

Discussion Paper Series – CRC TR 224

Discussion Paper No. 385
Project B 03

The Status Quo and Belief Polarization of Inattentive Agent: Theory and Experiment

Vladimir Novak ¹
Andrei Matveenko ²
Silvio Ravaoli ³

January 2023

¹ National Bank of Slovakia, Email: vladimir.novak@nbs.sk

² Department of Economics, University of Mannheim, Email : matveenko@uni-mannheim.de

³ Cornerstone Research, New York, Email: vio.ravaoli@cornerstone.com

Support by the CRC

The Status Quo and Belief Polarization of Inattentive Agents: Theory and Experiment*

Vladimír Novák^{†1}, Andrei Matveenko², and Silvio Ravaoli³

¹National Bank of Slovakia, Imricha Karvaša 1, 813 25 Bratislava, Slovakia

²Department of Economics, University of Mannheim, L7 3-5, 68 161 Mannheim, Germany

³Cornerstone Research, 599 Lexington Avenue, New York, USA

January, 2023

Abstract

We show that rational but inattentive agents can become polarized ex-ante. We present how optimal information acquisition, and subsequent belief formation, depend crucially on the agent-specific status quo valuation. Beliefs can systematically - in expectations over all possible signal realizations conditional on the state of the world - update away from the realized truth and even agents with the same initial beliefs might become polarized. We design a laboratory experiment to test the model's predictions. The results confirm our predictions about the mechanism (rational information acquisition), its effect on beliefs (systematic polarization) and provide general insights into demand for information.

Keywords: polarization, beliefs updating, rational inattention, status quo, experiment.

JEL codes: C92, D72, D83, D84, D91.

*An earlier version of this paper appeared in CERGE-EI Working Papers series under the same title. Many thanks to Filip Matějka for invaluable advice and crucial help in setting up the framework employed in this paper, and to Mark Dean for continuous support on this project. We are grateful for comments from Sandro Ambuehl, Niels Boissonnet, Andrew Caplin, Alessandra Casella, Yeon-Koo Che, Alexander Nesterov, Pietro Ortoleva, Jacopo Perego, Debraj Ray, Avner Shaked, Jakub Steiner, Michael Woodford, Jan Zápál, and numerous seminar participants for helpful comments and suggestions. This project has received funding from the European Research Council under the European Union's Horizon 2020 research and innovation programme (grant agreements No. 678081 and No. 740369) and we acknowledge the financial support of grants from the Columbia Experimental Laboratory for Social Sciences. Funding by the German Research Foundation (DFG) through CRC TR 224 (project B03) is gratefully acknowledged. Vladimír Novák was supported by the Charles University in Prague (GA UK projects No. 197216 and No. 616219) and by the H2020-MSCA-RISE project GEMCLIME-2020 GA. no 681228. The authors affirm no financial interest in the outcomes of the experiments detailed herein. All data were collected with the approval of the Columbia University Institutional Review Board (protocol AAAS5801). The views expressed in this paper are those of the authors and should not be attributed to the National Bank of Slovakia or Cornerstone Research.

[†]Corresponding author.

E-mail addresses: vladimir.novak@nbs.sk (V. Novák), matveenko@uni-mannheim.de (A. Matveenko), silvio.ravaoli@cornerstone.com (S. Ravaoli).

1 Introduction

Societies are polarized¹ not only in their beliefs about future policies but show significant disagreement even in their valuations of reality – the implemented policies (status quo) – that can be factually evaluated (e.g., Alesina, Miano and Stantcheva, 2020). Such heterogeneity in evaluations of reality leads to differences in perceived gains and losses associated with the adoption of a new policy and thus may have a significant influence on the demand for information, belief polarization, and eventually, on essential economic decisions.² Consider policies³ that aim to achieve a climate-neutrality by 2050⁴ as an illustrative example.⁵ The consequences of adopting, for instance, carbon taxes were uncertain at the moment of choice, whereas the opposite – to preserve the status quo – appeared to be more certain in its consequences. Public opinion surveys suggest that such binary policy choices are often associated with an increase in belief polarization in society, while misconceptions about reality play a crucial role.⁶ This raises the question: How do status quo valuations influence belief polarization and what important aspects of the environment determine the demand for information?

In this paper, we present a new model of belief polarization and a laboratory experiment that tests the model’s predictions and provides insights into determinants of the demand for information. The key theoretical mechanism is that when attention is scarce, the relative valuation of the status quo determines the choice of information structure, which may in turn lead to systematic polarization of beliefs even between rational agents with the same prior beliefs. In particular, unlike the existing literature that shows that there exist such signal realizations that lead to belief polarization (see e.g. Nimark, Sundaresan (2019)), we show

¹See, for instance Poole and Rosenthal (1984); McCarty, Poole and Rosenthal (2008); Gentzkow, Shapiro and Taddy (2016).

²For instance, Alesina, Stantcheva and Teso (2018) document that people that are pessimistic about the actual intergenerational mobility tend to be more favorable towards redistribution. Similar results were shown for immigration (Alesina, Miano and Stantcheva, 2018) and income inequality (Kuziemko et al., 2015).

³Examples of such policies include carbon taxes, bans on polluting vehicles, support of infrastructure, agriculture policies, and many others.

⁴Such target was declared by more than a hundred countries worldwide NPUC (2021).

⁵Other examples might include referendums like, for instance, the 2016 Brexit referendum.

⁶Dechezleprêtre et al. (2022) present how attitudes toward climate change policies differ across twenty countries, how they vary with knowledge of the present situation, expectations about their impact, and document impact of information treatments on the support of these policies. Douenne and Fabre (2022) show that misperceptions about the impact of carbon pricing policies were one of the drivers of opposition to these policies during the Yellow Vest movement in France. In general, the polarization of climate change news is documented, for instance, by Chinn, Hart and Soroka (2020); Leiserowitz et al. (2019). In the case of the Brexit referendum British society a few months after the Brexit referendum was even more polarized than on the referendum day (Smith, 2019).

that there exist states of the world in which the agents become polarized in their posterior expected value of the new policy on average over all possible signal realizations from the selected information structure. We label such polarization as *polarization ex-ante* because it can be identified even before signals are realized as it is driven by the agent’s selection of the information structure that is given by the agents’ objectives and primitives.

We design a laboratory experiment in which we manipulate the value of the status quo action in order to test the theoretical predictions. The results confirm both predictions about the key mechanism of information acquisition and the qualitative effect on beliefs, although the overall effect is mitigated by behavioral factors. The main mitigating channel is a preference for simple signal structures, a result that generalizes the well-known preference for certainty. Other factors, like risk preference and subjective beliefs, are not sufficient to explain the observed deviations from the predictions.

We begin our paper with a model that generalizes the intuition provided in a simple illustrative example which is below in the introduction. Specifically, we model the agent to be rationally inattentive, following Sims (1998, 2003), which allows us to account for endogenous information acquisition without imposing any restrictions or biases on the agent’s learning process. We consider a static decision problem with n states and two actions, in the manner of Matějka and McKay (2015). However, in contrast to Matějka and McKay (2015), our focus is on the evolution of beliefs. Since information is plentiful but attention is scarce, the agent chooses to learn the essential pieces of information for her decision problem, whether to adopt the new policy or to preserve the status quo. Thus, the agent endogenously partitions states into groups, separating the states in which payoffs of the new policy are higher or lower than under the status quo, and acquires costly information just to identify which of the two groups contains the realized state. We refer to such avoidance of redundant information about states associated with the same optimal action as *state pooling effect*. Our main theoretical result (Theorem 1) proves that two agents can become polarized ex-ante (Definition 1), that is, their expected posterior values of the new policy in expectation over all possible signal realizations evolve in opposite directions and are further apart than the prior expected values. In addition, we show that two agents might become polarized when they either differ in the valuation of the status quo or in their prior expected beliefs about the payoff of the new policy. We acknowledge, however, that the presented mechanism relies on agents’ ability to influence the information-gathering process and thus does not aim to capture situations in which all agents are presented with the same selective evidence, with no power over the information source, and are asked questions about an emotionally charged

topics (e.g., death penalty) as it is often the case in many classical psychological experiments (see, e.g., Lord, Ross and Lepper, 1979).

The theoretical results presented rely on several assumptions about a decision maker’s preferences and updating process that have been challenged by previous experimental findings.⁷ In Section 3, we introduce an experiment designed to test our theoretical predictions, in particular the state pooling effect and the presence of belief polarization ex-ante. The experiment’s setup is restricted with respect to the general theoretical framework, focusing on testing predicted behavior and not necessarily rational inattention per se. In the main task, we focus on the choice between signal structures that might differ in their internal mental processing costs but without explicitly stated cost, removing the assumption on a specific functional form for the attention cost. The subjects are presented with a binary choice and can acquire instrumentally valuable information from advisors (signal structures) before making their decision. In the main task, participants make an information choice followed by an action choice. For the information choice, they are presented with a pair of advisors and can select only one of them. After that, they indicate the chosen action (status quo or new policy) conditional on the observed signal. Our key manipulation consists in varying the value of the status quo, and we expect this to be sufficient to revert the choice of optimal advisor. In two separate tasks we elicit, for each participant, the subjective beliefs about the likelihood of a signal realization, and each state (conditional on the realized signal).

Section 4 presents our main experimental results. Participants do react to the value of the status quo, as predicted by our theoretical model, and display a preference for state pooling information structures. Importantly, we document the belief polarization ex-ante in our laboratory setting. The magnitude of the polarization is mitigated with respect to the predicted magnitude, and we investigate the main behavioral channels that can explain the deviation. We show that the evaluation of information structures is affected by non-instrumental characteristics; in particular, preference for advisors that are characterized by certainty (degenerate posteriors) or simplicity (fewer possible outcomes). In Section 5 we discuss the robustness of our experimental findings and Section 6 concludes the paper with a discussion of the limits of our model and experiment and directions for future research.

⁷The experimental literature contains systematic evidence of deviation from theoretical predictions in several domains. See, for example Holt and Laury (2002) (risk preferences), Kareev, Arnon and Horwitz-Zeliger (2002) (subjective beliefs), and Ambuehl and Li (2018) (preference over signal structures).

1.1 Example and policy implications

In this section, we provide the intuition for our result using a simple highly restricted example that, however, still allows us to demonstrate the main mechanism, its connection with the real world and possible policy implications. Furthermore, the reader interested purely in the experimental investigation should be able to skip over the details of the general model and immediately move to the experimental part, after reading this subsection.

Intuition. Assume the state of the world $v \sim U[0, 1]$. There are two risk-neutral payoff maximizing agents (A and B) facing a binary action $a \in \{0, 1\}$, referring to preserving the status quo and adoption of a new policy, respectively. Agent A prefers option $a = 1$ if $v \geq R_A$; and agent B prefers option $a = 1$ if $v \geq R_B$, where $R_i \in (0, 1) \forall i \in \{A, B\}$.⁸ For simplicity let us assume that $R_B < R_A$ and that both agents have the same uninformative prior beliefs. If information acquisition is costly, agents will demand the most instrumental signal structure, that is, agent A will ask whether $v \geq R_A$ and agent B will ask whether $v \geq R_B$, but none of the agents would care about the exact value of v . The fact that the agents do not distinguish some states of the world after the optimal acquisition of information, and pool states associated with the same action together, is referred to as a *state pooling effect*. Consequently, when the true state of the world $v \in (R_A, R_B)$ the agents would receive opposite signals with respect to whether they should adopt the new policy, given the assumption they receive noiseless truthful signals. Therefore, agents' posterior expected values from the new policy would get polarized, i.e. move in opposite directions (towards opposite extremes) and further apart as they were.

In section 2 we show that a similar pattern of behavior can be observed in a much richer setting for rationally inattentive agents. Crucially, the full-fledged model allows us to show in addition that the agents would not only get polarized in their posterior expected conditional on a true state for a particular realized signal but in expectation over all possible signal realizations from the selected information structure. We denote such divergence of beliefs as polarization ex-ante, as it allows us to claim that agents will polarize even before signals are received (or observed) purely on knowing their valuations of the status quo policy and prior beliefs.

Connection with real-world examples. Our framework allows interpreting considered policies quite broadly. Consider the adoption of some particular climate change mitigating policies, especially the carbon tax policy, as an example. The crucial impact of

⁸We thank a referee for suggesting this simpler version of the example to communicate the intuition of the model.

the expectations and status quo valuations on people’s disagreement about measures that should be adopted to mitigate climate change is documented by Dechezleprêtre et al. (2022), who presents how attitudes toward climate change policies differ across twenty countries. The specific example is provided by a series of papers by Douenne and Fabre (2022) and Douenne and Fabre (2020). The focus of their papers is on aversion towards the carbon tax in France, especially after the yellow vests movement in France. The mapping to our setting is as follows. The prior dispersion in the status quo valuations can be represented by households’ purchasing power. Simultaneously, they document dispersion in prior beliefs (10% approving the reform vs. 70% opposing it). Importantly, for the connection with our paper, they are able to determine the true effect of the reform, v in our case. Specifically, 70% of households are supposed to benefit from the reform and in connection with their energy usage, they are able to compute a respondent-specific estimate of the tax incidence on their purchasing power. They document that misperceptions about the impact of carbon tax policies were one of the drivers of opposition to these policies during the Yellow Vest movement in France. Interestingly, the authors also provide the information treatment⁹ to reverse the pessimistic beliefs. They observe that generally, the information they provided fails to reverse pessimistic beliefs, but they update their posterior expected valuation of the reform in an asymmetric way - that is, they overweight negative information, showing that they would slightly lose from the reform.¹⁰ Such observation is completely in line with our predictions and thus provides a possible mechanism explaining the observation of Douenne and Fabre (2022).

Policy implications. The endogenously arising state pooling effect represents a new mechanism that can generate belief polarization. It differs from other mechanisms, for example, confirmatory learning and information misunderstanding, which assume biases in the processing of new information and to which the observed polarization is often attributed (see, e.g., Fryer Jr, Harms and Jackson, 2019). It is typically considered desirable to mitigate polarization and design policies that aim to reach this effect. In order to be effective in doing so, it is essential to understand which mechanism is responsible for the observed polarization and under what circumstances.

First, note that observed behavior resulting from the state pooling might be empirically

⁹Specifically, respondents randomly receive a piece of information about the progressivity and/or the effectiveness of the policy as well as customized information derived from our respondent-specific estimation on whether their household is expected to win or lose from the policy.

¹⁰In terms of the full model with three states, this corresponds to the setting when $v_1 < \mathbb{E}v < v_{s^*=2} < R < v_3$, the impact of the policy is slightly negative in comparison with the current situation $R > v_{s^*=2}$ and as a consequence $m^j(s^*) < \mathbb{E}v$.

reminiscent of behavior generated by confirmatory learning. This remark gives reason for caution when inferring biases from observed beliefs. Second, our mechanism highlights that interventions affecting the status quo valuation might dramatically change the information acquisition strategy in our setting. In comparison, the agent with a confirmation bias (acquiring information that supports the agent’s prior belief) should not alter the information acquisition strategy when the valuation of the status quo changes, but prior belief stays the same.

1.2 Related literature

This paper contributes to the theoretical and experimental literature on information acquisition, as well as to the broader literature on belief polarization. We highlight the role of rational endogenous information acquisition about instrumental choice in belief polarization. Furthermore, our experimental study provides important insights into what aspects of information structure are important in determining a demand for information.¹¹

In our model, we consider a multiple states environment in which rationally inattentive agents can choose any distribution of signals about the value of the new policy subject to the cost of information and their valuation of the status quo policy. Modeling the agent as rationally inattentive relieves us of the need to assume exogenously given biases¹² or bimodality of preferences (Dixit and Weibull, 2007), which are common in the preceding literature. This in turn allows us to move in a new direction, away from findings that the beliefs of Bayesian agents would converge over time and that they will almost surely assign probability 1 to a true state (Savage, 1954; Blackwell and Dubins, 1962); and allowed us to contribute to the polarization literature also by providing rational decision theory mechanism that is not driven by the media competition (Perego and Yuksel, 2018; Bernhardt, Krassa and Polborn, 2008) or necessarily partisan conflict (Prior, 2013; Oliveros and Várdy, 2015; Gentzkow, Shapiro and Taddy, 2016).

The question of how inattentiveness can lead to persistent belief polarization was previously studied by Nimark and Sundaresan (2019), however, there are several significant differences that set our reasoning apart. Nimark and Sundaresan (2019) study binary state environments in which an agent that wants to determine the state of the world can pay a cost to influence the noise of binary signal about the state. They document a confirmation effect, that is, the agent selects to receive less noisy information about the state that is more

¹¹We thank a referee for noting that our results contribute to this more general point.

¹²Gerber and Green (1999) review the literature that invokes some biases in learning or perception in order to modify Bayesian updating.

likely a priori. Consequently, two agents with the same prior belief permanently disagree, in their setting only when they first observe different signal realizations which are in the following periods reinforced by the confirmation effect. This is due to the fact that both agents select initially the same information structure. In contrast, the agent in our model endogenously decides to pool states into binary categories, thus the binary signal form is a result, not an assumption. Even more importantly, we show that agents become polarized ex-ante conditional on the realized state, i.e., on average over all possible signal realizations. Such polarization as we show occurs for intermediate states and thus cannot occur in the two-state environment. Our conclusions also do not depend on the fact that two agents would first observe different signal realizations and thus they lead to different policy implications.

The state pooling effect that emerges as a result of the model can be viewed as an endogenously arising categories formation mechanism, a specific decision-making heuristic (Gigerenzer and Gaissmaier, 2011), and can further provide a foundation for the settings which assume at most two states, i.e., setting that is often used in dynamic settings (see, e.g. Che and Mierendorff, 2019; Nimark and Sundaresan, 2019). Endogeneity of the coarse reasoning differentiates us also from the previous literature that assumes exogenously given categories (see, e.g., Suen, 2004; Manzini and Mariotti, 2012; Maltz, 2020). Most closely related paper from those considering exogenously given information coarsening, is work by Suen (2004). It investigates how information coarsening can lead to self-perpetuation of biased beliefs and polarization of opinions. Specifically, the agents face a choice among advisors who observe continuous signals about the binary state of the world. The signals are modeled as a random draw from one of two state-dependent distributions with continuously differentiable density functions. Each advisor, by assumption, is coarsening (pooling) signals into binary recommendations, based on their advisor specific threshold. They show that an agent with biased belief prefers to receive information from an advisor that conforms their existing beliefs and hence two agents with different beliefs interpret same evidence differently, due to different advisor selection. Focus of our paper is, on endogenous information acquisition rather than on interpretation of evidence. Hence in our model, the agent is pooling together the states and thus selecting different information structure, rather than pooling signals together from the same underlying information structure. Consequently, the agent neither needs to have biased beliefs nor necessarily prefers conforming information structures.

Recently, two contemporaneous papers Bloedel and Segal (2021) and Hu, Li and Segal

(2022) study the optimality of binary signals when information is costly and its effect on policy polarization, respectively. Bloedel and Segal (2021) consider a model of persuasion of a rationally inattentive agent and show that the sender can pool states together to attract the receiver’s attention. In our work, we show that state pooling can appear without strategic interactions, but the exact way how the state are pooled differs among the papers. In Bloedel and Segal (2021) the receiver pools together states with intermediate and high-stakes and reveals perfectly low-stake states. In the current work the agent pools states differently (e.g., low-stake states are never revealed perfectly), because the agent’s objective is not to manipulate its attention and state pooling is driven by the status quo valuation. Hu, Li and Segal (2022) study a two-candidate electoral competition model, with attention-maximizing infomediary that aggregates information about candidates’ valence and investigate implications of personalized news vs. providing same news for everyone. Similarly as us they show that optimal personalized signal for any voter is binary. Our paper differs from that of Hu, Li and Segal (2022) in the driving force of the state pooling, the mechanism behind the polarization, and the very form polarization that we study (policy polarization vs. belief polarization).

Naturally, this paper also adds to the rational inattention literature¹³ by studying the evolution of beliefs alongside the manifestation of the crucial implications of incorporating the safe option into the choice set. The main mechanism of our paper - the state pooling effect - is connected with the presence of the status quo policy, a particular form of reference point. However, in contrast with the rich theories of individually determined reference points (see, e.g., Kahneman and Tversky, 1979; Köszegi and Rabin, 2006; Guney, Richter and Tsur, 2018), the status quo is exogenously given. At the same time, we neither assume the status quo bias (Samuelson and Zeckhauser, 1988) nor aim to study the formation of such bias (Ortoleva, 2010; Masatlioglu and Ok, 2014), but our focus is on the evolution of beliefs and not necessarily the chosen actions. Hence, our paper complements the studies documenting how current economic standing, in our setting represented by the status quo policy valuation, influences the action taken by citizens¹⁴.

Finally, the paper adds to the experimental literature. First, in contrast to Charness,

¹³A survey of literature on rational inattention is provided in Maćkowiak, Matějka and Wiederholt (2020). For a posterior-based approach see Caplin and Dean (2015) and a dynamic discrete choice model is presented in Steiner, Stewart and Matějka (2017).

¹⁴For instance, consider Fetzer (2019) who claims that economic losers were more likely to vote for Brexit, Dal Bó et al. (2018) who connect economic losers with the rise of the Swedish radical right, and Alesina, Stantcheva and Teso (2018) who show that people that view social mobility more pessimistically are more favorable towards redistribution.

Oprea and Yuksel (2021) we focus on variation in the value of the status quo and not on the variation in prior beliefs. Second, our findings give reason for caution for empirical and experimental work inferring preference for information. In particular, this paper provides a disciplined model that suggests how the preference for skewed information might crucially depend on the value of the status quo and thus provides an important channel that is missing in the research on whether people prefer negatively or positively skewed information, e.g. Masatlioglu, Orhun and Raymond (2017). Third, we replicate and extend previous evidence of certainty preference (Ambuehl and Li, 2018) in a three states environment.

2 The model

In this section, we describe the general case of the agents' decision problem, introduce a methodology for assessment of beliefs evolution, and present the main theoretical results. The structure of this section is as follows. In the subsections 2.1 and 2.2, we describe the choice problem faced by agents, which is a special case of the agent's problem from Matějka and McKay (2015). Subsection 2.3 discusses the evolution of the agents' beliefs, provides a definition of polarization ex-ante for rationally inattentive agents, and demonstrates how such polarization can take place.

2.1 Description of the setup

There are two agents $j \in \{A, B\}$. Each agent independently faces a problem of discrete choice between two options. The first option, which we refer to as *a new policy*, provides a common payoff $v_s \in \mathbb{R}$ that depends on the realized state of the world $s \in S = \{1, \dots, n\}$, where $n \in \mathbb{N}$, $n \geq 3$. When we say that the state is realized, we mean that the potential outcome of the new policy has become possible to be evaluated. The states are labeled in ascending order $v_1 < v_2 < \dots < v_n$. The second option, which we refer to as *a status quo policy*, yields a known fixed agent-specific payoff: $R^j \in \mathbb{R}$, $\forall j$ respectively. We assume that $v_1 < R^j < v_n$, $\forall j$ in order to exclude trivial cases.¹⁵

The agents are uncertain which state of the world s is going to be realized and we denote each agent's j prior belief as a vector of probabilities $\mathbf{g}^j = [g_1^j \ g_2^j \ \dots \ g_n^j]^T$, where $\mathcal{P}^j(s = l) = g_l^j$, $\forall l \in S$; $\sum_{l=1}^n g_l^j = 1$ and $g_l^j > 0$, $\forall l \in S$. We model the agents to be

¹⁵If $R^j \leq v_1$, the safe option is weakly dominated by the risky option, and if $R^j \geq v_n$ the risky option is weakly dominated by the safe option. In both of these cases, the agent j does not have incentive to acquire information about the realization of the state of the world.

rationally inattentive in the fashion of Sims (1998, 2003). Prior to making their decisions, the agents have a possibility to acquire some information about the actual value of the new policy, which is modeled as receiving a signal $x^j \in \mathbb{R}$. The distribution of the signals, $f^j(x^j, s) \in P(\mathbb{R} \times S)$, where $P(\mathbb{R} \times S)$ is the set of all probability distributions on $\mathbb{R} \times S$, is subject to the agent's j choice. Upon receiving a signal, each agent updates her belief using Bayes rule. However, observing a signal is costly and we assume the cost κ^j to be proportional to the expected reduction in entropy¹⁶ between the agent's j prior and posterior beliefs.

Upon receiving a signal, each agent chooses an action, and her choice rule is modeled as $\sigma^j(x^j) : x^j \rightarrow \{\text{new policy, status quo}\}$. The agent's j objective is to maximize the expected value of the chosen policy less the cost of information. Given the updated belief, the agent j chooses the action with the highest expected payoff. The timing of the decision problem is depicted in Figure 1.

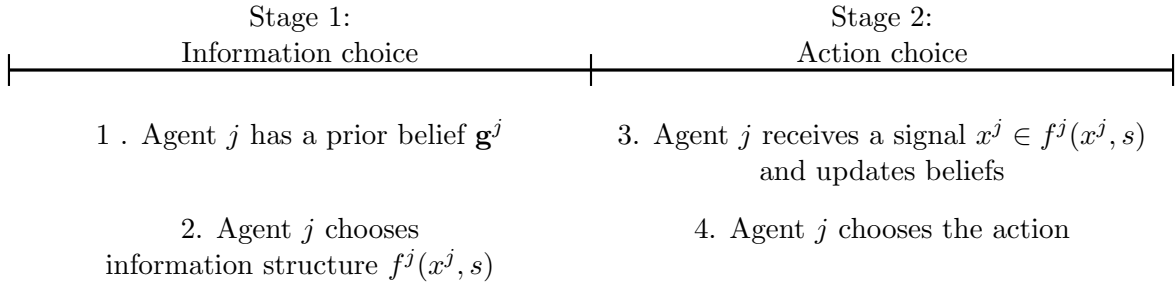


Figure 1: Timing of the events in the problem. The decision problem consists of two stages: an information strategy selection stage and a standard choice under uncertainty stage.

2.2 Agent's problem

The information strategy of agent j is characterized by the collection of conditional probabilities of choosing option i in state of the world s : $\mathcal{P}^j = \{\mathcal{P}^j(i|s) \mid i = 1, 2; s \in S\}$, where $i \in \{\text{new policy, status quo}\} = \{1, 2\}$ denotes the option.¹⁷ Each agent $j \in \{A, B\}$ independently solves :

$$\max_{\{\mathcal{P}^j(i|s) \mid i=1,2; s \in S\}} \left\{ \sum_{s=1}^n (v_s \mathcal{P}^j(i=1|s) + R^j \mathcal{P}^j(i=2|s)) g_s^j - \lambda \kappa^j \right\}, \quad (1)$$

¹⁶The entropy $H(Z)$ of a discrete random variable Z with support $\mathcal{Z}$ and probability mass function $p(z) = Pr\{Z = z\}, z \in \mathcal{Z}$ is defined by $H(Z) = -\sum_{z \in \mathcal{Z}} p(z) \log p(z)$. For detailed treatment of entropy, see, for example, Cover and Thomas (2012).

¹⁷According to Lemma 1 from Matějka and McKay (2015), the choice behavior of the rationally inattentive agent can be found as a solution to a simpler maximization problem that is stated in terms of state-contingent choice probabilities alone.

subject to

$$\forall i : \mathcal{P}^j(i|s) \geq 0 \quad \forall s \in S, \quad (2)$$

$$\sum_{i=1}^2 \mathcal{P}^j(i|s) = 1 \quad \forall s \in S, \quad (3)$$

and where

$$\kappa^j = - \underbrace{\sum_{i=1}^2 \mathcal{P}^j(i) \log \mathcal{P}^j(i)}_{\text{prior uncertainty}} - \sum_{s=1}^n \left(- \underbrace{\left(\sum_{i=1}^2 \mathcal{P}^j(i|s) \log \mathcal{P}^j(i|s) \right)}_{\text{posterior uncertainty in state } s} g_s^j \right). \quad (4)$$

$\mathcal{P}^j(i)$ is the unconditional probability that option i will be chosen and is defined as

$$\mathcal{P}^j(i) = \sum_{s=1}^n \mathcal{P}^j(i|s) g_s^j, \quad i = 1, 2.$$

Here κ^j denotes the expected reduction in entropy between the prior and the posterior beliefs about the choice outcome, $\lambda > 0$ is the unit cost of information and thus, $\lambda \kappa^j$ reflects the cost of generating signals with different precision.

2.3 Description of beliefs evolution

The main aim of this paper is to describe the evolution of the agents' beliefs, represented by the expected payoff of the new policy, in each state of the world. In order to exclude situations in which agents decide not to acquire information we assume that $0 < \mathcal{P}^j(i = 1) < 1$ ¹⁸, $\forall j \in \{A, B\}$. Let us first introduce the main objects used in our analyses.

Prior expected value. The uncertainty in this model is about the realized state of the world and thus about the actual payoff of the new policy. Without the information acquisition stage of the problem the agent would choose the option based on the comparison of the status quo payoff R with the agent's j *prior expected value of the new policy* being

¹⁸See Caplin, Dean and Leahy (2019) for the characterization of the necessary and sufficient conditions for solution of the discrete rational inattention problems.

$$\mu^j = \mathbb{E}^j v = \sum_{s=1}^n v_s g_s^j.$$

Posterior conditional expected value. In order to judge how the expected payoff from the new policy changes after the signal is received and the option is chosen, we take the position of an external observer. The observer knows that a realized state of the world is s^* and is interested in the agent's posterior expected belief about the payoff of the new policy v given the realized state s^* . Note that the agent's posterior belief is given by the signal she receives and thus the observer not only wants to know what the expected posterior belief is for a given signal, but is interested in the expected posterior belief about the new policy on average across all possible signals the agent may receive. Since there is a one to one mapping between the selected information structure and consequently chosen action, the posterior expected belief of interest is

$$m^j(s^*) = \mathbb{E}_i^j [\mathbb{E}^j(v|i)|s^*] = \sum_{i=1}^2 \left(\sum_{s=1}^n v_s \mathcal{P}^j(s|i) \right) \mathcal{P}^j(i|s^*),$$

where option $i \in \{1, 2\} = \{\text{new policy, status quo}\}$.

Polarization ex-ante. Our objective is to study whether the agents can become polarized after receiving new information and how the agents' disagreement evolves. Let us thus first denote the difference between prior and posterior beliefs, represented by expected values, for agent j , $j \in \{A, B\}$ in the realized state $s^* \in S$ as

$$\Delta^j(s^*) = m^j(s^*) - \mu^j.$$

Definition 1. We say that two agents $j \in \{A, B\}$, who are characterized by the pair $(R^j, \mathbf{g}^j)$ and are choosing between actions $i = \{1, 2\}$, become *polarized ex-ante* in the state $s^* \in S$ when the following two conditions are satisfied

1. $|m^A(s^*) - m^B(s^*)| > |\mu^A - \mu^B|$.
2. $\Delta^A(s^*) \cdot \Delta^B(s^*) < 0$.

The condition 1 secures that the expected posterior beliefs in the state s^* of two agents

are further apart than the prior beliefs are ¹⁹ whereas the condition 2 ensures that the agents update their beliefs in the opposite directions in the state s^* . Importantly, we denote such polarization as polarization ex-ante because the measure $m^j(s^*)$, $\forall s^* \in S$ takes the expectation of the expected value conditional on the realized state over all possible signals and as it was shown in the rational inattention literature²⁰ there is a one-to-one mapping between the selected information structure and chosen action i . Consequently, once the agents select the joined distribution of the signals and states, we can say that the agents become polarized in expectation even before the signal is realized.

In the following theorem, we provide conditions for the presence of the states of the world in which the agents become polarized.

Theorem 1. *Two agents $j \in \{A, B\}$, who are characterized by the pair $(R^j, \mathbf{g}^j)$ and for whom one of the following conditions holds:*

$$(a) \mu^A \neq \mu^B \text{ and } (\mu^A - \mu^B)(v_{s^*} - R^A) > 0 \wedge (\mu^A - \mu^B)(v_{s^*} - R^B) < 0;$$

$$(b) \mu^A = \mu^B \text{ and } (v_{s^*} - R^A)(v_{s^*} - R^B) < 0;$$

become polarized ex-ante in the state of the world $s^ \in S$.*

Intuitively, the condition $(\mu^A - \mu^B)(v_{s^*} - R^A) > 0$ means that if, for example, agent A has a higher prior belief about the payoff of the risky option than agent B ($\mu^A - \mu^B > 0$), then her belief should go up after information acquisition ($v_{s^*} - R^A > 0$). If the belief of agent B goes down at the same time ($v_{s^*} - R^B < 0$), then the two agents become polarized in their beliefs. Due to costly information acquisition, the rationally inattentive agent chooses only the necessary information in order to disentangle whether to select the status quo or the new policy. This leads to the *state pooling effect*, when the agent endogenously divides the states into two categories (the states in which the payoff of the new policy is higher than the payoff of the status quo and the states in which it is lower) and chooses only information that helps to disentangle which of these two categories the realized state s^* is from. It is also worth noting that the set of states where the agents become polarized are those states where payoffs are neither very high nor very low.²¹

¹⁹The same measure of disagreement is used in the analyses of Kartik, Lee and Suen (2021). For axiomatic foundations of disagreement measures see Zanardo (2017).

²⁰See, e.g., Matějka and McKay (2015).

²¹For instance, it is clear that in case (b) of Theorem 1 the agents become polarized in the states of the world in which $v_{s^*} \in (R_A, R_B)$. Given the structure of payoffs, such states s^* are intermediate.

The rationally inattentive agents do not always diverge in their beliefs, but their beliefs in some situations can converge or they can become further apart in their beliefs while updating towards the same extreme state. We describe all these scenarios in Appendix B.

3 Experimental design

Our theoretical results show that the agent-specific value of the status quo determines the information structure selected. As a consequence, two agents with different values of the status quo might become polarized in expectations. In particular, the rationally inattentive agent chooses to learn whether the outcome of the new policy is better or worse than the status quo, and not to learn the exact state-dependent outcome of the new policy. We have denoted this information strategy as state pooling behavior.

The main results of the model rely on several assumptions about a decision maker’s preferences (risk neutrality), ability to estimate probabilities (by correctly updating beliefs), and motivation (information has a purely instrumental value). The experimental literature reports a large amount of evidence that casts doubts on human ability to perform these tasks as accurately as the theory requires, and highlights that belief divergence could be mitigated or enhanced by human biases. We are interested in testing whether belief divergence in expectations, the main result of our model, can occur in a lab setting and whether behavioral components enhance or mitigate its magnitude.

In the following sections of the paper we investigate whether our normative model is also accurate in describing human behavior. We do so by running a lab experiment in which participants are allowed to collect information before making choices under uncertainty. We collect actions and beliefs separately, and combine them to compute a cardinal indicator for beliefs divergence and to compare human behavior and theoretical predictions. Our design allows us to test whether a manipulation of the status quo affects the choice of information source and generates belief polarization in expectations.

The experiment is quite demanding in terms of complexity of the setting, structure of the decision process, and length of the instructions. We take several steps in order to improve the subjects’ experience and reduce potential sources of confusion (details in Appendix C). In particular, before the main task – task 2 – we use a simplified training task – task 1 – to make sure subjects familiarize themselves with the environment and the interface. Our main exposition and results refer to task 2 (and subsequent control tasks). Additional insights from task 1 results are presented at the end of the results section and in Appendix I.

3.1 Overview of the experimental design

The experiment comprises four tasks and a final questionnaire. In the first and second tasks, subjects face a binary choice between the opaque box (*risky action*), which contains a single “color ball,” the value of which depends on the unknown color, and the transparent box (*safe action*), containing a single ball whose value is known. The color ball is randomly drawn from a box containing three balls (*states*) with different colors (red, yellow, blue) with uniform probability of being selected. The two tasks differ in the way we provide interim information about the color ball. In task 1 four possible advisors (representing degenerate signal structures) are evaluated separately and subjects report their willingness to accept (Becker-DeGroot-Marschak method,²² BDM thereafter) renunciation of the signal. In each trial for task 2 only two advisors are displayed and the subject makes a binary choice between them. In tasks 3 and 4 we elicit unconditional and conditional beliefs for different advisors, assigned exogenously. We ensure incentive compatibility by paying subjects for a single decision randomly selected from the entire experiment. Subjects never receive feedback about their decisions until the very end of the experiment. Each subject participates in all of the following tasks, in the order listed below.

3.2 Task 1 - Colorblind advisor game

In each round of task 1 subjects (i) choose an action contingent on the advisor and signal received (Figure 2a) and then (ii) indicate for each advisor the willingness to accept renunciation of its signal (Figure 2b). This task is presented at the beginning of the session to familiarize participants with the choice environment and the notion of the advisor.

Subjects play 10 rounds with the same four advisors and different lottery return values. Three of the advisors in this game (named Red, Yellow, and Blue) are described as colorblind to all colors except the subject’s own. They are able to observe the color ball and report truthfully whether it matches her own name’s color.²³ The fourth advisor is named Rainbow and reports every color accurately, without uncertainty. In each round, the subject chooses which box she would pick in each hypothetical advisor/answer scenario (strategy method). Then, the subjects fill out a multiple choice list for each of the four advisors, choosing between pairs of options: “Choice with the X advisor” (X is replaced with the advisor’s name) or “Choice without advisor + w extra points,” for w between 0 and 20 points, in 2

²²Becker, DeGroot and Marschak (1964).

²³For example, the Red advisor returns a signal RED or NOT RED, which is easy to interpret.

(a) Task 1, Screen 1: Action choice

(b) Task 1, Screen 2: WTA for each advisor

Figure 2: Task 1 - Colorblind advisor game. Left: Subjects choose an action (box) contingent on the advisor and signal received. The possible values of each action are indicated on the top of the screen. Each state (ball color) is equally likely to occur. Right: Subjects indicate for each advisor the willingness to accept renunciation of its signal in a series of binary choices (BDM method). At most one switch is allowed. Action choices selected in the previous stage are reported on the bottom of the screen.

points intervals.²⁴

3.3 Task 2 - Imprecise advisor game

In each round of task 2 subjects (i) choose one advisor between the two options available (Figure 3a) and then (ii) indicate the signal-contingent action for each signal (Figure 3b).

(a) Task 2, Screen 1: Advisor choice

(b) Task 2, Screen 2: Action choice

Figure 3: Task 2 - Imprecise advisor game. Left: Subjects choose one signal structure (advisor) between the two options available. Each advisor is a triplet of state-contingent signal probabilities. Right: Subjects indicate the signal-contingent action for each signal (strategy method).

²⁴The value w at which a subject i switches from preferring the former to the latter option reveals her subjective valuation w_i^i .

Subjects play 40 rounds with different pairs of advisors and values for the ball in the transparent box.²⁵ Each round comprises two parts. First, subjects observe a pair of advisors and make a binary choice to select which advisor they want to consult. Subsequently, only the selected advisor is consulted, a signal-contingent binary choice is implemented, and the participant chooses one box based on the signal received. Each advisor is defined as a triplet of state-contingent conditional probabilities of providing a binary signal.

The 40 rounds are designed as a combination of 20 advisor pairs and two values for the ball in the transparent box (safe option). The advisors are selected in order to examine preference over sources of information and formulate predictions about the effect of the safe option on information collection and posterior beliefs. In particular, 11 out of 20 pairs of advisors are designed such that a Bayesian agent would pick different advisors by changing the safe option.

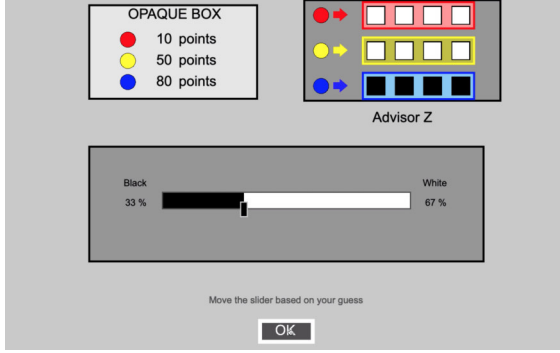
3.4 Task 3 - Card color prediction game

In each round of task 3 subjects indicate the likelihood of observing each signal for a given advisor (Figure 4a). We elicit the subjects' signal probability beliefs for each of the 20 advisors using a single slider with sensibility to the unitary percentage level. Each round contains a single advisor from those used in task 2 and subjects are asked to report the probability of a black or white card being shown. We incentivize accurate and truthful reporting by using the quadratic loss scoring rule with payoffs expressed in probability points (likelihood of winning the bonus prize).

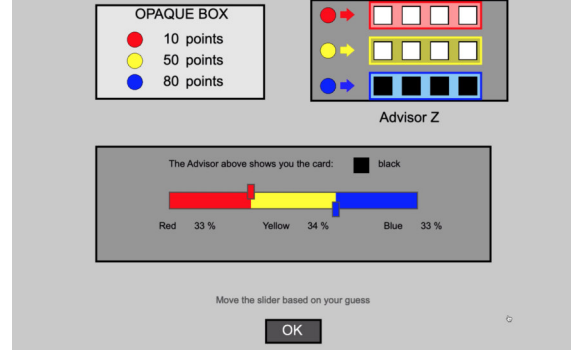
3.5 Task 4 - Ball color prediction game

In each round of task 4 subjects indicate the likelihood of each state given an advisor and signal (Figure 4b). We elicit the subjects' posterior probability beliefs for each of the 20 advisors and for each possible signal realization, using a double slider with sensibility to the unitary percentage level. Each round contains a single advisor from those used in task 2 and one realized signal (black or white card). The subject is asked to report the probability of a red, yellow, or blue ball being in the opaque box after observing the card color. We incentivize accurate and truthful reporting by using the quadratic loss scoring rule with payoffs expressed in probability points.

²⁵The ball in the transparent box can take two values (30 and 65 points). The values for the balls in the opaque box are unchanged during the task (10, 50, and 80 points, uniform probability of being drawn). Details on the pairs of advisors used in the 40 rounds can be found in Appendix E.



(a) Task 3: Beliefs over signal likelihood



(b) Task 4: Beliefs over state likelihood

Figure 4: Left: Task 3 (Card color prediction game). Subjects indicate the likelihood of observing each signal (card color) for the given advisor. Right: Task 4 (Ball color prediction game). Subjects indicate the likelihood of each state (ball color) given an advisor and signal. In both tasks subjects move the slider(s) and receive a number of probability points (chance of winning the bonus prize) according to the quadratic loss scoring rule described in the instructions.

3.6 Questionnaire

The final part of the experiment is a questionnaire designed to collect demographic variables (including field of study and familiarity with Bayes' rule), psychological measures (Life Orientation Test - Revisited,²⁶ LOT-R hereafter), risk attitude (Holt-Laury risk elicitation method with multiple price list, Holt and Laury 2002, HL hereafter), cognitive ability (five questions from the Raven Progressive Matrices Test, Raven hereafter), as well as questions on the subjects' strategy in the first and second task.²⁷

3.7 Procedure

The experiment was run in the CELSS (Columbia Experimental Laboratory for Social Sciences) between August and September 2019.²⁸ The experiment was coded in MATLAB (Release 2018b) using Psychtoolbox 3 (Psychophysics Toolbox Version 3). Eighty-five volunteers were recruited using the platform ORSEE²⁹ (Online Recruitment System for Economic Experiments) and were naive to the main purpose of the study. All subjects provided written, informed consent. The whole experiment took on average 85 minutes, including instructions and payment. On completion of the experiment, the subjects received payment

²⁶See Scheier, Carver and Bridges (1994).

²⁷We did not collect questionnaire information for the first 22 participants to the experiment.

²⁸The experimental protocol was approved by the Columbia Institutional Review Board, protocol AAAS5801.

²⁹See Greiner (2015).

in cash according to task performance. Each subject received a \$10 show-up fee, and played for a bonus prize of \$15. In addition, a subject could earn between \$0.10 and \$4 in the risk elicitation task and \$0.50 for each question of the cognitive test, up to \$5. The average payment was \$25. A small number of participants were recruited for each laboratory session (6 on average) in order to facilitate clarification questions during the experiment.

3.8 Hypotheses

Our experiment allows us to test the main predictions of the model, as well as disentangle the possible factors that mitigate or enhance the results with respect to the behavior of an optimal decision maker. The first hypothesis refers to the crucial effect of the status quo on information acquisition.

Hypothesis 1. *A change in the status quo (safe option) generates a reversal in the choice between advisors when such a reversal is optimal.*

We test this hypothesis by collecting choices over information structures under different values of the safe option. The optimal advisor choice would not be sufficient to generate belief polarization. The second hypothesis refers to the result that appears in the title of the paper.

Hypothesis 2. *A change in the status quo (safe option) generates beliefs polarization in expectations.*

In order to test this hypothesis we need to collect additional data, on top of advisor choices and final actions. For this reason, we collect subjective beliefs after receiving a signal.³⁰ There are three main channels that may mitigate or enhance the results with respect to the behavior of an optimal decision maker: non-standard preferences over realized states, biased beliefs, and non-standard preferences for information structures. We summarize these potential confounding factors in three channels.

Channel 1. *Subjects evaluate information structures based on non-standard preferences over the realized state, e.g. risk aversion or risk seeking.*

Channel 2. *Subjects evaluate information structures based on subjective beliefs about the likelihood of signal and state realizations.*

³⁰An alternative approach to construct polarization would be to elicit subjective evaluations of the risky action after receiving a signal. The elicitation of subjective beliefs has the advantage of providing additional information on the belief channel, one of the possible confounding factors for the emergence of polarization.

Channel 3. *Subjects evaluate information structures based on non-instrumental characteristics, including ease of processing of the signals (certainty, few possible outcomes).*

Our setup allows us to test the main hypotheses and study the three behavioral channels by collecting detailed data about willingness to pay and binary choices between information structures.

3.9 Experimental investigation

We consider a setting with three possible states and two actions that generate state-contingent payoffs. The actions represent two policies: the current policy (the status quo), whose return $R \in \mathbb{R}$ is known and independent of the state, and a new policy, whose return $v_s \in \mathbb{R}$ is uncertain. The state of the world $s \in S = \{r, y, b\}$ is represented by a color associated with the deterministic return for the uncertain policy: r (red, low return), y (yellow, intermediate return), or b (blue, high return), with $0 < v_r < v_y < v_b < 100$ and $v_r < R < v_b$. An agent with a correct uniform prior belief $\mathcal{P}(s) = \frac{1}{3}$, $\forall s$ observes an informative signal about the state and selects one of the two policies. The return of own choice depends on the selected action and the realized state, and represents the probability, expressed in percents, of receiving a fixed prize k (\$15 in our laboratory experiment).

Information is valuable because it informs the subsequent binary choice between policies. We let $\sigma \in \{0, 1\}$ denote the realization of a stochastic signal that the subject may observe.³¹ Since we have three states and two possible signal realizations, a signal structure is a triplet of state-dependent probabilities $\mathcal{P}_I(\sigma = 1|s)$.³² We will refer to such a triplet I as an information source or advisor.³³

The Bayesian agent represents a natural benchmark to consider the objective value of information in this environment. Let $V(I)$ denote the bonus that renders the agent indifferent between playing the game without additional signals³⁴ (but receiving additional $V(I)$ “tickets”) and playing the game with the signal structure I . The valuation of the information

³¹In the main task of the study, all advisors have two signal realizations. In the first task, the fully revealing (*rainbow*) advisor has three possible signal realizations $\sigma \in \{0, 1, 2\}$.

³²Notice that the signal is deterministic in the case of a triplet of degenerate probabilities.

³³Notice that even though the three states are equally likely, the two signals need not be equally likely.

³⁴Playing without any additional information is, from a theoretical perspective, equivalent to playing with a purely noisy signal. We prefer to refer to the former case for the sake of clarity.

structure I is given by a chosen lottery and by the observed signal

$$V(I) = \sum_{\sigma \in \{0,1\}} \max \left\{ \underbrace{\sum_{s \in S} v_s \mathcal{P}_I(s|\sigma), R}_{=V(\sigma|I)} \right\} \mathcal{P}_I(\sigma) - \max \left\{ \underbrace{\sum_{s \in S} v_s \mathcal{P}(s), R}_{=V(\emptyset)} \right\}, \quad (5)$$

where $V(\emptyset)$ is the expected value of the action chosen without observing any signal and $V(\sigma|I)$ is the expected value of the action chosen after receiving signal σ from advisor I .

We can generalize the subjective valuation in order to include non-instrumental preference over information. A decision maker i has a subjective valuation $V^i(I)$ of the signal structure I that depends both on the instrumental value $V(I)$ and other characteristics of I ; for example the type of “optimistic/pessimistic” information that it provides. We postpone further discussion about possible differences between Bayesian and subjective valuation of information to the results section.

4 Experimental results

This section contains the main results of the experimental investigation. We report aggregate choices between sources of information (Section 4.1) and provide evidence that (i) subjects do react to the value of the status quo as predicted by the theoretical model, (ii) the variation in the status quo leads qualitatively to the belief polarization in our laboratory setting and (iii) we observe preference for state pooling information structures.

In section 4.2, we discuss to what extent the main behavioral channels introduced in section 3.8 cause deviations of the participants’ behavior from optimality. We verify that subjects’ actions are consistent with the optimal behavior of a risk-neutral agent and that their beliefs about the likelihood of signal and state realizations are close to the optimal ones. Consequently, we identify the evaluation of information structures based on non-instrumental characteristics as a major driver of the observed deviations from optimality.

Finally, we analyze the willingness to pay (WTP) for information structures in the first task (Section 4.3), where we observe that (iv) subjects display compression in their WTP and (v) are willing to pay higher amounts for information about the most desirable state.

4.1 Beliefs polarization in the laboratory experiment

Our model predicts that a change in the status quo creates belief polarization because of the endogenous choice of information structures (advisors). In our experimental design

this means that, given the true state, the same decision maker will have different beliefs (conditional on the state realization, but before the signal realization) based on her status quo value. The experiment contains 11 pairs of trials that can be used to test whether such polarization occurs. We combine the data collected for the binary advisor choice (task 2) with the subjective beliefs about posterior distribution (task 4) to calculate the magnitude of the observed polarization in the laboratory experiment.

The first hypothesis is that a change in the status quo generates a reversal in the advisor choice when such a reversal is optimal. Figure 5a shows that the hypothesis is confirmed in the trials from Task 2 where we expect to observe the reversal. Advisor I_1 , defined here as the best one under $R = 30$, is chosen 66% of the times when it is optimal to do so, and only 30% when $R = 65$. The difference between the two treatments is large and significant, and confirms our first hypothesis. In Figure 5b we depict the probability of choosing the advisor (out of the pair of advisors labeled I_1 and I_2) as a function of the difference in the instrumental values of the two presented advisors, with all the values computed as in equation 5. The probability of selecting advisor I_1 increases with the difference between the instrumental values of advisor I_1 and advisor I_2 . All trials but one lie in the first and the third quadrant; in lay terms, the advisor with the highest instrumental value is chosen more often. Choice probabilities increase almost linearly and not stepwise near the zero, suggesting that subjects do not respond only to the sign of the difference in the instrumental values between the advisors, but to the actual value of the difference.

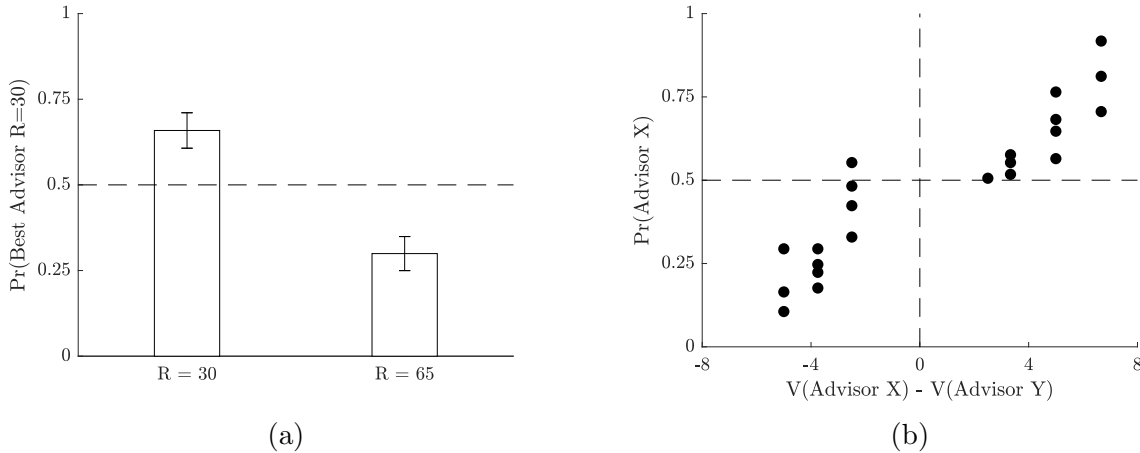


Figure 5: Advisor selection probability in Task 2. Left: Probability of selecting I_1 , defined here as the best advisor under $R = 30$, for each of the two treatments (11 trials per treatment, 935 observations per treatment). Right: Probability of selecting advisor I_1 based on the difference between instrumental values (22 trials, 85 observations per trial).

Experimental Result 1. *Subjects systematically react to the value of the status quo and choose the optimal advisor (information structure).*

The advisor choice reversal is necessary but not sufficient to generate polarization. If participants are systematically biased in the computation of probabilities, and react less (more) to the new information provided by the signal, this will reduce (increase) the magnitude of the polarization. We elicit subjective beliefs with a strategy-proof mechanism and we verify that the participants hold, on average, very accurate beliefs. In the fourth task of the session we elicit the posterior likelihood of each state, given an advisor and a signal realization. On average, we observe accurate probability estimates, close to the predictions of an optimal Bayesian agent. Results are displayed in Appendix H and show that 1) participants are on average accurate in the estimate of probabilities, 2) we do not observe a systematic difference between estimates involving different states (i.e., we do not have evidence of motivated beliefs, Bénabou, 2015), and 3) the results show mild evidence of conservatism (central tendency of judgement), as vastly reported in experiments with subjective estimates (Hollingworth, 1910; Anobile, Cicchini and Burr, 2012). A linear fit of the average subjective estimates over the unbiased ones confirms the mild conservatism (slope $\beta = 0.825$) and the overall good fit ($R^2 = 0.993$).

By combining advisor choices and posterior beliefs we can finally calculate the magnitude of polarization of beliefs. We consider separately each of the 11 pairs of trials in which the agents should switch advisor because of the status quo manipulation. For every trial and state, we calculate the ex-ante expected value for the risky action (conditional on the state, but not conditional on the signal realization). The *predicted polarization* is computed based on the Bayesian agent’s behavior. It is the absolute difference between the two ex-ante expected values: one is generated under the low status quo, the other under the high status quo. The *realized polarization* (Figure 6a) is calculated in the same way but it includes the subjects’ behavior in two stages. First, we replace the posterior beliefs (Bayesian) with the average subjective ones from task 4. Second, we replace the advisor choice probabilities (deterministic) with the observed ones from task 2. Participants switch advisors less than predicted, and update their beliefs slightly less than predicted, so both the components reduce the magnitude of the realized polarization. A linear fit of the distribution shows that the realized polarization is, on average, 32% of the predicted one, with little dispersion across pairs of trials, as confirmed by the high $R^2 = 0.892$.

Experimental Result 2. *Variations in the value of the status quo generate ex-ante belief divergence (before the signal realization, and after controlling for the true state) qualitatively*

analogous to those predicted by the model, but with smaller magnitude.

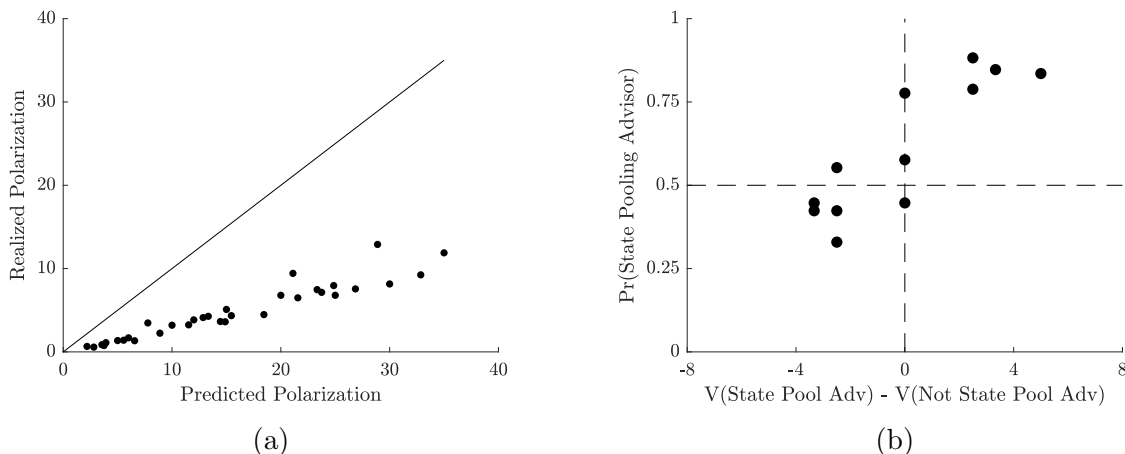


Figure 6: Left: Polarization. Predicted polarization (for the Bayesian decision maker) and realized polarization (based on subjects' responses, $n=85$) in the 11 pairs of trials with predicted advisor switches (3 states per pair of trials). Right: Probability of choosing the state pooling advisor over the alternative one in the trials with exactly one state pooling advisor. (12 trials, 85 observations per trial).

State pooling is the key mechanism for our model that determines the advisor switch. We define state pooler advisors as follows.

Definition 2. State pooler advisor. An advisor with information structure I is a *state pooler* under status quo value R when it can provide a signal σ such that the posterior belief for the agent is either $\mathcal{P}(\pi_s \geq R|\sigma) = 0$ or $\mathcal{P}(\pi_s \geq R|\sigma) = 1$.

When we investigate how the probability of choosing the state pooling advisor depends on the instrumental value of such an advisor (see Figure 6b), we can notice that the probability is greater than $1/2$ when it is optimal to select the state pooler, and otherwise it is below $1/2$. However, we can also notice that the probability of selecting the state pooler increases with the instrumental value. When the state pooling and non state pooling advisors both have the same instrumental value (0 on x-axis), subjects strictly prefer the state pooling advisor in comparison to a non state pooling advisor, even though it is not more informative.

In the next section we further discuss the quantitative departure from the theoretical prediction and we analyze separately the three main behavioral channels that represent potential confounding factors.

4.2 Behavioral channels of departure from theoretical predictions

Possible deviations from the optimal choice between alternative sources of information in task 2 (imprecise advisor game) can be rationalized by risk preferences, systematically biased beliefs, and non-instrumental preferences over information structures, as discussed in Section 3. We now consider each of these channels separately and discuss whether they provide an explanation for the advisor choices observed.

Risk preferences represent the first behavioral channel. Possible deviations from the optimal choice between alternative sources of information in the imprecise advisor game (task 2) can be consistent with risk aversion, as this would affect the evaluation of each advisor. We design our experiment in order to minimize this concern, by using probability points as prizes (see, e.g., Harrison, Martínez-Correa and Swarthout, 2013) and assigning probabilities between 10% and 90% (to minimize concerns about extreme probability events).

Our analysis confirms that risk aversion does not represent a driver of the participants' departure from the predictions: we summarize below the main results, with more analysis available in Appendix G. In task 2, given an advisor and a signal realization, participants were asked to choose one action corresponding to the risky lottery (opaque box) or the safe option (transparent box). If we assume that the decision maker adopts a CRRA utility function $u_\alpha(x) = \frac{x^{1-\alpha}}{1-\alpha}$, and we use maximum likelihood method to estimate the concavity of the utility function, represented by the risk aversion coefficient α , we obtain $\hat{\alpha} = 0.34$, which indicates a small risk aversion, compared to the null hypothesis $\alpha = 0$ (risk neutral agent). We test the statistical significance of adding the risk aversion coefficient in the action choice model using the likelihood ratio test, and we reject the null hypothesis ($p < 0.001$). The use of risk preferences slightly improves the predictive power, which we measure using the likelihood ratio (Cohen et al., 2013). This measure, representing an R squared statistics for logistic regressions, is $R^2_{\text{risk neutral}} = 0.382$ using risk neutrality and $R^2_{\text{risk averse}} = 0.422$ using risk aversion. Nevertheless, the magnitude of the deviation from risk neutrality is modest and unable to explain a choice reversal for the advisor selection.

Experimental Result 3. *Participants behave similarly to a risk neutral agent, and risk preferences represent a small driver of deviation from optimality in the choice between actions.*

Biased beliefs represent the second behavioral channel. We usually assume that the instrumental value of each advisor is determined using the exact probabilities for signal and state realizations. We relax this assumption, and consider the behavior of an agent who

selects the advisors based on the subjective value instead. We compute the subjective advisor value by using the subjective beliefs about signal likelihood and state likelihood (conditional on the observed signal) that we elicit in tasks 3 and 4, respectively. The average responses for these tasks are shown in Appendix F, and both distributions display small conservatism in the estimates.

In Figure 7a we depict the advisor choice probability (from a pair of advisors labeled I_1 and I_2) as a function of the difference in the instrumental values of the two presented advisors, with all the values computed as in equation 5 and using unbiased beliefs. This is similar to Figure 5b, but now we display all 40 trials in the experiment, including those where we do not expect to observe the advisor switch. We compare this result with Figure 7b, where the instrumental value of each advisor is computed using the average subjective beliefs ($V_{\text{subjective}}$). The use of subjective beliefs reduces the predictive power of the model, which we measure using the likelihood ratio (pseudo R squared, as above): $R_{\text{unbiased}}^2 = 0.563$ using unbiased beliefs and $R_{\text{subjective}}^2 = 0.225$ using subjective beliefs. Results from tasks 3 and 4 show that subjects’ beliefs are on average accurate (see Appendix H). The lower predictability of advisor choices suggests that there is no systematic departure from unbiased beliefs (that would affect both beliefs advisor evaluation). Instead, the evidence points towards independent noise in the responses in the different tasks.³⁵

Experimental Result 4. *Participants behave similarly to an agent with correct understanding of probabilities. Subjective beliefs about the likelihood of signal and state realizations represent a small driver of deviation from optimality in the choice between advisors.*

Non-instrumental preferences over information structures represent the third behavioral channel. Differently from the previous two approaches, we now assume that the decision makers evaluate some non-instrumental features of the advisor (e.g. simplicity and certainty) and chooses the preferred advisor based on the tradeoff between instrumental and non-instrumental characteristics. This is equivalent to the addition of a cognitive cost for each of the signal structures. Ambuehl and Li (2018) show that, in a two-states, two-signals setting, subjects “disproportionately prefer information structures that may perfectly reveal the state of the world” (certainty effect). Since we are considering a more complex three-states, two-signals setting we separate this effect using two separate measures: certainty and

³⁵A possible explanation for the result shown in Figure 7 is that the elicitation of beliefs is noisy and that this noise is not systematic across the main task (advisor choice) and the elicitation task. If the beliefs are noisy yet unbiased (or, more generally, if the systematic bias in beliefs is small enough), then the adoption of subjective beliefs is adding noise in the estimation of the advisor evaluation and decreasing the predictive power.

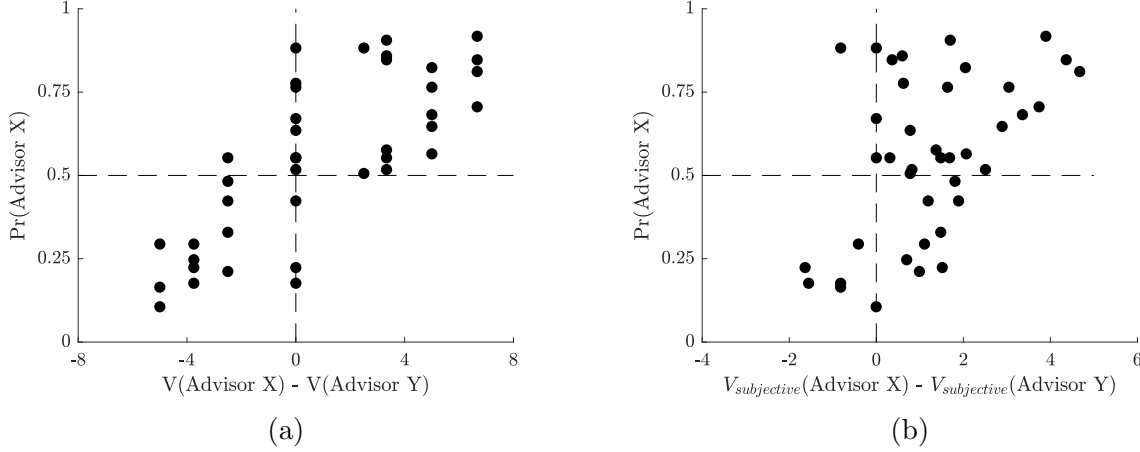


Figure 7: Task 2: Advisor selection probability (all trials). Left: Probability of selecting I_1 based on the difference between instrumental values (40 trials, 85 observations per trial). Right: Instrumental values of the advisors are calculated using the average subjective beliefs elicited in tasks 3 and 4.

simplicity. According to the definitions we introduce below, a *certain* advisor always removes uncertainty about a specific state, whereas a *simple* advisor refine the posterior beliefs by removing impossible states.

We now define formally these concepts and show how they relate to observed behavior.

Definition 3. Certain advisor. An advisor with information structure I is *certain* when there exists a state s and a signal σ_s such that $\mathcal{P}(s|\sigma_s) = 1$ and $\mathcal{P}(s|\sigma) = 0 \quad \forall \sigma \neq \sigma_s$.

In lay terms, a certain advisor answers a yes-no question about state s . It always provides certainty that state s is, or is not, the actual state. Figure 8a plots the probability of choosing the certain advisor given the difference between the values of the certain and uncertain advisors. It is apparent that the probability does not increase monotonically with the informativeness as in the previous figures. Once it is optimal to select the certain advisor, the probability of choosing it jumps to 86%, on average. On the other hand, when it is optimal to choose the uncertain advisor, the probability of selecting the certain advisor is, on average, only 46%.

In order to extend the effect to a larger set of conditions, we introduce here c_I as a discrete measure of complexity (or cost) for the signal structure I :

$$c_I = \sum_{\sigma} \left(\sum_s \mathbb{1}(\mathcal{P}(s|\sigma) > 0) - 1 \right). \quad (6)$$

We use this measure to compare advisors based on their complexity.³⁶

Definition 4. Simpler advisor. An advisor with information structure I is *simpler* than an advisor with information structure J if $c_I < c_J$.

In lay terms, a simpler advisor provides messages that can exclude more states. This index counts, for every distinct signal realization, how many states receive positive probability in the posterior beliefs generated by that signal. In task 2 of the study, it takes values between 1 (simplest) and 4 (most complex).³⁷ Figure 8b shows that participants tend to prefer advisors with a low index, even after controlling for the advisors’ value difference.

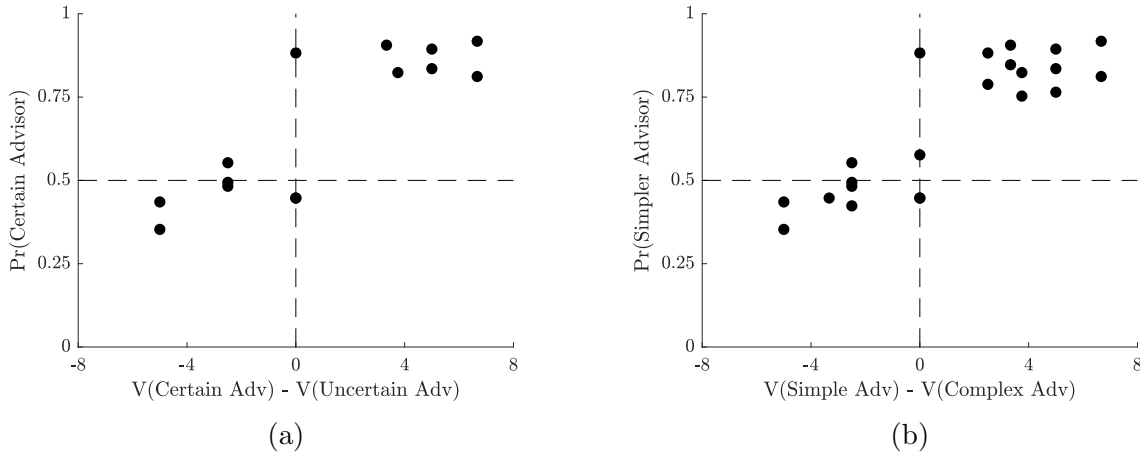


Figure 8: Preference for certainty and simplicity. Left: Certainty effect. Probability of choosing a certain advisor in the trials containing a certain and an uncertain advisor (14 trials, 85 observations per trial). Right: Simplicity effect. Probability of choosing the simplest advisor (as defined in equation 6) in the trials containing two advisors with different complexity scores (21 trials, 85 observations per trial).

We run a series of logit regressions aimed at measuring the relative importance of various features of the advisors. In addition to the instrumental value $V^{\text{Bayes}}(I)$, we use the complexity index c_I and three binary variables to indicate whether the advisor 1) is the best in the pair, 2) provides certainty, and/or 3) provides state pooling. The results from the regression

³⁶Examples of advisors with different complexity scores are as follows. Complexity 1 (lowest): Certainty advisor: 1-0-0 answers the question “is the state Red?” and always gives the correct answer. Complexity 2: Advisor 0.5-0-0 answers the question “is the state Red?”, if the state is not Red the advisor always replies “no”. If the state is Red, the advisor replies “yes” half of the time. Complexity 3: Advisor 1-0.5-0.5 answers the question “is the state Red?”, if the state is not Red, the advisor replies “no” half of the time, if the state is Red, the advisor always replies “yes”. Complexity 4: Advisor 0.25-0.5-0.75 returns the message “Blue” with a higher probability when the state is blue and with a lower probability when the state is red.

³⁷In task 1, the complexity score c_I takes values 0 for the fully revealing (rainbow) advisor and 1 for all the other (colorblind) advisors.

are presented in Table 1, and are consistent with the patterns previously discussed. The instrumental value of the advisor $V^{\text{Bayes}}(I)$ has a significant effect on the probability of the advisor being selected. However, certainty and state pooling are also significant.

Method: Logit, Dependent variable: Advisor choice				
	(1)	(2)	(3)	(4)
Value $V^{\text{Bayes}}(I)$	0.246*** (0.018)	0.217*** (0.011)	0.235*** (0.011)	0.232*** (0.019)
Best Advisor (dummy)	-0.084 (0.096)			-0.007 (0.102)
Complexity c_I		-0.359*** (0.037)		-0.074 (0.076)
Certainty (dummy)			0.511*** (0.069)	0.428*** (0.110)
State Pooling (dummy)			0.404*** (0.069)	0.330** (0.102)
Trials	All	All	All	All
Observations	3,400	3,400	3,400	3,400

Table 1: Advisor choice across all the trials in task 2. The coefficients refer to the differences between advisor I and the alternative advisor J . Statistical significance: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Experimental Result 5. *Participants significantly prefer simple advisors, and advisors providing certainty in particular. This result is robust even after controlling for the instrumental value of the advisors available. This is a major driver of the observed deviation from optimality in the choice between advisors.*

In order to determine the extent to which the subjects’ deviations from Bayesianism determine the demand for certain and simple information, we run a second set of regressions where we control for the subjective valuation of information sources $V^{\text{Subjective}}(I)$ that reflect the subjects’ deviations from Bayesianism. These results are presented in Table 2. While the deviations from Bayesianism do mute the observed polarization, they do not play such a significant role in affecting the demand for information. The general measure of complexity and the certainty dummy prove significant even after we control for non-Bayesianism. Thus we believe our results suggest an intrinsic preference for certainty/simplicity.³⁸

³⁸Note, however, that dummy for certainty and complexity measure are highly correlated, thus not being simultaneously significant if they are both included.

Method: Logit, Dependent variable: Advisor choice					
	(1)	(2)	(3)	(4)	(5)
Value $V^{\text{Bayes}}(I)$					0.232*** (0.019)
Value $V^{\text{Subjective}}(I)$	0.054* (0.029)	0.252*** (0.020)	0.250*** (0.020)	-0.034 (0.035)	-0.020 (0.036)
Best Advisor (dummy)	0.871*** (0.089)			1.085*** (0.112)	0.043 (0.136)
Complexity c_I		-0.434*** (0.035)		-0.113 (0.081)	-0.053 (0.085)
Certainty (dummy)			0.508*** (0.065)	0.571*** (0.118)	0.462*** (0.125)
State Pooling (dummy)			0.293*** (0.067)	0.218** (0.106)	0.355*** (0.112)
Trials	All	All	All	All	All
Observations	3,400	3,400	3,400	3,400	3,400

Table 2: Advisor choice across all the trials in task 2. The coefficients refer to the differences between advisor I and the alternative advisor J . Statistical significance: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

4.3 Willingness to pay for signal structures

The Colorblind advisor game introduced in task 1 provides a different dataset that we can compare with the results from the other tasks presented in a previous section. In each of the ten trials we collect signal-contingent actions (risky or safe options) as well as the subjective willingness to pay (WTP) in order to observe a certain signal structure. More precisely, we elicit the willingness to accept (WTA), expressed in probability points of winning the bonus, in exchange for the opportunity of playing the game without the advisor.³⁹ For each of the four advisors in the game we elicit subjects' valuation of the advisor using multiple price lists, an incentive-compatible implementation of the BDM mechanism. The red, yellow, and blue advisor provide a binary message, whereas the rainbow advisor fully reveals the true state. Figure 9 shows, for each advisor type, the relation between subjective (averaged across participants) and theoretical advisor value (for an optimal decision maker). We notice that for all the advisors the subjective evaluation tends to exceed the theoretical one (positive

³⁹Standard theory predicts that the difference between WTA and WTP is negligible when income effects are small. We implemented several measures in order to maintain this difference as negligible: i) we used probability points to minimize risk aversion concerns, ii) we maintain the gain domain during the whole session and iii) we avoid framing the payment in task 1 as a cost/loss. For an exact relation of WTA and WTP see Weber (2003).

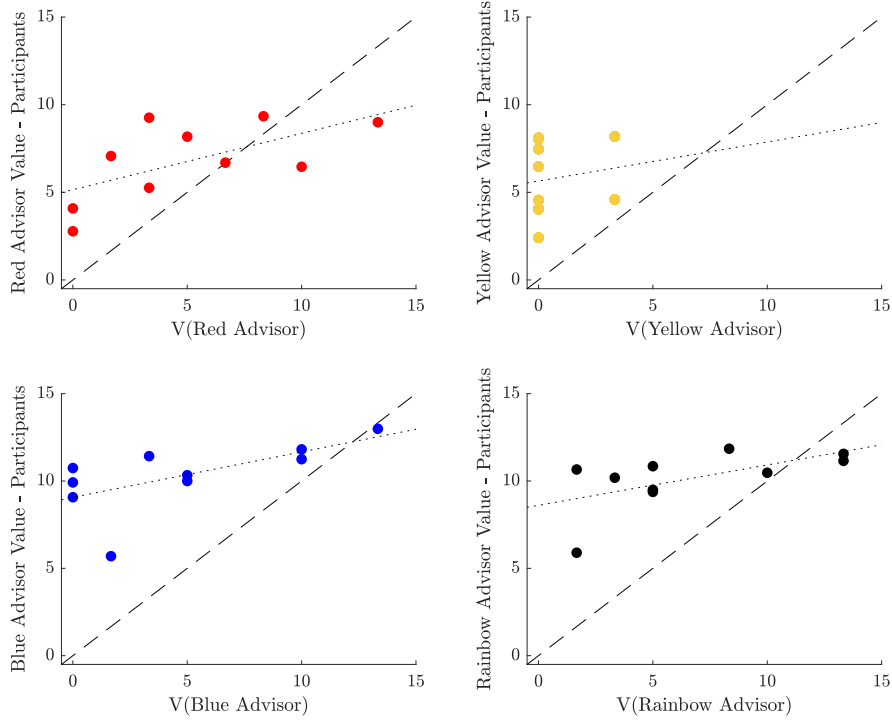


Figure 9: Willingness to pay for information (task 1). Comparison between the average subjective valuations of advisors across participants and the valuations for the optimal decision maker. The color used in each panel indicates the type of advisor: red, yellow, blue, and rainbow advisors (ordered from top-left to bottom-right). Optimal valuation (dashed lines) and linear regression estimates (dotted lines) are shown for comparison. $n=85$ observations per point in each panel (10 trials, 4 advisors per trial).

intercept) and there is a general positive relation between the two, with subjective values increasing with the theoretical ones, but not as much as the latter (slope coefficients lower than 1). This pattern is known as the compression effect and is well known in experiments with explicit elicitation of WTP for sources of information (e.g. Ambuehl and Li, 2018). We have several cases of advisors whose theoretical value is equal to zero (including for most of the yellow advisors): the observation of a signal from them is not pivotal for the chosen action with respect to the decision without an advisor, yet the subjects invest a significant amount of points to receive this piece of information.

Experimental Result 6. *Participants display a compression effect in their willingness to pay for information structures. They tend to overpay for advisors with low and even zero instrumental value, and their subjective WTP increases with the theoretical values, but with*

a slope smaller than one.

The comparison of the plots of different advisors highlights a consistent pattern. The compression effect appears similarly for all four advisors, with a similar slope of the linear regression between observed and theoretical values. At the same time, intercepts are significantly different, with similar values for red and yellow advisors, and much higher levels for blue and rainbow advisors. This difference is aligned with preference for information structure biased in favor of the most desirable (blue) state.

This result is confirmed by running a simple OLS regression of the subjective advisor value using the theoretical value as the regressor. Table 3 shows that the slope is positive but lower than one (compression effect) and the intercept is positive and significantly different from zero (a result analogous to the conservative probability estimates observed in tasks 3 and 4). When we allow the intercept to differ across advisors, we notice that they are not different between the red and yellow advisor, whereas the blue and rainbow advisors receive significantly higher WTPs. This result is consistent with those observed in environments with non-instrumental information (Masatlioglu, Orhun and Raymond, 2017) in which subjects display wishful thinking and desire to observe signals that are more accurate about the positive outcomes. The blue state represents the most desirable outcome in our setting, and it is fully revealed by consulting either the blue or the rainbow advisors. The slope of the curves is not significantly different across advisors (not reported in the table, see Appendix I) confirming that the effect does not arise from a different sensitivity to instrumental value. Instead, it provides evidence in favor of intrinsic (non-instrumental) preference for information structures, similarly to the previously discussed preferences in favor of advisors providing certainty or state pooling in the posterior beliefs.⁴⁰

Experimental Result 7. *Participants are willing to pay significantly higher amounts for advisors that provide evidence in favor of the most desirable state, as well as for advisors that fully reveal the true state.*

The result is qualitatively robust to the separate analysis of trials with high or low status quo. Columns 3 and 4 contain the regressions run independently, with the two parts of the dataset, divided based on the relative value of the status quo R with respect to the intermediate payoff v_y . Although the signs and significance of the estimates are unchanged,

⁴⁰In Appendix I we provide a further analysis of willingness to accept. Importantly, we present there also evidence that the participant demands the simplest information source that is sufficient for quaranteeing the optimal action, as is predicted by the theoretical model. We do so by comparing the WTP for the fully revealing signal and the highest WTP among the other advisors.

Method: OLS, Dependent var: $V^i(I)$				
	(1)	(2)	(3)	(4)
Constant	6.66*** (0.166)	5.62*** (0.231)	3.65*** (0.299)	7.58*** (0.339)
Value $V^{\text{Bayes}}(I)$	0.372*** (0.027)	0.269*** (0.030)	0.430*** (0.042)	0.116** (0.047)
Red advisor		-0.19 (0.353)	-0.31 (0.443)	0.26 (0.546)
Blue advisor		3.41*** (0.349)	3.51*** (0.496)	2.59*** (0.486)
Rainbow advisor		2.74*** (0.373)	2.70*** (0.496)	2.62*** (0.546)
Trials	All	All	$R > v_y$	$R < v_y$
Observations	2520	2520	1260	1260

Table 3: Aggregate valuations of information structures in task 1. Statistical significance: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

we observe different magnitudes. When the value of the status quo is higher than the intermediate state (column 3) the intercept is lower and the slope steeper (less compression). In this case the decision maker faces a safer choice problem and she reacts more to the incentive represented by the instrumental value of the advisor.

5 Heterogeneity across subjects

In this section, we present how the results of the main analysis are robust across individuals, and whether observable characteristics, such as mathematical literacy and risk attitude, can predict the heterogeneity in their behavior. We showed in Figure 6a that at the aggregate level we observe 32% of the predicted polarization. We replicate the analysis for every participant: Figure 10a shows the estimated polarization coefficient $\hat{p}_i$ for every subject, with 0 for no polarization and 1 indicating the magnitude of polarization predicted by our model. We observe substantial heterogeneity, with the full range of possible values, and an average polarization equal to 53% of the prediction. We even encounter a few values above 1; this is possible when subjective beliefs are characterized by over-reaction to evidence (instead of conservatism, as we observe at the aggregate level). In appendix K, we investigate what role in over-reaction to evidence is played by two channels influencing the polarization. We show that both non-Bayesian updating and ii) demand for information systematically decrease the empirical polarization, with the belief channel increasing empirical polarization

in only a few instances.

We previously noted that the polarization coefficient depends on advisor choices (task 2) and posterior beliefs (task 4). We separate the two components: even if subjective beliefs represent, on average, a small deviation from the predictions, the result changes at the individual level, and the average score jumps from 53% to 71%, reducing the missing polarization by over one-third.

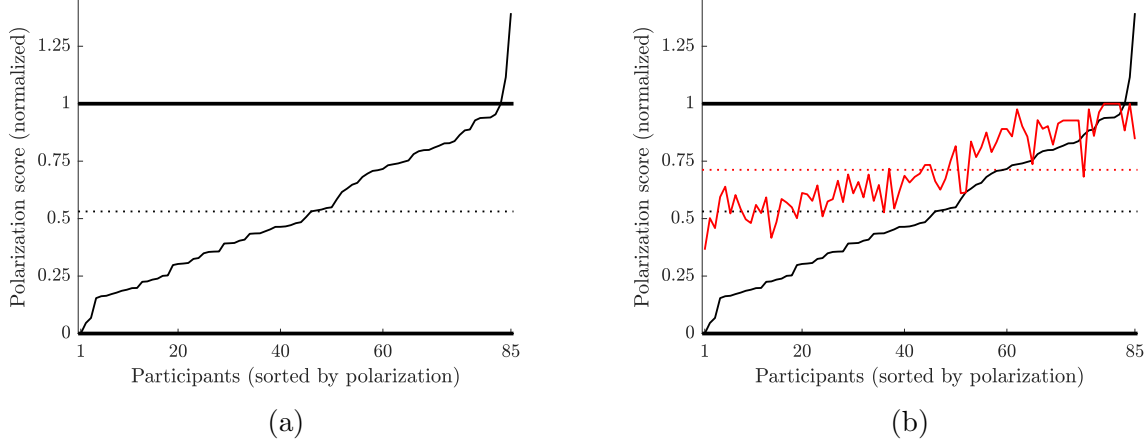


Figure 10: Estimated polarization coefficient $\hat{p}_i$ by subject. Left: Distribution of coefficients, subjects ordered by $\hat{p}_i$. Right: Decomposition of the missing polarization, replacing subjective beliefs with unbiased ones.

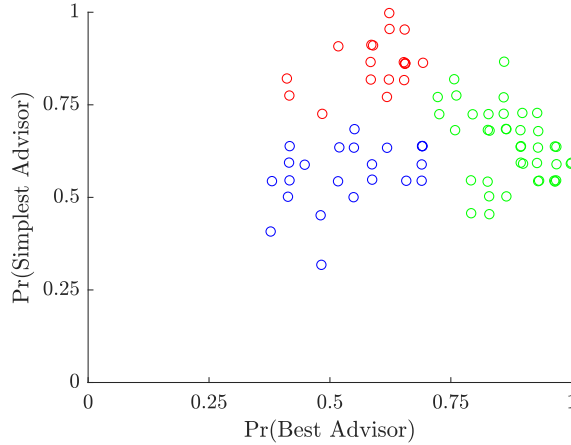


Figure 11: Clustered distribution of participants' advisor choices: the probability of choosing the best advisor (based on instrumental value) and simplest advisor (based on the complexity score).

Another way to recognize the vast heterogeneity in the subjects' behavior is to consider two dimensions of the advisor choice: the choice of the best advisor in terms of the instru-

mental value, and the choice of the simplest advisor using the complexity index introduced in Section 4.2. It could be the case that most participants suffer from a minor bias in favor of simple advisors, or that fewer subjects have strong preferences for them. Figure 11 provides some evidence in favor of the second explanation and suggests that the participants can be categorized into three broad groups based on these two dimensions.⁴¹ A cluster (green) of accurate participants that display little or no bias, a group (red) of simplicity-driven participants consistently selecting the advisor with lower complexity, and a group (blue) of participants whose advisor choices are close to random. Finally, we use the observable characteristics of the participants to predict the heterogeneity in the polarization score presented in Figure 10b .

6 Discussion and Conclusion

Opinions about proposed policies and pertinent issues often become polarized. The literature provides several explanations of the phenomenon, including (among others) preference for information which confirms existing beliefs, imperfect memory, and interpretation of ambiguous evidence as confirming existing beliefs. In this paper we explore a new source of belief polarization that arises as a consequence of the state-pooling effect when information is costly to acquire.

We find that the valuation of the status quo plays an important role in determining the direction of belief updating, as it directly affects the information acquisition strategy. In our interpretation, the agent partitions the states of the world into categories. This partition into categories is determined exactly by the valuation of the status quo. If the two agents have different valuations of the status quo, their information acquisition is such that they might diverge in their opinions *ex ante* (before the realization of the private signal).

The large number of assumptions required by the model may cast doubts on whether belief divergence can emerge from human behavior. We introduce an experiment in which we manipulate the value of the status quo, and we observe that this exogenous variation is able to generate belief polarization. We qualitatively replicate the model’s prediction, and observe that the magnitude of the polarization is lower than predicted. We explore the possible drivers of this difference and conclude that intrinsic (non-instrumental) preferences for information comprise the leading factor.

Our paper sheds new light on the problem of opinion polarization in society. Although

⁴¹See Appendix K for a more detailed analysis.

our analysis focuses on beliefs, the implications of our results easily extend to other variables of interest, including actions and ability to infer agents’ preferences from search behavior. In terms of inference of the agents’ type from search behavior, the reader can find more details in Appendix L. We consider a platform, similar to Facebook, that can access both actions (likes) and information collection (search) of its users, and we show that the search behavior can be a powerful predictor of the agent’s private type.

We acknowledge the limitations of our model and experiment and encourage further exploration of this important phenomenon in several directions. Our model considers individual decision making, and there is space for extensions in different strategic environments, including strategic voting and team coordination. Another limit of our analysis lies in the restriction to the binary action space, and we encourage exploration of the problem with larger action and state spaces: this feature would allow the creation of several endogenous categories and provide a connection with models of categorical thinking. On the experimental side, our design is restricted to binary information decisions (pick one out of two advisors, or buy-no-buy choices), despite the model being much more flexible. We encourage extension of the paradigm to different settings in which a larger signal space is associated with a cost for the signal accuracy (either explicit, in experimental currency, or implicit, for example time required to process the information available). The experimental estimates of beliefs polarization combine advisor choices and beliefs over states, implicitly relying on risk-neutral preferences. This assumption is justified by the fact that payoffs are expressed in terms of lottery points. We also show that both risk preferences and probability update are not major drivers of the milder polarization, and instead noninstrumental features of information play an important role. We encourage future investigations to test the robustness of the effect to alternative measures of polarization based on direct elicitations of the value of the risky option. Finally, on the empirical side, we encourage future research to test the implications of our model on referendum data, either with direct intervention on the sources of information available, or using identification strategies that capture different ease of access to media.

References

- Alesina, Alberto, Armando Miano, and Stefanie Stantcheva. 2018. “Immigration and redistribution.” *National Bureau of Economic Research*, , (w24733).
- Alesina, Alberto, Armando Miano, and Stefanie Stantcheva. 2020. “The polarization of reality.” *AEA Papers and Proceedings*, 110.

- Alesina, Alberto, Stefanie Stantcheva, and Edoardo Teso.** 2018. "Intergenerational mobility and preferences for redistribution." *American Economic Review*, 108(2): 521–54.
- Ambuehl, Sandro, and Shengwu Li.** 2018. "Belief updating and the demand for information." *Games and Economic Behavior*, 109(21-39).
- Anobile, Giovanni, Guido Marco Cicchini, and David C. Burr.** 2012. "Linear Mapping of Numbers onto Space Requires Attention." *Cognition*, 122(3): 454–459.
- Becker, Gordon M., Morris H. DeGroot, and Jacob Marschak.** 1964. "Measuring utility by a single-response sequential method." *Behavioral science*, 9(3): 226–232.
- Bénabou, Roland.** 2015. "The economics of motivated beliefs." *Revue d'économie politique*, 125(5): 665–685.
- Bernhardt, Dan, Stefan Krasa, and Mattias Polborn.** 2008. "Political polarization and the electoral effects of media bias." *Journal of Public Economics*, 92(5-6): 1092–1104.
- Blackwell, David, and Lester Dubins.** 1962. "Merging of Opinions with Increasing Information." *The Annals of Mathematical Statistics*, 33(3): 882–886.
- Bloedel, Alexander W., and Ilya Segal.** 2021. "Persuading a Rationally Inattentive Agent." *Working Paper*.
- Caplin, Andrew, and Mark Dean.** 2015. "Revealed preference, rational inattention, and costly information acquisition." *American Economic Review*, 105(7): 2183–2203.
- Caplin, Andrew, Mark Dean, and John Leahy.** 2019. "Rational inattention, optimal consideration sets and stochastic choice." *The Review of Economic Studies*, 86(3): 1061–1094.
- Charness, Gary, Ryan Oprea, and Sevgi Yuksel.** 2021. "How do people choose between biased information sources? Evidence from a laboratory experiment." *Journal of the European Economic Association*, 19(3): 1656–1691.
- Che, Yeon-Koo, and Konrad Mierendorff.** 2019. "Optimal dynamic allocation of attention." *American Economic Review*, 109(8): 2993–3029.
- Chinn, Sedona, P. Sol Hart, and Stuart Soroka.** 2020. "Politicization and Polarization in Climate Change News Content, 1985-2017." *Science Communication*, 42(1): 112–129.
- Cohen, Jacob, Patricia Cohen, Stephen G. West, and Leona S. Aiken.** 2013. *Applied multiple regression/correlation analysis for the behavioral sciences*. Routledge.
- Cover, Thomas, and Joy Thomas.** 2012. *Elements of information theory*. John Wiley & Sons.
- Dal Bó, Ernesto, Frederico Finan, Olle Folke, Torsten Persson, and Johanna Rickne.** 2018. "Economic losers and political winners: Sweden's radical right." *Unpublished manuscript, Department of Political Science, UC Berkeley*.

- Dechezleprêtre, A., A. Fabre, T. Kruse, B. Planterose, A. S. Chico, and S. Stantcheva. 2022. “Fighting climate change: International attitudes toward climate policies.” *National Bureau of Economic Research*, , (w30265).
- Dixit, Avinash K., and Jörgen W Weibull. 2007. “Political polarization.” *Proceedings of the National Academy of Sciences*, 104(18): 7351–7356.
- Douenne, Thomas, and Adrien Fabre. 2020. “French attitudes on climate change, carbon taxation and other climate policies.” *Ecological Economics*, 169.
- Douenne, Thomas, and Adrien Fabre. 2022. “Yellow vests, pessimistic beliefs, and carbon tax aversion.” *American Economic Journal: Economic Policy*, 14(1): 81–110.
- Fetzer, Thiemo. 2019. “Did austerity cause Brexit?” *American Economic Review*, 109(11): 3849–86.
- Fryer Jr, Roland G., Philipp Harms, and Matthew O. Jackson. 2019. “Updating beliefs when evidence is open to interpretation: Implications for bias and polarization.” *Journal of the European Economic Association*, 17(5): 1470–1501.
- Gentzkow, Matthew, Jesse Shapiro, and Matt Taddy. 2016. “Measuring Polarization in High-Dimensional Data: Method and Application to Congressional Speech.” NBER Working Paper Series.
- Gerber, Alan, and Donald Green. 1999. “Misperceptions about perceptual bias.” *Annual Review of Political Science*, 2(1): 189–210.
- Gigerenzer, Gerd, and Wolfgang Gaissmaier. 2011. “Heuristic decision making.” *Annual review of psychology*, 62: 451–482.
- Greiner, Ben. 2015. “Subject pool recruitment procedures: organizing experiments with ORSEE.” *Journal of the Economic Science Association*, 1(1): 114–125.
- Guney, B., M. Richter, and M. Tsur. 2018. “Aspiration-based choice.” *Journal of Economic Theory*, 176: 935–956.
- Harrison, Glenn W., Jimmy Martínez-Correa, and J. Todd Swarthout. 2013. “Inducing risk neutral preferences with binary lotteries: A reconsideration.” *Journal of Economic Behavior & Organization*, 94: 145–159.
- Hollingworth, Harry Levi. 1910. “The Central Tendency of Judgment.” *The Journal of Philosophy, Psychology and Scientific Methods*, 7(17): 461–469.
- Holt, Charles A., and Susan K. Laury. 2002. “Risk aversion and incentive effects.” *American economic review*, 92(5): 1644–1655.
- Hu, Lin, Anqi Li, and Ilya Segal. 2022. “The Politics of News Personalization.” *Working paper*.

- Kahneman, D., and A. Tversky.** 1979. "Prospect Theory: An Analysis of Decision under Risk." *Econometrica*, 47: 263–291.
- Kareev, Yaakov, Sharon Arnon, and Reut Horwitz-Zeliger.** 2002. "On the misperception of variability." *Journal of Experimental Psychology: General*, 131(2): 287.
- Kartik, Navin, Frances Xu Lee, and Wing Suen.** 2021. "Information validates the prior: A theorem on bayesian updating and applications." *American Economic Review: Insights*, 3(2): 165–82.
- Khaw, Mel Win, Ziang Li, and Michael Woodford.** 2019. "Cognitive Imprecision and Small-Stakes Risk Aversion." *NBER Working Paper*, no. 24978.
- Kőszegi, B., and M. Rabin.** 2006. "A model of reference-dependent preferences." *The Quarterly Journal of Economics*, 121(4): 1133–1165.
- Kuziemko, I., M. I. Norton, E. Saez, and S. Stantcheva.** 2015. "How elastic are preferences for redistribution? Evidence from randomized survey experiments." *American Economic Review*, 105(4): 1478–1508.
- Leiserowitz, Anthony, Edward Maibach, Seth Rosenthal, John Kotcher, Matthew Ballew, Matthew Goldberg, Abel Gustafson, and Parrish Bergquist.** 2019. "Politics & Global Warming." *Yale University and George Mason University. New Haven, CT: Yale Program on Climate Change Communication.*
- Lord, Charles G., Lee Ross, and Mark R. Lepper.** 1979. "Biased assimilation and attitude polarization: The effects of prior theories on subsequently considered evidence." *Journal of personality and social psychology*, 37(11).
- Maćkowiak, Bartosz, Filip Matějka, and Mirko Wiederholt.** 2020. "Rational Inattention: A Review." *CEPR Discussion Papers (No. 15408).*
- Maltz, A.** 2020. "Exogenous Endowment—Endogenous Reference Point." *The Economic Journal*, 130(625): 160–182.
- Manzini, P., and M. Mariotti.** 2012. "Categorize then choose: Boundedly rational choice and welfare." *Journal of the European Economic Association*, 10(5): 1141–1165.
- Masatlioglu, Yusufcan, and Efe A. Ok.** 2014. "A Canonical Model of Choice with Initial Endowments." *The Review of Economic Studies*, 81(2): 851–883.
- Masatlioglu, Yusufcan, A. Yesim Orhun, and Collin Raymond.** 2017. "Intrinsic information preferences and skewness." *Ross School of Business Paper.*
- Matějka, F., and A. McKay.** 2015. "Rational Inattention to Discrete Choices: A New Foundation for the Multinomial Logit Model." *American Economic Review*, 105(1): 272–298.

- Matějka, Filip, and Alisdair McKay.** 2015. “Rational inattention to discrete choices: A new foundation for the multinomial logit model.” *American Economic Review*, 105(1): 272–298.
- McCarty, Nolan, Keith T. Poole, and Howard Rosenthal.** 2008. *Polarized America: The dance of ideology and unequal riches*. MIT Press Books.
- Mosteller, Frederick, and Philip Noguee.** 1951. “An Experimental Measurement of Utility.” *Journal of Political Economy*, 59: 371–404.
- Nimark, Kristoffer P., and Savitar Sundareshan.** 2019. “Inattention and Belief Polarization.” *Journal of Economic Theory*, 180: 203–228.
- NPUC.** 2021. “Race to Net Zero: Carbon Neutral Goals by Country.” <https://www.visualcapitalist.com/sp/race-to-net-zero-carbon-neutral-goals-by-country/>.
- Oliveros, Santiago, and Felix Várdy.** 2015. “Demand for slant: How abstention shapes voters’ choice of news media.” *The Economic Journal*, 125(587): 1327–1368.
- Ortoleva, Pietro.** 2010. “Status quo bias, multiple priors and uncertainty aversion.” *Games and Economic Behavior*, 69(2): 411–424.
- Perego, Jacopo, and Sevgi Yuksel.** 2018. “Media competition and social disagreement.” *Working Paper*.
- Poole, Keith T., and Howard Rosenthal.** 1984. “The polarization of American politics.” *The Journal of Politics*, 46(4): 1061–1079.
- Prior, Markus.** 2013. “Media and political polarization.” *Annual Review of Political Science*, 16: 101–127.
- Samuelson, William, and Richard Zeckhauser.** 1988. “Status Quo Bias in Decision Making.” *Journal of Risk and Uncertainty*, 1(1): 7–59.
- Savage, Leonard J.** 1954. *The foundations of statistics*. Dover Books.
- Scheier, Michael F., Charles S. Carver, and Michael W. Bridges.** 1994. “Distinguishing optimism from neuroticism (and trait anxiety, self-mastery, and self-esteem): a reevaluation of the Life Orientation Test.” *Journal of personality and social psychology*, 67(7): 1063.
- Sims, Christopher A.** 1998. “Stickiness.” *Carnegie-Rochester Conference Series on Public Policy*, 49: 317–356.
- Sims, Christopher A.** 2003. “Implications of Rational Inattention.” *Journal of Monetary Economics*, 50(3): 665–690.
- Smith, Matthew.** 2019. “Most voters have only become more sure about their EU referendum vote.” <https://yougov.co.uk/topics/politics/articles-reports/2019/04/16/most-voters-have-only-become-more-sure-about-their>.

- Steiner, Jakub, Colin Stewart, and Filip Matějka.** 2017. “Rational Inattention Dynamics: Inertia and Delay in Decision-Making.” *Econometrica*, 85(2): 521–553.
- Suen, Wing.** 2004. “The self-perpetuation of biased beliefs.” *The Economic Journal*, 114(495): 377–396.
- Weber, Thomas A.** 2003. “An exact relation between willingness to pay and willingness to accept.” *Economics Letters*, 80(3): 311–315.
- Zanardo, Enrico.** 2017. “Essays in Microeconomics.” *Doctoral dissertation, Columbia University*.

A Appendix: Proof of Theorem 1

We start the proof with several useful lemmas.

Lemma 1. *Conditional on the realized state of the world $s \in S$, the probability of choosing a new policy for $\lambda > 0$ is implicitly defined by:*

$$\mathcal{P}(i = 1|s) = \frac{\mathcal{P}(i = 1)e^{\frac{v_s}{\lambda}}}{\mathcal{P}(i = 1)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i = 1))e^{\frac{R}{\lambda}}},$$

the probability of choosing the status quo is

$$\mathcal{P}(i = 2|s) = \frac{(1 - \mathcal{P}(i = 1))e^{\frac{R}{\lambda}}}{\mathcal{P}(i = 1)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i = 1))e^{\frac{R}{\lambda}}},$$

where $\mathcal{P}(i = 1)$ is the unconditional probability of choosing a new policy.

Proof. Lemma 1 is a direct consequence of Theorem 1 from Matějka and McKay (2015). $\square$

Lemma 2. *The agent's posterior belief about the payoff of the new policy v , given the fixed state s^* is*

$$\mathbb{E}_i[\mathbb{E}(v|i)|s^*] = \sum_{s=1}^n v_s g_s \frac{\mathcal{P}(i = 1|s^*)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i = 1|s^*))e^{\frac{R}{\lambda}}}{\mathcal{P}(i = 1)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i = 1))e^{\frac{R}{\lambda}}}. \quad (7)$$

Proof.

$$\mathbb{E}_i[\mathbb{E}(v|i)|s^*] = \mathcal{P}(i = 1|s^*)\mathbb{E}(v|i = 1) + \mathcal{P}(i = 2|s^*)\mathbb{E}(v|i = 2).$$

Substituting for the conditional probabilities using lemma 1 and applying Bayes rule, we obtain

$$\begin{aligned} \mathbb{E}_i[\mathbb{E}(v|i)|s^*] &= \frac{\mathcal{P}(i = 1)e^{\frac{v_{s^*}}{\lambda}}}{\mathcal{P}(i = 1)e^{\frac{v_{s^*}}{\lambda}} + (1 - \mathcal{P}(i = 1))e^{\frac{R}{\lambda}}} \cdot \sum_{s=1}^n v_s g_s \frac{e^{\frac{v_s}{\lambda}}}{\mathcal{P}(i = 1)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i = 1))e^{\frac{R}{\lambda}}} + \\ &+ \frac{(1 - \mathcal{P}(i = 1))e^{\frac{R}{\lambda}}}{\mathcal{P}(i = 1)e^{\frac{v_{s^*}}{\lambda}} + (1 - \mathcal{P}(i = 1))e^{\frac{R}{\lambda}}} \cdot \sum_{s=1}^n v_s g_s \frac{e^{\frac{R}{\lambda}}}{\mathcal{P}(i = 1)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i = 1))e^{\frac{R}{\lambda}}}. \end{aligned}$$

Thus,

$$\mathbb{E}_i[\mathbb{E}(v|i)|s^*] = \sum_{s=1}^n v_s g_s \frac{\mathcal{P}(i=1|s^*)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i=1|s^*))e^{\frac{R}{\lambda}}}{\mathcal{P}(i=1)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i=1))e^{\frac{R}{\lambda}}}.$$

□

Lemma 3. *Relations $\mathcal{P}(i=1|s^*) \geq P(i=1)$ for $0 < P(i=1) < 1$ are equivalent to $v_{s^*} \geq R$. Relation $\mathcal{P}(i=1|s^*) = P(i=1)$ for $0 < P(i=1) < 1$ is equivalent to $v_{s^*} = R$.*

Proof. After substitution for the conditional probabilities, the conditions $\mathcal{P}(i=1|s^*) \geq P(i=1)$ can be rewritten as

$$\frac{\mathcal{P}(i=1)e^{\frac{v_{s^*}}{\lambda}}}{\mathcal{P}(i=1)e^{\frac{v_{s^*}}{\lambda}} + (1 - \mathcal{P}(i=1))e^{\frac{R}{\lambda}}} \geq \mathcal{P}(i=1),$$

which are equivalent to

$$(\mathcal{P}(i=1) - P(i=1)) \left(e^{\frac{v_{s^*}}{\lambda}} - e^{\frac{R}{\lambda}} \right) \geq 0.$$

For $0 < P(i=1) < 1$ the term in the first parenthesis is always positive. Therefore, the left hand side of the inequality is positive when $v_{s^*} > R$ and negative for $v_{s^*} < R$. The case when $\mathcal{P}(i=1|s^*) = P(i=1)$ can be considered analogously. □

Lemma 4. *Given that the realized state of the world is $s^* \in S$, the sign of the change in the mean of beliefs about the payoff of the new policy $\Delta(s^*) = \mathbb{E}_i[\mathbb{E}(v|i)|s^*] - \mathbb{E}v$ is the same as the sign of $(v_{s^*} - R)$. $\Delta(s^*) = 0$ if and only if $v_{s^*} - R = 0$.*

Proof. In order to solve the agent's problem given by equations (1) - (4) we need to find $\mathcal{P}(i=1)$ and $\mathcal{P}(i=2)$ defined as $\mathcal{P}(i=2) = 1 - \mathcal{P}(i=1)$. These probabilities have to be internally consistent, i.e., $\mathcal{P}(i) = \sum_{s=1}^n \mathcal{P}(i|s)g_s$. After dividing both sides of these conditions by $P(i)$ we obtain the following conditions

$$1 = \sum_{s=1}^n \frac{e^{\frac{v_s}{\lambda}}}{\mathcal{P}(i=1)e^{\frac{v_s}{\lambda}} + \mathcal{P}(i=2)e^{\frac{R}{\lambda}}} g_s, \quad \text{if } \mathcal{P}(i=1) > 0,$$

$$1 = \sum_{s=1}^n \frac{e^{\frac{R}{\lambda}}}{\mathcal{P}(i=1)e^{\frac{v_s}{\lambda}} + \mathcal{P}(i=2)e^{\frac{R}{\lambda}}} g_s, \quad \text{if } \mathcal{P}(i=2) > 0.$$

The difference of these two equations is

$$\sum_{s=1}^n \frac{e^{\frac{v_s}{\lambda}} - e^{\frac{R}{\lambda}}}{\mathcal{P}(i=1)e^{\frac{v_s}{\lambda}} + \mathcal{P}(i=2)e^{\frac{R}{\lambda}}} g_s = 0.$$

For k , which holds that $v_k \leq R \leq v_{k+1}$ we can further write the equation above as

$$\frac{e^{\frac{v_k}{\lambda}} - e^{\frac{R}{\lambda}}}{\mathcal{P}(i=1)e^{\frac{v_k}{\lambda}} + \mathcal{P}(i=2)e^{\frac{R}{\lambda}}} v_k g_k = - \sum_{s \neq k} \frac{e^{\frac{v_s}{\lambda}} - e^{\frac{R}{\lambda}}}{\mathcal{P}(i=1)e^{\frac{v_s}{\lambda}} + \mathcal{P}(i=2)e^{\frac{R}{\lambda}}} v_k g_s. \quad (8)$$

We use the last equation for determining the sign of $\Delta(s^*)$, which can be written as

$$\begin{aligned} \Delta(s^*) &= \sum_{s=1}^n v_s g_s \frac{\mathcal{P}(i=1|s^*)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i=1|s^*))e^{\frac{R}{\lambda}}}{\mathcal{P}(i=1)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i=1))e^{\frac{R}{\lambda}}} - \sum_{i=1}^n v_s g_s, \\ \Delta(s^*) &= \sum_{s=1}^n v_s g_s \frac{\mathcal{P}(i=1|s^*)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i=1|s^*))e^{\frac{R}{\lambda}}}{\mathcal{P}(i=1)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i=1))e^{\frac{R}{\lambda}}} - \sum_{i=1}^n v_s g_s \frac{\mathcal{P}(i=1)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i=1))e^{\frac{R}{\lambda}}}{\mathcal{P}(i=1)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i=1))e^{\frac{R}{\lambda}}}, \\ \Delta(s^*) &= \sum_{i=1}^n v_s g_s \frac{(\mathcal{P}(i=1|s^*) - \mathcal{P}(i=1))(e^{\frac{v_s}{\lambda}} - e^{\frac{R}{\lambda}})}{\mathcal{P}(i=1)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i=1))e^{\frac{R}{\lambda}}}, \\ \Delta(s^*) &= (\mathcal{P}(i=1|s^*) - \mathcal{P}(i=1)) \cdot \sum_{s=1}^n v_s g_s \frac{e^{\frac{v_s}{\lambda}} - e^{\frac{R}{\lambda}}}{\mathcal{P}(i=1)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i=1))e^{\frac{R}{\lambda}}}. \end{aligned}$$

Substituting equation (8) into the sum in the last equation we obtain

$$\Delta(s^*) = (\mathcal{P}(i=1|s^*) - \mathcal{P}(i=1)) \left[\sum_{s \neq k} (v_s - v_k) g_s \frac{e^{\frac{v_s}{\lambda}} - e^{\frac{R}{\lambda}}}{\mathcal{P}(i=1)e^{\frac{v_s}{\lambda}} + (1 - \mathcal{P}(i=1))e^{\frac{R}{\lambda}}} \right].$$

The expression in the square brackets is positive, because for the above-defined k the sign of $(v_s - v_k)$ and the sign of $e^{\frac{v_s}{\lambda}} - e^{\frac{R}{\lambda}}$ are the same. Hence, $\Delta(s^*)$ has the same sign as $(\mathcal{P}(i=1|s^*) - \mathcal{P}(i=1))$ that further, by Lemma 3, has the same sign as $(v_{s^*} - R)$. $\square$

Now we can get back to proving Theorem 1. Let us first consider the case $\mu^A \neq \mu^B$

Without loss of generality, let us assume that $\mu^A > \mu^B$. For the condition $(\mu^A - \mu^B)(v_{s^*} - R^A) > 0$ to be satisfied, it is necessary that $v_{s^*} > R^A$. Proposition 4 states that in this case $\Delta_A(s^*) > 0$. Analogously, the second condition $(\mu^A - \mu^B)(v_{s^*} - R^B) < 0$ holds when $v_{s^*} < R^B$, which further implies that $\Delta_B(s^*) < 0$. That is, two agents $j = 1, 2$ update in different directions and the expected posterior beliefs are farther away from each other than the priors are. Both conditions from Definition 1 are satisfied and the agents, indeed, become polarized in state s^* . Similarly, if $\mu^A = \mu^B$, the inequality $(v_{s^*} - R^A)(v_{s^*} - R^B) < 0$, due to Lemma 4 implies that $\Delta^A(s^*)\Delta^B(s^*) < 0$, and the conditions from Definition 1 are thus satisfied.

B Appendix: Belief divergence without polarization

In this appendix, we demonstrate a possibility of belief divergence when these beliefs are updated in the same direction, which, in our opinion, is an interesting feature of our model. Specifically, we explore the influence of the prior beliefs on the magnitude of the change in the mean of beliefs. To study this question, we take advantage of an example with three states and two actions. This problem is a simple benchmark and its solution exhibits the basic features of solutions to the problems with n states and 2 actions. The solution we analyze in this section is symbolic. Let us start with the definition of divergence of beliefs updated in the same direction.

Definition 5. We say that two agents $j \in \{A, B\}$, who are characterized by the pair $(R^j, \mathbf{g}^j)$ and are choosing between actions $i = \{1, 2\}$, *diverge in their belief updated in the same direction* when in the state $s^* \in S$ the following two conditions are satisfied

1. $|m^A(s^*) - m^B(s^*)| > |\mu^A - \mu^B|$.
2. $\Delta_A(s^*) \cdot \Delta_B(s^*) > 0$.

The parameter values that we use in this appendix are as follows: $v_1 = 0$, $v_2 = 1/2$, $v_3 = 1$, $g_1 = g \in (0, 2/3)$, $g_2 = 1/3$, $g_3 = 2/3 - g$, $R = 3/8$, $\lambda = 1/4$.

Note that keeping prior probability of state 2, g_2 , fixed, we can vary prior probability of state 1, g only between $(0, 2/3)$. Also, $\mathbb{E}v$ can vary only from $1/6$ to $5/6$. To solve the

problem (1)-(4) it is necessary to find the unconditional probabilities $\mathcal{P}(i = 1)$ and $\mathcal{P}(i = 0)$, which we then use for finding the conditional probabilities.

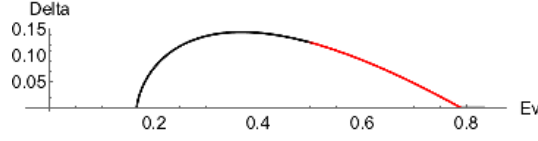


Figure 12: $\Delta(s^* = 2)$ as a function of $\mathbb{E}v$ for $R_1 = 3/8$ and $\lambda_2 = 1/4$. The red area depicts the region of updating in the opposite direction from the realized value.

Figure 12 provides an example in which two agents diverge in the beliefs while they update in the same direction. Since two agents differ only in their prior expectations about the new policy, it is sufficient to look at how a single agent's change in the mean of beliefs $\Delta(s^* = 2)$ depends on $\mathbb{E}v$. We are interested in finding two prior expected beliefs for which there is divergence of posterior beliefs. To do so, we need to find two points such that Δ for the left point is lower than Δ for the right point. In our example the red part of the plot is a decreasing function. This means that, in our example, two agents updating in the same direction with the same valuation of the status quo might diverge in their opinions only when they are updating towards the realized value. However, at the black part of the plot it is easy to find two points at which the agents diverge in their opinions.

C Appendix: Experimental Design and Procedure

C.1 Experimental Interface and Payment

Task 1 - Colorblind advisor game

If one round from this part is selected for the bonus payment, a subject receives the \$15 bonus with the percentage probability equal to the number of points that she collected in that round. Since each line counts as a separate decision, one of which might be randomly drawn for payment, truthful revelation is strictly optimal. We constrain subjects to have at most one switching point for every advisor.

Task 1, Screen 1: Action choice

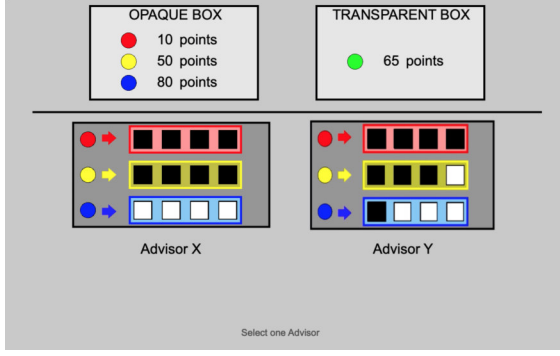
Task 2, Screen 2: WTA for each advisor

Figure 13: Task 1: Colorblind advisor game. Left: Subjects choose an action (box) contingent on the advisor and signal received. The possible values of each action are indicated on the top of the screen. Each state (ball color) is equally likely to occur. Right: Subjects indicate for each advisor the willingness to accept renunciation of its signal in a series of binary choices (BDM method). At most, one switch is allowed. Action choices selected in the previous stage are reported on the bottom of the screen.

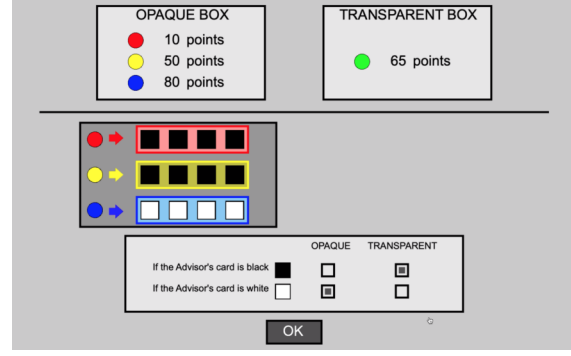
Task 2 - Imprecise advisor game If one round from this part is selected for the bonus payment, subjects receive the \$15 bonus with the percentage probability equal to the number of points that she collected in that round.

Task 3 - Card color prediction game If one round from this part is selected for the bonus payment, the computer randomly determines the state and realized signal, and subjects receive the \$15 bonus with the percentage probability determined by the quadratic loss scoring rule.

Task 4 - Ball color prediction game If one round from this part is selected for



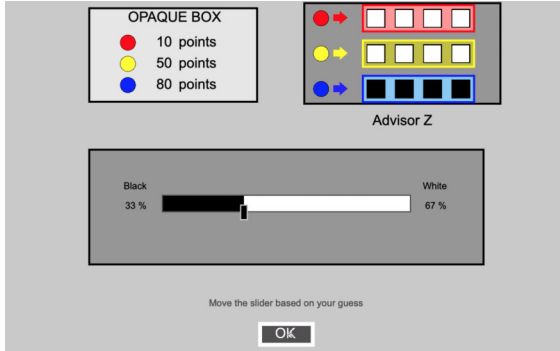
Task 1, Screen 1: Advisor choice



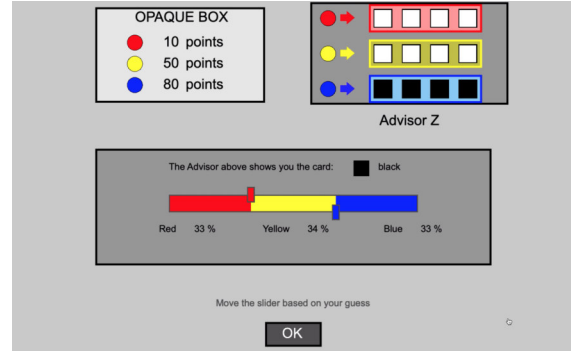
Task 2, Screen 2: Action choice

Figure 14: Task 2: Imprecise advisor game. Left: Subjects choose one signal structure (advisor) between the two options available. Each advisor is a triplet of state-contingent signal probabilities. Right: Subjects indicate the signal-contingent action for each signal (strategy method).

the bonus payment, the computer randomly determines the state and realized signal, and subjects receive the \$15 bonus with the percentage probability determined by the quadratic loss scoring rule.



Task 3: Beliefs over signal likelihood



Task 4: Beliefs over state likelihood

Figure 15: Left: Task 3 (Card color prediction game). Subjects indicate the likelihood of observing each signal (card color) for the given advisor. Right: Task 4 (Ball color prediction game). Subjects indicate the likelihood of each state (ball color) given an advisor and signal. In both tasks subjects move the slider(s) and receive a number of probability points according to the quadratic loss scoring rule described in the instructions.

C.2 Randomization

In all the tasks we randomize the order of the trials. For task 2 only, the first 3 trials are randomly drawn from the subset of trials where both the advisors provide certainty (in order

to facilitate the transition from task 1 to task 2).

In task 1 we randomize the order of the four hiring screens within each trial.

In task 2 we randomize the positions of the two advisors on the screen.

In tasks 2, 3, and 4 we randomize the advisors' card colors (black and white). This means that the signal-contingent choice in the second part of the round requires the subjects to analyze every advisor separately, since the colors do not convey any intrinsic message, and this procedure reduces the concern regarding inertia in the evaluation of the advisor and in actions.

C.3 Subject understanding

Instructions were provided on both the computer screen, as slides that can be browsed by each subject at the desired pace, and as a paper printout. The two versions of the instructions contained the same information verbatim. Before proceeding with every section of the experiment, subjects were required to correctly answer all the multiple-choice questions of the comprehension test to check understanding of the instructions. The number of questions ranged from two to four for every section, and subjects received a one-minute timeout before having a new attempt. Subjects were initially informed about the payment structure, the no-deception policy of the laboratory, and that choices in one section of the experiment did not affect any other section, or the questionnaire. A small number of subjects were recruited for each laboratory session (6 on average) in order to facilitate clarification of questions during the experiment.

D Appendix: Questionnaire

At the end of the four tasks, we have an additional section with a Holt and Laury test of risk aversion (Part 1), Raven matrices test of fluid intelligence (Part 2, five matrices of different difficulty), and a series of questions (Part 3), that we show here as they were presented to subjects.

Questionnaire - Part 3

* Required

Participant ID *

100

Next

Page 1 of 4

Never submit passwords through Google Forms.

This content is neither created nor endorsed by Google. [Report Abuse](#) · [Terms of Service](#) · [Privacy Policy](#)

Google Forms

Questionnaire - Part 3

* Required

In section 1 of the main experiment, you were asked to decide whether to make choices with or without the help of the advisors. Briefly describe your strategy in the task depicted in the figure. *

OPAQUE BOX

10 points

30 points

80 points

TRANSPARENT BOX

25 points

Choice with Red Advisor

Choice with Red Advisor

Choice with Red Advisor

Choice with Red Advisor

Choice with Red Advisor

Choice with Red Advisor

Choice with Red Advisor

Choice with Red Advisor

Choice with Red Advisor

Choice with Red Advisor

Choice without Advisor + 3 extra points

Choice without Advisor + 2 extra points

Choice without Advisor + 4 extra points

Choice without Advisor + 6 extra points

Choice without Advisor + 8 extra points

Choice without Advisor + 10 extra points

Choice without Advisor + 12 extra points

Choice without Advisor + 14 extra points

Choice without Advisor + 16 extra points

Choice without Advisor + 18 extra points

Choice without Advisor + 20 extra points

Transparent

Opaque

Make one choice for each line

Your answer

In section 2 of the main experiment, you were asked to select one advisor before choosing one of the boxes. Briefly describe your strategy in the task depicted in the figure. *

OPAQUE BOX

10 points

50 points

80 points

TRANSPARENT BOX

65 points

Advisor X

Advisor Y

Your answer

Back

Next

Page 2 of 4

Never submit passwords through Google Forms.

This content is neither created nor endorsed by Google. [Report Abuse](#) · [Terms of Service](#) · [Privacy Policy](#)

Google Forms

How much do you agree with the following statements? *

Try not to let your response to one statement influence your responses to other statements. There are no "correct" or "incorrect" answers. Answer according to your own feelings, rather than how you think "most people" would answer.

I agree a lot

I agree a little

I neither agree nor disagree

I disagree a little

I disagree a lot

Prefer not to say

I enjoy my friends a lot

I own an object that I feel brings me good luck during exams

I often take risk just for fun

I rarely count on good things happening to me

I am uncomfortable if I am assigned seat number 13 in a theatre

It's easy for me to relax

I refrain from taking risks when I feel I am having an unlucky day

I don't get upset too easily

Overall, I expect more good things to happen to me than bad

E Appendix: Advisor Pairs in the Main Task

In task 2, subjects play 40 rounds with different pairs of advisors and values for the ball in the transparent box. The ball in the transparent box can take two values: 30 points (low status quo) and 65 points (high status quo). The values for the balls in the opaque box are unchanged during the task (10, 50, and 80 points, with uniform probability of being drawn). The rounds are designed as a combination of 20 advisor pairs and two values for the ball in

51















Downloaded from <http://ajph.org/> on November 10, 2015

Downloaded from <https://www.cambridge.org/core>. University of Cambridge, on 01 Jun 2018 at 12:00:00, subject to the Cambridge Core terms of use, available at <https://www.cambridge.org/core/terms>. <https://doi.org/10.1017/9781315336435.008>

The table 4 shows the pairs of advisors, here labeled X and Y . Advisor names, positions on the screen, and signal colors were randomized at the subject level. Each advisor is presented as a triplet of conditional signal probabilities, conditional on the realized state (value of the risky action). For each pair of advisors, the table indicates which of them has the highest instrumental value under each status quo value (low or high), with $\sim$ to denote ties.

Pair	Advisor X			Advisor Y			Best Low R	Best High R
	x_{bad}	x_{med}	x_{good}	y_{bad}	y_{med}	y_{good}	($R = 30$)	($R = 65$)
1	0	1	1	0	0	1	X	Y
2	0.25	1	1	0	0	1	X	Y
3	0	1	0.75	0	0	1	X	Y
4	0	0.75	1	0	0	1	X	Y
5	0	1	1	0	0	0.5	X	Y
6	0	1	1	0	0.25	0.75	X	Y
7	0.25	1	1	0	0	0.75	X	Y
8	0.5	1	1	0	0	0.75	X	Y
9	0.25	1	1	0	0.25	0.75	X	Y
10	0.25	0.75	1	0	0	0.75	X	Y
11	0.25	0.75	1	0	0.25	0.75	X	Y
12	0	0.5	1	0.25	0.5	0.75	X	X
13	0	1	1	0	1	0	X	$\sim$
14	0.25	0.75	1	0.75	0.25	1	X	$\sim$
15	0	1	1	0.5	1	1	X	$\sim$
16	0	0	1	0	1	0	$\sim$	X
17	0	0	0.5	0.25	0.5	0.75	$\sim$	X
18	0	0.5	0.5	0	1	0	$\sim$	$\sim$
19	0.5	0	0.5	0.5	0.5	1	$\sim$	$\sim$
20	0	0.25	0.5	0	0.5	0.75	$\sim$	$\sim$

Table 4: Pairs of advisors used in Task 2. Each pair contains two different advisors (X and Y), with the triplet of signal probabilities $p_i = Pr(\sigma = 1 | s = i)$. The last two columns show the theoretical predictions for a Bayesian decision-maker. Each pair of advisors is presented with two different status quo values (low R and high R). For each of these values, we indicate which of the two advisors has the highest instrumental value, with $\sim$ in case of a tie.

The advisor pairs are selected in order to examine preference over sources of information and formulate predictions about the effect of the safe option on information collection and posterior beliefs. Based on the predicted behavior of a Bayesian agent, we can classify the pairs into the following groups:

- Pairs 1-11: pick different advisors by changing the safe option (strict preference);
- Pairs 12-16: pick different advisors by changing the safe option (weak preference);
- Pair 17: always pick advisor X regardless of the safe option (Blackwell ordered signals);
- Pairs 8-20: indifference between the advisors for both safe options.

The table 5 shows the pairs of advisors as in the table 4 extended for corresponding measures of simplicity. In particular, we include dummy for certainty, discrete complexity and the continuous measure of complexity. The discrete complexity measure is defined by equation 6 and the continuous measure of complexity for state s and a signal σ is defined by following equation

$$c_C = \sum_s \sum_{\sigma} \sqrt{\mathcal{P}(s|\sigma_s)} \quad (9)$$

Advisor	Advisor X			Value		Simplicity measures		
	x_{bad}	x_{med}	x_{good}	$(R = 30)$	$(R = 65)$	Certainty	Discrete complexity	Continuous complexity
1	0	0	1	46.667	70	1	1	2.4142
2	0	1	0	46.667	65	1	1	2.4142
3	0	1	1	53.333	65	1	1	2.4142
4	0	0	0.5	46.667	67.5	0	2	2.7121
5	0	0	0.75	46.667	68.75	0	2	2.6667
6	0	0.5	1	50	67.5	0	2	2.7877
7	0	0.75	1	51.667	66.25	0	2	2.7522
8	0	1	0.75	49.167	65	0	2	2.7522
9	0.25	1	1	51.667	65	0	2	2.6667
10	0.5	1	1	50	65	0	2	2.7121
11	0	0.25	0.5	46.667	66.25	0	3	3.1093
12	0	0.25	0.75	46.667	67.5	0	3	3.0391
13	0	0.5	0.5	46.667	65	0	3	3.1213
14	0	0.5	0.75	46.667	66.25	0	3	3.0755
15	0.25	0.75	1	50	65	0	3	3.0391
16	0.5	0	0.5	46.667	65	0	3	3.1213
17	0.5	0.5	1	46.667	65	0	3	3.1213
18	0.75	0.25	1	46.667	65	0	3	3.0391
19	0.25	0.5	0.75	46.667	65	0	4	3.3854
20	0.25	0.25	0.25	46.667	65	0	4	3.4641

Table 5: Pairs of advisors used in Task 2 and their certainty and complexity scores. Each pair contains two different advisors (X and Y), with the triplet of signal probabilities $p_i = Pr(\sigma = 1|s = i)$. The next two columns show the theoretical predictions for a Bayesian decision-maker. Each pair of advisors is presented with two different status quo values (low R and high R). For each of these values, we indicate which of the two advisors has the highest instrumental value, with $\sim$ in case of a tie. The last three columns consist of measures of simplicity: dummy for certainty, discrete complexity score from 1 to 4, and continuous complexity score, respectively.

F Appendix: Further Analysis on Advisor Choice

Certain advisors provide an answer to a question of the kind: “Is the state red(/yellow/blue)?” and allow the subject to learn with certainty if a particular state is realized (with probability one) or not realized (with probability zero).

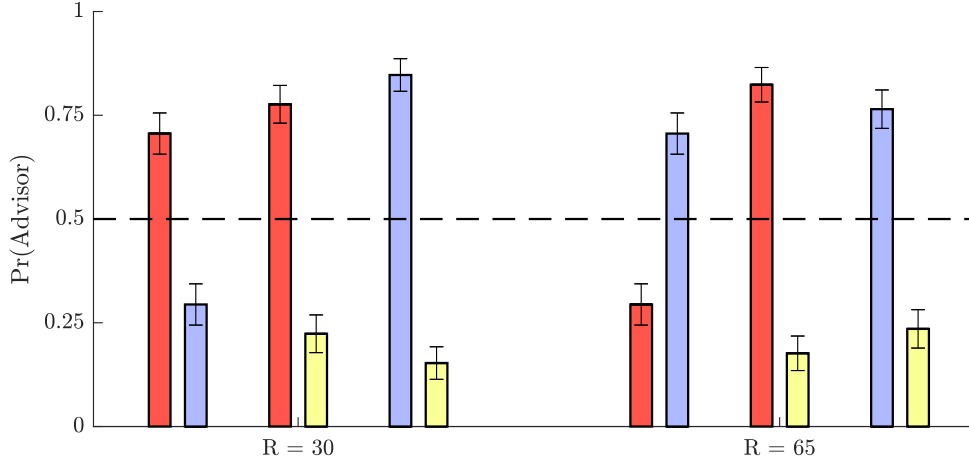


Figure 18: Comparison of advisors providing the ideal state pooling question: “Is the state red/yellow/blue ?” for two different status quo values. The color of the bar shows which state the question is about. The figure demonstrates the state pooling behavior, and also that participants do switch between advisors when it is valuable to do so.

Figure 18 shows advisor choice in the trials in which both advisors provide certainty. We display separately the trials with different status quo values. When the subjects have to choose between advisors that provide certainty and are also state poolers, that is, between an advisor providing information whether the state is blue and another advisor providing information whether the state is red (first couple of bars for $R = 30$ and $R = 65$), they significantly select the former for the high value of the status quo and latter for the low value of the status quo. This switch between advisors confirms our theoretically predicted state pooling effect. In particular, for a status quo value R the subject wants to learn whether the state-dependent payoff of the new policy is greater or lower than R . When subjects face a choice between a certainty state pooler⁴² and certainty advisors, they select on average the

⁴²A certainty state pooler advisor is certain (can fully reveal one state, as previously defined) and the revealed state is the singleton one from the state pooling.

certainty state pooler in 74% of the trials. This result appears at odds with the Experimental Result 7 (higher WTP for information on high-value states, in task 1). This suggests that, in case of conflict between state pooling and high-value preferences, the discrete choice task favors the first effect over the second one.

In the two scenarios of choice between the red-advisor and blue-advisor, we see that each option is chosen by more than one quarter of the participants, and this is at odds with the model’s prediction, especially in a trial with a relatively simple problem. This can be interpreted as a general signal of noise in the participants’ actions, or as a systematic preference towards information about low or high states. Figure 19 suggests that the latter interpretation can partially explain the pattern. Starting from all the trials, we calculate for every subject the probability of choosing the advisor that is best under low or high status quo value.

If choices are just noisy, we should observe most of the subjects to be clustered around the coordinates 0.5-0.5, which would be consistent both with optimal behavior (always pick the best advisor), and with completely erratic choices (pick randomly). If participants have non-instrumental preferences over skewed sources of information (as shown by Masatlioglu, Orhun and Raymond (2017)), and such preferences are heterogeneous, we would expect a distribution of subjects that systematically deviate towards 1-0 (reveal information about the low state) and 0-1 (about the high state). In fact, we observe that participants deviate in both directions, and some of them also deviate towards lower probabilities in both dimensions - this can be the case when the chosen advisor is the worse choice under both status quo scenarios.

G Appendix: Further Analysis on Risk Attitude

In Section 4.2 we discussed how risk preferences could represent an explanation for the subjects’ deviation from the model’s predictions. We provide here further details about how subjects’ actions are, on average, not characterized by a significant deviation from risk neutrality.

Figure 20a shows the realized probability of selecting the risky option as a function of

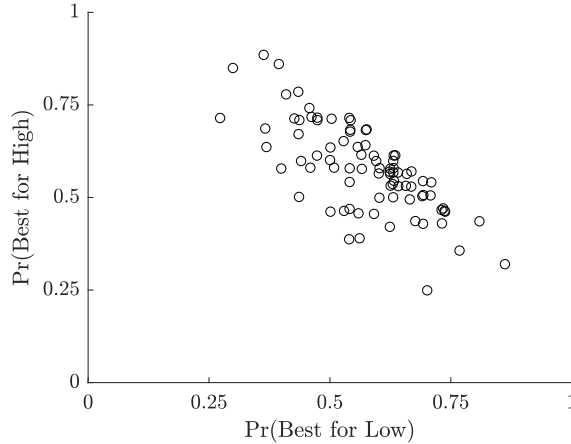


Figure 19: Distribution of participants’ advisor choices: probability of choosing the advisor that provides more information about the low or high state in different types of trials.

the difference in the EVs between the actions. Trials are grouped based on the x-axis value for visualization purposes. The optimal agent would have a sharp jump in probability from 0 (when the value difference is negative) to 1. We observe a smoother transition in our data, suggesting that action probability is modulated by the cost of mistakes, similarly to our discussion in Figure 5b about the choice between advisors. Such a sigmoid curve is normally found in experiments involving choice under risk (Mosteller and Nogee, 1951; Khaw, Li and Woodford, 2019). The indifference point appears close to the trials in which both actions have the same values, suggesting that the participants are overall close to risk neutrality.

We replicate the analysis for the choices made in task 1. An advantage of this dataset is that we observe two types of action choice scenario.

First, when the advisor confirms the color of the hidden ball, the decision maker faces a choice between two degenerate lotteries with different values, for example 80 points for sure (risky action if you know the color is blue) or 60 points for sure (safe action). In these cases, participants pick the best option 90% of the times, confirming the small amount of noise in the action implementation in these simple choices.

Second, all the participants encounter choices with full uncertainty about the color (decision without any hint) or with a hint about two possible colors (e.g. red or yellow with equal chance, but not blue). The participants pick the option with the highest expected value 84% of the times, and we encounter again the sigmoid curve discussed above. The MLE of the

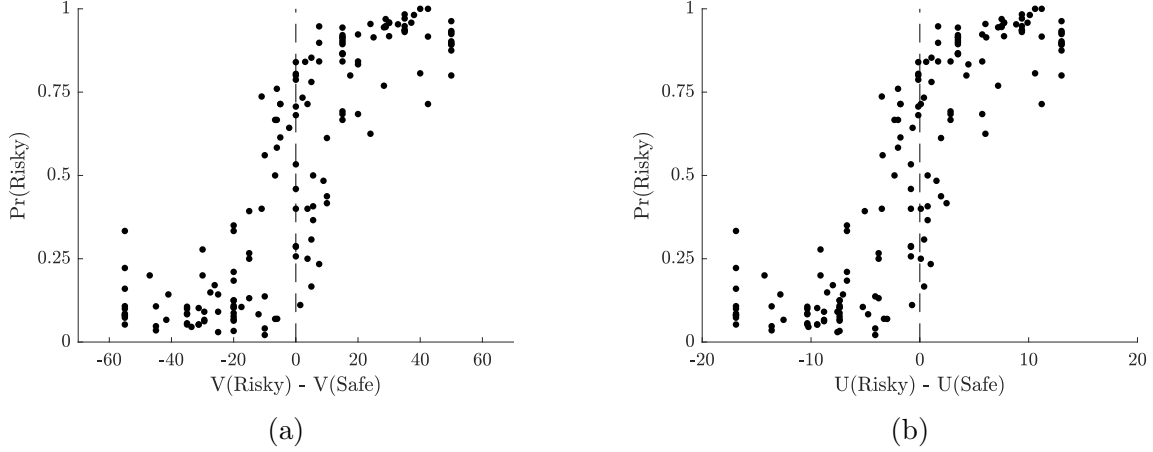


Figure 20: Task 2: Action selection probability. Left: Action choice under risk neutrality. Observed probability of choosing the risky action in task 2. 6,800 observations unequally divided across 160 cases (2 cases per trial, conditional on the advisor choice). Right: Action choice under risk aversion (best fit). The expected values for each action is replaced with the expected utility, with CRRA utility and the MLE coefficients $\hat{\alpha} = 0.34$ estimated from the dataset.

risk aversion parameter under CRRA utility is $\hat{\alpha} = 0.52$.

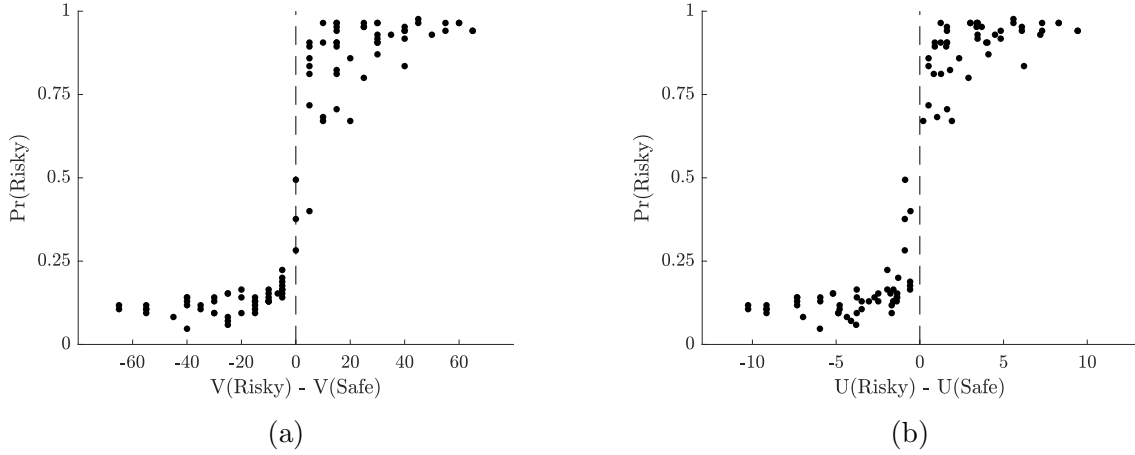


Figure 21: Task 1: Action selection probability. Left: Action choice under risk neutrality. Observed probability of choosing the risky action in task 1. 100 questions (10 questions for each of the 10 trials), 85 observations per question. Right: Action choice under risk aversion (best fit). The expected values for each action is replaced with the expected utility, with CRRA utility and the MLE coefficient $\hat{\alpha} = 0.52$ estimated from the dataset.

H Appendix: Further Analysis of Beliefs

In Section 4.2 we discussed how subjective beliefs could represent an explanation for the subjects’ deviation from the model’s predictions. We provide here further details about how subjects’ subjective beliefs elicited in tasks 3 and 4 display, on average, only mild evidence of conservatism.

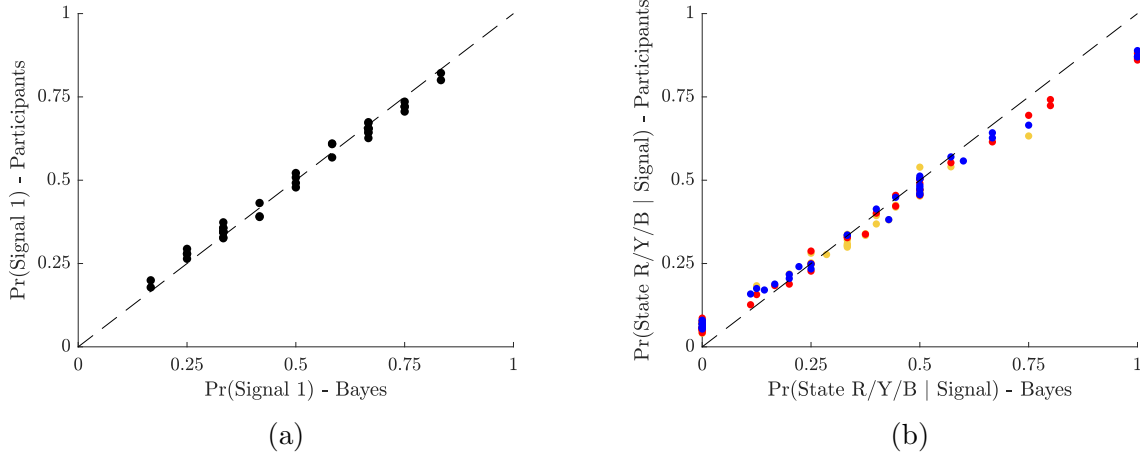


Figure 22: Average subjective beliefs. Left: Estimated probability of receiving a signal realization in Task 3. The plot compares the average of the subjective estimates collected with the optimal estimates of a Bayesian decision maker. 1,700 observations across 20 trials (85 observations per point). Right: Estimated posterior probability of each state in task 4 conditional on the realized signal. Colors indicate which state was estimated (red, yellow, blue). The plot compares the average of the subjective estimates collected with the optimal estimates of a Bayesian decision maker. 20,400 observations across 40 trials (6 observations per trial, 85 observations per point in the plot).

In both tasks we observe accurate probability estimates, close to the predictions of an optimal Bayesian agent. Figure 29a shows the subjective estimate of a signal realization (y-axis, averaged across participants) compared to the optimal estimates (x-axis). Similarly, Figure 29b shows the subjective estimate of each of the three possible states in the posterior compared to the unbiased posterior, with different colors in the figure matching the state. In both plots, the 45 degree lines represent our theoretical benchmark and we can see that 1) participants are on average accurate in the estimate of probabilities, 2) we do not observe a systematic difference between estimates involving different states (i.e., we do not have evidence of motivated beliefs, Bénabou (2015)), and 3) both tasks show mild

evidence of conservatism (central tendency of judgement), as vastly reported in experiments with subjective estimates (Hollingworth, 1910; Anobile, Cicchini and Burr, 2012).

For the signal probability (task 3) a linear fit of the subjective estimates $\hat{p}$ over the true probabilities p returns the coefficients $\hat{p} = 0.041 + 0.918 \cdot p$ with $R^2 = 0.991$. For the state probability (task 4) the linear fit for the whole dataset returns $\hat{p} = 0.058 + 0.825 \cdot p$ with $R^2 = 0.993$. The slopes are not significantly different across the three types of states: $\beta_{red} = 0.831$, $\beta_{yellow} = 0.805$, $\beta_{blue} = 0.827$.

I Appendix: Further Analysis of Willingness to Accept

In this section we add further results from the analysis of the willingness to accept renunciation of an advisor in task 1. We reported in Figure 9 and Table 3 that WTA is characterized by compression, a conservatism in the evaluation of the instrumental value of an advisor that leads to overpayment for the advisor with little or no informative value.

This result is robust across subjects, as displayed in Figure 23. For every subject, we estimate the sensitivity to the instrumental value by using a simple OLS regression of the subjective evaluation $V^j(i)$ over the instrumental value $V^{Bayes}(I)$. The graph shows the cumulative distribution of the fitted slopes, where 0 indicates no response to the true value and 1 indicates full alignment between the two variables. 82% of the participants show values between 0 and 1.

Another effect discussed in the paper is the excessive value assigned to blue advisors (that reveal the high-payoff state) and rainbow advisors (that provide full disclosure). Using Equation 5 we can easily recognize that the value of the rainbow advisor is equal to the highest value among the other three advisors. Participants do not seem to follow this rule, as they tend to pay much less for the rainbow advisor. Figure 24 shows the distribution of the differences, within each trial, between the WTA for the rainbow advisor and the maximum of the other three WTA. Participants are willing to pay strictly less 39% of the times, and they are willing to pay strictly more only 17% of the times.

Finally, we want to show that the extra WTA for blue and rainbow advisors is due to a fixed premium that participants are willing to add, and not because of different elasticity to

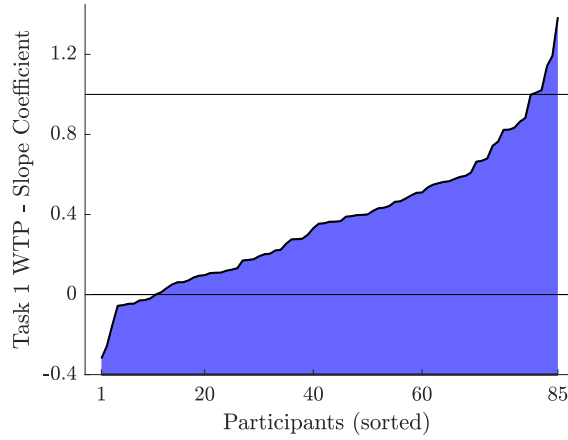


Figure 23: Distribution of participants' responses to the instrumental value of the advisors in Task 1.

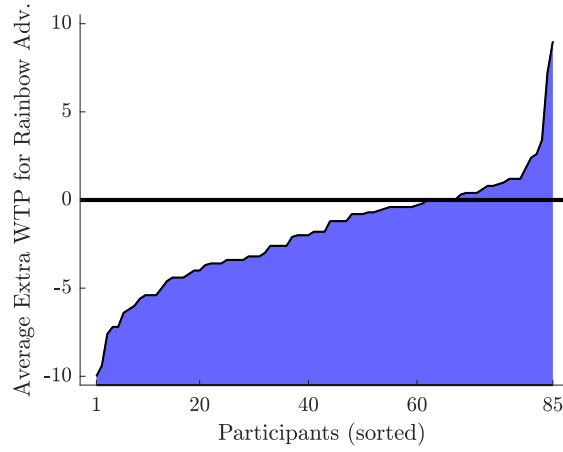


Figure 24: Distribution of participants' average excessive WTP for the rainbow advisor with respect to the highest WTP among the three simple advisors.

the instrumental values. In Table 3 we assumed that advisors' values can have different intercepts but share the same slope. We relax the assumption in a series of regressions displayed in Table 6. Compared to the benchmark model (column 1, same slope and intercepts for all), we notice a major improvement in the fit when we add the advisor-specific intercepts, which is not as much as for advisor-specific slopes.

Method: OLS, Dependent variable: $V^i(I)$				
	(1)	(2)	(3)	(4)
Constant	6.66*** (0.166)	5.62*** (0.231)	6.80*** (0.167)	5.65*** (0.257)
Red - constant		-0.193 (0.353)		-0.496 (0.447)
Blue - constant		3.41*** (0.349)		3.42*** (0.422)
Rainbow - constant		2.74*** (0.373)		2.96*** (0.505)
$V^{\text{Bayes}}(I)$ - slope	0.372*** (0.027)	0.269*** (0.030)	-0.123 (0.164)	0.222 (0.172)
Red - slope			0.252 (0.163)	0.099 (0.181)
Blue - slope			0.631*** (0.164)	0.038 (0.180)
Rainbow - slope			0.550*** (0.162)	0.009 (0.181)
Trials	All	All	All	All
Observations	2,520	2,520	2,520	2,520

Table 6: Aggregate valuations of information structures in task 1.

Statistical significance: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

J Appendix: Further Analysis of Questionnaire

We combine demographic information with two additional tasks (Holt-Laury test of risk attitude and Raven matrices as a measure of cognitive ability) and a final questionnaire with questions about the field of study, mathematical literacy, and other tests (Revised Life Orientation Test, LOT-R, as a measure of optimism, questions on superstition, and questions on risk attitude). Table 7 shows that neither of our measures of mathematical aptitude and cognitive style is significantly associated with our measure of within-subject polarization. Risk attitude, measured by the Holt and Laury test, shows a negative coefficient (a high risk seeking score is associated with a low polarization score), but the effect disappears once we introduce the demographic controls. Tables 8 - 12 show analogous analyzes for accuracy of advisor choice in task 2, probability of selecting the simple advisor in task 2, probability of selecting the risky action in task 1, accuracy in beliefs elicitation (tasks 3 and 4) and WTP slope in task 1, respectively.

Method: OLS, Dependent variable: Polarization score				
	Baseline	Full	Baseline	Full
	(1)	(2)	(3)	(4)
Risk attitude (Holt and Laury)	-0.52*** (0.16)	-0.50*** (0.16)	-0.27 (0.24)	-0.26 (0.25)
Fluid intelligence (Raven test)	0.13 (0.11)	0.10 (0.14)	0.20 (0.12)	0.07 (0.15)
Familiar with Bayes rule	0.03 (0.10)	0.02 (0.10)	0.10 (0.11)	0.12 (0.09)
Analytical studies	0.09 (0.09)	0.10 (0.10)	0.06 (0.10)	0.07 (0.11)
LOT-R scale		-0.03 (0.04)		-0.06 (0.05)
SUPERSTITION scale		-0.03 (0.04)		-0.01 (0.05)
RISK scale		-0.02 (0.04)		-0.07* (0.04)
Observations	63	63	63	63
Demographic Controls			✓	✓

Table 7: Polarization score. Statistical significance: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Method: OLS, Dependent variable: Accuracy of advisor choice in task 2				
	Baseline	Full	Baseline	Full
	(1)	(2)	(3)	(4)
Risk attitude (Holt and Laury)	-0.36*** (0.09)	-0.36*** (0.10)	-0.28* (0.15)	-0.28* (0.16)
Fluid intelligence (Raven test)	0.19** (0.07)	0.16 (0.10)	0.22*** (0.08)	0.13 (0.10)
Familiar with Bayes rule	0.09* (0.05)	0.11* (0.06)	0.17*** (0.06)	0.20*** (0.06)
Analytical studies	-0.03 (0.05)	-0.03 (0.05)	-0.06 (0.05)	-0.06 (0.06)
LOT-R scale		-0.01 (0.03)		-0.04 (0.03)
SUPERSTITION scale		0.02 (0.03)		0.03 (0.03)
RISK scale		-0.00 (0.03)		-0.03 (0.03)
Observations	63	63	63	63
Demographic Controls			✓	✓

Table 8: Accuracy of advisor choice in task 2. Statistical significance: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Method: OLS, Dependent variable: Probability of selecting the simple advisor in task 2				
	Baseline (1)	Full (2)	Baseline (3)	Full (4)
Risk attitude (Holt and Laury)	-0.07 (0.11)	-0.08 (0.11)	-0.17 (0.15)	-0.18 (0.15)
Fluid intelligence (Raven test)	-0.06 (0.08)	-0.08 (0.10)	-0.08 (0.09)	-0.11 (0.13)
Familiar with Bayes rule	-0.02 (0.04)	-0.01 (0.05)	-0.04 (0.05)	-0.03 (0.06)
Analytical studies	-0.01 (0.04)	-0.01 (0.04)	0.00 (0.04)	-0.00 (0.04)
LOT-R scale		0.01 (0.02)		0.01 (0.03)
SUPERSTITION scale		0.02 (0.02)		0.04 (0.03)
RISK scale		-0.01 (0.03)		-0.00 (0.03)
Observations	63	63	63	63
Demographic Controls			✓	✓

Table 9: Probability of selecting the simple advisor in task 2. Statistical significance: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Method: OLS, Dependent variable: Probability of selecting the risky action in task 1				
	Baseline (1)	Full (2)	Baseline (3)	Full (4)
Risk attitude (Holt and Laury)	0.03 (0.08)	0.03 (0.08)	-0.00 (0.10)	0.00 (0.10)
Fluid intelligence (Raven test)	0.05 (0.05)	0.06 (0.07)	0.03 (0.06)	0.05 (0.07)
Familiar with Bayes rule	-0.00 (0.04)	-0.02 (0.04)	-0.05 (0.04)	-0.06 (0.04)
Analytical studies	-0.01 (0.03)	0.00 (0.03)	0.00 (0.03)	0.01 (0.03)
LOT-R scale		0.01 (0.02)		0.02 (0.02)
SUPERSTITION scale		-0.02 (0.02)		-0.02 (0.02)
RISK scale		-0.01 (0.02)		-0.01 (0.02)
Observations	63	63	63	63
Demographic Controls			✓	✓

Table 10: Probability of selecting the risky action in task 1. Statistical significance: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Method: OLS, Dependent variable: Accuracy in beliefs elicitation (tasks 3 and 4)				
	Baseline (1)	Full (2)	Baseline (3)	Full (4)
Risk attitude (Holt and Laury)	−0.59 (0.68)	−0.60 (0.63)	−0.40 (0.50)	−0.51 (0.53)
Fluid intelligence (Raven test)	0.45 (0.39)	0.45 (0.41)	0.34 (0.35)	0.48 (0.35)
Familiar with Bayes rule	0.26 (0.20)	0.35* (0.20)	0.37 (0.26)	0.42 (0.30)
Analytical studies	−0.15 (0.19)	−0.21 (0.19)	−0.30 (0.20)	−0.30 (0.20)
LOT-R scale		−0.03 (0.13)		0.07 (0.12)
SUPERSTITION scale		0.13 (0.14)		0.11 (0.15)
RISK scale		0.10 (0.09)		0.14 (0.11)
Observations	63	63	63	63
Demographic Controls			✓	✓

Table 11: Accuracy in beliefs elicitation (tasks 3 and 4). Statistical significance: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Method: OLS, Dependent variable: WTP slope in task 1				
	Baseline (1)	Full (2)	Baseline (3)	Full (4)
Risk attitude (Holt and Laury)	−0.42** (0.20)	−0.39* (0.21)	−0.12 (0.27)	−0.15 (0.28)
Fluid intelligence (Raven test)	0.27 (0.18)	0.31 (0.24)	0.31* (0.16)	0.33 (0.21)
Familiar with Bayes rule	0.20* (0.10)	0.17 (0.11)	0.14 (0.12)	0.15 (0.13)
Analytical studies	−0.20** (0.10)	−0.16 (0.11)	−0.20** (0.09)	−0.16* (0.10)
LOT-R scale		0.02 (0.07)		0.03 (0.06)
SUPERSTITION scale		−0.06 (0.05)		−0.03 (0.06)
RISK scale		−0.02 (0.05)		−0.01 (0.06)
Observations	63	63	63	63
Demographic Controls			✓	✓

Table 12: WTP slope in task 1. Statistical significance: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

K Appendix: Further Analysis of Heterogeneity Across Subjects

We conducted a cluster analysis of the participants’ advisor choices using two approaches. In the first approach, we use simple clustering based on two dimensions depicted in Figure 11, i.e., the probability that the best advisor is selected and the probability of selecting the simplest advisor.

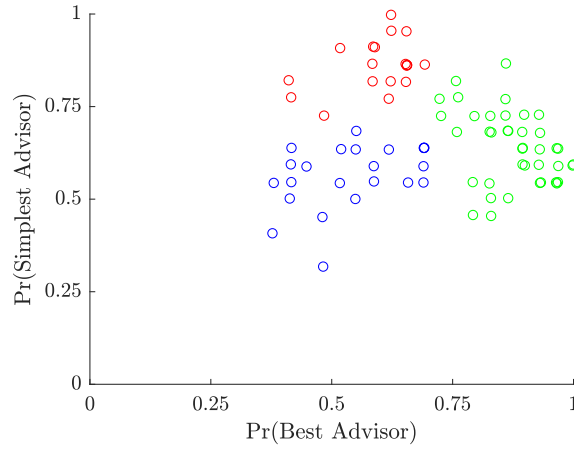


Figure 25: Distribution of participants’ advisor choices: the probability of choosing the best advisor (based on instrumental value) and simplest advisor (based on the complexity score). Clustered based on the approach 1.

	Population	Cluster		
		Red	Blue	Green
N	85	20	23	42
% best advisor	0.721	0.593	0.54	0.881
% simple advisor	0.671	0.868	0.563	0.636
% polarization (with Subjective beliefs)	0.531	0.302	0.366	0.731
% muted polarization (with Bayesian beliefs)	0.712	0.576	0.596	0.841
Avg beliefs slope in task 4 (1=Bayesian)	0.872	0.801	0.821	0.933
Avg Raven score	0.424	0.33	0.383	0.49

Table 13: Summary statistics for each cluster based on the approach 1.

As we describe in the main text (Section 6: Heterogeneity across subjects), the participants can be categorized into three broad groups based on these two dimensions. A cluster of accurate participants that display little or no bias on the right side - Green cluster, a group of simplicity-driven participants consistently selecting the advisor with lower complexity on

the top - Red cluster, and a smaller group of participants whose advisor choices are close to random - Blue cluster.

We further explore to what extent they show signs of updating beliefs in a Bayesian fashion and at the same time how much the polarization is mitigated for the participants corresponding to these three clusters.

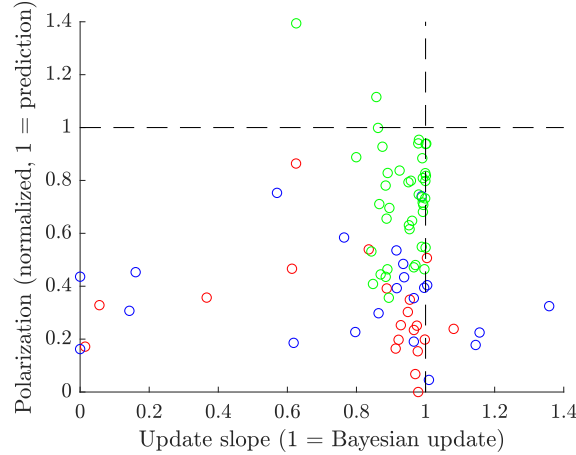


Figure 26: Distribution of participants' polarization and belief updating. Clustered according to participants' advisor choices - approach 1

Figure 26 and table 13 indicate that participants corresponding to all three groups are updating beliefs close to the Bayesian fashion, but those from the blue and the red clusters deviate slightly more from Bayesian updating than the participants from the green cluster. This together with advisor choices of each cluster provides an explanation for the degree of polarization mitigation. In particular, the green cluster's polarization is closest to the predicted one and the red cluster's polarization is the most mitigated. This is due to the fact that the green cluster is best in selecting the best advisor and preference for the simplest advisor is a little bit above one-half. The blue cluster on one side has almost no preference for the simplest advisor but identifies the best advisor almost randomly. In addition, deviates non-negligibly from Bayesian updating. Thus, the blue cluster ranks second in the degree of mitigated polarization. The red cluster has the most mitigated polarization, as it demonstrates the strongest preference for the simplest advisor while being only slightly above the coin flip from identifying the best advisor and being non-Bayesian on a similar level as the blue cluster. In table 13 we report how much the mitigated polarization changes when

subjective beliefs are replaced by Bayesian beliefs. Importantly, the green cluster accounts for almost half of all participants. These participants also show a marginally higher average Raven score.

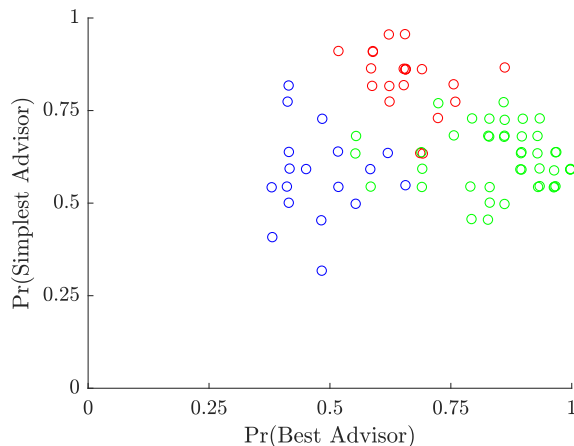


Figure 27: Distribution of participants' advisor choices: the probability of choosing the best advisor (based on instrumental value) and simplest advisor (based on the complexity score). Clustered based on the approach 2.

	Population	Cluster		
		Red	Blue	Green
N	85	23	18	44
% best advisor	0.721	0.654	0.477	0.856
% simple advisor	0.671	0.848	0.576	0.618
% polarization (with Subjective beliefs)	0.531	0.341	0.337	0.71
% muted polarization (with Bayesian beliefs)	0.712	0.616	0.58	0.817
Avg beliefs slope in task 4 (1=Bayesian)	0.872	0.868	0.739	0.928
Avg Raven score	0.424	0.4	0.378	0.455

Table 14: Summary statistics for each cluster based on the approach 2.

In the second approach, we do clustering based on the vector of all the choices for each subject. Each subject's vector of all the choices consists of 40 choices, where each is 0 or 1 (referring to the advisor selected).

We report the same figures as in the previous case (Figure 27-28) and summary statistics in table 14. The main difference between approaches 1 and 2 is in the probability of selecting the best advisor for clusters and that the red cluster - with a high preference for the simplest advisor is more numerous.

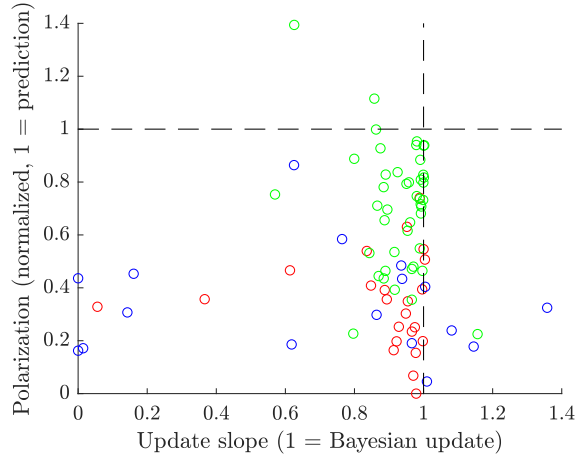


Figure 28: Distribution of participants' polarization and belief updating. Clustered according to participants' advisor choices - approach 2

In the rest of this appendix, we explore participants' over-reaction to the evidence as we documented by polarization exceeding the predicted values (above 1 on the normalized score) in Figure 10a. Specifically, if non-Bayesian agents over-react to the evidence received with the signal, they can reach more extreme posterior beliefs than the ones of a Bayesian agent. In this case, we would observe an actual magnitude of polarization that is larger than the predicted one. From figure 26, we observe that most of the participants show some conservatism in their response, and over-polarization occurs in participants with conservative beliefs. The two channels that can affect (reduce/increase) polarization are 1) non-Bayesian update (yet here it typically reduces polarization due to conservatism), and 2) demand for information. Figures 29 below show the decomposition of these two effects.

Both in the aggregate and at the subject level, non-Bayesian beliefs reduce the magnitude of polarization (with few exceptions, associated with conservatism). This is still possible because, despite the beliefs being on average conservative, there is dispersion around the regression curve and few participants over-react only to some evidence, but not in a systematic way.

Changes coming from the change due to information demand always reduce the magnitude of polarization, by construction (these are the trials in which the model always predicts the agents should switch, so there is no way for this channel to augment polarization).

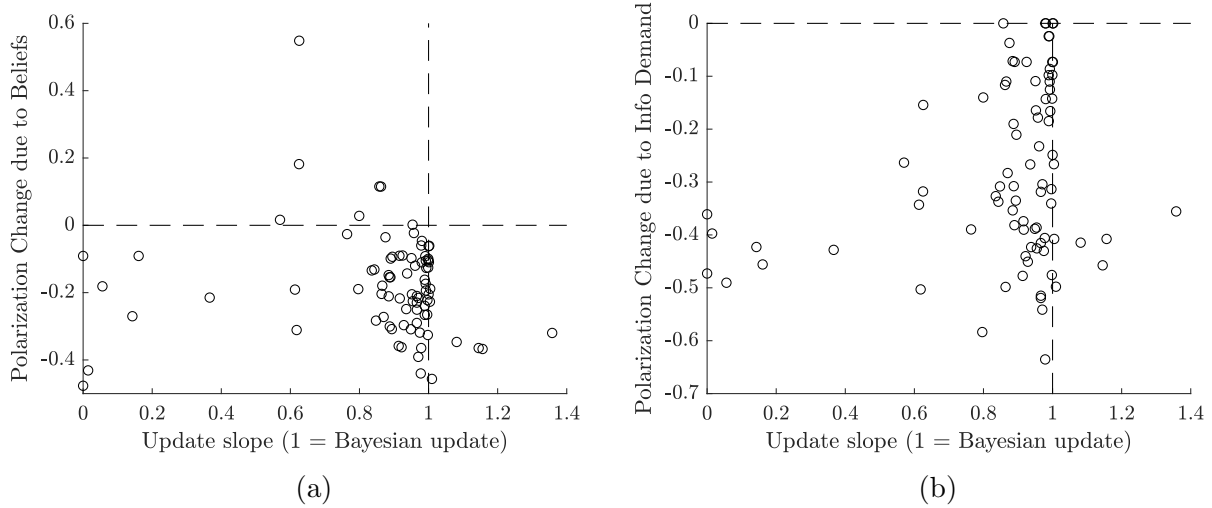


Figure 29: Decomposition of polarization according to two channels. Left: Change in polarization due to beliefs. The plot shows the difference between polarization with subjective beliefs and polarization with Bayesian beliefs, both normalized. Right: Change in polarization due to information demand. The plot shows the difference between polarization with Bayesian beliefs and model predictions, both normalized.

L Appendix: Prediction of the status quo type

We now look at our framework from the perspective of a platform that wants to infer the type (status quo value) of the decision maker and has access to a dataset with some observable activities. We can divide these activities into two groups: final actions, like *voting* or the choice between the status quo and a new policy, and information acquisition, like *reading newspapers* or selecting an advisor in our design. For a more concrete example, imagine a social media platform like Facebook or Twitter, that has access to a dataset of actions performed by its users. These actions include publicly observable actions (likes, list of friends or followers), but also a series of additional actions (clicks, searches) that involve the process of information acquisition. Are these search activities helpful in improving the prediction of the type of the user, on top of the observable actions?

We consider separate scenarios in which the platform has access to choices only (opaque or transparent box), searches only (advisor X or advisor Y), or both, under the assumptions of our model (rational decision makers) and in the dataset collected in the laboratory experiment. Table 15 shows the results of this exercise: having access to the search data guarantees a much higher accuracy with respect to the action data, with minor improve-

	Prediction	Data
No information	50.0%	50.0%
Choice only	69.7%	62.6%
Search only	100.0%	68.0%
Search+Choice	100.0%	68.4%
Search+Signal+Choice	100.0%	72.9%

Table 15: Inference of the agent’s status quo: predicted and realized accuracy (pairs of trials with expected advisor switch only). The table indicates the accuracy of the prediction of the type (status quo) of the decision maker based on the data available. The model’s predictions are based on rational and unbiased agents. The accuracy realized refers to the dataset collected in the laboratory experiment.

ments when both datasets are available. When we consider the trials in which we expect to observe an advisor switch (column 2), the type prediction accuracy with advisor choices is 68%, but it is only 62.6% when we observe only the final actions. When both datasets are available, the accuracy increases marginally to 68.4%, with a further improvement if the signal realization (that occurs between search and choice) is also observed.

We can conclude that, in this simple setup, the data about the choice over sources of information is more valuable than the final action from the perspective of an observer who wants to infer the type (status quo value) of the decision maker.

M Appendix: Timeline of the Problem

Our experimental design allows us to estimate how agents evaluate an informative signal structure (advisor), and measure how the subjective evaluation depends on the properties of the signal structure, including instrumental value (expected improvement in the choice process) and non-instrumental properties (ease of interpretation).

The timing of the problem (as in task 2) can be summarized as follows:

1. The agent is informed of the prior $\mathcal{P}(s) = \frac{1}{3} \forall s$ and the state-contingent returns $\{v_s\}_S, R$.
2. One state is realized, but the agent is unaware of it.
3. The agent is offered two sources of information (advisors) I_1 and I_2 .

4. The agent chooses one advisor and discards the other.
5. The selected advisor observes the realized state (ball in the opaque box).
6. The selected advisor returns a binary signal, whose likelihood depends on the realized state.
7. The agent observes the realized signal.
8. The agent chooses one action (opaque or transparent box) and receives the payoff π .
9. The agent plays a lottery and receives the final prize k with probability $\frac{\pi}{100}$.

The problem presented in task 1 is similar up to a change in steps 3 and 4:

- 3'. The agent is offered one single source of information (advisor) I
- 4'. The agent indicates how much she is willing to accept renunciation of the advisor.

In the *Colorblind advisor game* (task 1), we elicit the probability w_I such that the agent is indifferent between making a choice after observing the realization of a known signal structure I , and choosing without additional signals but receiving additional w_I tickets to win the prize. In the *Imprecise advisor game* (task 2), we offer pairs of signal structures, and collect binary choices between advisors. If the valuation and choices differ from those a Bayesian expected utility maximizer would display, we would like to pinpoint the source of the deviation. For this reason, we add two control tasks to elicit a subjective signal of beliefs' realization (*Card color prediction game*, task 3) and subjective posterior beliefs (*Ball color prediction game*, task 4). We collect posteriors only after eliciting preference over advisors, so we do not nudge the subjects towards thinking about information valuation in a specific fashion.