

Discussion Paper Series – CRC TR 224

Discussion Paper No. 296
Project B 02

Bayesian Persuasion With Costly Information Acquisition

Ludmila Matysková¹
Alfonso Montes²

April 2021

¹ University of Bonn, Institute for Microeconomics, Adenaueralle 24-42, 53113, Bonn, Germany;
email: imatysko@uni-bonn.de (corresponding author)

² Université Cergy-Pontoise, Laboratoire THEMA Chenes 1 - C145 33, Boulevard du Port, 95011, Cergy-Pontoise Cedex,
France; email: alfonso.montes-sanchez@cyu.fr

Funding by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation)
through CRC TR 224 is gratefully acknowledged.

Bayesian Persuasion With Costly Information Acquisition*

Ludmila Matysková[†]

University of Bonn

Alfonso Montes[‡]

Université Cergy-Pontoise

April 22, 2021

Abstract

A sender choosing a signal to be disclosed to a receiver can often influence the receiver’s actions. Is persuasion harder when the receiver has additional information sources? Does the receiver benefit from having them? We extend Bayesian persuasion to a receiver’s acquisition of costly information. The game can be solved as a standard Bayesian persuasion under an additional constraint: the receiver never learns. The ‘threat’ of learning hurts the sender. However, the outcome can also be worse for the receiver, in which case the receiver’s possibility to gather additional information decreases social welfare. Furthermore, we propose a new solution method that does not rely directly on concavification, which is also applicable to standard Bayesian persuasion.

Keywords: Bayesian persuasion, Rational inattention, Costly information acquisition, Information design

JEL classification: D72, D81, D82, D83

*For valuable comments, we thank: Andrew Caplin, Yeon-Koo Che, Mark Dean, Olivier Gossner, Navin Kartik, Frédéric Koessler, Andrei Maatvenko, Helene Mass, Filip Matějka, Laurent Mathevet, Margaret Meyer, Konrad Mierendorff, Stephen Morris, Vladimír Novák, Pietro Ortoleva, Avner Shaked, Vasiliki Skreta, Jakub Steiner, Dezső Szalay, Ina Taneva, Tristan Tomala, Michael Woodford, and Jan Zápál. Matysková gratefully acknowledges funding by the CRC TR 224 (project B02) provided by the German Research Foundation (DFG), by the H2020-MSCA-RISE (GEMCLIME-2020 GA No. 681228), and by the Charles University (GA UK no. 178715). Montes gratefully acknowledges funding by the Ecole Polytechnique, by the National Commission of Scientific and Technological Research of the Government of Chile, Ref. no. 72180147, and by the Labex MME-DII.

[†]University of Bonn, Institute for Microeconomics, Adenaueralle 24-42, 53113, Bonn, Germany; email: lmatysko@uni-bonn.de (corresponding author)

[‡]Laboratoire THEMA Chenes 1 - C145 33, Boulevard du Port, 95011, Cergy-Pontoise Cedex, France; email: alfonso.montes-sanchez@cyu.fr

1 Introduction

Decision makers often rely on free information provided by an interested party. Lobbying groups, such as tobacco and pharmaceutical companies, commission research activities with the goal of influencing politicians. Firms try to make sells by providing information to their potential customers; car dealers allow for limited inspection of their cars, and software developers offer trial versions of their products. The decision makers, however, may also be able to obtain independent information at a costly effort. Politicians can carry out their own research on tobacco and drugs, and customers can find online information about the characteristics of the product. In this paper, we study the welfare effects of access to independent information in such settings.

We consider a Bayesian persuasion model (Kamenica and Gentzkow 2011, henceforth KG) in which the receiver has access to additional information sources. As in KG, a sender chooses a signal to disclose information to a receiver, the decision maker. Unlike in KG, however, the receiver further chooses her own signal at a cost to acquire more information before taking an action. Within this framework, we address the following questions. Does the receiver benefit from having access to additional information? Is persuasion harder for the sender in such a case?

We solve the model and describe comparative statics with respect to the receiver's cost of information. We show that the possibility of additional learning reduces the sender's persuasive power, thus lowering his expected equilibrium utility. On the other hand, the effect on the receiver's expected equilibrium utility is positive in a binary state, binary action setting. However, in more general settings, the possibility of additional learning can be detrimental *not only* to the sender, *but also* to the receiver. If there is significant disagreement about preferred actions under full information, the sender optimally garbles a fully informative signal in order to prevent the receiver from taking certain actions. Whenever the receiver can also obtain further information, the sender might provide even less information. In turn, this can lead to less overall information being disclosed in the equilibrium. In such

cases, the receiver's access to additional information unambiguously decreases the social welfare.

We demonstrate the possibility of such detrimental effect in an application on *difference of opinions: extremism versus conservatism*. In the application, the payoffs are given by quadratic loss functions arising from the distance between the receiver's action and state. The players agree on the type of the action (negative vs positive), but they differ in their opinions on the level of the action (extreme vs conservative) as new information is generated. When the *receiver* is the *conservative player*, she under reacts to new information. She can never be persuaded to consider extreme levels, not even under full information. Then, the sender optimally provides full information so as to (at least) minimize mistakes in the type of actions. This effect is independent of how costly the receiver's information is, and hence, *a conservative receiver is not affected by changes in her cost of information*.

When the *receiver* is the *extreme player*, she overreacts to new information. Then, the sender optimally garbles a fully informative signal in order to prevent the receiver from choosing too extreme levels. The sender faces a trade-off between preventing extreme levels that would be taken if he provides too much information (or preventing further learning, which could also possibly lead to an extreme level) and minimizing the probability of mistakes in the type of action if he provides too little information. This tension becomes greater as the receiver's information becomes cheaper, leading the sender to provide less information (in Blackwell sense). In equilibrium, the receiver does not learn, and thus *an extreme receiver becomes worse off as her cost of information decreases*.

Our general model assumes finitely many actions and states, general preferences and uniformly posterior-separable cost function of receiver's information. To solve the model, we develop an *extreme-point solution method* that simplifies the task of finding equilibria. We fully characterize the solution for a specific, entropy-based cost, commonly used in models of rational inattention (Sims 2003). This characterization is applicable also at the limiting case of our model, the standard Bayesian persua-

sion, and it is thus complementary to the geometrical methods of concavification used in the Bayesian persuasion literature.¹

The extreme-point method relies on two results. First, Proposition 1 (Non-Learning Equilibrium), the main simplification step, shows that the game can be solved as a standard Bayesian persuasion under an additional constraint: the receiver never acquires additional information in equilibrium. The reason is that all the information that the receiver would have gathered following the sender’s signal can always be provided on her behalf by the sender. Therefore, once we partition the belief space into learning and non-learning regions (a non-learning region of a certain action is defined as all beliefs at which receiver does not learn and takes that action right away), the sender’s strategy optimally induces only posterior beliefs in the non-learning regions. To find a solution, it is thus sufficient to characterize the non-learning regions. This simplifies the analysis, because otherwise we need to identify the receiver’s optimal learning strategy for every belief in the learning regions—and there are infinitely many such beliefs—which becomes intractable when the state space is high dimensional.

Second, Proposition 2 (Extreme-Point Equilibrium) states that there always exists an optimal sender’s strategy that is supported only on the extreme points of the non-learning regions. For the entropy-based cost, we further provide a full characterization of the extreme points (Lemma 4), which implies that there are only finitely many of them (Corollary 1). Hence, we only need to consider finitely many sender’s strategies to find a solution, which generally may not be true under an arbitrary uniformly posterior-separable costs. We describe the steps how to find a solution with the entropy-based cost in *Extreme-point solution algorithm* in Section 3.3.1. We use this algorithm when solving the application.

In a binary action-state setting, there are two non-learning regions. Each is an interval with two extreme points. Extreme-point method then implies that only very few

¹Originally proposed by Aumann, Maschler, and Stearns (1995), concavification methods provides the analyst geometrical tools to solve the model. KG introduced this method in models of Bayesian persuasion, and Caplin and Dean (2013) introduced this method in decision problems with the entropy-based cost.

sender’s strategies must be considered for an optimum. It turns out that the sender optimally provides (i) no information (optimal when players disagree on preferred action in each state, in which case the sender does not benefit from persuasion), (ii) full information (optimal under aligned preferences), or (iii) information that maximizes the probability of one of the actions such that additional learning was prevented (optimal when the sender has a dominant action). In the binary-state setting, where the belief simplex is a line, extreme-point method then gives us an easy way to compare how the solution changes in the Blackwell order as the receiver’s cost parameter varies: we only need to determine whether the extreme points that are induced by the optimal strategy are moving apart from or towards each other. We use this observation when establishing our comparative statics results in a binary action-state setting and in the application.

Related literature. This paper lies at the intersection of the literature on Bayesian persuasion and on costly information acquisition. We extend the standard Bayesian persuasion model to consider an endogenously privately informed receiver.² Our model is closest in spirit to Bizzotto, Rüdiger, and Vigier (2020). In a binary action, binary state framework, they consider a receiver who, after receiving the sender’s information, can additionally obtain a binary signal of fixed precision by paying a fixed fee.³ Both Bizzotto, Rüdiger, and Vigier (2020) and we derive a similar result showing that the receiver may prefer commitment to worse information technology. However, the mechanisms beyond this result differ. In Bizzotto, Rüdiger, and Vigier (2020), the result stems from rigidity of receiver’s information acquisition technology, which the sender takes advantage of. In contrast, we show that when the receiver’s information technology is flexible, and thus this channel is not present, this result

²Extensions with an exogenously privately informed receiver have already been examined in Kolotilin, Mylovanov, Zapechelnnyuk, and Li (2017) and Kolotilin (2018). Other extensions of the standard Bayesian persuasion model that consider some natural constraints for the sender range from multiple senders (Gentzkow and Kamenica 2017) to a limited communication capacity (Le Treust and Tomala 2019), heterogeneous priors (Alonso and Camara 2016) and limited commitment power for the sender (Nguyen and Tan 2021). See Kamenica (2019) for a recent survey on Bayesian persuasion literature.

³We independently explore a similar framework in Example B3 in Appendix B, used to demonstrate failure of Proposition 1. At the time this example was framed, we were not aware of the existence of the above mentioned paper.

no longer holds in the binary action, binary state setting. It exists once a more general setup is considered. In our setting, the result is driven by the sender’s desire to prevent the receiver from additional learning, which, in a more general model, can have an unfavorable effect on a receiver with better information technology.⁴

Our modelling choice of the flexible information acquisition assumes that the receiver faces a uniformly posterior-separable (UPS) cost function.⁵ Other contemporaneous papers on information provision consider a similar modelling choice. However, the game specification differs from ours. Bloedel and Segal (2018) consider a receiver who needs to pay an attention, entropy-based, cost to process the sender’s messages. Lipnowski, Mathevet, and Wei (2020a) and Lipnowski, Mathevet, and Wei (2020b) consider players with aligned preferences, where the receiver acquires information about the state at a UPS cost, provided it can only be as informative (in Blackwell sense) as an upper bound disclosed by the sender. Subsequently, Wei (2020) extends this setup to binary action, binary state with misaligned preferences. He shows, similar to us, that the receiver’s equilibrium payoff is non-monotone in her cost parameter. The reason is that when the cost parameter is low, so that the receiver processes almost any information disclosed by the sender regardless of its information value, the sender provides very little useful information for the receiver.

Our paper’s technical contribution is on solution methods for standard and extended Bayesian persuasion problems. Our extreme-point solution method complements the results of Lipnowski and Mathevet (2017) and Lipnowski and Mathevet (2018), who independently show, in a standard Bayesian persuasion model and in a model with psychological preferences with aligned interest, that it is sufficient to focus on extreme-points of sets of beliefs on which the sender’s value function is (weakly) convex. They provide general abstract conditions to characterize those

⁴The notion that an agent in a strategic setting can be hurt by having access to better information technology is not unique to Bayesian persuasion setting, e.g., see Roesler and Szentes (2017) and Kessler (1998) for such a case in a contracting environment.

⁵Gentzkow and Kamenica (2014) introduce costly information acquisition into the model by assuming that the sender faces a cost of disclosing information to the receiver. They provide a class of cost functions (including entropy-based cost) that are compatible with the concavification approach.

sets. Complementary to them, we provide specific linear conditions to characterize the extreme-points when assuming the entropy-based cost, which can be used to solve standard Bayesian persuasion models as the cost parameter goes to infinity. For a special setting with continuous state and a sender with state-independent preferences, Arieli, Babichenko, Smorodinsky, and Yamashita (2019) and Kleiner, Moldovanu, and Strack (2020) characterize optimal sender’s strategies using the observation that a solution of this specific problem is an extreme-point of the set of all sender’s strategies. However, the notion of the extreme point is different from ours. In our setting, an optimal sender’s strategy *is supported* on the extreme-points of certain sets of beliefs, while in their setting, an optimal sender strategy *itself* is an extreme point (of a set of sender’s information strategies).⁶

By the modelling choice of the cost function, we also contribute to the growing literature that applies rational inattention (Sims 2003) framework to strategic setting, such as Martin (2017), Matějka and Tabellini (2016), Montes (2020), Ravid (2020), Yang (2015) and Yang (2020). We build on insights from single-agent decision problems in Matějka and McKay (2015), Caplin and Dean (2013), and Caplin, Dean, and Leahy (2017) when solving for the receiver’s maximization problem.

The rest of the paper is organized as follows. Section 2 sets the model up. Section 3 states the main simplification result (the Never-Learning Equilibrium), describes the solution method (the Extreme-Point Equilibrium) and provides a characterization of the solution for an entropy-based cost (Extreme-point solution algorithm). Section 4 provides comparative statics. Section 5 presents an application on difference of opinions. Section 6 is a conclusion.

⁶Other alternative method for persuasion problems in which the sender’s utility depends only on the expected state is introduced by Dworczak and Martini (2019).

2 Model

There are two players, a sender (he) and a receiver (she). There is an unknown payoff-relevant state ω drawn from a finite set Ω according to a prior $\mu_0 \in \text{int}(\Delta(\Omega))$.⁷ The receiver chooses an action a from a finite set A . The payoffs of the sender and the receiver are $u, v : A \times \Omega \rightarrow \mathbb{R}$ respectively.

After the state is drawn, the sender generates public information about the realized state, inducing the receiver to rationally update her prior belief μ_0 to an interim belief $\mu \in \Delta(\Omega)$. A sender's (information) strategy is a choice of a distribution $\tau \in \Delta(\Delta(\Omega))$ over the (updated) interim beliefs such that the martingale property holds: $\mathbb{E}_\tau[\mu] = \mu_0$. We assume that the sender can choose any information strategy at zero cost. After observing the sender's information and updating her beliefs to a particular interim belief μ , the receiver decides whether to acquire additional, costly, information about the realized state, inducing her to further rationally update her interim belief μ to a posterior belief $\gamma \in \Delta(\Omega)$. The receiver's (information) strategy is a choice of a distribution $\phi \in \Delta(\Delta(\Omega))$ over the (further updated) posterior beliefs γ such that the martingale property holds: $\mathbb{E}_\phi[\gamma] = \mu$. Finally, once interim beliefs are updated to a particular posterior belief γ , the receiver chooses an action $\sigma^*(\gamma)$ where $\sigma^* : \Delta(\Omega) \rightarrow A$ is defined such as⁸

$$\sigma^*(\gamma) \in \arg \max_{a' \in A} \sum_{\omega} u(a', \omega) \gamma(\omega) \quad \forall \gamma \in \Delta(\Omega) \quad (1)$$

where $\gamma(\omega)$ denotes the probability of state ω at posterior γ . If, for a given γ , there are multiple actions satisfying (1), we assume that $\sigma^*(\gamma)$ is the action that is (weakly) preferred by the sender.

⁷ $\text{int}(S)$ denotes interior of set S , and $\Delta(S)$ denotes the set of all probability distributions on S .

⁸As the focus of this paper is on information strategies, for the sake of notation, we do not specify the action strategies in the text. Instead, we assume that at the last stage of the game, the receiver is always automatically choosing the action that maximizes her expected utility.

Let us denote the receiver's highest expected payoff at posterior belief γ as

$$U(\gamma) := \sum_{\omega} u(\sigma^*(\gamma), \omega) \gamma(\omega) \quad (2)$$

While the sender's information has no cost, the receiver's information is costly. We assume that the receiver's cost of information is uniformly posterior-separable (UPS).⁹ That is, given a bounded and strictly concave function $F : \Delta(\Omega) \rightarrow \mathbb{R}_+$ and an interim belief μ , the cost of information strategy ϕ is given by

$$c(\phi; \mu, \lambda) = \lambda (F(\mu) - \mathbb{E}_{\phi}[F(\gamma)]) \quad (3)$$

where $\lambda > 0$ and the expectation is over posteriors γ distributed according to ϕ .

A leading example of such a cost function used in the literature is based on Shannon entropy. We say that the cost function is *Shannon* if

$$F(\gamma) = - \sum_{\omega \in \Omega} \gamma(\omega) \ln \gamma(\omega) \quad (4)$$

(with $0 \log 0 = 0$ by convention). This function underpinned important developments in the literature of rational inattention, from Sims (2003), Sims (2006) to Matějka and McKay (2015) and Steiner, Stewart, and Matějka (2017).

Definition 1. Given an interim belief μ , the *receiver's maximization problem* is

$$\begin{aligned} \max_{\phi \in \Delta(\Delta(\Omega))} \quad & \mathbb{E}_{\phi}[U(\gamma)] - c(\phi; \mu, \lambda) \\ \text{s.t.} \quad & \mathbb{E}_{\phi}[\gamma] = \mu, \end{aligned} \quad (5)$$

where $U(\gamma)$ is defined by (2) and the expectation is over posterior beliefs γ distributed according to the information strategy ϕ .

Lemma 1. The receiver's maximization problem (5) has a solution.

Throughout the paper, for a probability distribution π over a set S , we denote

⁹See Caplin, Dean, and Leahy (2017).

its support by $\text{supp}(\pi)$.¹⁰ We assume that if at some μ there is more than one optimal information strategy, the receiver chooses one that is (weakly) preferred by the sender. Hence, we focus on sender-preferred equilibria in both action and information strategies. For a given interim belief μ , let ϕ_μ^* denote the solution to problem (5) which is (weakly) preferred by the sender. We say that *the receiver does not learn at μ* if $\text{supp}(\phi_\mu^*) = \{\mu\}$, so that no new information is gathered and the posterior beliefs remain at the interim belief. We use ϕ_μ^N to denote such non-informative receiver's information strategies. Otherwise, we say that *the receiver learns at μ* .

When a belief μ is drawn from the information strategy τ , the receiver gathers information according to ϕ_μ^* . Finally, if a final posterior γ is drawn from ϕ_μ^* , the action $\sigma^*(\gamma)$ is taken. Applying backward induction, we can thus express the sender's conditional expected payoff for each interim belief μ as

$$\hat{v}(\mu) := \mathbb{E}_{\phi_\mu^*} \left[\sum_{\omega \in \Omega} v(\sigma^*(\gamma), \omega) \gamma(\omega) \right] \quad (6)$$

where the expectation is over posterior beliefs γ distributed according to the receiver's optimal information strategy ϕ_μ^* . The function $\hat{v}(\mu)$ is the sender's expected payoff at interim belief μ given the optimal continuation play of the receiver at μ .

Definition 2. Given prior μ_0 , the *sender's maximization problem* is

$$\begin{aligned} \max_{\tau \in \Delta(\Delta(\Omega))} \quad & \mathbb{E}_\tau[\hat{v}(\mu)] \\ \text{s.t.} \quad & \mathbb{E}_\tau[\mu] = \mu_0, \end{aligned} \quad (7)$$

where $\hat{v}(\mu)$ is given by (6) and the expectation is over interim beliefs μ distributed according to τ .

Lemma 2. The sender's maximization problem (7) has a solution.

¹⁰Hence, $\text{supp}(\pi) = \{s \in S : \pi(s) > 0\}$.

Definition 3. A (sender-preferred subgame perfect) *equilibrium* of the game is a pair $(\tau^*, (\phi_\mu^*)_\mu)$ such that τ^* solves (7) and for every $\mu \in \Delta(\Omega)$, ϕ_μ^* solves (5).¹¹

We define the *sender's equilibrium value* v^* as his expected payoff obtained under an equilibrium profile. We say that *the sender benefits from persuasion* when v^* is strictly larger than his expected payoff obtained under an information strategy τ^N : $\text{supp}(\tau^N) = \{\mu_0\}$, a strategy under which he generates no new information. Note that by Lemmas 1 and 2, an equilibrium always exists. We say that *the receiver never learns in equilibrium* $(\tau^*, (\phi_\mu^*)_\mu)$ if for all $\mu \in \text{supp}(\tau^*)$, $\phi_\mu^* = \phi_\mu^N$, i.e., the receiver does not learn at any μ in the support of the sender's optimal strategy.

3 Solution

A common solution method used in Bayesian persuasion is a *concavification approach*. Geometrically, one builds a concave closure of $\hat{v}(\cdot)$, which we denote by $\text{cav}(\hat{v})(\cdot)$.¹² The support of an optimal sender's strategy can then be read off from the graph.¹³ The concavification approach, sometimes also called a *posterior-based approach*, can also be used when solving the receiver's maximization problem (see Caplin and Dean (2013)).

For a binary state, applying the concavification method to both the receiver's and the sender's maximization problems is a tractable way to solve the model, because the receiver's solution implies $\hat{v}(\mu)$ is a piecewise linear function. The procedure is illustrated on Example B1 in Appendix B. Nevertheless, a double concavification approach becomes intractable with three or more states of nature, because the solution to the receiver's concavification problem does not imply a simple functional form

¹¹For the sake of notation, we do not include an action profile as the part of the definition of the equilibrium. Note, however, that the definition implicitly assumes that $\sigma(\gamma) = \sigma^*(\gamma)$ for all $\gamma \in \Delta(\Omega)$.

¹²Concave closure of $\hat{v}(\cdot)$ is the smallest concave function that is everywhere weakly greater than $\hat{v}(\cdot)$.

¹³The support of the optimal sender's strategy is the set of the interim beliefs that support the tangent hyperplane to the lower epigraph of the concavification above the prior belief.

for $\hat{v}(\mu)$. Then, it requires to determine $\hat{v}(\mu)$ for each belief separately and there is infinitely many such beliefs. The next section presents the main simplification step that allows us to solve the model without having to specify $\hat{v}(\mu)$ for all the beliefs.

3.1 Persuasion With No Learning

When the receiver is fairly uncertain about what the right thing to do is, she decides to gather some more information before taking an action. In turn, when her interim belief is precise enough, she does not learn and takes an action right away. The latter set of interim beliefs is important for our further analysis.

Definition 4. A *non-learning region* of action $a \in A$ is

$$NL^a := \left\{ \mu \in \Delta(\Omega) : \phi_\mu^N \text{ solves (5) and } a \in \arg \max_{a' \in A} \sum_{\omega} u(a', \omega) \gamma(\omega) \right\}. \quad (8)$$

The non-learning region of some action a is the set of all the interim beliefs at which taking that action right away instead of learning is optimal from the receiver's point of view. The set $\cup_a NL^a$ is non-empty as the vertices of the belief space always belong to $\cup_a NL^a$.¹⁴ Our first result states that there exists an equilibrium in which the receiver never learns. We call such an equilibrium a *non-learning* equilibrium.

Proposition 1 (Non-Learning Equilibrium). There exists an equilibrium in which the receiver never learns. That is, there exists $(\tau^*, (\phi_\mu^*)_\mu)$ such that $\forall \mu \in \text{supp}(\tau^*)$, $\phi_\mu^* = \phi_\mu^N$.

The intuition for Proposition 1 is the following: since the sender faces no cost of information, he can incorporate, at no cost, any subsequent receiver's learning strategy and save the learning part of the receiver. This new equilibrium preserves

¹⁴Furthermore, although it can be true that at some belief $\mu \in \cup_a NL^a$: $\phi_\mu^N \neq \phi_\mu^*$ —when there are multiple solutions to the receiver's maximization problem (5) and the sender prefers the learning strategy—it is also always true that $\phi_\mu^N = \phi_\mu^*$ at the vertices of the belief simplex, because no other strategy is available. Hence, the set of $\cup_a NL^a$ at which the receiver does not learn is also non-empty.

the final distribution of posteriors because the receiver will not engage in further learning, and so the expected payoff of the sender remains unchanged. Note that the receiver’s expected payoff may differ between an arbitrary equilibrium and a non-learning equilibrium, since the receiver may be undertaking costly learning in an arbitrary equilibrium. Hence, non-learning equilibria are Pareto dominant. Finally, note that whenever the equilibrium is unique, it is a non-learning equilibrium.

Proposition 1 is the main simplification step in our analysis, and it relies on several key properties of UPS cost functions. First, the learning technology does not vary with the interim belief and has an additive form. This implies that the support of any receiver’s optimal information strategy at any interim belief consists only of beliefs belonging to some non-learning regions (see Lemma 5 in Appendix A). That is, once the receiver learns, she does not wish to learn further even when given an extra chance to do so. Then, if the sender provides the desired information on the receiver’s behalf, she does not engage in any further learning afterwards, because her optimal learning behavior at a certain belief is independent of how she arrived at the belief. On the contrary, this property of receiver’s behavior can fail under a (not uniform) posterior-separable (PS) cost function, where the learning technology—the function F or the scaling parameter λ in (3)—is interim-specific. For instance, when the receiver’s learning technology becomes more efficient as her interim beliefs become more precise, provision of information on receiver’s behalf can change her subsequent behavior: the receiver would sometimes follow up with additional learning that would not have occurred otherwise.¹⁵ If this additional learning—which would not have happened if the sender were to leave the receiver to learn by herself—is not convenient for the sender, a non-learning equilibrium may fail to exist (see Example B2 in Appendix B).

Second, UPS cost functions do not restrict the set of receiver’s information strategies beyond Bayes’ law so that the receiver’s information technology is as flexible as the sender’s information technology. On the other hand, when the set of available

¹⁵One can think that when the interim belief is more precise, it may be easier for the agent to learn because she already knows where to find the information or whom to ask for it.

information strategies is constrained, Proposition 1 can fail. The sender can take advantage of the rigidity of the receiver's available strategies and optimally induce the receiver to learn in equilibrium (see Example B3 in Appendix B, when the receiver has available a binary signal of fixed precision by paying a small fixed fee c).

3.2 Extreme-Point Solution Method

Next, we show that there exists an equilibrium in which the posteriors induced by sender's strategy actually consist of (at most $|\Omega|$) *extreme points* of the (convex hull of the) non-learning regions.¹⁶ We call such an equilibrium an *extreme-point equilibrium*. Formally, for each $a \in A$, let EP^a denote the set of extreme points of the convex hull of NL^a .

Proposition 2 (Extreme-Point Equilibrium). There exists a non-learning equilibrium in which the sender's strategy is supported in (at most $|\Omega|$) extreme points of the convex hull of the non-learning regions. That is, there exists $(\tau^*, (\phi_\mu^*)_\mu)$ such that $|\text{supp}(\tau^*)| \leq |\Omega|$ and for all $\mu \in \text{supp}(\tau^*)$ we have $\mu \in \cup_a EP^a$ and $\phi_\mu^* = \phi_\mu^N$.

The rationale for Proposition 2 is the following. First, Proposition 1 implies that we can focus on non-learning equilibria. Then, if the sender induces an interim belief that lies in a non-learning region, but it is not an extreme point, it can instead be replaced by a distribution over a set of extreme points (of the convex hull) of the particular non-learning region. Doing this is feasible, because every belief in the non-learning region can be expressed as a convex combination of extreme points of its convex hull. It does not change the sender's expected payoff because the action that the receiver takes at the newly induced beliefs remains the same. Finally, the restriction on the size of the support of the optimal sender's strategy follows from the Carathéodory theorem. Note that, just as in Proposition 1, the expected payoff

¹⁶Recall that an extreme point of a convex set B is a point in the boundary of B which does not lie in any open line segment joining two points of B .

of the receiver may differ between an arbitrary non-learning equilibrium and an extreme-point equilibrium.

Proposition 2 provides an algorithm that can be used to find an equilibrium. We illustrate the algorithm in the next section when applied to the Shannon cost.

3.3 Shannon Model

In this subsection, we provide a characterization of a solution with the Shannon cost. The following two lemmas are based on known results in the literature of optimal behavior in decision problems with Shannon cost. Lemma 3, which uses equation (5) of Proposition 2 of Caplin, Dean, and Leahy (2019), gives a specific set of *linear* inequalities for an interim belief μ to be in a non-learning region. Lemma 4 then states that $|\Omega| - 1$ of these inequalities must be binding for the interim belief μ to also be an extreme point of the non-learning region.

Lemma 3. When the cost is Shannon, a non-learning region of action a is given by

$$NL^a = \left\{ \mu \in \Delta(\Omega) : \sum_{\omega \in \Omega} \mu(\omega) e^{d_{\omega}^{\lambda}(a, a')} \leq 1 \quad \forall a' \neq a \right\} \quad (9)$$

where $d_{\omega}^{\lambda}(a, a') = \frac{u(a', \omega) - u(a, \omega)}{\lambda}$.

Lemma 3 implies that for the Shannon cost, the non-learning regions are defined by polytopes, which in turn guarantees that the set of extreme points is always finite.

Corollary 1. When the cost is Shannon, the set $\cup_a EP^a$ is finite.

Corollary 1 is important, because it implies that only finitely many sender's strategies need to be considered in order to find a solution. This is not the case if $\cup_a EP^a$ is infinite, which could occur in the case of a general UPS cost function.

We now turn to the characterization of the extreme points of the non-learning re-

gions. Lemma 3 endows us with $2 \times |\Omega| + |A| \times (|A| - 1)$ inequalities, and an extreme point is uniquely identified when $|\Omega| - 1$ affine independent constraints are binding.

Lemma 4. When the cost is Shannon, a vector $\mu \in \mathbb{R}^{|\Omega|}$ is an extreme point of NL^a if and only if $\sum_{\omega} \mu(\omega) = 1$ and

- (i) $\mu(\omega) \leq 1$ for all $\omega \in \Omega$
- (ii) $\mu(\omega) \geq 0$ for all $\omega \in \Omega$
- (iii) $\sum_{\omega \in \Omega} \mu(\omega) e^{d_{\omega}^{\lambda}(a, a')} \leq 1$ for all $a' \neq a$, where $d_{\omega}^{\lambda}(a, a') = \frac{u(a', \omega) - u(a, \omega)}{\lambda}$

where $|\Omega| - 1$ affine independent constraints from among (i), (ii), (iii) are binding.

3.3.1 Extreme-Point Solution Algorithm With Shannon Cost

Using the previous results, we obtain the following algorithm to find an equilibrium with the Shannon cost.

Extreme-Point Solution Algorithm. 1. For every action $a \in A$, find all the extreme points of its non-learning regions, EP^a , using Lemma 4.

2. Consider all the sets of beliefs in $\cup_a EP^a$ whose size is at most $|\Omega|$ and for which the prior lies in the convex hull of the set. These sets are the supports of the sender's candidate strategies. Note that since $\cup_a EP^a$ is finite (Corollary 1), there are only finitely many such sets.

3. For each support that we consider, pin the distribution of each candidate strategy down using Bayes' law: $\mathbb{E}_{\tau}[\mu] = \mu_0$. Note that since $\text{supp}(\tau) \leq |\Omega|$, the distribution is unique.

4. For each candidate strategy τ , determine the value of $\hat{v}(\mu)$ on its support by noting that for each $a \in A$, if $\mu \in EP^a$, then $\hat{v}(\mu) = \mathbb{E}_{\mu}[v(a, \omega)]$. If two non-

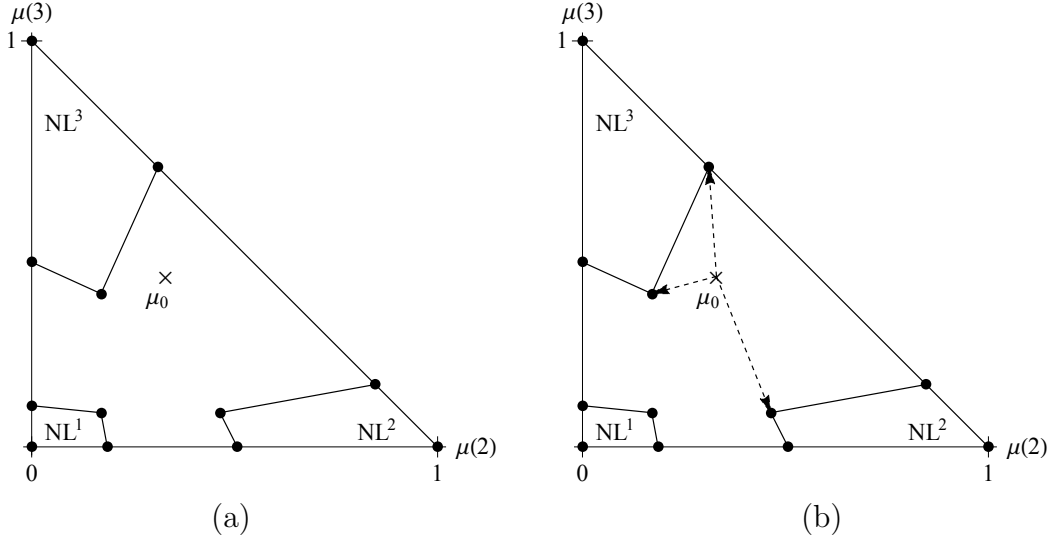


Figure 1: Model with Shannon cost, where $a, \omega \in \{1, 2, 3\}$, $u(a, \omega) = a$ if $a = \omega$ and 0 otherwise, $\lambda = 2$: (a) Non-learning regions and their extreme points; (b) A candidate sender's optimal strategy.

learning regions share the same extreme point (for instance, when $\lambda \rightarrow \infty$), use the action that is preferred by the sender at that belief.¹⁷

5. For each candidate strategy τ , compute the sender's expected utility $\mathbb{E}_\tau[\hat{v}(\mu)]$. The candidate strategy that yields the highest sender's expected payoff is an optimal one.

See Figure 1 for an illustration of the non-learning regions, their extreme points and a candidate optimal strategy for the sender in an example with three actions, three states, and the Shannon cost.

Note that Proposition 2 is applicable also to the standard Bayesian persuasion model by considering the limiting case $\lambda \rightarrow \infty$. In that case, the whole (interim) belief space is partitioned into non-learning regions, each of them with finitely many extreme points. The provided algorithm can thus be used to solve standard Bayesian persuasion models by using Lemma 4 for $\lambda \rightarrow \infty$. We demonstrate the partitioning

¹⁷For some $\mu \in EP^a$ for some a , the solution to (5) may not be unique and then ϕ_μ^* can be different from ϕ_μ^N (if the sender strictly prefers the learning strategy). In that case, $\hat{v}(\mu) \neq \mathbb{E}_\mu[v(a, \omega)]$. However, Proposition 2 implies that treating every extreme point as if $\phi_\mu^* = \phi_\mu^N$ is without loss of generality to find at least one solution of the model.

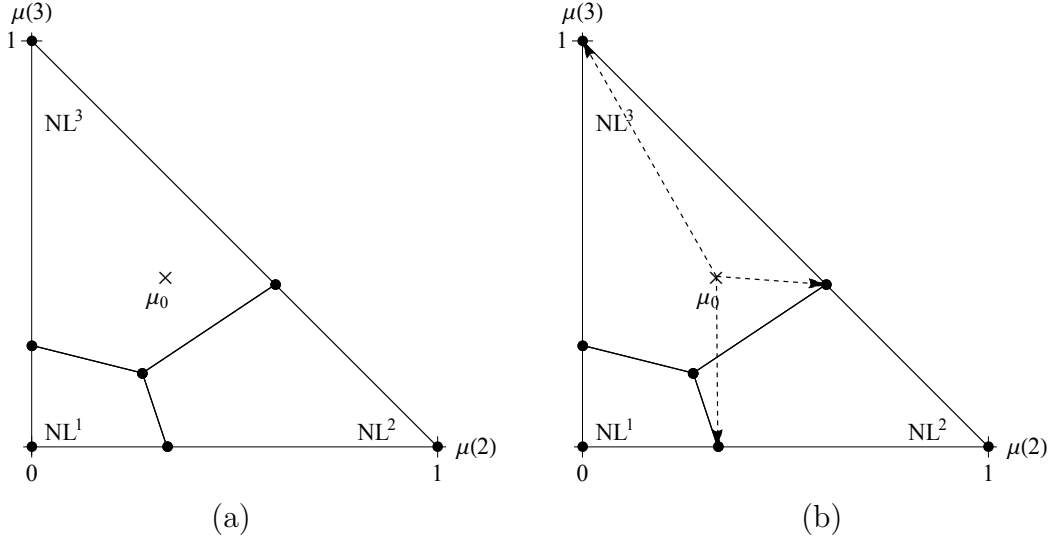


Figure 2: Limiting case of model with Shannon cost, where $\lambda \rightarrow \infty$, $a, \omega \in \{1, 2, 3\}$, $u(a, \omega) = a$ if $a = \omega$ and 0 otherwise: (a) Non-learning regions and their extreme points; (b) A candidate sender's optimal strategy.

of the belief simplex together with a candidate optimal strategy for the sender in such a limiting case in Figure 2.

4 Comparative Statics

In this section, we examine the relationship between the agents' expected equilibrium payoffs and the receiver's information cost parameter λ . First, we show that the sender cannot be better off if the receiver becomes better at gathering information. Recall that a higher λ means that the receiver is worse at gathering information.

Proposition 3. The sender's expected equilibrium payoff is (weakly) increasing in λ .

The intuition of Proposition 3 is straightforward. We show that the non-learning regions—and hence the set of sender's strategies under which the receiver never learns—do not shrink as λ increases. Therefore, increasing λ can only make the sender better off, as he can choose from a larger set of strategies.

The relationship between the receiver’s expected equilibrium payoff and her information cost parameter λ is less trivial. While one might expect that a lower information cost should benefit the receiver, the opposite can also be true, and the receiver’s expected equilibrium payoff can be non-monotone in λ .

Proposition 4. In general, the receiver’s expected equilibrium payoff is not necessarily (weakly) decreasing in λ .

We prove the statement by giving an example where a non-monotone relationship arises (see the next section).

In general, it is hard to provide necessary and sufficient conditions under which the receiver is always (weakly) better off as her cost of information becomes cheaper. Three actions and binary state are enough to build examples with non-monotonicity of the receiver’s payoff, even when assuming that the sender has state-independent preferences. Nonetheless, we present a setting where a monotone relationship is attained: when the state and action are binary—which is a setting used in a number of applied papers in Bayesian persuasion¹⁸—the receiver’s expected payoff is always (weakly) decreasing in her cost parameter λ . First, we state a restriction on the sender’s preferences that is needed for unique equilibrium in this setting.

Assumption 1. There exists no action $a \in A$ s.t. (i) $\forall \mu \in \Delta(\Omega) : \hat{v}(\mu) \leq \sum_{\omega \in \Omega} \mu(\omega)v(a, \omega)$, and (ii) $\exists \mu \in NL^{a'}$ where $a' \neq a$ and $\hat{v}(\mu) = \sum_{\omega \in \Omega} \mu(\omega)v(a, \omega)$.

Assumption A1 takes care of possible cases of indifference on the sender’s side, which could otherwise lead to multiplicity of equilibria. Lemma 8 in Appendix A shows that, for a given λ , when the sender benefits from persuasion, A1 is a sufficient condition for a unique equilibrium in a binary action-state setup.

Proposition 5. Assume $|A| = |\Omega| = 2$, and that for all $\lambda > 0$: A1 holds and the sender benefits from persuasion. Then the receiver’s expected equilibrium utility (weakly) *decreases* in λ .

¹⁸Standard Bayesian persuasion was applied to bank regulation (Gick and Pausch 2012), electoral manipulation (Gehlbach and Simpson 2015), investment decision (Bizzotto, Rüdiger, and Vigier 2021), and forecasting of disasters (Aoyagi 2014).

The proof takes advantage of the structure of the extreme-point method. In a binary action-state setting, there are two non-learning regions. Each is an interval with two extreme points. Extreme-point method then implies that only very few sender's strategies must be considered for an optimum. It turns out that optimal sender's strategies are: (i) providing no information (This is optimal when the players disagree on a preferred action in each state. Then, the game resembles a zero-sum game and the sender does not benefit from persuasion. This specification is thus excluded from Proposition 5.); (ii) providing full information (optimal under aligned preferences), which is independent of λ , and (iii) providing partial information in order to maximize the probability of sender's preferred action, given that additional learning was prevented (optimal when the sender has a dominant action). He does so by providing as little information as possible that convinces the receiver to take the dominant action upon favorable evidence and prevents her from additional learning. As λ increases, the receiver is willing to take the dominant action without additional learning even when less convinced about its optimality. The sender thus needs to provide less convincing information, making the receiver worse off.

Formally, in the binary-state setting, where the belief simplex is a line, the extreme-point method gives us an easy way to understand how the solution changes in the Blackwell order as the receiver's cost parameter varies: we only need to determine whether the extreme points induced by the optimal strategy are moving apart from or towards each other. To determine this, we further use the notion that the non-learning regions do not shrink as λ increases. For binary action, this gives us a clear direction of how the extreme points that are induced by the optimal strategy move.

5 Application: Extremism vs Conservatism

Here, we present a discretized version of a model with quadratic loss functions, a common specification used in the literature on communication games. Consider a binary state, $\omega \in \{-1, 1\}$. In this section, we use μ_0, μ, γ to denote the probability

of the state $\omega = 1$ (i.e., $\Pr[\omega = 1]$) at the prior, the interim, and the posterior belief respectively. Let $\mu_0 = 0.5$. Consider five actions, $a \in \{-2, -1, 0, 1, 2\}$. Call the actions $\{-1, 1\}$ *conservative* and the actions $\{-2, 2\}$ *extreme*. The payoff functions of the receiver and the sender are $u(a, \omega) = -(\alpha\omega - a)^2$ and $v(a, \omega) = -(\beta\omega - a)^2$, respectively. The cost is Shannon.

Under the prior, there is no disagreement among the players—the preferred action is 0 for both of them. The parameters α and β capture the players’ extremism—how much their preferred action changes with the precision of their beliefs. The higher the parameter, the more extreme they are, since their preferred action in any given state becomes (weakly) more extreme. We consider two cases: (i) $\alpha = 1$, $\beta = 2$, and (ii) $\alpha = 2$, $\beta = 1$. In (i), the *receiver* is the *conservative player* (knowing the state, the receiver chooses a conservative action, but the sender prefers an extreme action); in (ii), the opinions are reversed and the *receiver* is the *extreme player*.

The sender’s optimal information strategy, and the way in which it varies with the receiver’s cost parameter λ , is different in each case. In Appendix B, we analytically solve the application by applying the Extreme-point solution algorithm with Shannon cost developed in Section 3.3.1. Again, since we assume binary-state setting, where the belief simplex is a line, extreme-point method then gives us an easy way to compare how the solution changes in the Blackwell order with λ : we only need to determine whether the extreme points that are induced by the optimal strategy are moving apart from or towards each other. The following proposition captures the qualitative properties of the solution.

Proposition 6. For any $\lambda > 0$, the application has a unique equilibrium. Moreover:

- (i) When the *receiver* is the *conservative player*, the sender optimally provides full information for all values $\lambda > 0$.
- (ii) When the *receiver* is the *extreme player*, there exist a threshold value $\bar{\lambda} \in (0, \infty)$ such that
 - a) If $\lambda < \bar{\lambda}$, the sender optimally provides full information.

- b) If $\lambda \geq \bar{\lambda}$, the sender optimally provides partial information. Then, the receiver's expected equilibrium utility is *strictly increasing* in λ .

When the *receiver* is the *conservative player*, an extreme sender gains nothing in providing less than full information. Since the receiver is conservative, the best the sender can do is to fully reveal the state so as to minimize the probability of mistakes in conservative actions. Full revelation is thus always optimal, independently of the receiver's cost parameter λ . The receiver's expected equilibrium utility is (weakly) monotone in λ .

On the other hand, when the *receiver* is the *extreme player*, full revelation leads to extreme actions. A conservative sender wants to prevent that. He thus optimally garbles a fully informative signal leaving the receiver uncertain enough about the state so that she chooses conservative rather than extreme action and she does not undergo additional learning. Thus, the sender faces a trade-off between preventing the receiver from taking extreme actions (or learning further) if too much information is provided and minimizing the probability of mistakes in conservative actions if too little information is provided.

The sender's optimal information strategy is then no longer independent of the receiver's cost parameter λ . As λ decreases, the sender becomes more constrained in how much information he can provide and still guarantee that the receiver is not going to undergo additional costly learning. Hence, as λ decreases, the sender optimally generates strictly less information (in Blackwell sense). Since the receiver never learns in equilibrium, she then becomes strictly worse off. Note that the non-learning regions shrink as λ decreases. At the threshold value $\bar{\lambda}$, the non-learning regions of conservative actions are no longer intervals, but singletons. For any $\lambda < \bar{\lambda}$, only the non-learning regions of extreme actions are not empty sets. The cost of information is so cheap that whenever the receiver decides to learn, she acquires enough information to be confident enough about the state so that she always chooses an extreme action. Thus, the sender can never persuade her to consider any action but an extreme one. Then, the best he can do is to minimize the

probability of mistakes in extreme actions. Full revelation is then optimal. Hence, when the receiver is the extreme player, the receiver's expected equilibrium utility is non-monotone in λ .

5.1 A Graphical Illustration

Figure 3 demonstrates the solution for both cases for $\lambda \rightarrow \infty$ (the receiver cannot acquire additional information) and for $\lambda = 15$ (the receiver can costly acquire additional information). We depict the solution using the concavification approach, since, in a binary state setting, this approach offers an easily understandable graphical form. One expresses the sender's expected payoff at interim belief given the optimal continuation play $\hat{v}(\mu)$ —the black function—and take its concave closure $\text{cav}(\hat{v})(\mu)$ —the red function. The support of the sender's optimal information strategy is the set of those interim beliefs that support the tangent of $\text{cav}(\hat{v})(\mu)$ at the prior μ_0 , and for which $\text{cav}(\hat{v})(\mu) = \hat{v}(\mu)$.

The left column depicts the sender's optimal strategy when the receiver is conservative. In this case, the support of the sender's optimal strategy remains the same for both values of λ : the state is always revealed in equilibrium. Let τ_λ^* denote the sender's optimal strategy when the cost parameter is λ . Then $\text{supp}(\tau_\infty^*) = \{\mu^*, \mu'^*\} = \{0, 1\}$ in (a), and $\text{supp}(\tau_{15}^*) = \{\tilde{\mu}^*, \tilde{\mu}'^*\} = \{0, 1\}$ in (b). The right column depicts the sender's optimal strategy when the receiver is extreme. Here, the optimal strategy changes with λ . As λ decreases, the sender provides less information in Blackwell sense:¹⁹ $\text{supp}(\tau_\infty^*) = \{\mu^*, \mu'^*\}$ in (c), and $\text{supp}(\tau_{15}^*) = \{\tilde{\mu}^*, \tilde{\mu}'^*\}$ in (d), where $\text{supp}(\tau_{15}^*)$ lies inside a convex hull of $\text{supp}(\tau_\infty^*)$.

¹⁹An information strategy τ is more Blackwell-informative than τ' if and only if obtaining information via τ is preferred to information via τ' by all expected utility maximizers. Equivalently, τ is more Blackwell-informative than τ' if and only if $\text{supp}(\tau')$ lie inside the convex hull of $\text{supp}(\tau)$.

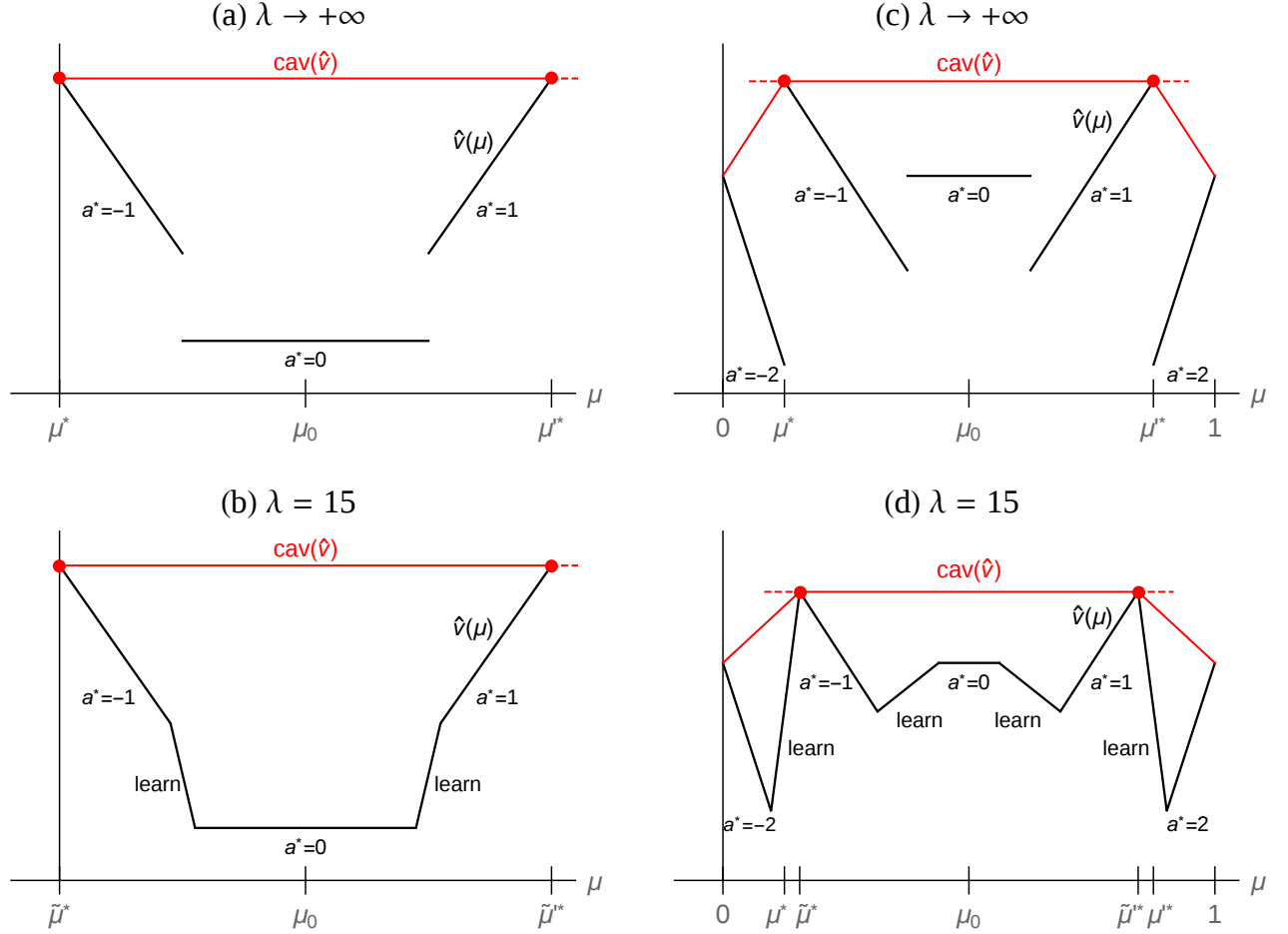


Figure 3: Sender's value function $\hat{v}(\mu)$, its concavification $cav(\hat{v})(\mu)$ and his optimal information strategy. With conservative receiver—cases (a) and (b)—the sender fully reveals the state for both values of λ : $\{\mu^*, \mu'^*\} = \{\tilde{\mu}^*, \tilde{\mu}'^*\} = \{0, 1\}$. With extreme receiver—cases (c) and (d)—the sender provides Blackwell less information for lower λ : $\{\tilde{\mu}^*, \tilde{\mu}'^*\}$ lie in the convex hull of $\{\mu^*, \mu'^*\}$.

6 Conclusion

We extend a model of Bayesian persuasion by allowing for additional costly information acquisition by a receiver under a uniformly posterior-separable cost function. We exploit common features of Bayesian persuasion and decision problems with this type of cost, resulting in a tractable model which can be used as a building block for applied problems. Using the characterization of the receiver’s optimal behavior, we propose a solution method that generates a smaller set of posterior beliefs on which at least one optimal strategy of the sender must be supported. Under an entropy-based cost function, as in rational inattention, such an optimal strategy is characterized by a finite series of specific linear conditions. This method, which does not rely on standard concavification method, is also applicable to the standard Bayesian persuasion model and can be used to find a solution. We further show that having additional information sources can be detrimental to the receiver, who could then prefer to commit to not having any such sources.

This result illustrates that the ability to gather information is not always desirable in strategic environments. It is well-known that the value of information may be negative when decisions are taken strategically. That is, in a game, equilibrium payoffs may be lower when more information is provided to the players.²⁰ This paper shows that a similar counter-intuitive phenomenon can occur in the model of Bayesian persuasion for the ability to gather information. In other words, even though a better ability to gather information is always desirable in a non-strategic environment, the reverse may be true in strategic environments.

References

Alonso, R. and O. Camara (2016). Bayesian persuasion with heterogeneous priors. *Journal of Economic Theory* 165, 672–706.

²⁰See for instance Bassan, Gossner, Scarsini, and Zamir (2003), Gossner (2000), and Neyman (1991).

- Aoyagi, M. (2014). Strategic obscurity in the forecasting of disasters. *Games and Economic Behavior* 87, 485–496.
- Arieli, I., Y. Babichenko, R. Smorodinsky, and T. Yamashita (2019). Optimal persuasion via bi-pooling. *Available at SSRN 3511516*.
- Aumann, R. J., M. Maschler, and R. E. Stearns (1995). *Repeated games with incomplete information*. MIT press.
- Bassan, B., O. Gossner, M. Scarsini, and S. Zamir (2003). Positive value of information in games. *International Journal of Game Theory* 32(1), 17–31.
- Bizzotto, J., J. Rüdiger, and A. Vigier (2020). Testing, disclosure and approval. *Journal of Economic Theory* 187, 105002.
- Bizzotto, J., J. Rüdiger, and A. Vigier (2021). Dynamic persuasion with outside information. *American Economic Journal: Microeconomics* 13(1), 179–94.
- Bloedel, A. W. and I. R. Segal (2018). Persuasion with rational inattention. *Available at SSRN 3164033*.
- Caplin, A. and M. Dean (2013). Behavioral implications of rational inattention with shannon entropy. Technical report, National Bureau of Economic Research.
- Caplin, A., M. Dean, and J. Leahy (2017). Rationally inattentive behavior: Characterizing and generalizing shannon entropy. Technical report, National Bureau of Economic Research.
- Caplin, A., M. Dean, and J. Leahy (2019). Rational inattention, optimal consideration sets, and stochastic choice. *The Review of Economic Studies* 86(3), 1061–1094.
- Dworczak, P. and G. Martini (2019). The simple economics of optimal persuasion. *Journal of Political Economy* 127(5), 1993–2048.
- Gehlbach, S. and A. Simpser (2015). Electoral manipulation as bureaucratic control. *American Journal of Political Science* 59(1), 212–224.

- Gentzkow, M. and E. Kamenica (2014). Costly persuasion. *The American Economic Review* 104(5), 457–462.
- Gentzkow, M. and E. Kamenica (2017). Bayesian persuasion with multiple senders and rich signal spaces. *Games and Economic Behavior* 104, 411–429.
- Gick, W. and T. Pausch (2012). Persuasion by stress testing: Optimal disclosure of supervisory information in the banking sector.
- Gossner, O. (2000). Comparison of information structures. *Games and Economic Behavior* 30(1), 44–63.
- Kamenica, E. (2019). Bayesian persuasion and information design. *Annual Review of Economics* 11, 249–272.
- Kamenica, E. and M. Gentzkow (2011). Bayesian persuasion. *American Economic Review* 101(6), 2590–2615.
- Kessler, A. S. (1998). The value of ignorance. *The Rand Journal of Economics*, 339–354.
- Kleiner, A., B. Moldovanu, and P. Strack (2020). Extreme points and majorization: Economic applications. *Available at SSRN*.
- Kolotilin, A. (2018). Optimal information disclosure: A linear programming approach. *Theoretical Economics* 13(2), 607–635.
- Kolotilin, A., T. Mylovannov, A. Zapechelnnyuk, and M. Li (2017). Persuasion of a privately informed receiver. *Econometrica* 85(6), 1949–1964.
- Le Treust, M. and T. Tomala (2019). Persuasion with limited communication capacity. *Journal of Economic Theory* 184, 104940.
- Lipnowski, E. and L. Mathevet (2017). Simplifying bayesian persuasion.
- Lipnowski, E. and L. Mathevet (2018). Disclosure to a psychological audience. *American Economic Journal: Microeconomics* 10(4), 67–93.
- Lipnowski, E., L. Mathevet, and D. Wei (2020a). Attention management. *American Economic Review: Insights* 2(1), 17–32.

- Lipnowski, E., L. Mathevet, and D. Wei (2020b). Optimal attention management: A tractable framework. *arXiv preprint arXiv:2006.07729*.
- Martin, D. (2017). Strategic pricing with rational inattention to quality. *Games and Economic Behavior* 104, 131–145.
- Matějka, F. and A. McKay (2015). Rational inattention to discrete choices: A new foundation for the multinomial logit model. *American Economic Review* 105(1), 272–98.
- Matějka, F. and G. Tabellini (2016). Electoral competition with rationally inattentive voters. *Journal of the European Economic Association*.
- Montes, A. (2020). Canonical equilibrium in games with flexible information acquisition. *Working Paper*.
- Neyman, A. (1991). The positive value of information. *Games and Economic Behavior* 3(3), 350–355.
- Nguyen, A. and T. Y. Tan (2021). Bayesian persuasion with costly messages. *Journal of Economic Theory* 193, 105212.
- Ok, E. A. (2007). *Real analysis with economic applications*, Volume 10. Princeton University Press.
- Ravid, D. (2020). Ultimatum bargaining with rational inattention. *American Economic Review* 110(9), 2948–63.
- Rockafellar, R. (1997). *Convex Analysis*. Princeton Landmarks in Mathematics and Physics. Princeton University Press.
- Roesler, A.-K. and B. Szentes (2017). Buyer-optimal learning and monopoly pricing. *American Economic Review* 107(7), 2072–80.
- Simon, B. (2011). *Convexity: an analytic viewpoint*, Volume 187. Cambridge University Press.
- Sims, C. A. (2003). Implications of rational inattention. *Journal of monetary Economics* 50(3), 665–690.

- Sims, C. A. (2006). Rational inattention: Beyond the linear-quadratic case. *The American economic review* 96(2), 158–163.
- Steiner, J., C. Stewart, and F. Matějka (2017). Rational inattention dynamics: Inertia and delay in decision-making. *Econometrica* 85(2), 521–553.
- Wei, D. (2020). Persuasion under costly learning. *Journal of Mathematical Economics*.
- Yang, M. (2015). Coordination with flexible information acquisition. *Journal of Economic Theory* 158, 721–738.
- Yang, M. (2020). Optimality of debt under flexible information acquisition. *The Review of Economic Studies* 87(1), 487–536.

7 Appendix A: Proofs

7.1 Proofs for Section 2

Lemma 1: The receiver’s maximization problem (5) has a solution.

We prove a stronger statement: *The solution correspondence of the receiver’s maximization problem (5) is non-empty, compact valued and upper hemicontinuous in μ .*

Proof. Since the cost function is strictly increasing in Blackwell informativeness, under optimal receiver’s strategy (if it exists), each action is selected in at most one posterior. Inducing distinct posteriors that lead to the same action is inefficient as information is acquired but not acted upon. Hence, the support of the receiver’s optimal strategy contains at most as many posteriors as $|A|$ and the receiver then chooses different action at each of the posteriors. The proof of this statement is shown in proof of Lemma 1 in Matějka and McKay (2015) used for a model of Shannon cost. Their proof is derived for Shannon entropy, where only the fact that

it is a strictly concave function, not its functional form per se, was used and thus can be applied in our setting as well.

We can thus rephrase the receiver's maximization problem (5) as choosing *conditional action probabilities*, $\pi(a|\omega)$. That is, the problem can be rephrased as choosing a family of probability distributions $\{\pi(\cdot|\omega)\}_{\omega \in \Omega}$, where $\pi(\cdot|\omega) \in \Delta(A)$ for each ω .

Definition. 1': Given an interim belief μ , the receiver's maximization problem is

$$\begin{aligned} \max_{\pi=\{\pi(a|\omega)\}_{a,\omega}} \quad & \sum_{a,\omega} u(a,\omega)\pi(a|\omega)\mu(\omega) - \tilde{c}(\pi;\mu,\lambda) \\ \text{s.t.} \quad & \forall a \quad \forall \omega : \pi(a|\omega) \geq 0 \\ & \forall \omega : \sum_a \pi(a|\omega) = 1 \end{aligned} \tag{10}$$

where $\tilde{c}(\pi;\mu,\lambda) = \lambda (F(\mu) - \sum_a F(\gamma_a^\pi) \tilde{\pi}(a))$, and where for all a , $\tilde{\pi}(a) = \sum_\omega \pi(a|\omega)\mu(\omega)$ is the unconditional probability of choosing action a under strategy π , and for all ω and for all a ,

$$\gamma_a^\pi(\omega) = \begin{cases} \frac{\pi(a|\omega)\mu(\omega)}{\tilde{\pi}(a)} & \text{if } \tilde{\pi}(a) \neq 0 \\ p(\omega) & \text{otherwise} \end{cases} \tag{11}$$

is the posterior probability of state ω when action a is taken under strategy π , where $p \in \Delta(\Omega)$ is an arbitrary probability distribution over states.²¹ The posterior probability depends on the strategy π and, for notation purposes, such dependence is reflected in the upper-script.

Next, we need to show that the cost function $\tilde{c}(\pi;\mu,\lambda)$ is continuous in π . Function (11) is well-defined, but it is not continuous. To show that $\tilde{c}(\pi;\mu,\lambda)$ is still continuous in π , the only points of concern are strategies under which there exist an action a which is never chosen, $\tilde{\pi}(a) = 0$. Let $\{\pi_n\}_n$ be a converging sequence of

²¹Since the cost function of the original problem is defined in terms of beliefs, we need to map all possible strategies π into well-defined beliefs, which are described in equation (11). We do it by assigning an arbitrary belief p to the beliefs that are not well-defined (posterior beliefs associated with actions which are never chosen under π).

receiver's strategies with a limit $\{\pi_n\}_n \rightarrow \hat{\pi}$, where $\hat{\pi}$ is a strategy that assigns zero unconditional probability to choosing action a' , $\hat{\pi}(a') = 0$. Then

$$\begin{aligned}
\lim_{\{\pi_n\}_n \rightarrow \hat{\pi}} \tilde{c}(\pi_n; \mu, \lambda) &= \lim_{\{\pi_n\}_n \rightarrow \hat{\pi}} \lambda \left(F(\mu) - \sum_a F(\gamma_a^{\pi_n}) \tilde{\pi}_n(a) \right) \\
&= \lambda \lim_{\{\pi_n\}_n \rightarrow \hat{\pi}} \left(F(\mu) - \sum_{a \neq a'} F(\gamma_a^{\pi_n}) \tilde{\pi}_n(a) - F(\gamma_{a'}^{\pi_n}) \tilde{\pi}_n(a') \right) \\
&= \lambda \left(F(\mu) - \sum_{a \neq a'} F(\gamma_a^{\hat{\pi}}) \hat{\pi}(a) \right) \\
&= \tilde{c}(\hat{\pi}; \mu, \lambda)
\end{aligned}$$

The third equality follows because (i) $F(\gamma_{a'}^{\pi_n}) \tilde{\pi}_n(a')$ goes to zero, as $F(\cdot)$ is bounded and $\tilde{\pi}_n(a')$ approaches zero, and (ii) $F(\cdot)$ is continuous, so $F(\gamma_a^{\pi_n}) \rightarrow F(\gamma_a^{\hat{\pi}})$. Hence, we conclude that $\tilde{c}$ is continuous on π .

Finally, since the objective function in the receiver's problem is continuous in π and the set of all possible strategies π in $\Delta(A)^{|\Omega|}$ is compact, Berge Maximum Theorem implies that the solution correspondence is non-empty, compact valued and upper hemicontinuous in μ .

□

Lemma 2: The sender's maximization problem (7) has a solution.

Proof. As in KG, the objective of the sender is to build the concave closure of the function $\hat{v}(\mu)$. The sufficient condition for that is when the function $\hat{v}(\mu)$ is upper semi-continuous. To show it is true, for now, we will drop the assumption of sender-preferred equilibrium in both the receiver's information strategies and her chosen actions, and consider the conditional expected value of the sender as a correspondence. We then show that this correspondence is upper hemicontinuous with closed graph. This in turn implies that $\hat{v}(\mu)$ is upper semi-continuous.

Formally, let us drop the sender-preferred equilibrium assumption and let us consider

the receiver's maximization problem (10) as it is redefined in terms of conditional action probabilities in Definition 1' in proof of Lemma 1. Let π_μ^* be the solution correspondence to the problem (10) for a given interim belief μ and $\gamma_a^\pi(\omega)$ the posterior probability of state ω when action a is chosen under strategy π defined as in (11). Then, let $\tilde{v} : \Delta(\Omega) \rightrightarrows \mathbb{R}$ be the correspondence that assigns to every μ the set of payoffs

$$\bigcup_{\pi \in \pi_\mu^*} \left(\sum_{a, \omega} v(a, \omega) \gamma_a^\pi(\omega) \tilde{\pi}(a) \right) \quad (12)$$

Now we show that correspondence $\tilde{v}$ is compact valued and upper hemicontinuous in μ , which in turn implies that $\tilde{v}$ has closed graph.²²

To prove that $\tilde{v}$ is upper hemicontinuous in μ , we use the sequential characterization of upper hemicontinuous correspondences as formulated in Ok (2007). Formally, take $\{\mu_n\}_n$ converging to μ , and $\{z_n\}_n$ with $z_n \in \tilde{v}(\mu_n)$ for all n . We want to show that there exists a sub-sequence $\{z_{n_k}\}_k$ such that it converges to $z \in \tilde{v}(\mu)$. Since π_μ^* is upper hemicontinuous (as shown in the proof of Lemma 1), we know that for given $\{\mu_n\}_n$ converging to μ , and $\{\pi_n\}_n$ with $\pi_n \in \pi_{\mu_n}^*$ for all n , there exists a sub-sequence $\{\pi_{n_k}\}_k$ such that it converges to $\pi' \in \pi_\mu^*$. Since $\{\pi_{n_k}\}_k$ converges to π'

$$\left\{ \sum_{a, \omega} v(a, \omega) \gamma_a^{\pi_{n_k}}(\omega) \tilde{\pi}_{n_k}(a) \right\}_{n_k} \longrightarrow \sum_{a, \omega} v(a, \omega) \gamma_a^{\pi'}(\omega) \tilde{\pi}'(a) \quad (13)$$

which is an element in $\tilde{v}(\mu)$. Note that if π' is such that $\tilde{\pi}'(a) = 0$ for some a , the term $\sum_{\omega} v(a, \omega) \gamma_a^{\pi_{n_k}}(\omega) \tilde{\pi}_{n_k}(a) \rightarrow 0$ since $v(a, \omega)$ is bounded, $\gamma_a^{\pi_{n_k}}(\omega)$ is bounded and $\tilde{\pi}_{n_k}(a) \rightarrow 0$. For any other actions with $\tilde{\pi}'(a) > 0$, the convergence in (13) follows from the fact that $\sum_{\omega} v(a, \omega) \gamma_a^{\pi_{n_k}}(\omega) \tilde{\pi}_{n_k}(a)$ is then continuous on π . Hence, $\tilde{v}$ is upper hemicontinuous.

Finally, since π_μ^* is a compact set (as shown in the proof of Lemma 1), its image through a continuous function, that is, $\{\sum_{a, \omega} v(a, \omega) \gamma_a^{\pi'}(\omega) \pi'(a) | \pi' \in \pi_\mu^*\}$, is also

²²See Ok (2007), Chapter E, Continuity II.

compact. Therefore, $\tilde{v}$ is also compact valued. Since $\hat{v}$ is compact valued and upper hemicontinuous in μ , $\tilde{v}$ has a closed graph, and a solution to sender's problem (7) always exists. The reason is that, since $\tilde{v}$ is compact valued, the function $\hat{v}(\mu) := \max_{\pi \in \pi_\mu^*} \left(\sum_{a, \omega} v(a, \omega) \gamma_a^\pi(\omega) \tilde{\pi}(a) \right)$ is well defined for every μ , and since $\tilde{v}$ is upper hemicontinuous, $\hat{v}$ is an upper semi-continuous function. $\square$

7.2 Proofs for Section 3.1

The next lemma is used to prove Proposition 1. It states that whenever the receiver learns at μ , she induces beliefs that all lie in some non-learning regions.

Lemma 5. For any μ , the support of the receiver's optimal strategy lies in non-learning regions. That is, $\forall \mu \in \Delta(\Omega)$, we have $\text{supp}(\phi_\mu^*) \in \cup_a NL^a$.

Proof. Let ϕ_μ^* be receiver's optimal strategy at μ and suppose, contrary to the statement, that $\exists \tilde{\gamma} \in \text{supp}(\phi_\mu^*)$ for which $\tilde{\gamma} \notin \cup_a NL^a$. What it means is that, if $\tilde{\gamma}$ was an interim belief, then the receiver's optimal strategy at the interim belief $\tilde{\gamma}$ includes learning. That is, there exists a distribution of posterior beliefs $\phi_{\tilde{\gamma}}^L \neq \phi_{\tilde{\gamma}}^N$ with $\mathbb{E}_{\phi_{\tilde{\gamma}}^L}[\gamma] = \tilde{\gamma}$ (an optimal receiver's strategy at $\tilde{\gamma}$) which yields a strictly higher receiver's expected payoff at $\tilde{\gamma}$ than no learning, that is,

$$\begin{aligned} \mathbb{E}_{\phi_{\tilde{\gamma}}^L}[U(\gamma)] - c(\phi_{\tilde{\gamma}}^L; \tilde{\gamma}, \lambda) &= \mathbb{E}_{\phi_{\tilde{\gamma}}^L}[U(\gamma) + \lambda F(\gamma)] - \lambda F(\tilde{\gamma}) > U(\tilde{\gamma}) = \mathbb{E}_{\phi_{\tilde{\gamma}}^N}[U(\gamma)] - c(\phi_{\tilde{\gamma}}^N; \tilde{\gamma}, \lambda) \\ \mathbb{E}_{\phi_{\tilde{\gamma}}^L}[U(\gamma) + \lambda F(\gamma)] &> U(\tilde{\gamma}) + \lambda F(\tilde{\gamma}) \end{aligned} \quad (14)$$

where we used the fact that no-learning comes at zero cost.

We will show that there exist another receiver's strategy, ϕ'_μ , that yields a strictly higher expected utility than ϕ_μ^* . For the initial receiver's problem—a re-

ceiver who faces interim belief μ —consider a different receiver’s strategy ϕ'_μ satisfying

$$\phi'_\mu(\gamma) = \begin{cases} \phi_\mu^*(\gamma) & \text{if } \gamma \notin \{\tilde{\gamma}\} \cup \text{supp}(\phi_{\tilde{\gamma}}^L) \\ \phi_\mu^*(\gamma) + \phi_\mu^*(\tilde{\gamma})\phi_{\tilde{\gamma}}^L(\gamma) & \text{if } \gamma \in \text{supp}(\phi_{\tilde{\gamma}}^L) \\ 0 & \text{if } \gamma = \tilde{\gamma} \end{cases}$$

Note that for $\gamma \notin \text{supp}(\phi_\mu^*) \cup \text{supp}(\phi_{\tilde{\gamma}}^L)$, we have $\phi_\mu^*(\gamma) = \phi'_\mu(\gamma) = 0$ and thus $\text{supp}(\phi'_\mu) = \text{supp}(\phi_\mu^*) \cup \text{supp}(\phi_{\tilde{\gamma}}^L) \setminus \{\tilde{\gamma}\}$. Note that the optimality of ϕ_μ^* and $\phi_{\tilde{\gamma}}^L$ implies that their supports are of finite size (as discussed in the proof of Lemma 1) and thus the support of ϕ'_μ is also finite.

However, then, the strategy ϕ'_μ yields a strictly higher receiver’s expected utility than ϕ_μ^* , contradicting the optimality of ϕ_μ^* . To see that, let us compare the receiver’s expected utility under both of the strategies. We have $\mathbb{E}_{\phi'_\mu}[U(\gamma)] - c(\phi'_\mu; \mu, \lambda) > \mathbb{E}_{\phi_\mu^*}[U(\gamma)] - c(\phi_\mu^*; \mu, \lambda)$ whenever the following inequality holds

$$\begin{aligned} \sum_{\gamma \in \text{supp}(\phi'_\mu)} [U(\gamma) + \lambda F(\gamma)] \phi'_\mu(\gamma) - \lambda F(\mu) &> \sum_{\gamma \in \text{supp}(\phi_\mu^*)} [U(\gamma) + \lambda F(\gamma)] \phi_\mu^*(\gamma) - \lambda F(\mu) \\ \sum_{\gamma \in \{\tilde{\gamma}\} \cup \text{supp}(\phi_{\tilde{\gamma}}^L)} [U(\gamma) + \lambda F(\gamma)] \phi'_\mu(\gamma) &> \sum_{\gamma \in \{\tilde{\gamma}\} \cup \text{supp}(\phi_{\tilde{\gamma}}^L)} [U(\gamma) + \lambda F(\gamma)] \phi_\mu^*(\gamma) \\ \sum_{\gamma \in \text{supp}(\phi_{\tilde{\gamma}}^L)} [U(\gamma) + \lambda F(\gamma)] \phi_\mu^*(\gamma) &+ \sum_{\gamma \in \text{supp}(\phi_{\tilde{\gamma}}^L)} [U(\gamma) + \lambda F(\gamma)] \phi_{\tilde{\gamma}}^L(\gamma) \phi_\mu^*(\tilde{\gamma}) > \\ \sum_{\gamma \in \text{supp}(\phi_{\tilde{\gamma}}^L)} [U(\gamma) + \lambda F(\gamma)] \phi_\mu^*(\gamma) &+ [U(\tilde{\gamma}) + \lambda F(\tilde{\gamma})] \phi_\mu^*(\tilde{\gamma}) \\ \mathbb{E}_{\phi_{\tilde{\gamma}}^L} [U(\gamma) + \lambda F(\gamma)] \phi_\mu^*(\tilde{\gamma}) &> [U(\tilde{\gamma}) + \lambda F(\tilde{\gamma})] \phi_\mu^*(\tilde{\gamma}) \end{aligned}$$

The inequality holds, since we only made equivalent operations and the last inequality is true based on (14) and $\phi_\mu^*(\tilde{\gamma}) > 0$. The second line follows from $\phi'(\gamma) = \phi^*(\gamma)$ for all $\gamma \notin \{\tilde{\gamma}\} \cup \text{supp}(\phi_{\tilde{\gamma}}^L)$ and thus the terms in the sum for $\gamma \notin \{\tilde{\gamma}\} \cup \text{supp}(\phi_{\tilde{\gamma}}^L)$ cancel out on both sides. The third line uses $\phi'(\tilde{\gamma}) = 0$ and $\phi'(\gamma) = \phi_\mu^*(\gamma) + \phi_\mu^*(\tilde{\gamma})\phi_{\tilde{\gamma}}^L(\gamma)$ when $\gamma \in \text{supp}(\phi_{\tilde{\gamma}}^L)$.

□

Proposition 1: There exists an equilibrium in which the receiver never learns. That is, there exists $(\tau^*, (\phi_\mu^*)_\mu)$ such that $\forall \mu \in \text{supp}(\tau^*), \phi_\mu^* = \phi_\mu^N$.

Proof. By Lemmas 1 and 2, an equilibrium exists. Suppose that there exists an equilibrium in which the receiver learns. We will construct another equilibrium in which the receiver never learns. Suppose there exists an equilibrium $(\hat{\tau}, (\hat{\phi}_\mu)_\mu)$, such that for some $\mu' \in \text{supp}(\hat{\tau})$ we have $\hat{\phi}_{\mu'} \neq \phi_{\mu'}^N$. As discussed in proof of Lemma 1, for all μ , the size of the support of any optimal receiver's strategy at μ does not exceed $|A|$ and hence it is finite.

Consider τ' , another sender's strategy, defined by

$$\tau'(\mu) = \begin{cases} \hat{\tau}(\mu) & \text{if } \mu \notin \{\mu'\} \cup \text{supp}(\hat{\phi}_{\mu'}) \\ \hat{\tau}(\mu) + \hat{\tau}(\mu')\phi_{\mu'}(\mu) & \text{if } \mu \in \text{supp}(\hat{\phi}_{\mu'}) \\ 0 & \text{if } \mu = \mu' \end{cases} \quad (15)$$

Note that for $\mu \notin \text{supp}(\tau') \cup \text{supp}(\hat{\tau})$, we have $\tau'(\mu) = \hat{\tau}(\mu) = 0$. Note that, by Lemma 5, we have $\text{supp}(\hat{\phi}_\mu) \in \bigcup_a NL^a$ for all μ . Hence, the receiver's optimal equilibrium strategies $(\hat{\phi}_\mu)_\mu$ satisfy $\hat{\phi}_\mu = \phi_\mu^N$ for all $\mu \in \text{supp}(\hat{\phi}_{\mu'})$. That is, whenever the sender induces an interim belief in $\text{supp}(\hat{\phi}_{\mu'})$ —provides the information on receiver's behalf—the receiver does not learn at that interim belief and takes an action right away. Hence, the final distribution over posteriors and associated actions under τ' is the same as under $\hat{\tau}$ and the expected payoff of the sender remains the same. Thus, $(\tau', (\hat{\phi}_\mu)_\mu)$ is also an equilibrium. This process can be done for all $\mu \in \text{supp}(\hat{\tau})$ at which the receiver learns until we reach an equilibrium in which the receiver never learns. Note that generally, it can be true that at some belief $\tilde{\mu} \in \bigcup_a NL^a$, the solution to the receiver's maximization problem (5) may not be unique, in which case $\phi_{\tilde{\mu}}^*$ may be different from $\phi_{\tilde{\mu}}^N$ (whenever the sender strictly prefers the learning strategy). However, this cannot be the case for any belief in $\text{supp}(\hat{\phi}_\mu)$ and we indeed have $\hat{\phi}_\mu = \phi_\mu^N$. The reason is that if this were not the case, then τ' yields strictly higher payoff than $\hat{\tau}$, contradicting that $\hat{\tau}$ was an equilibrium strategy. $\square$

7.3 Proofs for Section 3.2

The next lemma is used to prove Proposition 2, Proposition 5 and Lemmas 7 and 8.

Lemma 6. i) If $\mu \in \text{supp}(\tau^*)$, then $\text{cav}(\hat{v})(\mu) = \hat{v}(\mu)$.

ii) The sender benefits from persuasion if and only if $\hat{v}(\mu_0) < \text{cav}(\hat{v})(\mu_0)$.

Lemma 6 is an analogy of Lemma 2 and Corollary 2 from KG, which is applicable to our setting when $\hat{v}(\mu)$ modified to our setting is considered.

Proposition 2: There exists a non-learning equilibrium in which the sender's strategy is supported in (at most $|\Omega|$) extreme points of the convex hull of the non-learning regions. That is, there exists $(\tau^*, (\phi_\mu^*)_\mu)$ such that $|\text{supp}(\tau^*)| \leq |\Omega|$ and for all $\mu \in \text{supp}(\tau^*)$ we have $\mu \in \cup_a EP^a$ and $\phi_\mu^* = \phi_\mu^N$.

Proof. By Proposition 1, there exists a non-learning equilibrium. We will construct an extreme-point equilibrium from an arbitrary never-learning equilibrium. Let $(\hat{\tau}, (\hat{\phi}_\mu)_\mu)$ be a non-learning equilibrium, and suppose it is not an extreme-point equilibrium, i.e., for some $\mu' \in \text{supp}(\hat{\tau})$, we have $\mu' \in NL^a \setminus EP^a$ for some a . By Corollary 18.3.1 in Rockafellar (1997), the extreme points of the convex hull of NL^a belong to NL^a . Then, by Minkowsky-Caratheodory Theorem (see Simon (2011), Theorem 8.11), there exists a collection $(x_i)_{i=1}^{|\Omega|}$ with $x_i \in EP^a$ for all i and a list $(\alpha_i)_{i=1}^{|\Omega|}$ with $\alpha_i \geq 0 \ \forall i$ and $\sum_i \alpha_i = 1$ such that $\mu' = \sum_i \alpha_i x_i$.

Consider another sender's strategy τ' defined as

$$\tau'(\mu) = \begin{cases} \hat{\tau}(\mu) & \text{if } \mu \notin \{\mu'\} \cup (x_i)_{i=1}^{|\Omega|} \\ \hat{\tau}(\mu) + \hat{\tau}(\mu')\alpha_i & \text{if } \mu = x_i, \ i = 1, \dots, |\Omega| \\ 0 & \text{if } \mu = \mu' \end{cases} \quad (16)$$

Note that for $\mu \notin \text{supp}(\hat{\tau}) \cup (x_i)_{i=1}^{|\Omega|}$, we have $\hat{\tau}(\mu) = \tau'(\mu) = 0$ and thus $\text{supp}(\tau') = \text{supp}(\hat{\tau}) \cup (x_i)_{i=1}^{|\Omega|} \setminus \{\mu'\}$. Note that, as $x_i \in EP^a$, for all i , the receiver takes the

same action at x_i that he takes at μ' . If that were not the case—which could happen when multiple convex hulls of non-learning regions share the same extreme point—then the sender could only benefit from inducing the collection $(x_i)_i$ instead of μ' by the assumption of the sender preferred equilibrium. However, this would contradict that $(\hat{\tau}, (\hat{\phi}_\mu)_\mu)$ is an equilibrium, because then there existed another sender's strategy under which the sender would have a strictly higher expected payoff. Hence, since for all i , the receiver takes the same action at x_i that he takes at μ' , the expected payoff of the sender is the same under τ' as under $\hat{\tau}$, since it is linear on the non-learning regions and $\mu' = \sum_i \alpha_i x_i$. Hence, $(\tau', (\hat{\phi}_\mu)_\mu)$ is also an equilibrium. Proceeding like this for every μ' that induces a posterior that is not an extreme point of some non-learning regions, we can construct a corresponding extreme-point equilibrium. As in proof of Proposition 1, note that, generally, it can be true that at some belief $\tilde{\mu} \in \cup_a EP^a$, the solution to the receiver's maximization problem (5) may not be unique, in which case ϕ_μ^* may be different from ϕ_μ^N (whenever the sender strictly prefers the learning strategy). However, this cannot be the case for any belief in the constructed extreme-point equilibrium and we indeed have that the receiver does not learn in equilibrium. The reason is that if this were not the case, then τ' yields strictly higher payoff than $\hat{\tau}$, contradicting that $\hat{\tau}$ was an equilibrium strategy.

Finally, let $(\tau', (\hat{\phi}_\mu)_\mu)$ be an equilibrium where $\mu \in \cup_a EP^a$ for all $\mu \in \text{supp}(\tau')$ and suppose that $|\text{supp}(\tau')| > |\Omega|$. Since τ' is sender's optimal strategy, $\text{supp}(\tau')$ supports a tangent hyperplane to $\text{cav}(\hat{v})$ at the prior μ_0 , and $\text{cav}(\hat{v})(\mu) = \hat{v}(\mu)$ for all $\mu \in \text{supp}(\tau')$ (Lemma 6). Caratheodory Theorem implies that there exists a subset $C \subset \text{supp}(\tau')$ with $|C| \leq |\Omega|$ such that the prior μ_0 lies in the convex hull of C . Hence, there exists a sender's strategy $\tilde{\tau}$ where $\text{supp}(\tilde{\tau}) = C$. Since $\text{cav}(\hat{v})(\mu) = \hat{v}(\mu)$ for all $\mu \in C$, $C \subset \text{supp}(\tau')$ and $\text{supp}(\tau')$ supports the tangent hyperplane to $\text{cav}(\hat{v})$ at the prior μ_0 , the strategy $\tilde{\tau}$ also supports the tangent hyperplane to $\text{cav}(\hat{v})$ at the prior μ_0 . Then $(\tilde{\tau}, (\hat{\phi}_\mu)_\mu)$ is an extreme-point equilibrium with $|\text{supp}(\tilde{\tau})| \leq |\Omega|$. $\square$

7.4 Proofs for Section 4

Proposition 3: The sender's expected equilibrium payoff is weakly increasing in λ .

Proof. By Proposition 1, we know that a sender's optimal strategy is contained in the set of strategies under which the receiver never learns. We show that the non-learning regions do not shrink as λ increases. As the sender can choose from the same (or possibly even bigger) set of strategies, he never becomes strictly worse off as λ increases.

Let $\lambda \geq 0$ and let us denote $\phi_{\mu,\lambda}^*$ an optimal strategy at μ when the cost parameter is λ . Suppose $\phi_{\mu,\lambda}^* = \phi_\mu^N$. Let ϕ_μ^L be an arbitrary receiver's strategy with strictly positive learning at μ , i.e., $\phi_\mu^L \neq \phi_\mu^N$. By optimality of ϕ_μ^N at μ when the cost parameter is λ , we have

$$\mathbb{E}_{\phi_\mu^N}[U(\gamma)] - c(\phi_\mu^N; \mu, \lambda) \geq \mathbb{E}_{\phi_\mu^L}[U(\gamma)] - c(\phi_\mu^L; \mu, \lambda)$$

Let $\lambda' > \lambda$. Then $c(\phi_\mu^N; \mu, \lambda) = c(\phi_\mu^N; \mu, \lambda') = 0$ (no-learning costs zero) and $c(\phi_\mu^L; \mu, \lambda) < c(\phi_\mu^L; \mu, \lambda')$. Hence,

$$\begin{aligned} \mathbb{E}_{\phi_\mu^N}[U(\gamma)] - c(\phi_\mu^N; \mu, \lambda') &= \mathbb{E}_{\phi_\mu^N}[U(\gamma)] - c(\phi_\mu^N; \mu, \lambda) \geq \mathbb{E}_{\phi_\mu^L}[U(\gamma)] - c(\phi_\mu^L; \mu, \lambda) \\ &> \mathbb{E}_{\phi_\mu^L}[U(\gamma)] - c(\phi_\mu^L; \mu, \lambda') \end{aligned}$$

showing that no-learning strategy remains optimal at λ' : $\phi_{\mu,\lambda'}^* = \phi_\mu^N$.

□

Binary action-state setting

The next two lemmas are used to prove Proposition 5. They describe sufficient conditions for unique equilibrium in a binary action-state setting. Let us first state simplifying notation that we will use in proofs of Lemmas 7, 8, and Proposition 5.

Let $\Omega = \{\omega_0, \omega_1\}$, $A = \{a, b\}$. In this subsection of the proofs, we use μ_0, μ, γ to denote the probability of the state $\omega = 1$ (i.e., $\Pr[\omega = 1]$) at the prior, the interim, and the posterior belief respectively. Without loss of generality, assume that in each state, different action is uniquely optimal (if that is not satisfied, then the sender can never benefit from persuasion). Let $\sigma^*(\mu) = \{a\}$ when $\mu = 0$ and $\sigma^*(\mu) = \{b\}$ when $\mu = 1$. Then there exists two threshold beliefs $0 \leq \underline{\mu} \leq \bar{\mu} \leq 1$ such that $NL^a = [0, \underline{\mu}]$ and $NL^b = [\bar{\mu}, 1]$. Without loss of generality, we consider sender's strategies with no more than two elements in their support. Note that $\hat{v}(\mu)$ is a piecewise-linear, upper semi-continuous function with linear segments on $[0, \underline{\mu}]$, $[\underline{\mu}, \bar{\mu}]$, and $[\bar{\mu}, 1]$ (see Example B1 in Appendix B for explanation why $\hat{v}(\mu)$ is linear on $[\underline{\mu}, \bar{\mu}]$). Further, note that whenever $\underline{\mu} \neq \bar{\mu}$, the function $\hat{v}(\mu)$ is continuous in μ for all $\mu \in [0, 1]$. The concavification $cav(\hat{v})(\mu)$ of $\hat{v}(\mu)$ can attain four forms: (i) $\hat{v}(\mu) = cav(\hat{v})(\mu)$ iff $\mu \in [\bar{\mu}, 1] \cup \{0\}$; (ii) $\hat{v}(\mu) = cav(\hat{v})(\mu)$ iff $\mu \in [0, \underline{\mu}] \cup \{1\}$; (iii) $\hat{v}(\mu) = cav(\hat{v})(\mu)$ iff $\mu \in \{0\} \cup \{1\}$; (iv) $\hat{v}(\mu) = cav(\hat{v})(\mu)$ for all $\mu \in [0, 1]$. Note that in iv), the sender does not benefit from persuasion for any possible prior $\mu_0 \in (0, 1)$, since $\hat{v}(\mu_0) = cav(\hat{v})(\mu_0)$ holds (Lemma 6).

Lemma 7. Suppose $|A| = |\Omega| = 2$ and, for a given λ , the sender benefits from persuasion. Then only never-learning equilibria exist.

Proof. Suppose that the sender benefits from persuasion, but, contrary to the statement, there exists a sender's optimal strategy τ^* with $\tilde{\mu} \in \text{supp}(\tau^*)$ and $\tilde{\mu} \in (\underline{\mu}, \bar{\mu})$, i.e., the receiver learns at $\tilde{\mu}$. Then $\hat{v}(\tilde{\mu}) = cav(\hat{v})(\tilde{\mu})$ (Lemma 6). Since $\hat{v}(\mu)$ is linear over $[\underline{\mu}, \bar{\mu}]$, this implies that $\hat{v}(\mu) = cav(\hat{v})(\mu)$ for all $\mu \in [\underline{\mu}, \bar{\mu}]$. But then, only case iv) can happen. In particular, $\hat{v}(\mu_0) = cav(\hat{v})(\mu_0)$, which contradicts that the sender benefits from persuasion. $\square$

Lemma 8. Suppose $|A| = |\Omega| = 2$, and, for a given λ , that the sender benefits from persuasion and A1 holds. Then there exists a unique equilibrium.

Proof. First, we show that only extreme-point equilibria exist. We then show that this implies unique equilibrium. Suppose that the sender benefits from persuasion

and A1 holds. By Lemma 7, only non-learning equilibria exist. Suppose there exist a non-learning equilibrium, which is not an extreme-point equilibrium. That is, suppose there exists $\mu' \in \text{supp}(\tau^*)$ with $\mu' \notin EP^a \cup EP^b = \{0, \underline{\mu}, \bar{\mu}, 1\}$. WLOG, suppose $\mu' \in (\bar{\mu}, 1)$. Then $\hat{v}(\mu') = \text{cav}(\hat{v})(\mu')$ (Lemma 6). Since $\hat{v}(\mu)$ is linear over $[\bar{\mu}, 1]$, it then implies that $\hat{v}(\mu) = \text{cav}(\hat{v})(\mu)$ for all $\mu \in [\bar{\mu}, 1]$. Then only case (i) or (iv) can happen. Under (iv), the sender does not benefit from persuasion, hence it must be the case (i). Furthermore, since μ' is in the support of the optimal sender's strategy, it supports the tangent hyperplane to the concavification above the prior. Hence, $\hat{v}(\mu)$ is also part of this tangent for all $\mu \in [\bar{\mu}, 1]$, as it is linear and coincides with $\text{cav}(\hat{v})(\mu)$. However, then action b violates the assumption A1. Hence, only extreme-point equilibria exist.

Note that, since the sender benefits from persuasion, $\text{supp}(\tau^*) \in \{\{0, \bar{\mu}\}, \{\underline{\mu}, 1\}, \{0, 1\}\}$. Any other combination of extreme points contradicts that the sender benefits from persuasion, because then, from the shape of $\hat{v}(\mu)$ and Lemma 6, it holds that $\hat{v}(\mu_0) = \text{cav}(\hat{v})(\mu_0)$. Suppose, contrary to the proposition, there are two different optimal sender's strategies. Then there is a non-learning region of (at least) one of the actions such that the two sender's strategies each induce different extreme point of the same non-learning region. But then a new strategy that would, instead, *ceteris paribus*, induce a convex combination of these two extreme points would also be optimal (since still the same action is taken under the convex combination; more specifically, both the extreme points support the tangent hyperplane to the concavification above prior. Since $\hat{v}(\mu)$ is linear over that whole non-learning region, it then follows that all beliefs in that non-learning region also support the tangent hyperplane to the concavification above the prior and can thus be part of an optimal strategy). However, such a new belief lie inside the non-learning region, which contradicts that there are only extreme-point equilibria. $\square$

Proposition 5. Assume $|A| = |\Omega| = 2$, and that for all $\lambda > 0$: A1 holds and the sender benefits from persuasion. Then the receiver's expected equilibrium utility (weakly) *decreases* in λ .

Proof. Suppose that, for all $\lambda > 0$, the sender benefits from persuasion and assumption A1 holds. Then, for all $\lambda > 0$, there always exist a unique equilibrium (Lemma 8). For each λ , the support of the unique optimal strategy lies in the set of extreme points of non-learning regions (Proposition 2). Consider an arbitrary value $\lambda > 0$ and the associated sender's optimal strategy τ_λ^* . Since the sender benefits from persuasion, i.e., $\hat{v}(\mu_0) < \text{cav}(\hat{v})(\mu_0)$ (Lemma 6), and since $\hat{v}(\mu)$ is a piecewise-linear function, with linear segments over $[0, \underline{\mu}]$, $[\underline{\mu}, \bar{\mu}]$ and $[\bar{\mu}, 1]$, the support of τ_λ^* is either (A) $\{0, 1\}$, (B) $\{0, \bar{\mu}\}$ or (C) $\{\underline{\mu}, 1\}$, as any other combination of extreme points contradicts that the sender benefits from persuasion. Consider $\lambda' > \lambda$. As shown in the proof of Proposition 3, the non-learning regions do not shrink as λ increases. Hence, the new non-learning regions are $NL^a = [0, \underline{\mu}']$ and $NL^b = [\bar{\mu}', 1]$, with $\underline{\mu} \leq \underline{\mu}' \leq \bar{\mu}' \leq \bar{\mu}$. It is direct to see that in all three cases (A), (B) and (C), the receiver is weakly better off (whenever the sender does not switch between the three cases as λ changes).

To complete the proof, we show that the sender does not switch between the three cases (A), (B), and (C) as λ changes. That is, we show that if inducing $\{0, \bar{\mu}\}$ is optimal at λ , then inducing $\{0, \bar{\mu}'\}$ is optimal at λ' (and not, for instance, inducing $\{\underline{\mu}', 1\}$). Let us consider all different possible specifications of sender's payoff:

(1.) First, let the sender prefer different actions in each of the state. (a) If he agrees with the receiver on the preferred actions under full information (he prefers action a when $\omega = \omega_0$ and action b when $\omega = \omega_1$), then $\hat{v}(\mu)$ is convex in μ for any λ and he optimally provides full information regardless of λ . We are in case (A). (b) If he disagrees with the receiver on the preferred actions under full information (he prefers action b when $\omega = \omega_0$ and action a when $\omega = \omega_1$), then, for low enough λ —so that the interim belief at which the sender is indifferent between the two actions, $\hat{\mu}$: $\mathbb{E}_{\hat{\mu}}[v(a, \omega)] = \mathbb{E}_{\hat{\mu}}[v(b, \omega)]$, lies in the learning region of the receiver— $\hat{v}(\mu)$ is concave in μ and the sender optimally provides zero information for any prior $\mu_0 \in (0, 1)$. Then, this case violates the assumption that the sender benefits from persuasion for all values of λ . Intuitively, it resembles a zero-sum game, in which case the sender cannot benefit from persuasion as any information leads the receiver to do

the opposite of what the sender prefers.

(2.) Second, let the sender prefer (at least weakly) one action over the other in both states. Without loss of generality, let b be such a dominant action. Let us consider an arbitrary $\lambda > 0$. Then the expected sender's payoff from action b satisfies $\mathbb{E}_\mu[v(b, \omega)] \geq \hat{v}(\mu)$ for all μ . More precisely, $\mathbb{E}_\mu[v(b, \omega)] = \hat{v}(\mu)$ when $\mu \in NL^b = [\bar{\mu}, 1]$ and $\mathbb{E}_\mu[v(b, \omega)] > \hat{v}(\mu)$ otherwise (the situation when $v(b, \omega_0) = v(a, \omega_0)$ leading to $\mathbb{E}_\mu[v(b, \omega)] = \hat{v}(\mu)$ when $\mu = 0$ is excluded by assumption A1). In this case, the concavification of $\hat{v}(\mu)$ is piecewise linear, with line connecting $v(a, \omega_1)$ at $\mu = 0$ with $\mathbb{E}_{\bar{\mu}}[v(b, \omega)]$ at $\mu = \bar{\mu}$ and a line coinciding with $\mathbb{E}_\mu[v(b, \omega)]$ for $\mu \in [\bar{\mu}, 1]$. Then, the only possible case under which the sender benefits from persuasion is when the prior μ_0 lies somewhere between $(0, \bar{\mu})$, because only for these beliefs the concavification is strictly above $\hat{v}(\mu)$ (Lemma 6). Then, the support of optimal sender's strategy is case (B), $\{0, \bar{\mu}\}$. As this holds for an arbitrary λ , no switch between the cases occurs as λ changes. Intuitively, when the sender has a (weakly) dominant action, his strategy is to maximize the probability of taking that action. This is done by inducing the closest possible belief to the prior at which the receiver already takes the dominant action without learning, $\mu = \bar{\mu}$, and inducing the farthest possible belief from the prior at which the dominated action is taken, $\mu = 0$.

□

8 Appendix B: Examples

8.1 Application on Extremism vs Conservatism: Analytical Solution

In this section, we use μ_0, μ, γ to denote the probability of the state $\omega = 1$ (i.e., $\Pr[\omega = 1]$) at the prior, the interim, and the posterior belief respectively. We use the Extreme-Point Solution Algorithm to obtain an analytical solution. Note that

in both cases (setting with the conservative receiver and with the extreme receiver), the shape of the receiver's utility implies that $\hat{v}(\mu)$ is symmetric around the prior $\mu_0 = 0.5$. Therefore, there is an optimal sender's strategy with a symmetric support around the prior. As such, for step 2 of the Algorithm, we can state an additional restriction to only consider such symmetric supports. Under this restriction, step 3 immediately determines that the candidate seller's strategies place equal probabilities on the beliefs in their support. If τ is then such a candidate sender's strategy, then $\hat{v}(\mu) = \hat{v}(\mu')$ for $\mu, \mu' \in \text{supp}(\tau)$, and $\mathbb{E}_\tau[\hat{v}(\mu)] = \hat{v}(\mu) = \hat{v}(\mu')$. Therefore, to compare the values of each candidate strategy, for each strategy, we only need to determine $\hat{v}(\mu)$ at one of the beliefs from its support and compare these values of different strategies against each other. Further, note that in this setting, a necessary and sufficient condition for multiple equilibria is that there are multiple extreme-point strategies with symmetric support. This follows from the fact that $\hat{v}(\mu)$ is linear on a learning region between two non-learning regions (as explained in Example B1 in this Appendix). Therefore, the value of any sender's strategy under which the receiver learns can be obtained when used a convex combination of two different extreme-point strategies. The same is true for any non-learning strategy which is not an extreme-point strategy. Therefore, a necessary condition for multiple equilibria is existence of two different optimal symmetric extreme-point strategies.

8.1.1 The Conservative Receiver

Note that for any posterior belief $\gamma \in [0, 1]$, the extreme actions $\{-2, 2\} \notin \sigma^*(\gamma)$. Therefore, for any value of $\lambda > 0$: $NL^{-2} = NL^2 = \emptyset$. Further, note that there is a threshold value $\hat{\lambda} \in (0, \infty)$, such that for all $\lambda \geq \hat{\lambda}$: $NL^{-1}, NL^0, NL^1 \neq \emptyset$, and for all $\lambda < \hat{\lambda}$: $NL^0 = \emptyset$ and $NL^{-1}, NL^1 \neq \emptyset$.

Consider $\lambda \geq \hat{\lambda}$. Step 1: Using Lemma 4, equation (iii), we obtain the sets of extreme points of each non-learning regions: $EP^{-1} = \left\{0, \frac{1}{1+e^{\frac{1}{\lambda}}+e^{\frac{2}{\lambda}}+e^{\frac{3}{\lambda}}}\right\} := \{0, \mu'\}$, $EP^0 = \left\{\frac{e^{\frac{3}{\lambda}}}{1+e^{\frac{1}{\lambda}}+e^{\frac{2}{\lambda}}+e^{\frac{3}{\lambda}}}, 1 - \frac{e^{\frac{3}{\lambda}}}{1+e^{\frac{1}{\lambda}}+e^{\frac{2}{\lambda}}+e^{\frac{3}{\lambda}}}\right\} := \{\mu'', 1 - \mu''\}$ and $EP^1 = \left\{1 - \frac{1}{1+e^{\frac{1}{\lambda}}+e^{\frac{2}{\lambda}}+e^{\frac{3}{\lambda}}}, 1\right\} := \{1 - \mu', 1\}$. Step 2: the supports of symmetric candidate sender's strategies are:

$\text{supp}(\tau_1) = \{0, 1\}$ (i.e., full information), $\text{supp}(\tau_2) = \{\mu', 1 - \mu'\}$ and $\text{supp}(\tau_3) = \{\mu'', 1 - \mu''\}$. Step 3: as already noted, each candidate strategy places equal probability at the beliefs in its support. Step 4 & 5: Note that since action 0 is optimal at the prior, $\hat{v}(\mu'') = \hat{v}(1 - \mu'') = \hat{v}(\mu_0)$, the sender does not benefit from persuasion under τ_3 and the value of the strategy is $\mathbb{E}_{\tau_3}[\hat{v}(\mu)] = \mathbb{E}_{\mu_0}[v(0, \omega)] = \hat{v}(\mu_0)$. Further, note that beliefs 0 and μ' belong to the non-learning region of action $a = -2$. Noting that $\mathbb{E}_{\mu}[v(-2, \omega)]$ is strictly decreasing in μ , we have $\hat{v}(0) > \hat{v}(\mu')$ (similarly, $\hat{v}(1 - \mu') < \hat{v}(1)$) for any specification of λ . Therefore, τ_1 is strictly better than τ_2 . Noting that $\hat{v}(0) = \hat{v}(1) > \hat{v}(\mu_0)$, we get that τ_1 is strictly better than τ_3 . Full revelation is strictly optimal.

The threshold value $\hat{\lambda}$ is the value at which NL^0 is a singleton: $\mu'' = 1 - \mu''$ yielding $\hat{\lambda} = 1.64102$. Consider $\lambda < \hat{\lambda}$. Using Lemma 4, equation (iii), we obtain $EP^{-1} = \left\{0, \frac{1}{1+e^{\frac{1}{\lambda}}}\right\}$ and $EP^1 = \left\{1 - \frac{1}{1+e^{\frac{1}{\lambda}}}, 1\right\}$. By the same logic as when comparing τ_1 and τ_2 in previous paragraph, we conclude that full revelation is the unique solution.

Therefore, for any $\lambda > 0$, full revelation is uniquely optimal.

8.1.2 The Extreme Receiver

Note that there are two threshold values $\tilde{\lambda} > \bar{\lambda} > 0$ such that: for all $\lambda \geq \tilde{\lambda}$, all non-learning regions are nonempty, for all $\lambda \in [\bar{\lambda}, \tilde{\lambda})$, $NL^0 = \emptyset$ and all others are nonempty, and for all $\lambda < \bar{\lambda}$, $NL^{-2}, NL^2 \neq \emptyset$ and all others are empty sets.

Consider $\lambda \geq \hat{\lambda}$. Step 1: Using Lemma 4, equation (iii), we obtain the sets of extreme points of each non-learning regions: $EP^{-2} = \left\{0, \frac{-1+e^{\frac{1}{\lambda}}}{-1+e^{\frac{8}{\lambda}}}\right\} := \{0, \mu'\}$, $EP^{-1} = \left\{-\frac{e^{\frac{7}{\lambda}}(1-e^{\frac{1}{\lambda}})}{-1+e^{\frac{8}{\lambda}}}, \frac{-1+e^{\frac{3}{\lambda}}}{-1+e^{\frac{8}{\lambda}}}\right\} := \{\mu'', \mu'''\}$, $EP^0 = \left\{-\frac{e^{\frac{5}{\lambda}}(1-e^{\frac{3}{\lambda}})}{-1+e^{\frac{8}{\lambda}}}, 1 - \frac{e^{\frac{5}{\lambda}}(1-e^{\frac{3}{\lambda}})}{-1+e^{\frac{8}{\lambda}}}\right\} := \{\mu'''', 1 - \mu''''\}$, $EP^1 = \{1 - \mu''', 1 - \mu''\}$, $EP^2 = \{1 - \mu', 1\}$. Step 2: The candidate sender's strategies with a symmetric support around the prior are strategies $\text{supp}(\tau_1) = \{0, 1\}$, $\text{supp}(\tau_2) = \{\mu', 1 - \mu'\}$, $\text{supp}(\tau_3) = \{\mu'', 1 - \mu''\}$, $\text{supp}(\tau_4) = \{\mu''', 1 - \mu'''\}$, and $\text{supp}(\tau_5) = \{\mu'''', 1 - \mu''''\}$. Step 3: As already noted, each candidate strat-

egy places equal probabilities at the beliefs in its support. Step 4 & 5: Noting that $\mathbb{E}_\mu[v(-2, \omega)]$ and $\mathbb{E}_\mu[v(-1, \omega)]$ are strictly decreasing in μ (and $\mathbb{E}_\mu[v(2, \omega)]$ and $\mathbb{E}_\mu[v(1, \omega)]$ are strictly increasing in μ), strategy τ_1 is strictly better than τ_2 and strategy τ_3 is strictly better than τ_4 , $\forall \lambda \geq \hat{\lambda}$. Further, note that $\hat{v}(0) = \hat{v}(1) = -1 = \hat{v}(\mu''''') = \hat{v}(\mu''''') = \hat{v}(\mu_0)$. Therefore, strategy τ_1 and τ_5 yields the same expected payoff and, furthermore, the sender does not benefit from persuasion. Noting that $\hat{v}(\mu'') = \hat{v}(1 - \mu'') > \hat{v}(\mu_0)$ for all $\lambda \geq \hat{\lambda}$, we conclude that strategy τ_3 is optimal. The threshold $\hat{\lambda}$ is a value at which the non-learning region NL^0 is a singleton, i.e., at which $\mu'''' = 1 - \mu''''$, yielding $\hat{\lambda} = 7.66092$

Consider $\lambda \in [\bar{\lambda}, \hat{\lambda}]$. Step 1: Using Lemma 4, equation (iii), we obtain the sets of extreme points of each non-learning regions. They have the same form as in previous paragraph, except for these changes: $EP^0 = \emptyset$, $EP^{-1} = \left\{ \mu'', \frac{1}{1+e^{\frac{8}{\lambda}}} \right\}$ and $EP^1 = \left\{ 1 - \frac{1}{1+e^{\frac{8}{\lambda}}}, 1 - \mu'' \right\}$. By the same logic as before, since $\hat{v}(\mu'') = \hat{v}(\mu') > \hat{v}(0) = \hat{v}(1)$ for all values of $\lambda \in [\bar{\lambda}, \hat{\lambda}]$, and the optimal strategy is thus τ_3 . The threshold value $\bar{\lambda}$ is a value at which the non-learning regions NL^{-1} and NL^1 are singletons, i.e., at which $\mu'' = \frac{1}{1+e^{\frac{8}{\lambda}}}$, which yields $\bar{\lambda} = 5.99735$.

Consider $\lambda < \bar{\lambda}$. Then $EP^{-1} = EP^0 = EP^1 = \emptyset$ and $EP^{-2} = \left\{ 0, \frac{1}{1+e^{\frac{16}{\lambda}}} \right\}$ and $EP^2 = \left\{ 1 - \frac{1}{1+e^{\frac{16}{\lambda}}}, 1 \right\}$. Since $\mathbb{E}_\mu[v(-2, \omega)]$ is strictly decreasing in μ (and $\mathbb{E}_\mu[v(2, \omega)]$ is strictly increasing in μ), strategy τ_1 is strictly optimal for all $0 < \lambda < \bar{\lambda}$.

Therefore, for $0 < \lambda < \bar{\lambda}$, full revelation is uniquely optimal. For $\lambda \geq \bar{\lambda}$, full revelation is not optimal and the unique optimal strategy is strategy τ_3 where $\text{supp}(\tau_3) = \left\{ -\frac{e^{\frac{7}{\lambda}}(1-e^{\frac{1}{\lambda}})}{-1+e^{\frac{8}{\lambda}}}, 1 - \frac{e^{\frac{7}{\lambda}}(1-e^{\frac{1}{\lambda}})}{-1+e^{\frac{8}{\lambda}}} \right\}$. The derivatives are $\frac{e^{\frac{7}{\lambda}}(-7+8e^{\frac{1}{\lambda}-e^{\frac{8}{\lambda}}})}{(-1+e^{\frac{8}{\lambda}})^2 \lambda^2} > 0$ and $-\frac{e^{\frac{7}{\lambda}}(-7+8e^{\frac{1}{\lambda}-e^{\frac{8}{\lambda}}})}{(-1+e^{\frac{8}{\lambda}})^2 \lambda^2} < 0$. The extreme points that are induced by strategy τ_3 are moving strictly apart from each other as λ increases. Hence, the sender provides strictly more information in Blackwell sense as λ increases.

8.2 Example B1: Double Concavification Approach

A seller (he) is persuading a buyer (she) to buy his product (e.g. a music record). The product is either a good ($\omega = 1$) or a bad ($\omega = 0$) match. The buyer can either buy ($a = 1$) or not buy ($a = 0$). In this section, we use μ_0, μ, γ to denote the probability of the state $\omega = 1$ (i.e., $\Pr[\omega = 1]$) at the prior, the interim, and the posterior belief respectively. Let $\mu_0 < 0.5$. Consider Shannon cost. Tables 1 and 2 depict the buyer's and the seller's payoffs.

$u(a, \omega)$	$\omega = 0$	$\omega = 1$
$a = 0$	0	0
$a = 1$	-1	1

Table 1: Buyer's Payoff

$v(a, \omega)$	$\omega = 0$	$\omega = 1$
$a = 0$	0	0
$a = 1$	1	1

Table 2: Seller's Payoff

We first solve the buyer's maximization problem. Given μ , the buyer maximizes

$$\begin{aligned} \max_{\phi \in \Delta([0,1])} \quad & \mathbb{E}_\phi[B(\gamma)] - \lambda(F(\mu) - \mathbb{E}_\phi[F(\gamma)]) \\ \text{s.t.} \quad & \mathbb{E}_\phi[\gamma] = \mu, \end{aligned} \tag{17}$$

where the expectation is taken over the posterior beliefs induced by ϕ and $B(\gamma) = \max\{0, 2\gamma - 1\}$ is the buyer's gross expected utility at posterior γ under optimal action. We follow Caplin and Dean (2013) when solving for the buyer's optimal behavior. They show that the buyer's optimal behavior can be read off from a geometric interpretation of the problem (17), which can be rewritten as

$$\begin{aligned} \max_{\phi \in \Delta([0,1])} \quad & \mathbb{E}_\phi[\hat{u}(\gamma)] - \underbrace{\lambda F(\mu)}_{=const.} \\ \text{s.t.} \quad & \mathbb{E}_\phi[\gamma] = \mu, \end{aligned} \tag{18}$$

where $\hat{u}(\gamma) = B(\gamma) + \lambda F(\gamma)$ is the buyer's value function at posterior γ . Let $cav(\hat{u})(\gamma)$ denote the concavification of $\hat{u}(\gamma)$ defined as the smallest concave function that is everywhere weakly greater than $\hat{u}(\gamma)$. Then the support of the receiver's optimal information strategy, $\text{supp}(\phi_\mu^*)$, are those posterior beliefs that support the tangent hyperplane to the lower epigraph of the concavification above the interim

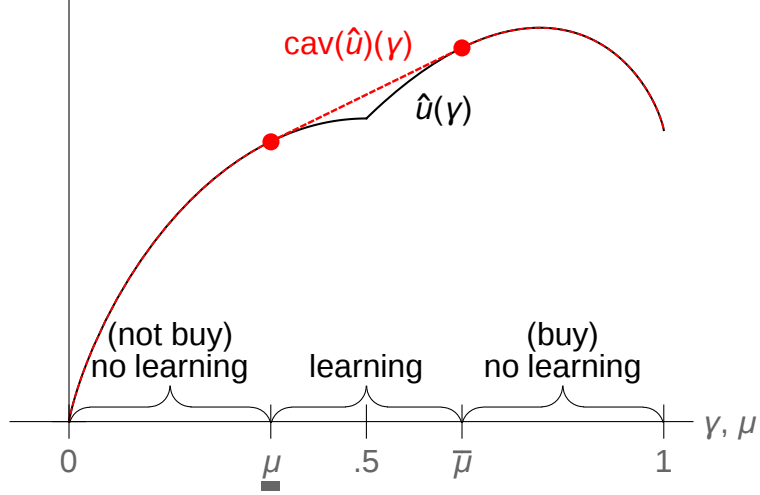


Figure 4: Buyer's value function $\hat{u}(\gamma)$ and its concavification $\text{cav}(\hat{u})(\gamma)$. When $\mu \in [0, \underline{\mu}] \cup [\bar{\mu}, 1]$, the buyer does not learn at μ . When $\mu \in (\underline{\mu}, \bar{\mu})$, the buyer learns at μ and the support of the buyer's optimal strategy is $\text{supp}(\phi_\mu^*) = \{\underline{\mu}, \bar{\mu}\}$.

belief μ . Whenever $\text{cav}(\hat{u})(\gamma) = \hat{u}(\gamma)$, the receiver does not learn at μ , and whenever $\text{cav}(\hat{u})(\gamma) > \hat{u}(\gamma)$, she learns at μ , where $\text{supp}(\phi_\mu^*) = \{\underline{\mu}, \bar{\mu}\}$ is the support of the optimal information strategy, see Figure 4.²³ Note that the buyer's optimal strategy satisfies a “locally invariant posteriors” property, which states that the support of the optimal strategy is invariant to local changes in the interim belief (see Caplin and Dean (2013)). That is, $\text{supp}(\phi_\mu^*) = \text{supp}(\phi_{\mu'}^*)$ for every μ, μ' that is in the convex hull of $\text{supp}(\phi_\mu^*)$. In binary state, this property implies that $\hat{v}(\mu)$ is linear on a particular learning region.

Next, let us turn to the seller's maximization problem. Given μ_0 , the seller maximizes

$$\begin{aligned} \max_{\tau \in \Delta([0,1])} \quad & \mathbb{E}_\tau[\hat{v}(\mu)] \\ \text{s.t.} \quad & \mathbb{E}_\tau[\mu] = \mu_0, \end{aligned} \tag{19}$$

where the expectation is taken over interim beliefs induced by τ and $\hat{v}(\mu)$ is the

²³Using equation (iii) from Lemma 4, we obtain $\underline{\mu} = \frac{1}{1+e^{\frac{1}{\lambda}}}$ and $\bar{\mu} = \frac{e^{\frac{1}{\lambda}}}{1+e^{\frac{1}{\lambda}}}$.

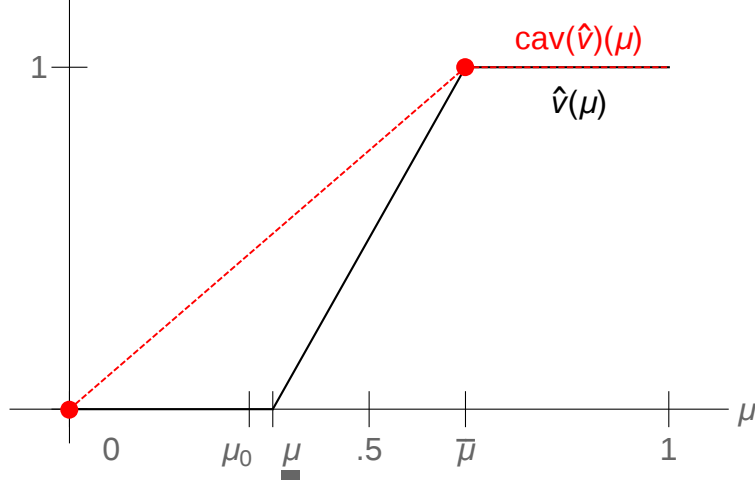


Figure 5: Seller's conditional expected payoff $\hat{v}(\mu)$, its concavification $\text{cav}(\hat{v})(\mu)$ and support of the sender's optimal strategy $\text{supp}(\tau^*) = \{0, \bar{\mu}\}$.

seller's expected utility at μ under optimal continuation play of the buyer at μ .²⁴ Recall, that $\text{cav}(\hat{v})(\mu)$ denotes the concavification of $\hat{v}(\mu)$. Determining the seller's optimal behavior can be done in a similar fashion as we did it for the buyer. The seller's expected equilibrium utility is the concavification evaluated at the prior, $\text{cav}(\hat{v})(\mu_0)$, and $\text{supp}(\tau^*) = \{0, \bar{\mu}\}$ is the support of the optimal sender's strategy, see Figure 5.

8.3 Example B2: Failure of Proposition 1 Under a Non-Uniformly Posterior-Separable Cost Function

Let $\Omega = \{\omega_0, \omega_1\}$, $A = \{a, b, c\}$. Let the sender's and receiver's payoffs be given by Tables 3 and 4. In this section, we use μ_0, μ, γ to denote the probability of the state $\omega = 1$ (i.e., $\Pr[\omega = 1]$) at the prior, the interim, and the posterior belief respectively. Let $\mu_0 = 0.5$.

²⁴Hence, $\hat{v}(\mu) = 0$ if $\mu \leq \underline{\mu}$ (the buyer does not learn and does not buy), $\hat{v}(\mu) = 1$ if $\mu \geq \bar{\mu}$ (the buyer does not learn and buys), and $\hat{v}(\mu) = \frac{1}{\bar{\mu} - \underline{\mu}}\mu - \frac{\underline{\mu}}{\bar{\mu} - \underline{\mu}}$ for $\mu \in (\underline{\mu}, \bar{\mu})$.

$u(a, \omega)$	ω_0	ω_1
a	1	-2
b	0	0
c	-2	1

Table 3: Receiver's Payoff

$v(a, \omega)$	ω_0	ω_1
a	0	0
b	0	1.5
c	2	2

Table 4: Sender's Payoff

8.3.1 Receiver's Optimal Behavior Under Shannon Cost

Let us first solve the model under the Shannon cost. Note that the receiver wants to take action a in state ω_0 , action c in state ω_1 , and action b when uncertain about the state. For high enough λ , there are three non-learning regions: $NL^a = [0, \mu_a(\lambda)]$, $NL^b = [\underline{\mu}_b(\lambda), \bar{\mu}_b(\lambda)]$, and $NL^c = [\mu_c(\lambda), 1]$, where $0 < \mu_a(\lambda) < \underline{\mu}_b(\lambda) \leq \bar{\mu}_b(\lambda) < \mu_c(\lambda) < 1$. For low enough λ , the non-learning region of action b disappears: whenever the receiver acquires information, she does so to decide on taking either action a or c . Then, there are only two non-learning regions: $NL^a = [0, \mu'_a(\lambda)]$, and $NL^c = [\mu'_c(\lambda), 1]$, where $0 \leq \mu'_a(\lambda) < \mu'_c(\lambda) \leq 1$.

Using equation (iii) in Lemma 4, for high enough λ , we obtain

$$\mu_a(\lambda) = \frac{e^{\frac{1}{\lambda}} - 1}{e^{\frac{3}{\lambda}} - 1}, \quad \bar{\mu}_b(\lambda) = \frac{e^{\frac{2}{\lambda}} - 1}{e^{\frac{3}{\lambda}} - 1}, \quad \underline{\mu}_b(\lambda) = \frac{1 - e^{-\frac{1}{\lambda}}}{1 - e^{-\frac{3}{\lambda}}}, \quad \mu_c(\lambda) = \frac{1 - e^{-\frac{2}{\lambda}}}{1 - e^{-\frac{3}{\lambda}}} \quad (20)$$

Analogously, for low enough λ , we obtain

$$\mu'_a(\lambda) = \frac{e^{\frac{3}{\lambda}} - 1}{e^{\frac{6}{\lambda}} - 1}, \quad \mu'_c(\lambda) = \frac{1 - e^{-\frac{3}{\lambda}}}{1 - e^{-\frac{6}{\lambda}}} \quad (21)$$

8.3.2 Non-Uniformly Posterior-Separable Cost Function: Setup

Now, let us consider a particular case of posterior-separable cost functions that are not UPS. In particular, we consider a function, in which the scaling parameter decreases in μ in the interval $[\mu_0, 1]$ (as opposed to UPS cost, where λ is independent of μ). Intuitively, the scaling parameter decreasing (increasing) in μ when $\mu > \mu_0$ ($\mu < \mu_0$) means that information acquisition technology is more efficient (cheaper)

	$\text{supp}(\phi_\mu^*)$	$Pr[a]$	$Pr[b]$	$Pr[c]$
$\mu \in [0, \mu_a(\bar{\kappa})]$	$\{\mu\}$	1	0	0
$\mu \in (\mu_a(\bar{\kappa}), \underline{\mu}_b(\bar{\kappa}))$	$\{\mu_a(\bar{\kappa}), \underline{\mu}_b(\bar{\kappa})\}$	$\frac{\underline{\mu}_b(\bar{\kappa}) - \mu}{\underline{\mu}_b(\bar{\kappa}) - \mu_a(\bar{\kappa})}$	$\frac{\mu - \mu_a(\bar{\kappa})}{\underline{\mu}_b(\bar{\kappa}) - \mu_a(\bar{\kappa})}$	0
$\mu \in [\underline{\mu}_b(\bar{\kappa}), \mu^b]$	$\{\mu\}$	0	1	0
$\mu \in (\mu^b, \hat{\mu})$	$\{\bar{\mu}_b(\tilde{\kappa}(\mu)), \mu_c(\tilde{\kappa}(\mu))\}$	0	$\frac{\mu_c(\tilde{\kappa}(\mu)) - \mu}{\mu_c(\tilde{\kappa}(\mu)) - \bar{\mu}_b(\tilde{\kappa}(\mu))}$	$\frac{\mu - \bar{\mu}_b(\tilde{\kappa}(\mu))}{\mu_c(\tilde{\kappa}(\mu)) - \bar{\mu}_b(\tilde{\kappa}(\mu))}$
$\mu \in [\hat{\mu}, \mu^c]$	$\{\mu'_a(\tilde{\kappa}(\mu)), \mu'_c(\tilde{\kappa}(\mu))\}$	$\frac{\mu'_c(\tilde{\kappa}(\mu)) - \mu}{\mu'_c(\tilde{\kappa}(\mu)) - \mu'_a(\tilde{\kappa}(\mu))}$	0	$\frac{\mu - \mu'_a(\tilde{\kappa}(\mu))}{\mu'_c(\tilde{\kappa}(\mu)) - \mu'_a(\tilde{\kappa}(\mu))}$
$\mu \in [\mu^c, 1]$	$\{\mu\}$	0	0	1

Table 5: Receiver's optimal behavior under a non-uniformly posterior-separable cost function

when the receiver is better informed about the state.²⁵

Formally, let $\bar{\kappa} = 3$, $\underline{\kappa} = 1.5$ and denote $\mu^b \equiv \bar{\mu}_b(\bar{\kappa})$ and $\mu^c \equiv \mu'_c(\underline{\kappa})$, where $\bar{\mu}_b(\cdot)$ is given by (20) for $\lambda = \bar{\kappa}$, and $\mu'_c(\cdot)$ is given by (21) for $\lambda = \underline{\kappa}$. That is, μ^b is the higher of the extreme points of NL^b when the receiver faces a Shannon cost with scaling cost parameter $\bar{\kappa}$, where the cost parameter is high enough so that there are three non-learning regions. Analogously, μ^c is the lower of the extreme points of NL^c when the receiver faces a Shannon cost with scaling cost parameter $\underline{\kappa}$, where the cost parameter is low enough so that there are only two non-learning regions.

Let us define the receiver's cost function as

$$c(\phi; \mu, \kappa) = \begin{cases} \bar{\kappa} [H(\mu) - \sum_{\gamma} H(\gamma)\phi(\gamma)] & \text{if } \mu \in [0, \mu^b] \\ \tilde{\kappa}(\mu) [H(\mu) - \sum_{\gamma} H(\gamma)\phi(\gamma)] & \text{if } \mu \in (\mu^b, \mu^c] \\ \underline{\kappa} [H(\mu) - \sum_{\gamma} H(\gamma)\phi(\gamma)] & \text{otherwise} \end{cases} \quad (22)$$

where $\tilde{\kappa}(\mu) = \bar{\kappa} \frac{\mu^c - \mu}{\mu^c - \mu^b} + \underline{\kappa} \frac{\mu - \mu^b}{\mu^c - \mu^b}$. Note that $\tilde{\kappa}(\mu)$ is a uniformly decreasing function on $[\mu^b, \mu^c]$ satisfying $\lim_{\mu \rightarrow \mu^b+} \tilde{\kappa}(\mu) = \bar{\kappa}$ and $\lim_{\mu \rightarrow \mu^c-} \tilde{\kappa}(\mu) = \underline{\kappa}$.

8.3.3 Non-Uniformly Posterior-Separable Cost Function: Solution

Table 5 captures how the receiver's optimal behavior and corresponding unconditional probabilities of taking each action vary with her interim belief. When $\mu \in [0, \mu_a(\bar{\kappa})]$, where $\mu_a(\cdot)$ is given by (20) for $\lambda = \bar{\kappa}$, the receiver does not learn and takes an action a . When $\mu \in (\mu_a(\bar{\kappa}), \underline{\mu}_b(\bar{\kappa}))$, where $\underline{\mu}_b(\bar{\kappa})$ is given by (20) for $\lambda = \bar{\kappa}$, the receiver learns in order to take either action a or b . In that case, the support of her optimal learning strategy are posteriors $\gamma_a = \mu_a(\bar{\kappa})$, under which she takes an action a , and $\gamma_b = \underline{\mu}_b(\bar{\kappa})$ under which she takes an action b . When $\mu \in [\underline{\mu}_b(\bar{\kappa}), \mu^b]$, the receiver does not learn and takes an action b . When $\mu \in (\mu^b, \hat{\mu})$, the receiver learns in order to take either action b or c . The support of her optimal learning strategy are posteriors

$$\gamma_b(\mu) = \frac{e^{\frac{2}{\bar{\kappa}(\mu)}} - 1}{e^{\frac{3}{\bar{\kappa}(\mu)}} - 1}, \quad \gamma_c(\mu) = \frac{1 - e^{-\frac{2}{\bar{\kappa}(\mu)}}}{1 - e^{-\frac{3}{\bar{\kappa}(\mu)}}}, \quad (23)$$

where she takes an action b if $\gamma_b(\mu)$ is induced and an action c otherwise. The belief $\hat{\mu}$ is a point of indifference, where the receiver is indifferent between learning in order to take either action b or c , or learning in order to take either action a or c . It is an interim belief, under which in a model with Shannon cost with $\lambda = \tilde{\kappa}(\hat{\mu})$, the non-learning region of action b is not an interval, but a singleton. That is, the following condition is satisfied:

$$\underline{\mu}_b(\tilde{\kappa}(\hat{\mu})) = \bar{\mu}_b(\tilde{\kappa}(\hat{\mu})) \quad (24)$$

where $\underline{\mu}_b(\cdot)$ and $\bar{\mu}_b(\cdot)$ are given by (20). When $\mu \in (\hat{\mu}, \mu^c)$, the receiver learns in order to take either action a or c . The support of her optimal learning strategy are posteriors

$$\hat{\gamma}_a(\mu) = \frac{e^{\frac{3}{\bar{\kappa}(\mu)}} - 1}{e^{\frac{6}{\bar{\kappa}(\mu)}} - 1}, \quad \hat{\gamma}_c(\mu) = \frac{1 - e^{-\frac{3}{\bar{\kappa}(\mu)}}}{1 - e^{-\frac{6}{\bar{\kappa}(\mu)}}}, \quad (25)$$

²⁵The failure of Proposition 1 holds when the scaling parameter decreases in μ for $\mu_0 < \mu$ and increases in μ when $\mu < \mu_0$ at the same time. To keep the analysis simple, we focus only on one side, when $\mu_0 < \mu$.

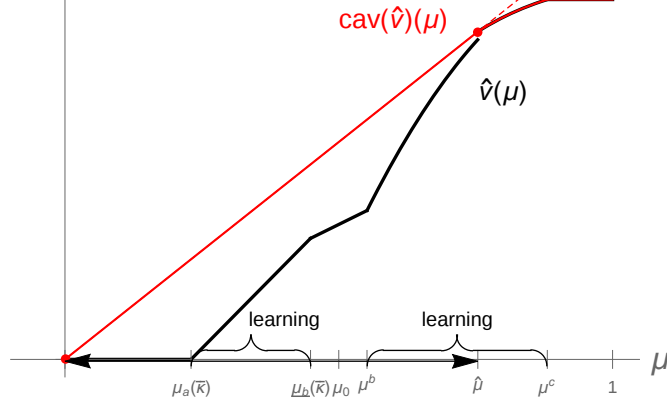


Figure 6: Non-existence of a non-learning equilibrium under a non-uniformly posterior-separable cost function (22): Sender's conditional expected payoff $\hat{v}(\mu)$, its concavification $\text{cav}(\hat{v})(\mu)$ and the support of sender's unique optimal strategy $\text{supp}(\tau^*) = \{0, \hat{\mu}\}$. The receiver learns at $\hat{\mu}$.

where she takes an action a if $\hat{\gamma}_a(\mu)$ is induced and an action c otherwise. When $\mu \in [\mu^c, 1]$, the receiver does not learn at μ and takes an action c . Knowing the support of the receiver's optimal information strategies for each μ , we can determine the unconditional probabilities for each of the action to be taken using Bayes' law.

Having characterized receiver's optimal behaviour, we can determine $\hat{v}(\mu)$, which takes the following form:

$$\hat{v}(\mu) = \begin{cases} 0 & \mu \in [0, \mu_a(\bar{\kappa})] \\ \frac{\mu_b(\bar{\kappa}) - \mu}{\mu_b(\bar{\kappa}) - \mu_a(\bar{\kappa})} \times 1.5\mu_b(\bar{\kappa}) & \mu \in (\mu_a(\bar{\kappa}), \mu_b(\bar{\kappa})) \\ 1.5\mu & \mu \in [\mu_b(\bar{\kappa}), \mu^b] \\ \frac{\mu - \bar{\mu}_b(\tilde{\kappa}(\mu))}{\mu_c(\tilde{\kappa}(\mu)) - \bar{\mu}_b(\tilde{\kappa}(\mu))} \times 2 + \frac{\mu_c(\tilde{\kappa}(\mu)) - \mu}{\mu_c(\tilde{\kappa}(\mu)) - \bar{\mu}_b(\tilde{\kappa}(\mu))} \times 1.5\bar{\mu}_b(\tilde{\kappa}(\mu)) & \mu \in (\mu^b, \hat{\mu}) \\ \frac{\mu - \mu'_a(\tilde{\kappa})}{\mu'_c(\tilde{\kappa}(\mu)) - \mu'_a(\tilde{\kappa}(\mu))} \times 2 & \mu \in [\hat{\mu}, \mu^c) \\ 2 & \mu \in [\mu^c, 1] \end{cases} \quad (26)$$

Figure 6 captures the sender's conditional expected payoff $\hat{v}(\mu)$, its concavification $\text{cav}(\hat{v})(\mu)$ and the support of sender's unique optimal strategy $\text{supp}(\tau^*) = \{0, \hat{\mu}\}$. Since $\hat{\mu} \in (\mu^b, \mu^c)$, the receiver learns at $\hat{\mu}$, and the Proposition 1 fails.

8.4 Example B3: Failure of Proposition 1 Under a Binary Signal Strategy with Exogenous Precision

In this section, we solve the buyer-seller setup used in Example B1 under a different cost function: we assume the receiver can obtain a partially revealing binary signal at a fixed cost $c \geq 0$.

Let $\Omega = \{0, 1\}$, $A = \{0, 1\}$, $v(a, \omega) = a$, $u(a, \omega) = 1$ if $a = \omega = 1$, $u(a, \omega) = -1$ if $a = 1$ and $\omega = 0$ and 0 otherwise. In this section, we use μ_0, μ, γ to denote the probability of the state $\omega = 1$ (i.e., $\Pr[\omega = 1]$) at the prior, the interim, and the posterior belief respectively. Given μ , the receiver can obtain a binary signal $s \in \{0, 1\}$ of precision $p := \Pr[s = \omega | \omega] > 0.5$ by paying $c \geq 0$. In this example, say *the receiver learns* if she pays c and gets the signal.

8.4.1 Receiver's Maximization Problem: Solution

Given μ , if the receiver learns, she updates her beliefs to a posterior $\gamma_s(\mu) := \Pr[\omega = 1 | s, \mu]$ with probability $\phi_s(\mu) := \Pr[s | \mu]$, where

$$\begin{aligned} \gamma_1(\mu) &= \frac{p\mu}{\phi_1}, & \phi_1(\mu) &= p\mu + (1-p)(1-\mu), \\ \gamma_0(\mu) &= \frac{(1-p)\mu}{\phi_0}, & \phi_0(\mu) &= (1-p)\mu + p(1-\mu). \end{aligned}$$

The receiver takes action $a = 1$ if $\gamma_s \geq 1/2$. Her expected utility from learning is

$$U^L(\mu) = \max\{0, 2\gamma_1(\mu) - 1\}\phi_1(\mu) + \max\{0, 2\gamma_0(\mu) - 1\}\phi_0(\mu) - c.$$

If she does not learn, she takes action $a = 1$ if and only if $\mu \geq 1/2$, obtaining expected utility

$$U^{NL}(\mu) = \max\{0, 2\mu - 1\}.$$

For sufficiently low cost c , there are two interim beliefs at which the receiver is indifferent between learning and not: $\underline{\mu} < 1/2$ such that upon $s = 1$, the receiver switches to action $a = 1$, but the expected marginal benefit is exactly c : $U^{NL}(\underline{\mu}|\underline{\mu} < 1/2) = U^L(\underline{\mu}|\underline{\mu} < 1/2, \gamma_1(\underline{\mu}) \geq 1/2)$; and $\bar{\mu} \geq 1/2$ such that upon $s = 0$, the receiver switches to action $a = 0$, but the expected marginal benefit is exactly c : $U^{NL}(\bar{\mu}|\bar{\mu} \geq 1/2) = U^L(\bar{\mu}|\bar{\mu} > 1/2, \gamma_0(\bar{\mu}) < 1/2)$. The first equation is $0 = (2\gamma_1(\underline{\mu}) - 1) \phi_1(\underline{\mu}) - c$ and the second equation is $2\bar{\mu} - 1 = (2\gamma_1(\bar{\mu}) - 1) \phi_1(\bar{\mu}) - c$, yielding

$$\underline{\mu} = 1 - p + c, \quad \bar{\mu} = p - c,$$

where $p - c \geq 1/2$ must hold.

Hence, if $c \leq p - 1/2$, there are two non-learning, $[0, \underline{\mu})$, $[\bar{\mu}, 1]$, and one learning, $[\underline{\mu}, \bar{\mu})$, regions. In contrast to the original model, a non-learning region need not be closed, as the sender-preferred equilibrium assumption puts the belief $\underline{\mu}$ to the learning region. If $c > p - 1/2$, the receiver never learns for any μ .

8.4.2 Sender's Maximization Problem: Solution

Suppose $c \leq p - 1/2$. A seller's conditional expected utility $\hat{v}(\mu)$ is

$$\hat{v}(\mu) = \begin{cases} 0 & 0 \leq \mu < \underline{\mu} \\ p\mu + (1-p)(1-\mu) & \underline{\mu} \leq \mu < \bar{\mu} \\ 1 & \bar{\mu} \leq \mu \leq 1 \end{cases}.$$

A sender's optimal strategy can be found by concavification $cav(\hat{v})(\mu)$ of $\hat{v}(\mu)$. Figure 7 depicts $\hat{v}(\mu)$ and an optimal sender's strategy when the precision of the receiver's signal is $p = 0.8$ and when (a) $c \rightarrow 0$ or (b) $c = 0.2$. Proposition 1 fails when $c \rightarrow 0$, since the (unique) sender's strategy targets learning of the receiver. Figure 8 then captures the player's expected equilibrium payoffs as a function of c when the precision of the signal is $p = 0.8$. Similarly to the application, the receiver's payoff is non-monotone in c , and in this example, the receiver prefers intermediate

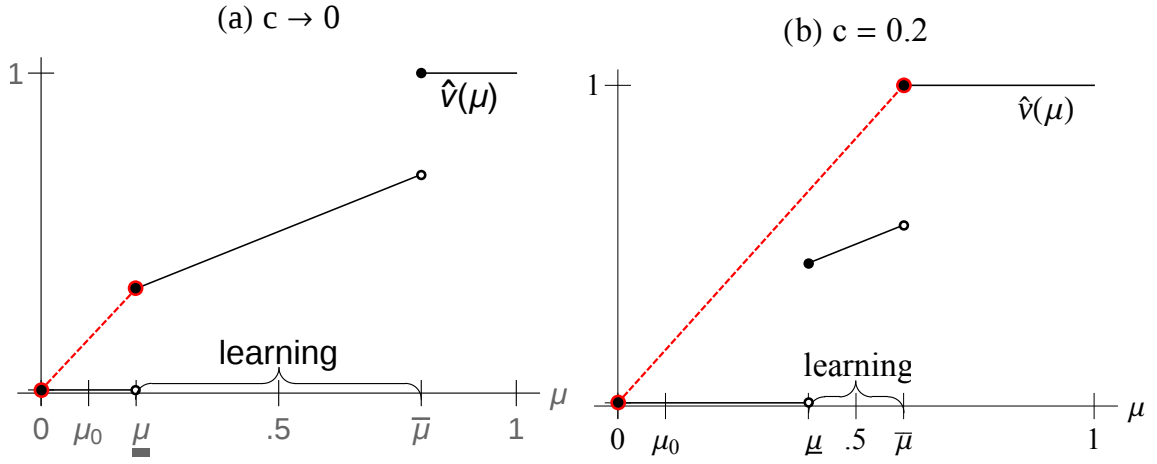


Figure 7: $\hat{v}(\mu)$ and the sender's optimal strategy with $p = 0.8$. The sender targets (a) learning (less information) and (b) no learning (more information).

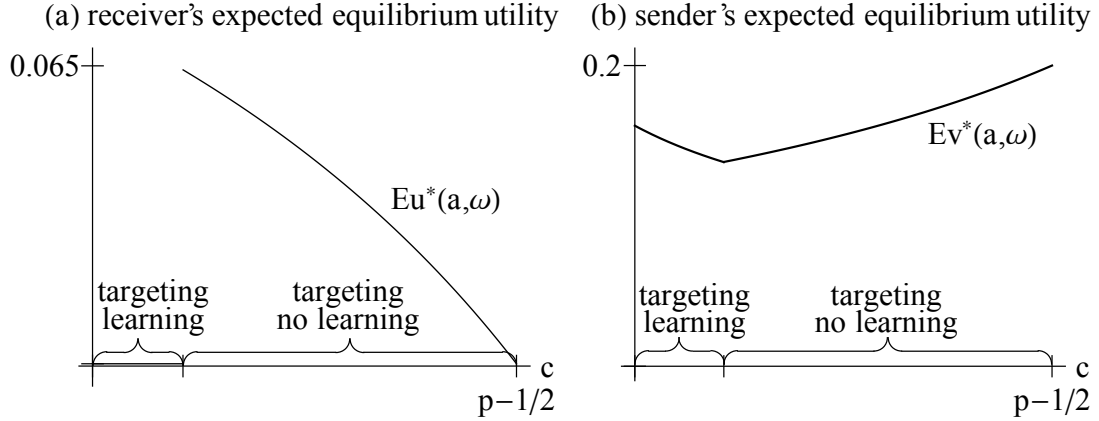


Figure 8: Equilibrium expected utilities as a function of c with $p = 0.8$. Both the sender's and the receiver's utilities are non-monotone in c .

to very low cost. In contrast to Proposition 3, however, the sender's payoff is also non-monotone in c : it is locally decreasing in c over the region in which the sender targets learning. The difference to our original model stems from the rigidity of receiver's information acquisition technology, which the sender takes advantage of. The sender optimally chooses between two types of strategies: providing enough information to prevent the receiver from additional learning, or providing less information and inducing her to learn the fixed amount of information that is available to her. The non-monotonicity results then stem from the interplay between these two strategies and how their optimality varies with changes in the cost c .