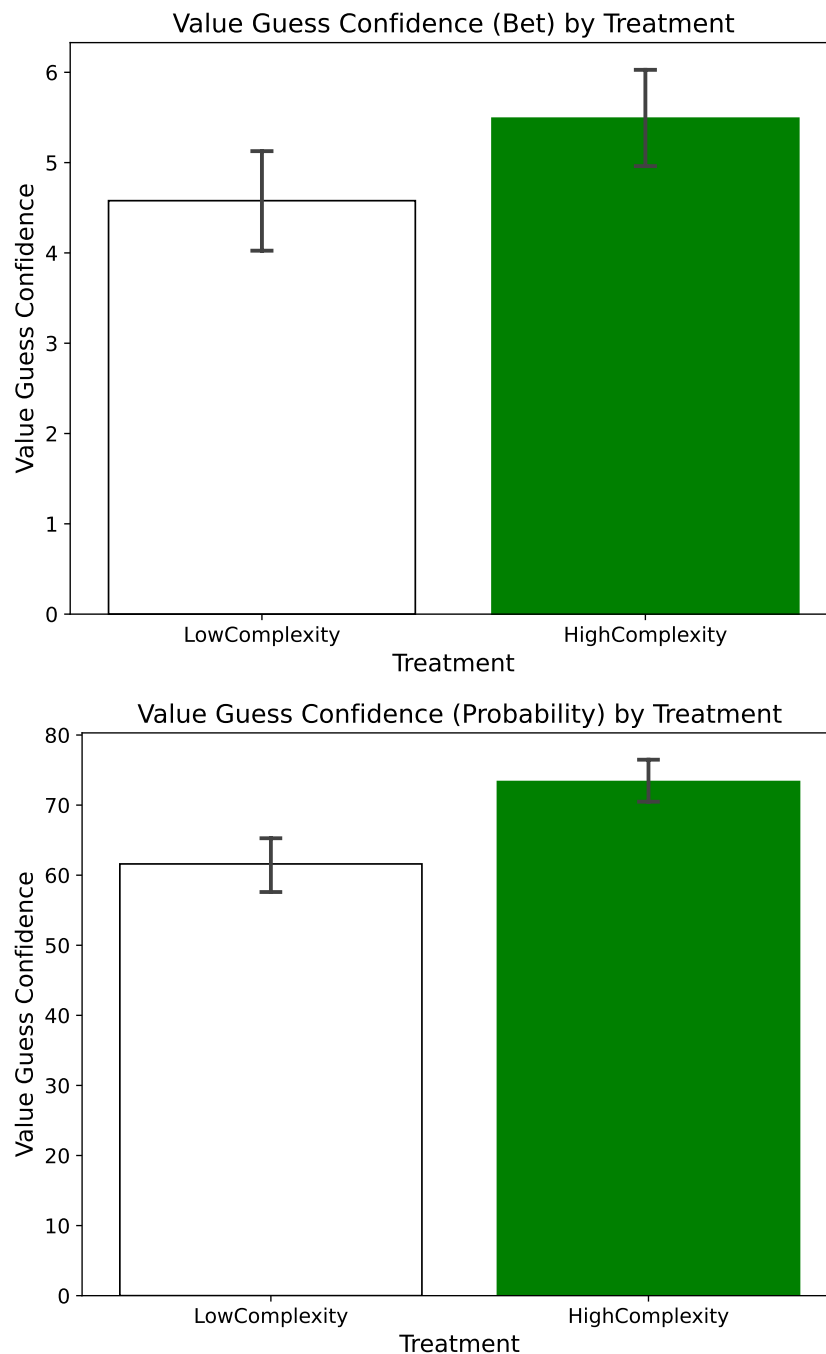
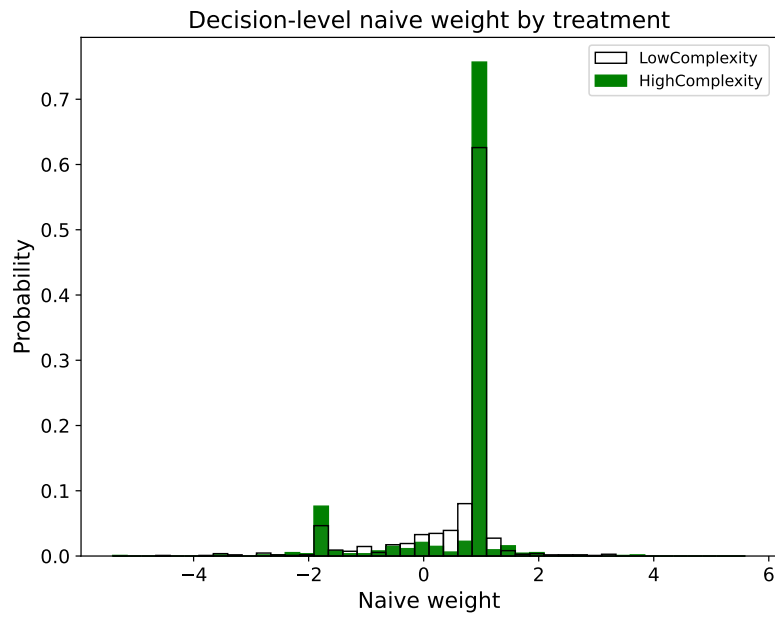
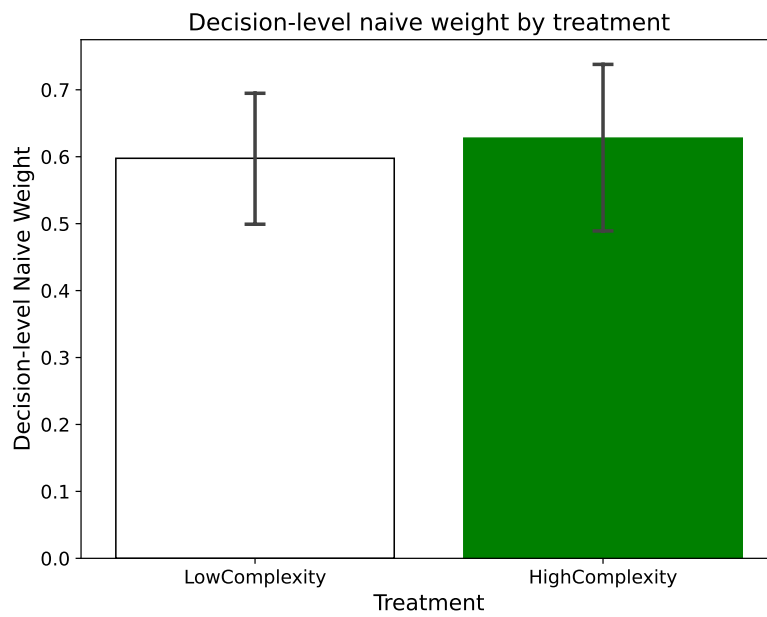


Figure B.26: Average value guess confidence in the lenient sample of the Incentivized Confidence Experiment. *The top figure plots the average incentivized confidence measure, while the bottom figure plots the non-incentivized measure, both using the lenient sample of the Incentivized Confidence study with 281 participants.*



(a) Distribution of decision-level naive weights



(b) Mean decision-level naive weights

Figure B.27: Decision-level naive weights in the lenient sample of the Incentivized Confidence Experiment. Panel (a) plots the distribution of naive weights λ calculated as specified in Equation 2, using the lenient sample of the Incentivized Confidence Experiment with 281 participants. Panel (b) plots average naive weights.

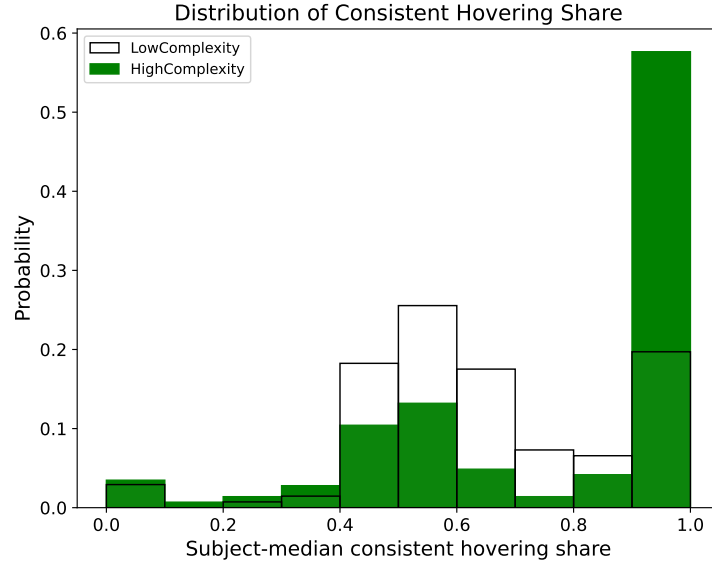
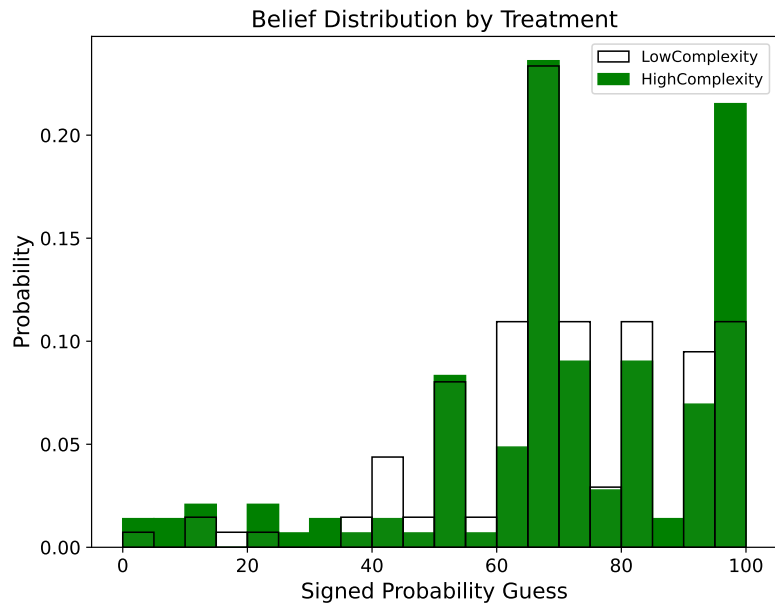


Figure B.28: Distribution of subject-medians of the consistent hovering shares in the lenient sample of the Incentivized Confidence Experiment. *The figure plots the distribution of the share of time that respondents spent looking at the values of the signal-consistent model, using the lenient sample of the Incentivized Confidence Experiment with 281 participants. Only the median consistent share for each participant is plotted.*

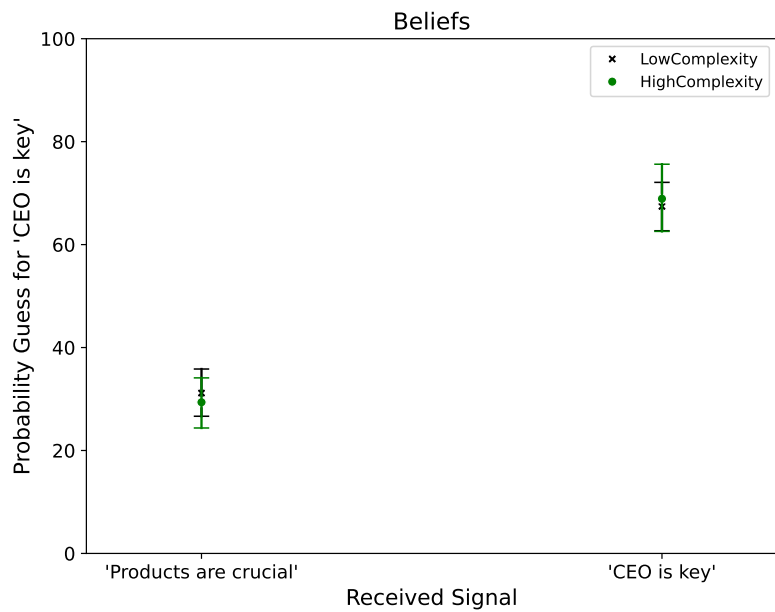
Table B.11: Company Value Guesses in the Lenient Sample of the Incentivized Confidence Experiment

<i>Dependent variable:</i>	Company Value Guess		
	LowComplexity (1)	HighComplexity (2)	Pooled (3)
Rational Benchmark	0.518*** (0.063)	0.451*** (0.071)	0.518*** (0.062)
Naive Benchmark	0.564*** (0.050)	0.591*** (0.062)	0.564*** (0.050)
Rational B. $\times$ HighComplexity			-0.067 (0.094)
Naive B. $\times$ HighComplexity			0.027 (0.079)
R^2	0.909	0.885	0.897
Observations	1096	1152	2248

The table presents OLS regressions of respondents' company value guesses on the rational and naive benchmarks as detailed in Section 3.1. The table is based on the lenient sample of the Incentivized Confidence Experiment. Column (1) uses observations from the LowComplexity treatment, column (2) from the HighComplexity treatment, and column (3) from both. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.



(a) Histogram of signed probability guesses.



(b) Mean probability guesses for "The CEO is key".

Figure B.29: Beliefs in the lenient sample of the Incentivized Confidence Experiment. In Panel (a) beliefs are converted into the direction of the more likely model, so that 65 corresponds to the Bayesian probability. Panel (b) plots mean probability guesses for the model "The CEO is key". The figure is based on the lenient sample of the Incentivized Confidence Experiment with 281 participants.

Table B.12: Beliefs and Recall in the Lenient Sample of the Incentivized Confidence Experiment

<i>Dependent variable:</i>	Probability Guess	Correct Recall
<i>Sample:</i>	Pooled (1)	Pooled (2)
Constant	68.139 ^{***} (1.676)	0.964 ^{***} (0.016)
HighComplexity	1.714 (2.626)	0.002 (0.022)
R^2	0.002	0.000
Observations	281	281

The table presents OLS regressions using the lenient sample of the Incentivized Confidence Experiment. In column (1), the dependent variable is the probability guess for the likely state. Column (2) uses a dummy for whether respondents correctly recall the more likely model. All columns use observations from both the HighComplexity and LowComplexity conditions. Stars highlight significant differences from 0 with * for $p < 0.10$, ** for $p < 0.05$, *** for $p < 0.01$. Standard errors are clustered on the subject level.

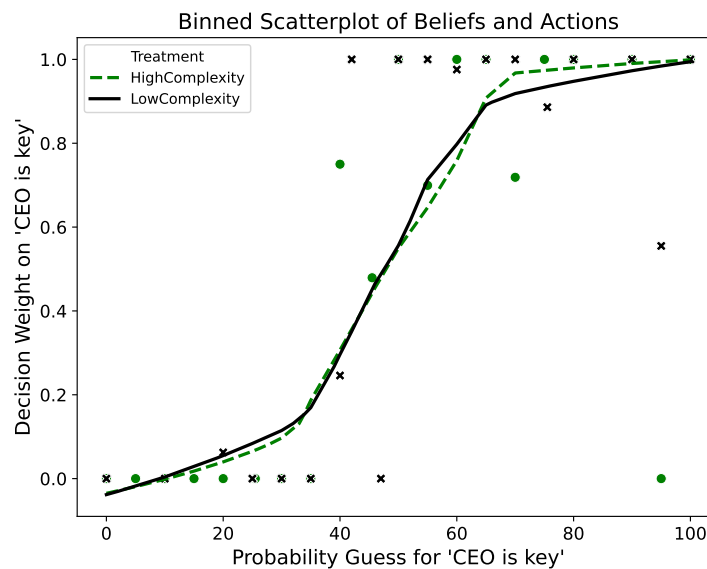


Figure B.30: The relationship between decisions and beliefs in the lenient sample of the Incentivized Confidence Experiment. The figure shows a binned scatterplot using the lenient sample of the Incentivized Confidence Experiment with 281 participants. It has the probability guess that "The CEO is key" is the more likely model on the horizontal axis, and the decision weight γ as defined in Equation 3 on the vertical axis. The lines show LOWESS regressions based on all datapoints.

C Model Framework

We present a simple model that can generate the key results of our experiments. The model is an augmented and simplified version of Bordalo, Gennaioli, Lanzani, and Shleifer (2025) and formalizes the construction of a mental representation of the decision problem. This process is shaped both by bottom-up and top-down attention. Upon presentation of a decision-problem, a bottom-up process of cue-dependent memory determines which of the currently stored mental representation is top of mind. Then, in a top-down process, the agent decides whether they want to further simplify this representation.

C.1 Setup

The model closely follows our experimental environment. There are two models of the world that can explain how company values are determined, i.e. $m \in \{A, B\}$. While only one model is correct, ex-ante there is a 50-50 chance of each model being correct. After observing a noisy but informative signal, the Bayesian posterior is given by $\pi = \Pr(m = A) = 0.65$, i.e. we assume without loss of generality that A is the model that is more likely to be correct.

C.2 Representations

When faced with a decision (action, confidence or belief elicitation), the agent forms a mental representation of decision problem. We assume that at any moment t the agent holds a database (a finite set) of representations

$$\mathbf{R}_t = \{ r_i = (\hat{\pi}(r_i), V_{r_i}^{\text{context}}) \}$$

The scalar $\hat{\pi}(r_i)$ stores the probability attached to model A , while $V_{r_i}^{\text{context}}$ collects the contextual features of the environment in which r was formed, including the type of decision (action, confidence or belief elicitation).

Given a decision problem, the agent forms a mental representation, which we model as a two stage process. The first stage is based on bottom-up attention and similarity based recall (Bordalo, Gennaioli, Lanzani, and Shleifer (2025)), while the second stage is a top-down attention decision of whether or not to simplify the representation that

was recalled in the first stage.

First Stage: Bottom-Up Similarity Based Recall. When a decision screen with cue vector ξ is displayed, the agent retrieves the stored representations based on the similarity of the contextual features V_r^{context} and ξ . In particular, the recall probability $p(r)$ for a given representation is given by the relative similarity of its contextual features to the cue:

$$p(r) = \frac{S(\xi, V_r^{\text{context}})}{\sum_{r' \in \mathbf{R}_t} S(\xi, V_{r'}^{\text{context}})}, \quad (5)$$

where $S(\cdot, \cdot)$ is the similarity kernel and $\mathbf{R}_t$ is the current set of stored representations.

Second Stage: Top-Down Simplification Decision. After retrieving a non-trivial representation r with $0.5 \leq \hat{\pi}(r) < 1$ the agent may keep it as it is or collapse it into a simplified version that sets $\hat{\pi} = 1$.

Optimal Action Given a Posterior π . For the case of actions (value estimates) let us first characterize the optimal estimate given some fixed representation. Under the binarised scoring rule with prize P the agent's utility is

$$u_m(a) = (1 - (\frac{a_m}{100} - \frac{a}{100})^2) \times u(P),$$

when the true model is $m \in \{A, B\}$ and a_m is the company value guess provided by model m . Hence the expected utility from guess a is

$$\mathbb{E}_{m \sim \pi}[u_m(a)] = (1 - \frac{1}{10000} [\pi(a_A - a)^2 + (1 - \pi)(a_B - a)^2]) \times u(P).$$

Because P and the constant are irrelevant for the maximisation, the agent chooses the a that minimises $\pi(a_A - a)^2 + (1 - \pi)(a_B - a)^2$, yielding

$$a^*(\pi) = \pi a_A + (1 - \pi) a_B.$$

Complexity and Cognitive Cost. Decisions differ in their level of computational complexity $c \in \{\emptyset, \text{low}, \text{high}\}$. Here, $\emptyset$ means negligible complexity. We set $c = \text{high}$ for company-value guesses in our high-complexity treatment, $c = \text{low}$ for company-value

guesses in the low-complexity treatment, and $c = \emptyset$ for the confidence and belief elicitations. Recall that we label the models so that the relevant posterior always satisfies $\pi \geq 0.5$ (i.e. model A is the more likely one), so that simplifying means collapsing to full certainty $\pi = 1$. Accordingly, we specify the cognitive cost of acting on belief π under complexity level c as

$$K(c, \pi) = (1 - \pi)\kappa_c$$

where we assume

$$\kappa_{\emptyset} = 0 < \kappa_{\text{low}} < \kappa_{\text{high}} \quad \text{and} \quad \kappa_{\text{low}} < \frac{\Delta U(0.65)}{1 - 0.65} < \kappa_{\text{high}},$$

where

$$\Delta U(\pi) = \mathbb{E}_{m \sim \pi}[u_m(a^*(\pi))] - \mathbb{E}_{m \sim \pi}[u_m(a^*(1))]$$

is the utility cost of simplifying the representation due to loss of precision in the company value guess if the belief for model A is π . Thus, for company-value guesses, a broad representation is more costly than a simplified one, with a larger gap under high than low complexity, while for confidence and belief screens with $c = \emptyset$ no cost difference arises.

Utility Comparison for Simplification. Given a stored representation r with implied belief $\hat{\pi}(r)$ and complexity level c , maintaining the broad representation yields

$$\mathbb{E}_{m \sim \hat{\pi}(r)}[u_m(a^*(\hat{\pi}(r)))] - K(c, \hat{\pi}(r)),$$

whereas collapsing to certainty ($\pi = 1$) yields

$$\mathbb{E}_{m \sim \hat{\pi}(r)}[u_m(a^*(1))] - K(c, 1).$$

Hence, the agent simplifies if and only if

$$K(c, \hat{\pi}(r)) - K(c, 1) > \mathbb{E}_{m \sim \hat{\pi}(r)}[u_m(a^*(\hat{\pi}(r)))] - \mathbb{E}_{m \sim \hat{\pi}(r)}[u_m(a^*(1))] = \Delta U(\hat{\pi}(r)),$$

i.e. the cognitive-cost saving $K(c, \hat{\pi}(r)) - K(c, 1)$ exceeds the expected-utility loss $\Delta U(\hat{\pi}(r))$. Because $K(\emptyset, \hat{\pi}(r)) - K(\emptyset, 1) = 0$, simplifying on the confidence and belief screens ($c = \emptyset$)

would reduce expected utility without any cost-saving. Hence the agent never simplifies under $c = \emptyset$.

Using and Storing Representations. Once a representation has been formed and used to guide a decision, it is added to the database of representations. $V_r^{context}$ of such a representation contains the contextual features of the decision problem in which it was used.

C.3 Timeline of the Experiment

Before stating our predictions, let us briefly recap the timeline of the experiment. Respondents first read the instructions and observe the signal about which model is more likely to be correct. We assume that this induces the default representation

$$r^{\text{initial}} = (0.65, V_{\text{initial}}^{\text{context}}), \quad \mathbf{R}_0 = \{r^{\text{initial}}\}.$$

We assume that $V_{\text{initial}}^{\text{context}}$ contains features relating to the general setup of the environment, in particular the two models, signal structure and received signal.

After going through the instructions, respondents provide the company value guesses, state their decision confidence and then finally their beliefs about which model is correct.

C.4 Predictions

C.4.1 Estimation of Company Values

Predictions. In the high-complexity condition, respondents will simplify the initial representation to $\hat{\pi} = 1$ and choose $a^* = a_A$. In the low-complexity condition, they will keep $\hat{\pi} = 0.65$ and choose $a^* = 0.65 a_A + 0.35 a_B$.

Proof. Since $R_0 = \{r^{\text{initial}}\}$, participants must decide whether or not to collapse that single representation. Hence they form

$$r^{\text{guess}} = \begin{cases} (1, V_{\text{guess}}^{\text{context}}) & \text{if } K(c, 0.65) - K(c, 1) > \Delta U(0.65), \\ (0.65, V_{\text{guess}}^{\text{context}}) & \text{otherwise.} \end{cases}$$

Substituting $K(c, \pi) = (1 - \pi)\kappa_c$ shows that simplification occurs exactly when $(1 - 0.65)\kappa_c > \Delta U(0.65)$. By assumption $(1 - 0.65)\kappa_{\text{low}} < \Delta U(0.65) < (1 - 0.65)\kappa_{\text{high}}$, hence only high-complexity agents collapse to $\pi = 1$, yielding the stated actions.

C.4.2 Confidence

Predictions. Assuming that the reported confidence for $\hat{\pi} = 1$ exceeds that for $\hat{\pi} = 0.65$, on average, respondents in the high-complexity condition will report higher confidence than those in the low-complexity condition.

Proof. Confidence is assessed within the formed mental representation. On the confidence screen, the set of available representations is $R_1 = \{r^{\text{initial}}, r^{\text{guess}}\}$. We assume the prompt ‘How certain are you that your answer is within 10 percentage points of the best possible guess?’ makes the guess context more salient, so

$$S(\xi_{\text{confidence}}, V_{\text{guess}}^{\text{context}}) > S(\xi_{\text{confidence}}, V_{\text{initial}}^{\text{context}}).$$

Therefore the most likely case is that r^{guess} is retrieved. Because $c = \emptyset$ on this screen, no further simplification occurs, so high-complexity subjects retain $\hat{\pi} = 1$ and low-complexity subjects retain $\hat{\pi} = 0.65$. By the assumption that higher stored $\hat{\pi}$ yields higher reported confidence, high-complexity subjects report greater confidence.

If r^{initial} is retrieved instead, once again no simplification occurs because of $c = \emptyset$ and participants in both conditions report a low confidence associated with $\hat{\pi} = 0.65$.

C.4.3 Belief

Predictions. Simplification in the action space might not carry over to the belief space. The most likely outcome is that respondents in both complexity conditions will state their belief as $\hat{\pi} = 0.65$. Respondents in the high-complexity condition may exhibit a greater tendency to state the simplified belief of $\hat{\pi} = 1$.

Proof. On the belief screen, $R_2 = \{r^{\text{initial}}, r^{\text{guess}}, r^{\text{confidence}}\}$. We assume the prompt ‘Which of the two rules generated correct company value estimates?’ explicitly echoes the orig-

inal signal description, so

$$S(\xi_{\text{belief}}, V_{\text{initial}}^{\text{context}}) > \max\{S(\xi_{\text{belief}}, V_{\text{guess}}^{\text{context}}), S(\xi_{\text{belief}}, V_{\text{confidence}}^{\text{context}})\}.$$

Thus the most likely case is that r^{initial} is retrieved. As $c = \emptyset$ here too, no collapse occurs and the stated belief is $\hat{\pi} = 0.65$ in both complexity conditions.

If r^{guess} or $r^{\text{confidence}}$ are retrieved instead, once again no simplification occurs because of $c = \emptyset$. For respondents in the high-complexity condition this implies $\hat{\pi}(r^{\text{guess}}) = \hat{\pi}(r^{\text{confidence}}) = 1$, while under low complexity $\hat{\pi}(r^{\text{guess}}) = \hat{\pi}(r^{\text{confidence}}) = 0.65$.

D Instructions

D.1 Introduction

D.1.1 Welcome Screen and Attention Check

Welcome!

This study is designed for **computer (PC or Mac) users only** (desktop, laptop, etc.). If you are accessing this study on a smartphone, a tablet or any other non-PC devices, please switch to PC and enter the study again, or return the submission on Prolific.

Please write **at least 15 words** describing your opinion about daylight savings time. Whether you are in favor or against daylight savings does not affect your eligibility to participate in this study. However, we ask that you write at least 15 words on your thoughts about this topic.



D.1.2 General Instructions

General Instructions

Thank you for participating in this study.

Today's survey will take approximately 20 minutes. You will earn a reward of \$4.00 for participating (implying an hourly wage of \$12/hr).

To earn your reward, you have to **read all instructions carefully and correctly answer the comprehension questions.**

To receive your payment, it is crucial that you pay attention throughout the whole study!

Every tenth participant has the chance to get an **additional bonus of \$10.**

Feel free to use a piece of paper and pen during the study.



D.1.3 Bonus Information

Bonus

Today's survey consists of **two parts.**

If you are selected for the bonus, **an answer in one of the two parts will be randomly selected to determine your bonus.**

Therefore, you should answer all questions carefully.



D.2 Company Value Guesses

D.2.1 Instructions

Part 1 - Instructions

Basic Setting

In this study, you will need to guess the value of 8 different hypothetical companies. All companies have a value between 10 and 100.

There are two rules that generate estimates for the value of the companies. One of these rules will generate the correct estimate, the other will generate an incorrect and uninformative estimate. Both rules take different variables as inputs from which the estimates can be calculated.

These two rules are:

1. **"The CEO is key"**

Under the "The CEO is key" rule, the CEO competence **C** and supporting staff **S** are the only variables that matter.

The estimate under this rule is given by $C \cdot S - C - S + 10$.

2. **"Products are crucial"**

Under the "Products are crucial" rule, the number of products **P** and the research cost **R** are the only variables that matter.

The estimate under this rule is given by $P \cdot (10 - R) + R - P$.

The realizations of the variables are different for each company and are always between 0 and 10. The company value estimates under the two rules will also differ between companies and are always between 10 and 100.

Indication of the Correct Rule

At the beginning of the study, the computer will randomly determine with a fair coin flip which of the two rules generates correct estimates. Therefore, only either the "The CEO is key" or "Products are crucial" rule will generate correct estimates for all 8 companies.

Importantly, you will not learn the outcome of the coin flip directly. Instead, we will provide you with an indication for which rule is the correct one. This indication will give you an idea for which rule is more likely to produce correct estimates, but you will not learn the correct rule with certainty.

Versions of the Study

There are two versions of the study. After you have read the instructions, the computer will randomly assign you to one of these versions.

If you are assigned to **study version 1**, you need to calculate the company value estimates under both rules yourself. For each company, you can learn about the realizations of the different variables by hovering over the variable names, which can be plugged into the formulas to obtain the company value estimates under each rule. Then you need to provide a guess about the value of that company.

If you are assigned to **study version 2**, the company value estimates under both rules will be calculated for you. For each company, you can learn about the company value estimates under both decision rules by hovering over the name of the respective rule. Then you need to provide a guess about the value of that company.

Example (Study Version 1)

The following is an example for the information you will see on the decision screen if you are assigned to study version 1.

Hover over the elements below to display the formulas:

- **"The CEO is key":**
- **"Products are crucial":**

Hover over the elements below to display the variables which can be used to calculate the company value estimates under both rules:

- **CEO competence C:**
- **Supporting staff S:**
- **Number of products P:**
- **Research cost R:**

Given these variable realizations, the company value estimates can be calculated under each decision rule:

- **Estimate for "The CEO is key":** $C \cdot S - C - S + 10 = 5 \cdot 4 - 5 - 4 + 10 = 21$
- **Estimate for "Products are crucial":** $P \cdot (10 - R) + R - P = 7 \cdot (10 - 2) + 2 - 7 = 51$

Example (Study Version 2)

The following is an example for the information you will see on the decision screen if you are assigned to study version 2.

Hover over the elements below to display the company value estimates under each decision rule:

- **Estimate for "The CEO is key":**
- **Estimate for "Products are crucial":**

Bonus Payment

We will randomly select one of your 8 answers. **The closer your guess is to the correct value of the company, the higher the likelihood that you receive the bonus of \$10.**

If you click on the triangle below, the precise formula will be displayed. While this formula might seem complicated, the underlying principle is very simple: the smaller the difference between your estimate and the truth, the higher the likelihood that you win \$10. It is hence in your best interest to simply state your best guess.

Importantly, it does not directly matter for your earnings whether the company value is high or low. All that matters is that you guess it correctly. Your chance of receiving the bonus only depends on how close your guess is to the actual company value, which is not necessarily the higher one.



Likelihood of winning the bonus (in %) = $100 - 100 \cdot (\text{Value} + 100 - \text{Guess} + 100)^2$

D.2.2 Comprehension Questions

Comprehension Questions

Please answer the following comprehension questions. You need to answer them correctly to continue the study.

For how many companies will you need to provide a guess?

8

12

10

Which one of the following statements is correct?

The rule that generates correct company value estimates might differ between the companies.

A fair coin flip will determine which rule generates correct company value estimates. This rule will then be relevant for all companies.

Which one of the following statements is correct?

The guesses I make in this study can affect my payoff. The study involves real stakes.

The guesses I make in this study will not affect my payoff. The study is purely hypothetical.

The computer will secretly determine which rule generates correct company value estimates. Which of the following is true?

The computer is more likely to select "The CEO is key".

The computer is more likely to select "Products are crucial".

Both rules are equally likely.

In the example of study version 1, what is the variable realization for **Number of products P**?

What is true regarding your chance of winning the 10\$ bonus?

The chance of getting the bonus is highest when I guess the larger of the two estimates proposed by the two rules.

The chance of getting the bonus is highest when I state the guess that is closest to the true company value, which is not necessarily the higher of the two estimates.

The chance of getting the bonus is highest when I guess the sum of the two estimates proposed by the two rules.



D.2.3 Example Decision Screens for Selecting Restricted Sample

You passed the comprehension questions!

Before the actual tasks start, we give you the opportunity to familiarize yourself with the calculations of study version 1 on the next two pages.

This is not an additional comprehension check, but you can use the opportunity to get familiar with the calculations.



Example for Study Version 1

Hover over the elements below to display the formulas:

- "The CEO is key": $C \cdot S - C - S + 10$
- "Products are crucial":

Hover over the elements below to display the variables which can be used to calculate the company value estimates under both rules:

- CEO competence **C**:
- Supporting staff **S**:
- Number of products **P**:
- Research cost **R**:

What is the company value estimate under the rule "The CEO is key"?

(Your response needs to lie between 10 and 100.)



Example for Study Version 1

Hover over the elements below to display the formulas:

- "The CEO is key":
- "Products are crucial":

Hover over the elements below to display the variables which can be used to calculate the company value estimates under both rules:

- CEO competence **C**:
- Supporting staff **S**:
- Number of products **P**:
- Research cost **R**: 5

What is the company value estimate under the rule "Products are crucial"?

(Your response needs to lie between 10 and 100.)



D.2.4 Treatment Assignment - HighComplexity

Study Version

The computer has assigned you to study version 1, which means you need to calculate the company value estimates under both decision rules yourself.

D.2.5 Treatment Assignment - LowComplexity

Study Version

The computer has assigned you to study version 2, which means the company value estimates under both decision rules will be calculated for you.



D.2.6 Determining the Correct Rule

Determination of the Rule

Next, the computer will flip a fair coin to secretly determine which of the two rules - "**The CEO is key**" or "**Products are crucial**" - will generate the correct value estimates for all 8 companies.

Please click on "Flip the coin!" to proceed.

Flip the coin!

D.2.7 Timeline - HighComplexity

Timeline

- The computer has now flipped a fair coin to determine which of the two rules - "The CEO is key" or "Products are crucial" - generates the correct value estimates. We will next provide you with an indication about which rule has been selected.
- For each of the 8 companies, you can see the realizations of the four variables - CEO competence C, supporting staff S, number of products P, research cost R - by hovering over the variables on the decision screen. You can also remind yourself of the formulas for both rules by hovering over them. You then need to provide a guess about the value of that company.



D.2.8 Timeline - LowComplexity

Timeline

- The computer has now flipped a fair coin to determine which of the two rules - "Products are crucial" or "The CEO is key" - generates the correct company value estimates. We will next provide you with an indication about which rule was selected.
- For each of the 8 companies, you can directly see the company value estimates under both rules - "Products are crucial" and "The CEO is key" - on the decision screen. You then need to provide a guess about the value of that company.



D.2.9 Drawing the Signal

Indication about the Rule

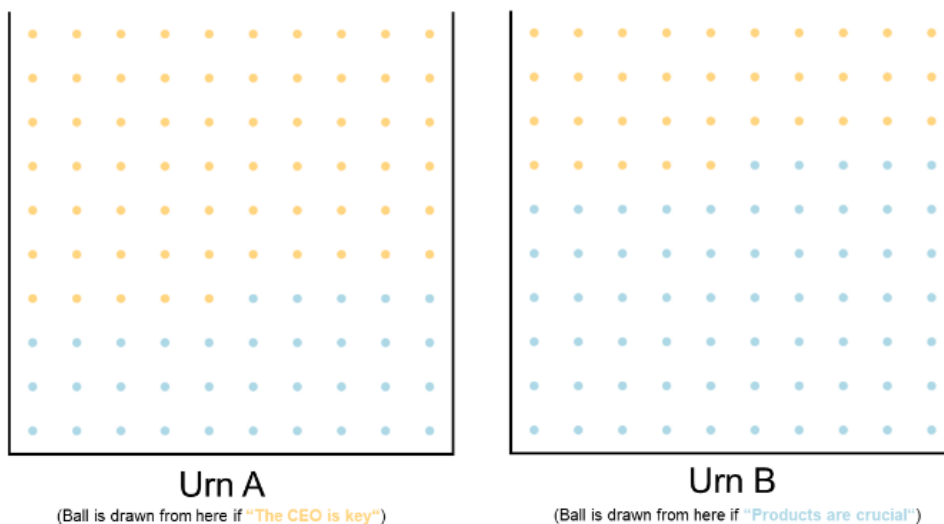
The computer has secretly flipped a coin to determine which of the two rules - **"The CEO is key"** or **"Products are crucial"** - generates the correct company value estimates.

On the next page, you will observe an indication for which of the two rules was randomly chosen.

You will get this information in the form of a ball that is drawn from one of the two urns displayed below.

- If **"The CEO is key"** generates correct estimates, the computer will draw a single ball from **Urn A**, which contains **65 orange balls** and **35 blue balls**.
- If **"Products are crucial"** generates correct estimates, the computer will draw a single ball from **Urn B**, which contains **35 orange balls** and **65 blue balls**.

This means that a **orange ball** is rather indicative of **"The CEO is key"**, and a **blue ball** is rather indicative of **"Products are crucial"**. Notice again that these are only indications, **you will not know with certainty which rule has been selected** by the computer.



Draw the ball!

D.2.10 Signal Screen - CEO is key

The drawn ball is orange, which means "**The CEO is key**" is more likely to generate correct estimates.



You can proceed to the next screen in 15 seconds.

D.2.11 Signal Screen - Products are crucial

The drawn ball is blue, which means "**Products are crucial**" is more likely to generate correct estimates.



You can proceed to the next screen in 15 seconds.

D.2.12 Signal Check

To verify that you have observed the indication, please choose the correct statement below.

The rule "**The CEO is key**" is more likely to generate correct estimates

The rule "**Products are crucial**" is more likely to generate correct estimates



D.2.13 Decision Screen - HighComplexity

Company 1/8

Hover over the elements below to display the formulas:

- "The CEO is key":
- "Products are crucial":

Hover over the elements below to display the variables which can be used to calculate the company value estimates under both rules:

- CEO competence **C**: 9
- Supporting staff **S**:
- Number of products **P**:
- Research cost **R**:

What is your guess for the value of the company?

(Your response needs to lie between 10 and 100.)



The formulas and variable realizations were displayed upon hovering over the respective text.

D.2.14 Decision Screen - LowComplexity

Company 1/8

Hover over the elements below to display the company value estimates under each decision rule:

- Estimate for "Products are crucial":
- Estimate for "The CEO is key": 25

What is your guess for the value of the company?

(Your response needs to lie between 10 and 100.)



The pre-calculated company value estimates were displayed upon hovering over the respective text.

D.3 Beliefs

D.3.1 Instructions

Part 2 - Instructions

You have just guessed the values of companies. Either the **"The CEO is key"** rule or the **"Products are crucial"** rule generated correct company value estimates.

In this part, we ask you to state which of the two rules generated the correct company value estimates. Specifically, we will **ask you for the probability (between 0 and 100%) that "The CEO is key" generated correct estimates**. The closer your guess is to the correct answer, the higher is your chance to win the \$10 bonus.

If you wish, you can click on the triangle to uncover the precise formula for the bonus implied by your probability guess. However, you only need to know that it's in your best interest to simply state your best guess.



Likelihood of winning the bonus (in %) = $100 - 100 \cdot (\text{Truth} - \text{Guess})^2$

"Truth" is 1 if **"The CEO is key"** is the correct rule, and 0 if **"Products are crucial"** is the correct rule.

"Guess" is your guess for the likelihood (in %) that **"The CEO is key"**.



D.3.2 Decision Screen

In the first part, which of the two rules - **"The CEO is key"** or **"Products are crucial"** - generated correct company value estimates?



"The CEO is key": 34 %.



"Products are crucial": 66%.



D.3.3 Confidence Elicitation

You guessed that the likelihood for **"The CEO is key"** is 34%.

How certain are you that this answer is within 10 percentage points of the best possible guess given the information you received in the first part?

(This question does not affect your bonus.)



D.3.4 Direct Recall

In the first part, before providing guesses for the company values, you received an indication for which rule was selected by the computer to generate correct company value estimates.

Do you recall which rule was indicated to be more likely to generate correct value estimates?

(This question does not affect your bonus.)

"The CEO is key"

"Products are crucial"

I don't remember

