

Rational Misspecification: Framework and Applications*

Ran Eilat[†] Kfir Eliaz[‡]

October 22, 2025

Abstract

This paper proposes a framework for assessing whether misspecified decision makers would be willing to pay for information that can potentially make them less misspecified. We introduce a prior-free approach, based on “constrained” maximal regret, to derive an upper bound on the subjective assessment of potential gains from acquiring a more accurate model. The constraint stems from the information currently available to the decision maker. We apply our approach to three prominent models of misspecified beliefs: coarse expectations, causal misperceptions and sampling equilibria.

Keywords: Misspecified beliefs; regret; value of information

JEL Classification Codes: D83; D84; D90

The literature on non-rational expectations, where decision-makers have misspecified beliefs about the steady-state mapping from their actions to consequences, offers several ways to model misspecifications and a variety of solution concepts for analyzing single-person and interactive decision-making under the various misspecifications. Notable examples include analogy-based expectation equilibrium ([Jehiel, 2005](#); [Jehiel, 2022](#)), sampling ([Osborne and Rubinstein, 1998](#); [Salant and Cherry, 2020](#)), causal misperceptions ([Spiegler, 2016](#); [Spiegler, 2020](#)), cursed equilibrium ([Eyster and Rabin, 2005](#); [Cohen and Li, 2023](#)), Berk-Nash ([Esponda and Pouzo, 2016](#)), and behavioral equilibrium ([Esponda, 2008](#)).

*We thank Sarah Auster, Geoffroy de Clippel, Ignacio Esponda, Kevin He, Philippe Jehiel, Margaret Meyer, Ran Spiegler, Jakub Steiner and Heidi Thysen for valuable comments. Financial support from BSF grant #2022294 is gratefully acknowledged.

[†]School of Economics, Tel Aviv University, eilatr@tauex.tau.ac.il

[‡]School of Economics, Tel Aviv University, kfire@tauex.tau.ac.il

A common feature of many models in this literature is that they typically take as given the decision-makers’ particular form of misspecification and its degree. For instance, in sampling equilibrium ([Osborne and Rubinstein, 1998](#)), the decision-maker may form his beliefs about the mapping from actions to consequences based on one, two, or any other number of samples of outcomes resulting from his past actions. Or, in analogy-based expectation equilibrium ([Jehiel, 2005](#)), a player may group her opponents’ types into one, two, or any other number of “analogy classes,” and form beliefs according to the average behavior observed within each class. However, in these cases and others, a natural question arises: how does the decision-maker arrive at these particular degrees of misspecifications to begin with? And, in particular, if a decision-maker suspects that his model of the environment is not perfectly accurate, why does he not attempt to reduce his error by acquiring more knowledge?

One possible answer is that more data is simply not available – if there exists only a single sample for the past consequence of each of the decision maker’s actions, then he must make a decision based solely on this information. Another answer might be that the decision-maker is just completely unaware of his misspecification. While this may be true in certain cases, oftentimes reality is more nuanced. Indeed, in many situations decision-makers are aware that they do not *fully* understand the relationship between actions and consequences in the environment they operate in, yet they still do not engage in improving their model, even when data can be collected.

A possible explanation for this behavior is that acquiring more data may be costly. In this case, if the costs exceed the benefits, it is a “rational” choice for the decision-maker to remain misspecified. However, this raises another conceptual question: how can a misspecified decision-maker compute the benefit of becoming less misspecified?¹ Our goal in this paper is to propose a framework for rational misspecification and to evaluate its implications in several settings.

To illustrate the challenge, consider a two-player Bayesian game between Alice and Bob, in which Bob forms expectations about Alice’s behavior using coarse data on her strategy, as in analogy-based expectation equilibrium ([Jehiel, 2005](#)). Specifically, suppose that although Alice’s strategy depends on her type, which is relevant for Bob’s payoff, Bob knows only the overall distribution of Alice’s

¹In the context of analogy-based expectation equilibrium, [Jehiel \(2022\)](#) highlighted the difficulty of employing a cost-benefit analysis to endogenize a decision-maker’s analogy partition, stating that “it is not clear how players would have the correct understanding about how their choices of analogy partitions translate into true payoffs.”

actions, irrespective of her type. Now, suppose Bob is given the opportunity to refine his data by learning the distribution of Alice’s actions conditional on her type belonging to each cell of some partition of her types. Clearly, this additional knowledge can only improve Bob’s decision-making. But how can he quantify the extent of this improvement?

Although the basic idea that individuals weigh costs and benefits when deciding whether to become more informed appears also in the literature on costly information acquisition, there is an inherent conceptual difference between the two problems. The distinction lies in the difficulty of the misspecified decision-maker to quantify the benefits of acquiring additional data. For example, in the rational inattention literature (e.g., [Sims \(2003\)](#) and [Maćkowiak et al. \(2023\)](#)) an uninformed decision-maker has to decide which information structure to acquire. To evaluate the expected benefit of any given information structure, the decision-maker relies on one of the framework’s primitives – the true prior distribution over the states – to form beliefs about the possible consequences of learning. In contrast, a misspecified decision-maker operates with an erroneous model of the steady-state and needs to form beliefs about the expected gain from employing a more precise model. Short of exogenously imposing some arbitrary prior beliefs on the set of correct models, there is no primitive of the environment to guide the decision-maker in forming these beliefs.

In the absence of an objective prior beliefs on the set of models, there are myriad ways to form beliefs about what one might learn from acquiring more data. In this paper we propose an *upper bound* on the subjective assessment of the expected gain from learning a more accurate model. This upper bound is computed by finding a new stochastic mapping from actions to consequences that satisfies the following properties: (i) it is consistent with the partial but correct information on the true mapping derived from the decision-maker’s current misspecified model, and (ii) it maximizes the difference between the expected payoff the decision-maker could achieve from the more accurate model and the expected payoff he would obtain if he remained with the action he planned to take given his current model, with the expectation taken with respect to the new stochastic mapping. We refer to this difference as the *maximal regret* from not acquiring the more accurate model, and interpret it as the decision-maker’s maximal willingness-to-pay for reducing his misspecification. Thus, when the cost of reducing the misspecification exceeds this maximal willingness-to-pay, the decision-maker can be said to be rationally misspecified.

Our approach is motivated by the economic literature on decision-making without priors (in particular, [Bergemann and Schlag \(2008\)](#) and [Bergemann and Schlag \(2011\)](#)) and by the phenomenon of “fear of missing out” or FOMO (see, e.g., [Milyavskaya et al. \(2018\)](#) and [Laurence and Temple \(2023\)](#)). The idea is that when a decision-maker is presented with the opportunity to acquire new knowledge that could make him better off, he considers what he might be giving up if he forgoes that opportunity. Specifically, the new knowledge might prompt him to take a different action and obtain a significantly higher payoff; it might also make him realize that the action he was planning to take without the new knowledge would result in a low payoff. The greater the difference between these two potential payoffs, the more valuable the new knowledge becomes. The worst-case, in terms of what the decision-maker would lose by not acquiring the new knowledge, is represented by the maximal value this difference in payoffs can take.

When assessing his potential regret, the decision-maker does not consider the entire set of (stochastic) mappings from actions to consequences. Instead, he focuses only on those mappings that are *consistent* with the partial information he already possesses. To illustrate this, consider the example of Alice and Bob presented above. Suppose that, at the outset, Bob knows that the steady-state distribution of Alice’s actions is uniform when her type lies in the interval $[0, 1]$. Now, suppose Bob is contemplating learning the conditional distributions of Alice’s actions when her type lies in the interval $[0, 0.5]$ and when it lies in the interval $[0.5, 1]$. In this case, when calculating his maximal regret, Bob considers only those conditional distributions that are consistent with the overall distribution of Alice’s actions remaining uniform on $[0, 1]$.

As a different example, consider a decision-maker who only knows the correlation between two pairs of variables, (x, y) and (y, z) . If this decision-maker can learn the true joint distribution over all three variables, consistency requires that this distribution must align with the pairwise correlations he already knows. Thus, a decision-maker’s maximal regret from not reducing his misspecification is constrained by this consistency requirement.

Our upper bound on the decision-maker’s willingness-to-pay for information is “conservative/permissive” in the sense that it is computed with respect to all the possible mappings from actions to consequences (“models”) that are relevant for the new data and are consistent with his current knowledge. We impose only minimal assumptions on how the decision maker aggregates the regret associated with not adopting each of the possible models. Thus, various ways to aggregate

the regret may lead to different levels of actual willingness-to-pay, but none will exceed our upper bound. Hence, if the cost of obtaining new information exceeds this bound, we can be certain that the decision-maker will not be willing to incur it, *no matter how he aggregates the potential regret associated with not adopting each of the possible models*. Despite this conservative approach, we will show that in some environments, the upper bound may in fact be zero.

We demonstrate the framework’s portability by applying it to a range of belief misspecification models: *coarse expectations* (as captured by [Jehiel \(2005\)](#) notion of Analogy-Based-Expectations-Equilibrium or ABEE), *causal misperceptions* (as captured by [Spiegler \(2016\)](#)’s Bayesian networks framework) and *sampling* (as captured by [Osborne and Rubinstein \(1998\)](#)’s notion of Sk equilibrium). We provide a detailed explanation of each model in the corresponding section below.

The remainder of the paper is organized as follows. Related literature is discussed immediately below. Section 1 formally introduces our approach. The next three sections analyze the three applications of our approach: coarse expectations in Section 2, causal misperceptions in Section 3 and sampling in Section 4.

Related literature. A number of alternative approaches have been proposed to endogenize decision-makers’ misspecified beliefs. [Gonçalves \(2023\)](#) considers normal form games where each player is endowed with some exogenous prior over the other players’ mixed strategies and decides sequentially whether to sample costly signals about these strategies. [Heller and Winter \(2020\)](#) study misspecified beliefs that constitute a fixed point: players best respond to their misspecified beliefs and these misspecified beliefs are best responses to each other. [He and Libgober \(2023\)](#) propose an evolutionary approach to define a notion of “stable misspecifications”. In the context of ABEE, [Jehiel and Weber \(2024\)](#) endogenize the composition of analogy partitions, by requiring them to satisfy a property that can be interpreted as minimizing prediction errors. Finally, there is a literature that takes a learning approach to justifying persistent misspecification. Notable examples include [Cho and Kasa \(2015\)](#) and more recently, [Ba \(2024\)](#).

A different approach to endogenizing misspecified beliefs is to consider an interested third party that strategically provides a decision-maker with a (possibly misspecified) model of the steady-state in order to persuade him to choose a particular action. Some recent examples include [Eliaz and Spiegler \(2020\)](#); [Eliaz et al. \(2021c,a,b\)](#); [Schwartzstein and Sunderam \(2021\)](#) and [Aina \(2024\)](#).

The problem of evaluating the impact of new information in the absence of ob-

jective priors naturally comes up in decision-making under ambiguity. Li (2020) assumes that the decision-maker uses the *same* model of ambiguity-aversion (e.g., max-min expected utility) to evaluate *both* his expected utility given the new information, *and* to assess his uncertainty about which information will realize. This approach is then shown to sometimes generate *negative* value of information.² In contrast, models of belief-misspecifications, which are the focus of this paper, do not provide guidance on how the decision-maker may evaluate information that reduces his misspecification. Consequently, the approach proposed by Li (2020) cannot be applied. This is where our framework, which bounds the willingness to pay for information with maximal regret, proves useful.

1 Framework

We present the framework in four steps. First, we define the objective (or “true”) environment in which the player operates. This environment is known to the modeler but not to the player. Next, we introduce the concept of a misspecified “type” and explain how a player’s type affects his decisions. We then define a player’s regret from *not* adopting an alternative model. Based on this, we derive an upper bound on the player’s willingness to pay for data that can lead him to adopt a new model, which is unknown at the time of acquiring the data. Finally, we say that a player rationally decides to remain misspecified if, among all models consistent with his type, there is no model for which the player’s willingness to pay exceeds the cost.

The objective environment. A player has to choose an action from a compact set A . Each action is stochastically mapped to a consequence from a set Y via a function $g : A \rightarrow \Delta(Y)$.³ The stochastic nature of this mapping could be due to an unknown state of nature, or because the consequence also depends on an unknown action by another player. We refer to g as the *true model of the environment*. The player’s preferences are defined over $A \times Y$ and are represented by a bounded and continuous utility function $u : A \times Y \rightarrow \mathbb{R}$.

Misspecified models. The player does not know g . Instead, he works according

²Kops and Pasichnichenko (2023) and Shishkin and Ortoleva (2023) experimentally study the relationship between ambiguity-aversion and negative value of information, finding mixed evidence.

³We assume that Y is a Polish space, and denote by $\Delta(Y)$ the set of probability distributions over Y endowed with the weak* topology. The function g is assumed to be both measurable and continuous.

to a (potentially) misspecified model, which we represent by his *type*. Let Θ denote the set of types. Each type $\theta \in \Theta$ possesses a subjective model $g_\theta : A \rightarrow \Delta(Y)$ from actions to consequences which guides his choice of actions.⁴ A model g_θ is considered *misspecified* if it differs from g . We assume that θ encapsulates all relevant information the player has about the environment. Thus, the optimal action for a player of type θ , denoted by a_θ , is given by:⁵

$$a_\theta = \arg \max_{a \in A} \int_{y \in Y} u(a, y) dg_\theta(y | a) \quad (1)$$

where $g_\theta(\cdot | a)$ is the probability measure over consequences generated by $g_\theta(a)$.

For example, consider a player of type θ who only knows, for each of his actions, the feasible outcomes, and a single observation of an outcome from a previous instance when the action was taken. A possible specification of g_θ is that each action leads with certainty to the outcome that was observed for that action.

Regret from *not* adopting a new model. Consider a player of type θ who possesses a subjective model g_θ that guides his choice of actions. Suppose this player encounters an alternative model $\tilde{g} : A \rightarrow \Delta(Y)$. The player is uncertain whether $\tilde{g}$ is the correct model, yet recognizes that ignoring it could result in regret.

We define *regret* as the difference between the expected payoff from the optimal action under the new model $\tilde{g}$, and the expected payoff from the original optimal action a_θ , with expectations about the relationship between actions and consequences evaluated according to $\tilde{g}$. Formally, the regret experienced by a player of type θ from continuing to operate under g_θ instead of adopting $\tilde{g}$ is given by:⁶

$$R_\theta(\tilde{g}) = \max_{a \in A} \int_{y \in Y} u(a, y) d\tilde{g}(y | a) - \int_{y \in Y} u(a_\theta, y) d\tilde{g}(y | a_\theta). \quad (\text{Regret})$$

where $\tilde{g}(\cdot | a)$ is the probability measure over consequences in Y that is generated by $\tilde{g}(a)$.

This approach to quantifying the player's regret is inspired by the common phenomenon of FOMO. That is, a player is concerned that if he were to pass on

⁴We assume that the function g_θ is both measurable and continuous.

⁵Existence of a maximum in Eq. (1) follows from the compactness of A and the continuity of u and g_θ . For simplicity, we assume that this maximizer is unique. In Section 3 we relax this assumption.

⁶When there is more than one optimal action according to g_θ , the regret may also depend on which of these actions is chosen. We demonstrate this in Section 3.

the opportunity to act according to a new model, he would not only miss out the chance to earn a significantly high payoff, but also could realize that his current action is truly suboptimal.

Rational misspecification. Suppose a player of type θ is offered the opportunity to pay a cost c to obtain *data* about the environment that would lead him to revise his model. We represent such data by a set of potential new models G_θ that are all *consistent* with the information encoded in the player’s type θ . The notion of consistency is context dependent and will be defined precisely for each of the applications in the subsequent sections. For now, we keep the definition abstract and illustrate it with the following example.

Consider the player with type θ described above. Suppose θ is offered access to new data that provides an additional observation of a past outcome for each action. This new data may lead type θ to adopt a new model. Crucially, the set of possible new models, namely G_θ , includes only those consistent with the original information that type θ has. Specifically, for each action, models in G_θ must assign a positive probability to the outcome that was originally observed.

Before acquiring the data, the decision maker does not know what his revised model will be. We assume that his willingness to pay for the data is determined by the regret he would experience from not adopting any of the models in G_θ . However, the precise way in which he aggregates these regret levels to form his valuation is unknown to the modeler. We impose only a weak assumption: The overall level of regret from not acquiring data that could lead to a model in G_θ cannot exceed the regret associated with not adopting *any* single model in G_θ .^{7,8} Hence, from the modeler’s perspective, an upper bound on the decision maker’s

⁷One family of aggregation rules consistent with this assumption is the following. Suppose that, given a set of consistent models G_θ , the decision maker aggregates the regret of not adopting any of them using a function $\rho(\{R_\theta(\tilde{g})\}_{\tilde{g} \in G_\theta})$, where $R_\theta(\tilde{g})$ denotes the regret from not adopting model $\tilde{g}$, as defined in (Regret). Assume that ρ is monotone – that is, (weakly) increasing in each argument $R_\theta(\tilde{g})$ – and satisfies the condition that if $R_\theta(\tilde{g}) = r$ for all $\tilde{g} \in G_\theta$ and some constant $r > 0$, then $\rho(G_\theta) = r$. Although the analyst does not observe the function ρ , she can conclude that $\rho(G_\theta) \leq \sup_{\tilde{g} \in G_\theta} R_\theta(\tilde{g})$.

⁸Our regret-based approach is rooted in the idea that the decision-maker is triggered to think about the models in G_θ only when confronted with the opportunity to acquire knowledge. We do not make any assumption on whether such opportunity may also trigger the decision-maker to change the action he planned on choosing absent any information in order to minimize his anticipated regret. Since this may only reduce his actual willingness-to-pay, he would still reject the offer if the cost is above our upper bound. In this sense, our upper bound may be “too permissive”.

regret from not adopting any model in G_θ is given by:

$$R^*(G_\theta) = \sup_{\tilde{g} \in G_\theta} R_\theta(\tilde{g}). \quad (\text{WTP})$$

We interpret $R^*(G_\theta)$ as an upper bound on what a player will pay for the data given by G_θ . To streamline the exposition, we will henceforth simply refer to $R^*(G_\theta)$ as the player's *willingness-to-pay (WTP)* for G_θ . Thus, if $c > R_\theta^*$ the player would pass on the opportunity to acquire the data and learn about a new model from actions to consequences. In this case, we say that the player is *rationally misspecified*.

Remark. Generally speaking, a player's type encodes many things: the information the agent possesses, how he completes missing information, his belief about how outcomes are determined, etc. However, only two components of a type are payoff relevant: (i) the mapping from actions to consequences, that is captured by g_θ , and (ii) the set of mappings from actions to consequences that the player considers possible when learning new information, which is captured by G_θ . Therefore, when defining types in the applications below we will focus on defining these two items, with the understanding that the full description of the type may potentially include other (non-payoff relevant) details as well.

In the following sections, we apply this framework to three forms of misspecifications that have been analyzed in the literature: coarse expectations, sampling, and causal misperceptions. For each case, we explain how it maps to the primitives defined in this section.

2 Coarse expectations

We apply our framework to misspecifications arising from coarse beliefs about the mapping from actions to consequences, based on coarse data on the equilibrium joint distribution over the action profile and states of nature. For instance, a player may only have access to the marginal distribution over another player's action and the marginal distribution over states, without data about the joint distribution. In this case, the player extrapolates and fills in the missing data to form a subjective belief about the joint distribution over actions and states.

The literature offers several approaches to how a player extrapolates from his coarse data. We adopt the Analogy Based Expectations Equilibrium (ABEE) ap-

proach originally proposed in [Jehiel \(2005\)](#) and later extended to Bayesian games in [Jehiel and Koessler \(2008\)](#). Rather than presenting this framework in its full generality, we apply it to a classic adverse selection setting. This setting was analyzed under various approaches to coarse expectations: “cursed equilibrium” in [Eyster and Rabin \(2005\)](#), “behavioral equilibrium” in [Esponda \(2008\)](#) and ABEE in [Jehiel and Koessler \(2008\)](#) and [Spiegler \(2011\)](#).

The setting. A seller owns an object with a privately observed quality ϕ , which is distributed according to a distribution F supported on $[0, 1]$ with density f . We assume that F is such that $\phi + F(\phi)/f(\phi)$ is increasing in ϕ , which means the seller’s “virtual value” is increasing in his type.

Trade occurs through a double auction protocol: The seller submits an ask price $x \in [0, 1]$ and the buyer submits a bid price $p \in [0, 1]$. A trade takes place at price p if $p \geq x$. The value of the object for the buyer is $v(\phi, b)$, where $v(\cdot, \cdot)$ is increasing in both parameters, and $v(\phi, b) \geq \phi$ for all $\phi \in [0, 1]$, ensuring there are always gains from trade. The parameter $b \in \mathbb{R}^+$ captures the gains from trade. The seller’s payoff is 0 if there is no trade, and $p - \phi$ otherwise. The buyer’s payoff is 0 if there is no trade, and $v(\phi, b) - p$, otherwise.

For expositional purposes, it is convenient to have a stark rational expectations benchmark in which the the market collapses due to adverse selection. We therefore make the following assumption:⁹

$$\mathbb{E}_{\phi \sim F} [v(\phi, b) | \phi < p] < p \quad \forall p \in (0, 1]. \quad (2)$$

In words, Eq. (2) states that, for any price p , the buyer’s expected value from trading with a seller who agrees to sell at price p is less than p .

For a seller with quality ϕ , submitting an ask price $x = \phi$ is a dominant strategy. Thus, in what follows, we focus on the buyer’s problem.

ABEE. We follow [Spiegler \(2011\)](#) in describing the notion of ABEE in the present context. Let $\sigma : [0, 1] \rightarrow \Delta([0, 1])$ be a seller’s strategy that maps each quality to a distribution over prices. As noted above, given the trading rule, the seller will use the deterministic mapping σ^* which equates the ask to the quality. A buyer with rational expectations would choose a bid that maximizes his expected payoff given a perfect perception of σ^* . In contrast, a misspecified buyer will maximize his expected payoff with respect to a coarse representation of σ^* . This representation

⁹For a simple example that satisfies this condition, see Example 1 below.

takes the following form. The buyer is endowed with an *analogy partition* $\mathcal{C} = (C_1, \dots, C_K)$ of $[0, 1]$, where each cell C_k is an interval. Let $C(\phi)$ denote the cell containing the seller's quality ϕ . The buyer's coarse representation of σ^* is a mixed strategy $\sigma^{\mathcal{C}}$ such that for every seller quality ϕ , the strategy $\sigma^{\mathcal{C}}$ mimics the price distribution in the entire cell $C(\phi)$ in the sense that for all $x \in [0, 1]$:¹⁰

$$\Pr[\sigma^{\mathcal{C}}(\phi) \leq x] = \Pr[\sigma^*(t) \leq x \mid t \in C(\phi)]$$

Next, we situate the above setting within the framework presented in Section 1. Readers who are not interested in the precise details of this mapping may skip ahead to Section 2.1 without loss of continuity.

Mapping the setting to our framework. To cast the setting within our framework, we proceed in two steps. First, we describe how a misspecified buyer computes the distribution over the possible consequences of submitting a bid (i.e. a model $g : A \rightarrow \Delta(Y)$). This computation is performed given that the buyer is misspecified in the ABEE sense and has certain beliefs about the marginal distributions over the set of seller's qualities and the marginal distribution of asks based on his analogy partition. Next, we explain how these marginal distributions are determined for the misspecified buyer, both when he possesses the initial analogy partition and when he considers which models are consistent with new data he might acquire.

Fix an analogy partition $\mathcal{C} = (C_1, \dots, C_K)$ of $[0, 1]$. Suppose the buyer believes that the marginal distribution over seller's quality is given by H_ϕ , which admits density h_ϕ . The buyer also believes that the marginal distribution over asks, conditional on the seller's quality being in the cell C_k , is given by H_x^k . These marginals can come either from the buyer's initial misspecified model, or these could be what he "imagines" the marginals would be when he is offered the opportunity to refine his original partition to $\mathcal{C}$.

For example, suppose the seller's quality is uniformly distributed over $[0, 1]$ and he plays his dominant strategy. If the buyer is initially fully coarse, i.e. his analogy partition $\mathcal{C} = (C_1 = [0, 1])$ has a single cell, then $H_\phi = H_x^1 =$

¹⁰As [Spiegler \(2011\)](#) explains, a possible interpretation of this notion of misspecification is the following: "when the buyer enters the market, he has access to records of all the ask prices that were previously submitted by the seller (or his previous incarnations), but he does not have access to the records of the valuations that lay behind these ask prices. Following the "Occam's razor" principle, the buyer adopts the simplest theory that is consistent with the historical records, where simplicity here means that the theory is not allowed to depend on unobservable variables as long as it is consistent with the data."

$U[0, 1]$. Alternatively, if the buyer initially has the analogy partition $\mathcal{C} = (C_1 = [0, 1/2], C_2 = [1/2, 1])$, then $H_\phi = U[0, 1]$, $H_x^1 = U[0, 1/2]$ and $H_x^2 = U[1/2, 1]$.

Trade occurs at the bid price, whenever it is higher than the ask. Thus, each bid price p induces a probability distribution over the set of consequences $Y = \{\emptyset\} \cup [0, 1]$, where the outcome $\{\emptyset\}$ is interpreted as “no-trade” and any outcome $\phi \in [0, 1]$ is interpreted as trade with a seller of quality ϕ . Therefore, for any bid price p , the buyer computes the (ex-ante) probability for no trade as follows:

$$\Pr(\{\emptyset\} \mid p) = \sum_{k=1}^K \Pr(\phi \in C_k) \Pr(x > p \mid \phi \in C_k) = \sum_{k=1}^K \left(H_\phi(\overline{C}_k) - H_\phi(\underline{C}_k) \right) \cdot \left(1 - H_x^k(p) \right), \quad (3)$$

where $\overline{C}_k$ and $\underline{C}_k$ denote the upper and lower boundaries of the cell C_k , respectively. For each k in the sum on the right-hand side of Eq. (3), the first multiplier is the probability that the seller’s quality ϕ is in the cell C_k , and the second multiplier is the probability that the ask is greater than the bid, conditional on the seller’s quality being in the cell C_k .

For each bid price p , the buyer can also compute the probability of trade with any set of qualities $\Phi \subseteq [0, 1]$. Under ABEE, his misspecification leads him to compute this probability using the marginals as follows:

$$\begin{aligned} \Pr(\text{trade with sellers in } \Phi \mid p) &= \sum_{k=1}^K \Pr(\phi \in C_k) \cdot \Pr(x < p \mid \phi \in C_k) \cdot \Pr(\phi \in \Phi \mid \phi \in C_k) \\ &= \sum_{k=1}^K \left(H_x^k(p) \cdot \int_{\phi \in \Phi \cap C_k} h_\phi(z) dz \right). \end{aligned} \quad (4)$$

For each k in the sum on the right-hand side of the first line of Eq. (4), the first multiplier denotes the probability that the seller’s quality ϕ lies in cell C_k ; the second multiplier is the probability that the ask is below the bid—so that trade occurs—conditional on the seller’s quality being in C_k ; and the third multiplier is the probability that the seller’s quality lies in the set Φ , conditional on being in C_k . Note that the player’s misspecification is reflected in his use of $\Pr(\phi \in \Phi \mid \phi \in C_k)$ instead of $\Pr(\phi \in \Phi \mid \phi \in C_k, x < p)$, as he fails to account for the dependence between the seller’s ask price and quality.

A buyer’s type is a partition of the interval $[0, 1]$, and represents the buyer’s analogy partition at the outset, before potentially acquiring new data. For simplicity, we restrict our attention to partitions with countably many elements. A type $\theta = \mathcal{C}^\theta = (C_1^\theta, \dots, C_K^\theta)$ correctly perceives the marginal distribution of the seller’s quality, and the marginal distribution of the seller’s asks, conditional on

the seller's quality being within any of the cells. Consequently, the buyer's type θ determines the marginal distributions $(H_\phi^\theta, H_x^{\theta,1}, \dots, H_x^{\theta,K})$ in which the buyer believes as follows:

$$H_\phi^\theta(\phi) = F(\phi) \quad \forall \phi \in [0, 1] \quad \text{and} \quad (5)$$

$$H_x^{\theta,k}(x) = \frac{F(x) - F(\underline{C}_k^\theta)}{F(\overline{C}_k^\theta) - F(\underline{C}_k^\theta)} \quad \forall k \in \{1, \dots, K\} \text{ and } \forall x \in C_k^\theta. \quad (6)$$

Thus, the misspecified model g_θ of type θ is determined by Eqs. (3) and (4), which are computed based on the marginal distributions in Eqs. (5) and (6).¹¹

Now, suppose the buyer is offered the opportunity to get access to a new partition $C = (C_1, \dots, C_M)$, which is a refinement of $\mathcal{C}$. The set of mappings that are feasible for θ (the set G_θ) includes all the mappings that are induced according to the marginals $(H_\phi, H_x^1, \dots, H_x^M)$, which satisfy the following:

$$H_\phi(\phi) = H_\phi^\theta(\phi) \quad \forall \phi \in [0, 1] \quad \text{and} \quad (7)$$

$$H_x^{\theta,k}(x) = \sum_{\ell: C_\ell \subseteq C_k^\theta} \left(F(\overline{C}_\ell) - F(\underline{C}_\ell) \right) H_x^\ell(x) \quad \forall k \in \{1, \dots, K\} \text{ and } \forall x \in C_k^\theta. \quad (8)$$

In words, Eq. (7) states that the marginal over the seller's quality is consistent with what the buyer's knowledge prior to obtaining the new partition. Equation (8) guarantees that for a cell that was refined, the new marginals over asks aggregate to the coarser marginal that the buyer started with. Note these constraints still leave the buyer with substantial freedom in imagining what the marginal over the asks may be in the new (refined) partition. In Section 2.2 we illustrate how to operationalize these constraints.

2.1 Two polar benchmarks

To illustrate the impact of the buyer's misspecification we begin by comparing the case of a correctly specified buyer with the case of a fully coarse one.

Rational expectations. Under rational expectations, the buyer has correct beliefs about the joint distribution of the object's quality and the seller's ask (i.e., he knows they are perfectly correlated, because it is a dominant strategy for the

¹¹Our analysis depends on the information encoded in the buyer's type θ , i.e. the marginals on the quality and on the asks as specified by Eqs. (5) and (6). In principle, these could be generated by a different market setting than the double auction we described above. That is, our analysis would continue to hold for any joint distribution as long as it induces these marginals.

seller to submit an ask that is equal to the quality). Hence, the buyer's problem is given by:

$$\max_p F(p) \cdot (\mathbb{E}_{\phi \sim F} [v(\phi, b) \mid \phi < p] - p).$$

By our assumption in Eq. (2), the optimal solution is obtained at $p = 0$. Thus, there is no trade in equilibrium.

Full coarseness. Next, consider a buyer whose analogy partition consists of a single cell, i.e., $\theta = (C_1^\theta = [0, 1])$. This buyer only knows the marginal distribution over the seller's quality and the overall marginal distribution over the seller's ask. We refer to this buyer as being *fully coarse*. According to Eqs. (5) and (6), we have $H_\phi^\theta = H_x^{\theta,1} = F$.

The problem of a fully coarse buyer is given by: $\max_p F(p) \cdot (\mathbb{E}_{\phi \sim F} [v(\phi, b)] - p)$. The optimal price satisfies:

$$\mathbb{E}_{\phi \sim F} [v(\phi, b)] = p + \frac{F(p)}{f(p)}. \quad (9)$$

There exists a unique price that solves this equation. Denote this solution by p_0 .

2.2 The willingness to pay of a fully coarse buyer

Suppose that a fully coarse buyer has the opportunity to refine his data by paying a fee to add a cell to his partition. Specifically, the buyer can acquire the analogy partition $(C_1 = [0, t], C_2 = [t, 1])$ for some $t \in (0, 1)$. This refinement allows the buyer to learn the marginal distribution of asks when the seller's quality is in $[0, t]$, denoted H_x^1 , and the marginal over seller's ask when the quality is in $[t, 1]$, denoted H_x^2 .

Before obtaining the new data, the buyer does not know what the marginal distributions H_x^1 and H_x^2 might be. However, since these distributions must be consistent with the data he already possesses (Eqs. (6) and (8)), he knows that:

$$F(t) \cdot H_x^1(x) + (1 - F(t)) \cdot H_x^2(x) = F(x) \quad \forall x \quad (10)$$

Denote by $W(p, H_x^1, H_x^2)$ the buyer's (misspecified) expected payoff from a bid p when the marginals over the seller's ask are given by H_x^1 and H_x^2 . We can then

compute this expected payoff as follows (see Eqs. (4)-(7) above):

$$\begin{aligned}
W(p, H_x^1, H_x^2) &= \Pr(\text{trade with sellers in } C_1 \mid \text{bid price is } p) \cdot \mathbb{E}_{\phi \sim H\phi} [v(\phi, b) - p \mid \phi \in C_1] \\
&\quad + \Pr(\text{trade with sellers in } C_2 \mid \text{bid price is } p) \cdot \mathbb{E}_{\phi \sim H\phi} [v(\phi, b) - p \mid \phi \in C_2] \\
&= F(t) \cdot H_x^1(p) \cdot (V_1 - p) + (1 - F(t)) \cdot H_x^2(p) \cdot (V_2 - p), \tag{11}
\end{aligned}$$

where $V_k \equiv \mathbb{E}_{\phi \sim F}(v(\phi, b) \mid \phi \in C_k)$ denotes the expected value of $v(\phi, b)$, conditional on the seller's quality being in C_k . Note that the buyer's misspecification is reflected in his use of the expected values V_1 and V_2 , which do not condition on the event that trade occurs. Thus, the buyer fails to account for the dependence between the object's quality and the seller's ask.

The buyer's WTP is determined by solving the following maximization problem:

$$\begin{aligned}
&\max_{H_x^1, H_x^2} \left\{ \left(\max_p W(p, H_x^1, H_x^2) \right) - W(p_0, H_x^1, H_x^2) \right\} \tag{OBJ} \\
&\text{s.t.} \quad F(t) \cdot H_x^1(x) + (1 - F(t)) \cdot H_x^2(x) = F(x) \quad \forall x
\end{aligned}$$

The first component in the objective function, $W(p, H_x^1, H_x^2)$, represents the expected payoff the buyer can achieve with the new data, where the maximum reflects the optimal price choice based on this refined information. The second component, $W(p_0, H_x^1, H_x^2)$, captures the expected payoff for the buyer from adhering to the original bid, with the expectation evaluated using the refined data. The constraint follows from Eq. (10) above.

Our first result characterizes the marginals, H_x^1 and H_x^2 , and the new bid p_1 that solve the problem presented in hat (OBJ).

Proposition 1. *Let p_1 be the price that solves the buyer's WTP, as determined by (OBJ). Then exactly one of the following statements holds:*

- i. *The price p_1 satisfies $V_1 = p_1 + F(p_1)/f(p_1)$, provided that $F(p_0) - F(p_1) \leq F(t)$. Otherwise, it satisfies $F(p_0) - F(p_1) = F(t)$.*
- ii. *The price p_1 satisfies $V_2 = p_1 + F(p_1)/f(p_1)$, provided that $F(p_1) - F(p_0) \leq 1 - F(t)$. Otherwise, it satisfies $F(p_1) - F(p_0) = 1 - F(t)$.*

Furthermore, the marginal distributions H_x^1 and H_x^2 that solve the buyer's WTP, as determined by (OBJ), satisfy the following condition: If the price p_1 is determined by (i) above, then $H_x^2(p_1) = H_x^2(p_0)$; if the price p_1 is determined by (ii) above, then $H_x^1(p_1) = H_x^1(p_0)$.

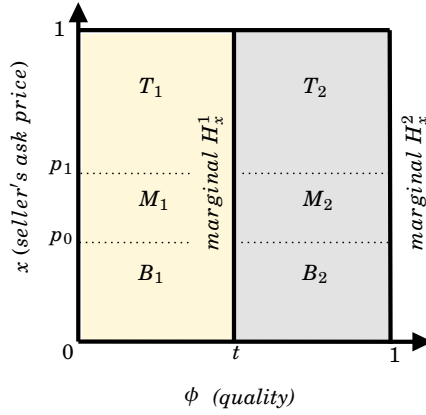


Figure 1: a joint probability distribution H over the space $[0, 1] \times [0, 1]$ of seller asks and qualities, partitioned into six regions under the assumption $p_1 > p_0$.

To interpret the conditions in Proposition 1, recall that a buyer faces two potential sources of regret from not refining his data. First, the new data might have enabled the buyer to place a more precise bid, potentially securing better quality. Second, the new data might have revealed that his original bid was too high relative to the actual quality. Proposition 1 states that the buyer experiences maximal regret when he focuses exclusively on one of these sources and envisions the worst-case scenario (in terms of regret) associated with that particular source.

To gain intuition for this result, note that the new price p_1 can be either above or below the original price p_0 . If $p_1 > p_0$, then p_1 will lead to trade with higher-quality sellers compared to p_0 if sellers who submit asks between the p_0 and p_1 are more likely to belong to the upper cell of the partition, i.e. $[t, 1]$. Consequently, by not refining his data, the buyer would miss an opportunity to trade with “good” sellers and secure a higher payoff. Indeed, $H_x^1(p_1) = H_x^1(p_0)$ indicates that the regret-maximizing distribution assigns zero probability to the event that a seller with an object of quality in the interval $[0, t]$ submits an ask between p_0 and p_1 . In other words, increasing the bid attracts only high-quality sellers.

On the other hand, if $p_1 < p_0$, then the original price p_0 was too high relative to the quality purchased. The loss from maintaining this price is accentuated if sellers who are willing to trade at p_0 , but not at p_1 , are more likely to belong to the lower cell of the partition, i.e. $[0, t]$. Indeed, $H_x^2(p_1) = H_x^2(p_0)$ implies that the regret-maximizing distribution assigns zero probability to the event that a seller with an object of quality in the interval $[t, 1]$ submits an ask between p_1 and p_0 .

The main idea behind the proof is to identify a *joint* distribution H over asks (x) and qualities (ϕ) that solves (OBJ). To provide a rough outline of the proof, we

focus on the case where the price p in the maximization problem is restricted to be greater than p_0 . In this case, the space $[0, 1] \times [0, 1]$ of asks and qualities – over which the joint distribution H is defined – can be partitioned into six regions, as illustrated in Figure 1.

With this partition of the space, we can express the objective function in (OBJ) in terms of the probability mass that the distribution H assigns to the regions M_1 and M_2 . The constraints on the marginal distributions pin down the total probability masses along each row and column of regions in Figure 1. This structure allows us to establish that, since $V_2 > V_1$, the distribution H that maximizes the buyer's WTP allocates as much probability as possible to the region M_2 , effectively shifting probability mass away from the region M_1 . As a result, the objective function simplifies significantly, and the solution follows directly from maximizing it.

The following example demonstrates how Proposition 1 can be operationalized:

Example 1: Suppose the seller's quality ϕ is uniformly distributed on the interval $[0, 1]$, i.e., $F(\phi) = \phi$. Additionally, assume that the $v(\phi, b) = \phi b$ for some $b \in (1, 2)$. Solving Eq. (9) shows that a coarse buyer sets the price $p_0 = b/4$. Now, suppose the buyer has the opportunity to refine his data according to the partition $([0, \frac{1}{2}], [\frac{1}{2}, 1])$. What is the buyer's WTP for this refined data?

Computation shows that $V_1 = b/4$ and $V_2 = 3b/4$. By Proposition 1, the bid p_1 that solves the problem presented in (OBJ) satisfies either $V_2 = 2p_1$ or $V_1 = 2p_1$. To determine the buyer's WTP for the refined data, we compare the two cases.

If $V_2 = 2p_1$, then $p_1 = 3b/8$. Indeed, since $F(3b/8) - F(b/4) = b/8 < \frac{1}{2} = 1 - F(\frac{1}{2})$, the price $p_1 = 3b/8$ is a candidate for maximizing the buyer's regret. From the proof of Proposition 1, we know that the buyer's regret from not acquiring the new data in this case is given by $F(p_1)(V_2 - p_1) - F(p_0)(V_2 - p_0)$, which simplifies to $b^2/64$.

If $V_1 = 2p_1$, then $p_1 = b/8$. Similarly, since $F(b/4) - F(b/8) = b/8 < \frac{1}{2} = 1 - F(\frac{1}{2})$, it follows that $p_1 = b/8$ is also a candidate for maximizing the buyer's regret. A similar computation shows that, in this case as well, the buyer's regret from not acquiring the refined data is $b^2/64$. Notably, the equality of regret in both cases is a consequence of this specific setup and does not hold in general.

Thus, the buyer's WTP for refining his data and obtaining a new partition is $b^2/64$. $\square$

Proposition 1 assumed that the refined partition available to the buyer was given. A natural question that arises is: which two-cell partition maximizes the buyer's WTP for data?

To address this question, we introduce the following notation. Given an analogy partition $([0, t], [t, 1])$ defined by a cutoff t , let $p_1^*(t)$ denote the price that solves the buyer's WTP to learn this analogy partition, as characterized in Proposition 1. We then have that:

Proposition 2. *The cutoff t^* that generates the partition $([0, t^*], [t^*, 1])$ for which the buyer's WTP is maximal satisfies*

- i. $F(p_0) - F(p_1^*(t^*)) = F(t^*)$ if $p_0 > p_1^*$, and
- ii. $F(p_1^*(t^*)) - F(p_0) = 1 - F(t^*)$ if $p_0 \leq p_1^*$ otherwise.

In terms of the explanation provided after Proposition 1, Proposition 2 implies that when $p_1^* > p_0$, the regret-maximizing distribution H assigns zero probability mass to regions M_1 , T_2 , and B_2 in Figure 1. This suggests that the buyer perceives the following: if he maintains the original bid p_0 , he will only receive quality below t , whereas raising his bid to p_1^* will allow him to trade with sellers offering quality above t . An analogous intuition applies for the case that $p_1^* < p_0$.

Thus, if the cost c of refining data and acquiring a new partition exceeds the buyer's WTP for the partition characterized in Proposition 2, the buyer will opt to retain his coarsest partition. That is, he will choose to remain *rationaly misspecified*. The next example illustrates this phenomenon in a simple setting. A noteworthy feature of this example is that even though the environment is symmetric, the partition associated with the highest WTP is highly skewed: the threshold t^* characterizing it lies above $\frac{2}{3}$.

Example 2. Consider the specification in Example 1. Under these parameters, the partition that maximizes the buyer's WTP, as characterized by Proposition 2, splits the interval $[0, 1]$ at $t^* = 4/(4 + b)$.¹² For brevity, we focus on the first partition in this example. Given the range of b , it follows that $t^* > \frac{2}{3}$.

To see how this conclusion is derived, consider the case where the solution satisfies (ii) in Proposition 1. From the proof of Proposition 1, the buyer's WTP is given by $F(p_1)(V_2 - p_1) - F(p_0)(V_2 - p_0)$. Recall that, under the parameters of this example, $p_0 = b/4$. Additionally, since $V_2 = (1 + t)b/2 = 2p_1$, it follows that

¹²An alternative partition with a threshold at $t^* = b/(b + 4)$ yields the same WTP, in which case $t^* < \frac{1}{3}$.

$p_1 = (1+t)b/4$. Applying Proposition 2, the threshold t^* that maximizes the buyer’s WTP satisfies

$$\frac{1}{4}(1+t^*)b - \frac{b}{4} = 1 - t^*$$

which implies to $t^* = 4/(4+b)$. The corresponding buyer’s WTP is $b^2/(b+4)^2$. Indeed, since $b < 2$, this value exceeds $b^2/64$, which was the buyer’s WTP for the partition $([0, \frac{1}{2}], [\frac{1}{2}, 1])$, as established in Example 1. The arguments for the case where the solution corresponds to case (i) in Proposition 1 are analogous.

Therefore, with the parameters of Example 1, if the cost c of acquiring the refined data exceeds $b^2/(b+4)^2$, the buyer will always prefer to remain fully coarse.

□

Extending the analysis to a buyer with an analogy partition containing n cells who considers paying for a refinement of that partition is not conceptually different but requires a more involved computation: there would be more ways to shift the admissible joint distribution in order to increase regret (a figure analogous to Figure 1 would have more regions to consider). However, the same type of reasoning would apply: regret would be maximized by either making it more likely that the new price attracts higher-quality sellers, or that the original price attracted very low-quality sellers.

3 Causal misperceptions

In this section we apply our approach to misspecifications that stem from causal misperceptions, where individuals misinterpret spurious correlations between variables as causal relations. To model such misperceptions, we adopt the framework introduced by Spiegler (2016), which employs tools from the Bayesian networks literature to represent individuals’ subjective causal perceptions using directed acyclic graphs (DAGs).

We consider agents who hold a misspecified causal model but have the option to pay a cost to learn the true model. Using our approach, we derive the agents’ willingness to pay for acquiring this knowledge. We then use this setting to illustrate that in a market for information, misspecified agents – who stand to benefit the most from learning the true model – may be excluded from access to information.

The setting. We demonstrate our approach in the context of the “dieter’s dilemma” analyzed in Spiegler (2016). This is a convenient setting for introduc-

ing the Bayesian network framework of causal misperceptions (for the general framework, see [Spiegler \(2016\)](#) and [Spiegler \(2020\)](#)). There is a unit measure of agents, who each needs to decide whether to take a dietary supplement. We let $a = 1$ denote the action of taking the supplement and $a = 0$ the action of not taking it. An agent cares about his health h , which is either good ($h = 1$) or bad ($h = 0$). His health h and supplement consumption a are potentially correlated with the level of some chemical c in his blood, which can be either normal ($c = 0$) or not ($c = 1$). An agent's payoff is equal to $h - ka$, where k is some positive constant.

The relations between these three variables are governed by a data generating process, which can be represented by a long-run joint distribution p over (a, c, h) .¹³ Let $p(a)$ and $p(h)$ denote the marginal distributions over the action and health outcome. Let $p(c | a, h)$ be the conditional probability of c , given a and h . The true objective joint distribution over (a, c, h) can be factorized as follows:

$$p(a, c, h) = p(a)p(h)p(c | a, h) \quad (12)$$

where $p(a)$ is determined endogenously by the agents' behavior according to the equilibrium notion defined below.

We assume that the true distribution p satisfies the following two properties, for all $a, h \in \{0, 1\}$. First,

$$p(h | a) = p(h) = \frac{1}{2}.$$

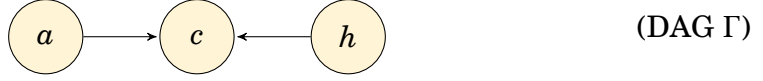
Namely, in reality, an agent's health state is independent of supplementation. Second,

$$p(c = 1 | a, h) = (1 - a)(1 - h).$$

That is, the agent's blood chemical level is abnormal if and only if he is unhealthy and have not consumed the supplement. Thus, given p , the rational decision is to *not consume* the supplement, i.e. choose $a = 0$.

The factorization presented in Eq. (12) can be depicted by a DAG, where each node corresponds to a variable, and a direct link from one variable to another indicates a causal relation – that is, the former causes the latter. Thus, the DAG Γ associated with p is

¹³One can think of this steady-state distribution as a giant excel sheet with infinite rows and three columns, one for each variable. Each entry in this sheet is a particular realization of the three variables, and the empirical frequencies are interpreted as the long-run probabilities.



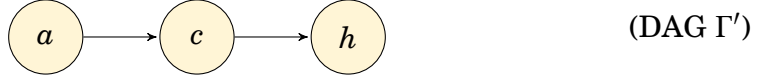
We refer to (DAG Γ) as the true DAG and to p as the true data generating process. Thus, according to p , the action a does not affect h , which is determined exogenously by nature. In addition, the chemical level c is merely a symptom, jointly influenced by the action a and the person's health h .

Remark. Given some DAG Γ' , any joint distribution p' on (a, c, h) that is consistent with Γ' can be factorized as follows:

$$p'(a, c, h) = p'(a)p'(c | \Gamma'(c))p'(h | \Gamma'(h))$$

where $\Gamma'(c)$ and $\Gamma'(h)$ are the variables that cause c and h , respectively (i.e., those variables with a link going into c and into h). If one of these variables, say $x \in \{c, h\}$, is exogenous (i.e., no links go into x), then $\Gamma'(x) = \emptyset$ and $p'(x | \Gamma'(x)) = p'(x)$.

Suppose the agents have a misspecified belief about the data generating process. Specifically, they believe that the joint distribution over (a, c, h) is consistent with the DAG Γ' given by:



Namely, the agents falsely believe that supplementation causes the health outcome via its effect on the chemical level. Thus, they believe that the joint distribution p' that is consistent with (DAG Γ') is factorized as follows:

$$p'(a, c, h) = p(a)p(c | a)p(h | c). \quad (13)$$

Suppose that, when deciding which action to take, agents do not have access to the true conditional probability $p(h | a)$. Instead, they must derive this quantity using their subjective model. That is, agents only have access to the probabilities in their factorization formula (13), namely: $p(a)$, $p(c | a)$ and $p(h | c)$. Using these, they compute their subjective belief about the probability of being healthy given an action, denoted by $p'(h = 1 | a)$, as follows:

$$p'(h = 1 | a) = p(c = 0 | a)p(h = 1 | c = 0) + p(c = 1 | a)p(h = 1 | c = 1).$$

Note that the misspecification lies in that the agent incorrectly uses $p(h = 1 \mid c = 0)$ (which is an expectation over all values of a) instead of $p(h = 1 \mid c = 0, a)$. Given our assumptions on p , these calculation yields:

$$p'(h = 1 \mid a) = \frac{1}{(2 - a)(1 + a)}$$

where a is the (endogenous) steady-state probability of choosing $a = 1$ (or equivalently, the steady-state proportion of agents who choose $a = 1$).

Mapping the setting to our framework. To describe this setting in the language of our framework, let $A = \{0, 1\}$ and $Y = \{0, 1\}$ (where consequences correspond to health status, i.e. $y = h$). The true mapping from actions to consequences $g(a)$ follows a uniform distribution over Y , for all $a \in A$. Consequently, the optimal action under rational expectations is $a = 0$.

The set of types Θ consists of all possible DAGs over the variables (a, c, y) such that a is an ancestral node. For a type $\theta = \Gamma'$, the agents' misspecified mapping g_θ is given by:

$$g_\theta(a)[1] = \begin{cases} \frac{1}{2(1+a)} & \text{if } a = 0 \\ \frac{1}{1+a} & \text{if } a = 1. \end{cases}$$

Note that the steady-state frequency of taking the supplement affects the agents' misspecified belief about the supplement's effect on health, which in turn affects the steady-state frequency of taking the supplement. Such feedback effect is common in models of misspecified beliefs.

Personal equilibrium. Because the agents' strategy affects their mapping from actions to consequences, their decisions are determined as a fixed point rather than as a solution to a maximization problem. [Spiegler \(2016\)](#) refers to this fixed point as a *personal equilibrium*. To present this definition, denote by $\mathbb{E}_{\tilde{g}}u(a, y)$ the expectation of $u(a, y)$ with respect to some mapping $\tilde{g} : A \rightarrow \Delta(Y)$.

Definition. Given a (possibly misspecified) mapping $\tilde{g} : A \rightarrow \Delta(Y)$, a probability $\alpha \in (0, 1)$ is an **ε -personal equilibrium**, if whenever $\alpha > \varepsilon$ (respectively, $1 - \alpha > \varepsilon$) then $\mathbb{E}_{\tilde{g}}u(1, y) \geq \mathbb{E}_{\tilde{g}}u(0, y)$ (respectively, $\mathbb{E}_{\tilde{g}}u(1, y) \leq \mathbb{E}_{\tilde{g}}u(0, y)$). A probability $\alpha \in [0, 1]$ is a **personal equilibrium** if it is the limit of ε -personal equilibria as $\varepsilon \rightarrow 0$.

Equipped with this definition, the following can be shown:

Proposition 3. [[Spiegler \(2016\)](#)] Suppose the agents' misspecified mapping is g_θ . If $k \in (\frac{1}{4}, \frac{1}{2})$ there is a personal equilibrium in which the proportion of agents who

take the supplement is given by

$$\alpha = \frac{1 - 2k}{2k}. \quad (14)$$

The willingness to pay for the true model. Suppose the agents could obtain access to the true probability $p(h | a)$ for a price. What would be their willingness to pay for this information in the above personal equilibrium?

According to our approach, the willingness to pay is given by the maximal expected regret from *not* obtaining the information. To determine this quantity, we seek the joint distribution q over (a, c, h) that maximizes the difference between: (i) the expected payoff from the action taken under q , and (ii) the expected payoff from the agent's equilibrium action, where: (a) this expectation is computed according to q , and (b) the joint distribution q must be consistent with the information that the agents use under their misspecified model:

$$\begin{aligned} q(h) = \frac{1}{2} \quad , \quad q(c = 0 | a = 1) = 1 \quad , \quad q(c = 0 | a = 0) = \frac{1}{2} \\ q(h = 1 | c = 1) = 0 \quad , \quad q(h = 1 | c = 0) = 2k. \end{aligned}$$

These constraints leave only one degree of freedom in the specification of q . Denoting $\beta := q(a = 1, c = 0, h = 0)$, the values of $q(a, c, h)$ are determined by β and α (where α is given by Eq. (14)) as follows:

	$a = 0$		$a = 1$	
	$c = 0$	$c = 1$	$c = 0$	$c = 1$
$h = 0$	$\frac{\alpha}{2} - \beta$	$\frac{1-\alpha}{2}$	β	0
$h = 1$	$\frac{1}{2} - \alpha + \beta$	0	$\alpha - \beta$	0

Table 1. the constrained joint distribution q

Denote by Q the set of probability distributions that are consistent with Table 1.

In contrast to the applications presented in Sections 2 and 4, here, the agent's regret depends not only on his type, but also on his equilibrium action. Let $R_a(q)$ denote the expected regret of agents who choose a and believe that q is the joint

distribution over (a, c, h) . Let:¹⁴

$$q_a^* \in \arg\max_{q \in Q} R_a(q) \quad \text{and} \quad R_a^* = R_a(q_a^*).$$

Denote by β_a^* the value of β that corresponds to q_a^* . Our next result characterizes these values:

Proposition 4. *The regret-maximizing distribution q_a^* satisfies $\beta_0^* = \max\{0, \frac{1-3k}{2k}\}$ and $\beta_1^* = \frac{1-2k}{4k}$.*

To gain intuition for this result, consider agents who choose $a = 0$ in equilibrium. The joint distribution q that maximizes their regret assigns a high probability to the outcome $h = 0$ when $a = 0$, and a high probability to $h = 1$ when $a = 1$. These objectives are achieved when β is minimized, subject to the constraint that all the entries in Table 1 remain non-negative. It can be shown that the effective constraint is $\frac{1}{2} - \alpha + \beta \geq 0$. This implies that the optimal value of β for agents choosing $a = 0$ is given by $\beta_0^* = \max\{0, \alpha - \frac{1}{2}\}$. Substituting for α , this result translates directly into the first condition stated in Proposition 4.

A similar argument applies to agents who choose $a = 1$ in equilibrium. In this case, the joint distribution q that maximizes regret assigns a high probability to $h = 1$ when $a = 0$, while simultaneously assigning a high probability to $h = 0$ when $a = 1$. These objectives are attained when β is maximized, provided that all the entries in Table 1 remain non-negative. The effective constraint in this case is $\alpha/2 - \beta \geq 0$. Hence, $\beta_1^* = \alpha/2$. This result translates directly into the second condition in Proposition 4.

Proposition 4 establishes a threshold price for learning the true causal model, above which, all agents will “rationally” choose to remain misspecified. This price is computed using the value of R_a^* , which is derived in the proof of Proposition 4. The following corollary formally describes this threshold:

Corollary 1. *No agent will pay to learn the true model if the price of learning exceeds*

$$\begin{array}{ll} \frac{2k^2}{1-2k} & \text{if } \frac{1}{4} < k \leq \frac{1}{3}, \\ \frac{2k(1-2k)}{4k-1} & \text{if } \frac{1}{3} < k \leq \frac{3}{8}, \\ k & \text{if } \frac{3}{8} < k < \frac{1}{2}. \end{array}$$

Corollary 1 leads to the following observation. Suppose there are consultants who can reveal the true causal model to the agents. However, the consultants

¹⁴Since $\beta \in [0, 1]$, this guarantees that a maximizer q_a^* exists.

have limited capacity, meaning they can serve at most a fraction $\lambda < 1 - \alpha$ of the agents population. In this market for consultants, the equilibrium prices will be determined by the WTP of the marginal agent, as defined above.

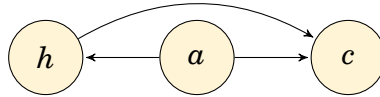
A key question arises: Will this market for consultants eliminate the misspecification and lead agents to stop consuming the supplement? The next observation demonstrates that the answer is negative when the supplement's cost falls below a certain threshold.

Corollary 2. *If $k \in (\frac{1}{4}, \frac{3}{8})$, then in the equilibrium of the market for consultants, only the agents who choose the rational action (not taking the supplement) will pay for consultants.*

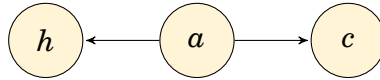
This result follows directly from the proof of Proposition 4 and Corollary 1. The WTP of agents who take the supplement is equal to k , which is lower than the WTP of agents who do not take the supplement when $k < \frac{3}{8}$. Consequently, there exists a range of parameter values in which *only agents who would not benefit from learning the true model choose to acquire the information*. Notably, when the agents who originally chose the rational action pay to learn $p(h | a)$, they simply confirm their initial choice was correct. As a result, the distribution over actions remains unchanged.

Our next result characterizes the causal models that are consistent with the joint distributions that maximize the agents' regret.

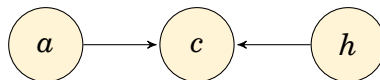
Proposition 5. (i) *If $k \geq \frac{1}{3}$, then q_0^* is consistent with a causal model represented by the DAG*



(ii) *if $k < \frac{1}{3}$, then q_0^* is consistent with a causal model represented by the DAG*



(iii) *q_1^* is consistent with the true casual model represented by the DAG*



Thus, an agent who chooses the irrational action $a = 1$ is willing to pay the most to learn the true model when he believes that this model is the one given by the objectively true DAG (i.e., (DAG Γ), where the action has no effect on the outcome, and both the action and the outcome influence the chemical level). Intuitively, such a model would make the agent realize that choosing $a = 1$ is a mistake.

On the other hand, an agent who is actually choosing the rational action (albeit for the wrong reason) is willing to pay the most to learn the true model when he believes that this model would reveal that the action affects both the outcome and the chemical level – and, for high k , that the health outcome also influences the chemical level. Intuitively, such a model would convince him to take the supplement, and hence, he would forego a high payoff if he didn't learn the truth.

4 Sampling

In this section we apply our framework to a strategic environment where players' beliefs about the mapping from actions to consequences are based on a sampling procedure proposed by [Osborne and Rubinstein \(1998\)](#). The key idea is that players are unaware of the equilibrium distribution of consequences conditional on their actions (moreover, a player may not even be aware that he is engaged in a game). Instead, each player uses the following sampling procedure: For each action, he draws K independent samples of realized payoffs generated by playing that action in the game; he then associates each action with the average payoff observed in the sample and chooses the action with the highest average payoff.

Suppose that each player receives the first sample for free, and has the option to pay a cost c to obtain a second sample for each action. What we have in mind is a setting where the player decides sequentially whether to pay for one additional sample and is myopic in this decision – when considering whether to pay for a second sample, he does not anticipate future opportunities to sample again. Thus, if the player declines the second sample, he is left with only the initial one. In light of this, would a player be willing to pay for this additional information? Our focus in this section is primarily on identifying cases in which the answer is negative, regardless of how small c is.

The setting. Following [Salant and Cherry \(2020\)](#) we focus on two-player symmetric binary action games (see [Salant and Cherry \(2020\)](#) for economic applications of sampling equilibrium in this class of games). Denote by $A = \{a, b\}$ the

	a	b
a	u_{aa}	u_{ab}
b	u_{ba}	u_{bb}

Table 2. Row player's payoffs in a symmetric game

set of actions for each player. A player who chooses action $x \in A$ while the other player chooses $y \in A$ receives a payoff of $u_{xy} = u(x, y) \geq 0$. We restrict attention to games without a dominant action. Table 2 depicts the row player's payoffs.

A *strategy* for a player is a probability distribution over actions in A . Since there are only two actions, such a distribution can be fully described by the probability with which the player chooses action a . We let α_i denote the probability that player i chooses action a .

At the outset, a player samples a realized payoff for each of his actions. This payoff is generated by drawing an action from his opponent's strategy. For example, when player i samples a payoff for action a , he draws u_{aa} with probability α_{-i} , and u_{ab} with probability $(1 - \alpha_{-i})$. A sample s_K is a list of K payoff pairs $((u_a^k, u_b^k))_{k=1, \dots, K}$, where u_a^k and u_b^k denote the realized payoff from actions a and b in the k^{th} sample, respectively.

Given a sample $s_K = ((u_a^k, u_b^k))_{k=1, \dots, K}$, define

$$\bar{u}_a(s_K) = \frac{1}{K} \sum_{k=1}^K u_a^k \quad \bar{u}_b(s_K) = \frac{1}{K} \sum_{k=1}^K u_b^k \quad (15)$$

In words, $\bar{u}_a(s_K)$ and $\bar{u}_b(s_K)$ are the average payoffs that were generated by actions a and b in the sample, respectively. A player who draws the sample s_K will prefer action a over b if and only if $\bar{u}_a(s_K) \geq \bar{u}_b(s_K)$ (with ties broken arbitrarily).

A steady-state – or a sampling equilibrium – is a probability distribution over actions that satisfies the following fixed point property: the probability that action a is played is equal to the probability that this action is associated with the highest average payoff in the sample.¹⁵ A single-sample (or S1) equilibrium is then defined as follows:

Definition. An S1 equilibrium is a distribution $(\alpha, 1 - \alpha)$ over the actions (a, b) ,

¹⁵In what follows, we focus on non-degenerate cases in which each action is played with positive probability. In such cases, the average payoff associated with each action is a random variable.

	a	b
a	2	4
b	3	1

Table 3. Payoffs in a symmetric game – an example

	a	b
a	2	2
b	3	1

Table 4. A symmetric game with an average-preserving-spread action

	a	b
a	4	2
b	1	3

Table 5. A symmetric game with main diagonal dominance

where α equals the probability of drawing a sample s_1 in which $\bar{u}_a(s_1) > \bar{u}_b(s_1)$, assuming that each player's strategy is given by $(\alpha, 1 - \alpha)$.

To illustrate this equilibrium concept, consider two players who play the binary symmetric game shown in Table 3, where the entries represent the row player's payoffs. This game appears as Example 1 in Osborne and Rubinstein (1998). In an $S1$ equilibrium, the probability α that a player chooses a is equal to the probability that $s_1 \in \{(2, 1), (4, 3), (4, 1)\}$. Thus, $\alpha = \alpha(1 - \alpha) + (1 - \alpha)$.¹⁶

Suppose that a player is approached with the opportunity to pay a cost c and obtain a second sample for each of his actions. To compute the maximal regret from *not* taking the second sample, the player considers all possible average empirical payoffs for each action following the new sample. To remain consistent with the information he already possesses, the player cannot "forget" the payoff he observed in the initial sample.

Mapping the setting to our framework. The set of actions is given by $A = \{a, b\}$ and the set of consequences is given by the set of payoffs $Y = \mathbb{R}$. The "true" mapping g faced by a player in an $S1$ equilibrium is given by

$$g(a)[u_{aa}] = g(b)[u_{ba}] = \alpha$$

where $g(a)[z]$ denotes the probability that the distribution $g(a)$ assigns to the outcome z .

A type θ is defined by the realized payoff the player observes for each of his actions in the initial sample $s = (u_a^1, u_b^1)$. Because the player associates each action with the average payoff it generates, a misspecified model simply maps each action to a degenerate distribution that assigns probability one to a single num-

¹⁶The first term on the right-hand side is the probability of drawing 2 when a is sampled and 1 when b is sampled, given that the opponent chooses his first and second actions with probability α and $1 - \alpha$, respectively. The second term is the probability of drawing 4 when a is sampled.

ber – the average payoff associated with that action. Thus, with a slight abuse of notation, we have $g_\theta(a) = u_a^1$ and $g_\theta(b) = u_b^1$. As noted above, the optimal action for a player of type θ is given by $a_\theta = a$ if and only if $u_a^1 \geq u_b^1$; otherwise, $a_\theta = b$.

Consider a player of type θ who is contemplating whether to take a second sample for each action. The set of models that type θ deems possible after obtaining a second sample, G_θ , includes all mappings from actions to empirical averages of payoffs that could possibly result from a second sample. To be consistent with his current knowledge, each average must incorporate the realized payoff from the initial sample. Thus,

$$\tilde{g} \in G_\theta \iff \begin{cases} \tilde{g}(a) \in \{\frac{1}{2}(u_a^1 + u_{aa}), \frac{1}{2}(u_a^1 + u_{ab})\}, \\ \tilde{g}(b) \in \{\frac{1}{2}(u_b^1 + u_{ba}), \frac{1}{2}(u_b^1 + u_{bb})\} \end{cases} \quad (16)$$

In words, a model $\tilde{g}$ is deemed possible for type $\theta = (u_a^1, u_b^1)$ if, after a second sample, action a is mapped either to the average of the observed payoff, u_a^1 , and the “new” payoff u_{aa} , i.e., $\frac{1}{2}(u_a^1 + u_{aa})$, or to the average of the observed payoff, u_a^1 , and the “new” payoff u_{ab} , i.e., $\frac{1}{2}(u_a^1 + u_{ab})$. An analogous requirement applies to the mapping of action b .

Non-willingness to pay for a second sample. Fix a type θ . For each model $\tilde{g} \in G_\theta$, if type θ has no strict incentive to switch from a_θ , then this type’s regret from not adopting $\tilde{g}$ is zero, i.e. $R_\theta(\tilde{g}) = 0$. If type θ has zero regret from not adopting any admissible model in G_θ , then this type’s maximal regret is also zero: $R^*(G_\theta) = 0$. If for every type θ the maximal regret is zero, then no player is willing to pay a positive amount for a second sample. In what follows we explore conditions under which players are not willing to pay for a second sample.

Consider the example depicted in Table 4. In this game, no type would be willing to pay a positive amount for a second sample. To see why, consider the types who observed a realized payoff of 3 when sampling action b . These types would choose action b if they do not sample again. However, even if they were to obtain a second sample for each action, the outcome would never provide a reason to switch actions. Hence, these types would not be willing to pay a positive amount for a second sample. By similar reasoning, the types for whom the realized payoff was 1 when they sampled the action b (and therefore, they would choose a without a second sample), would also not be willing to pay for a second sample.

Note that this game has the property that there is a “safe” action that delivers a constant payoff, and a “risky” action whose payoffs are an “average-preserving-

spread” of the safe action’s payoffs. Formally, we say that an action $y \in \{a, b\}$ is an *average-preserving-spread* of action $x \neq y$ if $u_{aa} + u_{ab} = u_{ba} + u_{bb}$ and $u_{xa} = u_{xb}$ while $u_{ya} \neq u_{yb}$. The next result shows that this property is both a necessary and sufficient condition for ensuring that no type is willing to pay for a second sample.

Proposition 6. *No type is willing to pay any positive amount for a second sample if and only if one action is an average-preserving-spread of the other.*

The proof of the “if” direction follows reasoning similar to that used in the above example. The “only if” direction identifies the types and second-sample realizations that would induce the player to switch actions in the absence of an average-preserving relationship between the actions.

Proposition 6 establishes that in the absence of an average-preserving action, there always exists a type for whom the maximal regret from not taking a second sample is strictly positive. However, we show that in a certain class of games, there exist *S1* equilibria in which the probability that such a type is realized is arbitrarily small.

We say that in a binary symmetric game, the main diagonal *dominates the off diagonal* if $\min\{u_{aa}, u_{bb}\} > \max\{u_{ab}, u_{ba}\}$. Table 5 depicts an example. Under this condition, we have the following result:

Proposition 7. *Suppose that the main diagonal dominates the off diagonal. Then, for any probability p there exists an *S1* equilibrium in which the probability that a player is not willing to pay for a second sample is higher than p .*

The idea behind the proof is to first identify types who are not willing to pay for a second sample. Then, when the main diagonal dominates the off-diagonal, it is possible to construct an *S1* equilibrium in which the probability of such types being realized is arbitrarily high.

Finally, we note that main diagonal dominance is not a necessary condition for the result. However, in its absence, one can construct games that have a unique *S1* equilibrium in which, for example, each player uniformly randomizes between the two actions. In such a game, the probability that a player is not willing to pay for a second sample cannot exceed 0.25. An example of such a game is given by the payoffs: $u_{aa} = 5$, $u_{ab} = 1$, $u_{ba} = 3$, and $u_{bb} = 2$.

Extending the above analysis to symmetric games with more than two actions would follow the same type of reasoning. However, the necessary and sufficient conditions for zero WTP may be more involved.

5 Proofs

Proof of Proposition 1

Instead of maximizing Eq. (OBJ) over the domain of marginals H_x^1 and H_x^2 that satisfy the constraint in Eq. (10), we solve the equivalent problem of maximizing the objective in Eq. (OBJ) over the domain of joint distributions over $[0, 1] \times [0, 1]$, whose marginals satisfy Eqs. (5) and (10). A solution to this problem is a pair, (p_1^*, H^*) , of a price and a joint distribution over the quality (ϕ) and seller's ask (x).

We start by considering the case where $p_1^* \geq p_0$ (recall that p_0 is the price that solves Eq. (9)). To solve this problem, it is useful to partition the space $[0, 1] \times [0, 1]$ into the following six subsets, as depicted in Figure 1:

$$\begin{aligned} B_i &\equiv \{(\phi, x) \mid \phi \in C_i \wedge x \in (0, p_0)\} \\ M_i &\equiv \{(\phi, x) \mid \phi \in C_i \wedge x \in (p_0, p)\} \\ T_i &\equiv \{(\phi, x) \mid \phi \in C_i \wedge x \in (p, 1)\} \end{aligned}$$

For any joint distribution H over $[0, 1] \times [0, 1]$, and for each $i = 1, 2$, denote by $\mu_H(B_i)$, $\mu_H(M_i)$, and $\mu_H(T_i)$ the probability mass of the sets B_i , M_i , and T_i , respectively, according to H . Hence, the conditional marginal distributions induced by H can be computed as follows:

$$H_x^i(p_0) = \frac{\mu_H(B_i)}{\mu_H(B_i) + \mu_H(M_i) + \mu_H(T_i)} \text{ and } H_x^i(p) = \frac{\mu_H(B_i) + \mu_H(M_i)}{\mu_H(B_i) + \mu_H(M_i) + \mu_H(T_i)}.$$

Moreover, in terms of these probability masses, the marginal constraint (5) implies that:

$$\mu_H(B_1) + \mu_H(M_1) + \mu_H(T_1) = F(t) \quad \text{and} \quad \mu_H(B_2) + \mu_H(M_2) + \mu_H(T_2) = 1 - F(t), \quad (17)$$

whereas the marginal constraint (10) implies that:

$$\mu_H(B_1) + \mu_H(B_2) = F(p_0) \quad \text{and} \quad \mu_H(M_1) + \mu_H(M_2) = F(p) - F(p_0). \quad (18)$$

Let $\mathcal{H}$ denote the set of joint distributions over qualities (ϕ) and asks (x), whose marginals satisfy Eqs. (17) and (18). We can then rewrite the maximization problem in Eq. (OBJ), restricted to the case that $p_1^* \geq p_0$, as follows:

$$\max_{p \geq p_0} \left[\max_{H \in \mathcal{H}} \left(V_1 \mu_H(M_1) + V_2 \mu_H(M_2) \right) - pF(p) \right] + F(p_0)p_0. \quad (19)$$

Recall that, by definition, $V_2 > V_1$. Moreover, given any price $p \geq p_0$, the constraint in Eq. (18) implies that the sum $\mu_H(M_1) + \mu_H(M_2)$ is fixed. Hence, the distribution H^* assigns as much probability as possible to $\mu_H(M_2)$ “at the expense” of the probability mass assigned to $\mu_H(M_1)$.

Consequently, the optimal solution (p_1^*, H^*) cannot satisfy $F(p_1^*) - F(p_0) > 1 - F(t)$. To see this, suppose that $F(p_1^*) - F(p_0) > 1 - F(t)$, and note that this implies $p_1^* > p_0$. Moreover, the constraints in Eqs. (17) and (18) imply that the highest probability mass that H^* can assign to the region M_2 cannot exceed $1 - F(t)$, and therefore $\mu_{H^*}(M_2) = 1 - F(t)$ and $\mu_{H^*}(M_1) = (F(p_1^*) - F(p_0)) - (1 - F(t))$.¹⁷ Plugging these into Eq. (19) we obtain that $p^* = \arg \max_{p \geq p_0} F(p)(V_1 - p)$. However, our assumption that $p + F(p)/f(p)$ is strictly increasing in p implies that the derivative $\frac{d}{dp} [F(p)(V_1 - p)] = f(p)[V_1 - (p + F(p)/f(p))]$ is negative for all $p \geq p_0$. This is a contradiction for the optimal price p^* being strictly above p_0 .

Therefore, the solution (p_1^*, H^*) must satisfy $F(p_1^*) - F(p_0) \leq 1 - F(t)$. The fact that $V_2 > V_1$ and the constraints (17) and (18) then imply that $\mu_{H^*}(M_1) = 0$ and $\mu_{H^*}(M_2) = F(p^*) - F(p_0)$.¹⁸ Substituting into Eq. (19), we obtain:

$$\begin{aligned} p_1^* &\in \arg \max_{p \geq p_0} F(p)(V_2 - p) \\ &\text{subject to } F(p) - F(p_0) \leq 1 - F(t) \end{aligned}$$

Our assumption that $p + F(p)/f(p)$ is strictly increasing in p implies that the function $F(p)(V_2 - p)$ has a unique extremum on $[p_0, 1]$, which occurs at $\tilde{p}$ that satisfies $V_2 = \tilde{p} + F(\tilde{p})/f(\tilde{p})$. Furthermore, this extremum is a maximum, and $\tilde{p} \geq p_0$. Therefore, $p_1^* = \tilde{p}$, provided that $F(\tilde{p}) - F(p_0) \leq 1 - F(t)$. Otherwise, p_1^* is equal to the price p that solves $F(p) - F(p_0) = 1 - F(t)$.

The proof in the case of $p_1^* < p_0$ is analogous and is therefore omitted.

Proof of Proposition 2

Suppose that the partition $((0, t), (t, 1))$ maximizes the buyer’s WTP. Let (p_1^*, H^*) be the solution to the maximization problem in Eq. (OBJ), subject to the constraint

¹⁷ Additionally, $\mu_{H^*}(T_2) = \mu_{H^*}(B_2) = 0$, $\mu_{H^*}(T_1) = 1 - F(p_1^*)$ and $\mu_{H^*}(B_1) = F(p_0^*)$

¹⁸ Additionally, $\mu_{H^*}(T_2)$ and $\mu_{H^*}(B_2)$ satisfy $0 \leq \mu_{H^*}(T_2) \leq 1 - F(p^*)$, $0 < \mu_{H^*}(B_2) < F(p_0)$, and $\mu_{H^*}(T_2) + \mu_{H^*}(B_2) = (1 - F(t)) - (F(p^*) - F(p_0))$.

in Eq. (10). Suppose, by contradiction, that $F(p_1^*) - F(p_0) < 1 - F(t)$. By Proposition 1, this implies that $V_2 = p_1^* + F(p_1^*)/f(p_1^*)$, and therefore $p_1^* > p_0$.

Because H^* is determined optimally given p_1^* , we know from the proof of Proposition 1 that the buyer's maximal WTP, as presented in Eq. (19), can be written as follows:

$$V_2(F(p_1^*) - F(p_0)) - p_1^*F(p_1^*) + F(p_0)p_0. \quad (20)$$

Recall that $V_2 \equiv \mathbb{E}_{\phi \sim F}(v(\phi, b) | \phi \in C_2)$, where $C_2 = (t, 1)$ is the second element in the analogy partition. Therefore, by slightly increasing the boundary between the partition elements from t to $t^+ > t$, so that $C_2 = (t^+, 1)$, we increase the value of V_2 . Holding (p_1^*, H^*) fixed, this change only increases the expression in Eq. (20), while the inequality $F(p_1^*) - F(p_0) < 1 - F(t^+)$ still holds. Note, however, that this new value of Eq. (20), which is computed when (p_1^*, H^*) is held fixed, is only a lower bound to the buyer's willingness to pay for the new analogy partition, which is computed by solving the maximization problem in Eq. (OBJ), subject to the constraint in Eq. (10), for the analogy partition $((0, t^+), (t^+, 1))$. This contradicts $((0, t), (t, 1))$ being the partition that maximizes the buyer's willingness to pay for knowledge. The proof of the second case is analogous and is omitted. ■

Proof of Proposition 4

For agents who choose $a = 0$, we have

$$R_0^* = \max_q [q(h = 1 | a = 1) - k - q(h = 1 | a = 0)]$$

which reduces to

$$R_0^* = \max_{\beta} \left(\frac{\alpha - \beta}{\alpha} - k - \frac{\frac{1}{2} - \alpha + \beta}{1 - \alpha} \right)$$

Since the R.H.S. *decreases* with β , regret is maximized for the *minimal* β that satisfies $\frac{1}{2} - \alpha + \beta \geq 0$. This yields that $\beta_0^* = \max\{0, \alpha - \frac{1}{2}\}$ and implies that

$$R_0^* = \begin{cases} \frac{2k(1-2k)}{4k-1} & \text{if } k \geq \frac{1}{3} \\ \frac{2k^2}{1-2k} & \text{if } k < \frac{1}{3} \end{cases} \quad (21)$$

(note that both values are positive since $k < \frac{1}{2}$).

For agents who choose $a = 1$, we have

$$R_1^* = \max_q [q(h = 1 | a = 0) - q(h = 1 | a = 0) - k]$$

which reduces to

$$R_1^* = \max_{\beta} \left(\frac{\frac{1}{2} - \alpha + \beta}{1 - \alpha} - \frac{\alpha - \beta}{\alpha} + k \right)$$

Since the R.H.S. *increases* with β , regret is maximized for the *maximal* β that satisfies $\alpha/2 - \beta \geq 0$ and $\frac{1}{2} - \alpha + \beta \leq 1$. Since $\alpha/2 < \alpha + \frac{1}{2}$, the solution is $\beta = \alpha/2$. It follows that $R_1^*(k) = k$. ■

Proof of Proposition 5

Proof of part (i). If $k \geq \frac{1}{3}$, then q_0^* satisfies

$$q_0^*(h = 1 | a = 1) = 1 > \frac{3k - 1}{4k - 1} = q_0^*(h = 1 | a = 0)$$

and

$$q_0^*(c = 1 | a = 0, h = 0) = \frac{4k - 1}{2k}$$

while

$$q_0^*(c = 1 | a = 0, h = 1) = q_0^*(c = 1 | a = 1, h = 0) = 0$$

Proof of part (ii). If $k < \frac{1}{3}$, then

$$\begin{aligned} q_0^*(h = 1 | a = 1) &= \frac{1}{2} > 0 = q_0^*(h = 1 | a = 0), \\ q_0^*(c = 1 | a = 1, h = 0) &= q_0^*(c = 1 | a = 1, h = 1) = 0, \text{ and} \\ q_0^*(c = 1 | a = 0, h = 0) &= q_0^*(c = 1 | a = 0) = \frac{1}{2}. \end{aligned}$$

(note that the event $(a = 0 \wedge h = 1)$ has zero probability).

Proof of part (iii). Note that

$$q_1^*(h = 1 | a = 1) = q_1^*(h = 1 | a = 0) = \frac{1}{2},$$

and that

$$q_1^*(c = 1 | a, h) = (1 - a)(1 - h).$$

■

Proof of Proposition 6

Sufficiency. Without loss of generality, assume that the safe action is a , yielding a payoff of z , i.e. $u_{aa} = u_{ab} = z$. The two payoffs from the risky action b can then be represented by $u_{ba} = z + \delta$ and $u_{bb} = z - \delta$. Without loss of generality we assume that $\delta \in (0, z]$.

In this setting, a type can be essentially characterized by the realized payoff from sampling b . Consider first the “type” $z + \delta$, who chooses b absent a second sample. The realization of a second sample which is “maximally in favor of switching to a ” is $(z, z - \delta)$. However, for this realization the average from both actions is z , hence there is no strict incentive to switch.

Consider next the “type” $z - \delta$, who chooses a absent a second sample. For this to be most inclined to switch to b , the second sample realization should be $(z, z + \delta)$. However, as before, even with this realization, the player has no strict incentive to switch actions.

Necessity. Suppose first that $u_{aa} + u_{ab} \neq u_{ba} + u_{bb}$ and assume, by contradiction, that no type is willing to pay for a second sample. This means that for each type there is no realization of the second sample that would give a strict incentive for that type to switch an action. Without loss of generality, assume that $u_{aa} > u_{ba}$. Since no action is dominated, this means that $u_{bb} > u_{ab}$. Consider types (u_{aa}, u_{ba}) and (u_{ab}, u_{bb}) , who would choose the actions a and b , respectively, absent a second sample. If for every realization of the second sample, these types have no strict incentive to switch an action, then $u_{aa} + u_{ab} \geq u_{ba} + u_{bb}$ and $u_{bb} + u_{ba} \geq u_{ab} + u_{aa}$. But these two inequalities cannot both hold.

Suppose next that $u_{aa} + u_{ab} = u_{ba} + u_{bb}$ but $u_{aa} \neq u_{ab}$ and $u_{ba} \neq u_{bb}$. Without loss of generality assume that u_{ba} is the maximal payoff. We consider two cases.

Case 1: $u_{aa} > u_{ab}$. Consider type (u_{aa}, u_{ba}) , who chooses action b absent a second sample. If this type would observe in the second sample the realization (u_{aa}, u_{bb}) , he would want to switch to action a because $2u_{aa} > u_{aa} + u_{ab} = u_{ba} + u_{bb}$.

Case 2: $u_{aa} < u_{ab}$. Consider type (u_{ab}, u_{ba}) , who chooses action b absent a second sample. If this type would observe in the second sample the realization (u_{ab}, u_{bb}) , he would want to switch to action a because $2u_{ab} > u_{aa} + u_{ab} = u_{ba} + u_{bb}$. ■

Proof of Proposition 7

Proof. Suppose that $\max\{u_{aa}, u_{bb}\} > \max\{u_{ab}, u_{ba}\}$. Without loss of generality assume $u_{aa} = \max\{u_{aa}, u_{bb}\}$. Since no action is dominated, $u_{ab} < u_{bb}$. We first show that any distribution over $\{a, b\}$ is an S1 equilibrium. There are two cases to consider. If $u_{ab} < u_{ba}$, then the only event in which a player chooses a is that when he sampled that action the realized payoff was u_{aa} . If $u_{ab} > u_{ba}$, then the only event in which a player chooses b is that when he sampled that action the realized payoff was u_{bb} . Both cases yield that any distribution over the two actions is an S1 equilibrium.

If $u_{aa} + u_{ab} \geq u_{ba} + u_{bb}$, then type (u_{aa}, u_{ba}) , who would choose a without a second sample, is not willing to pay for a second sample. In this case, there is an S1 equilibrium $\alpha > \sqrt{p}$ in which the probability that this type is realized is higher than p . If $u_{aa} + u_{ab} < u_{ba} + u_{bb}$, then type (u_{ab}, u_{bb}) , who would choose b without a second sample, is not willing to pay for a second sample. In this case there is an S1 equilibrium $\alpha < 1 - \sqrt{p}$ in which the probability that this type is realized is higher than p . ■

Bibliography

- Aina, Chiara.** 2024. “Tailored Stories.” Working paper.
- Ba, Cuimin.** 2024. “Robust Misspecified Models and Paradigm Shift.” Working paper.
- Bergemann, Dirk, and Karl H. Schlag.** 2008. “Pricing without Priors.” *Journal of the European Economic Association* 6 (2-3): 560–569. [10.1162/JEEA.2008.6.2-3.560](https://doi.org/10.1162/JEEA.2008.6.2-3.560).
- Bergemann, Dirk, and Karl H. Schlag.** 2011. “Robust monopoly pricing.” *Journal of Economic Theory* 146 (6): 2527–2543. [10.1016/j.jet.2011.10.018](https://doi.org/10.1016/j.jet.2011.10.018).
- Cho, In-Koo, and Kenneth Kasa.** 2015. “Learning and Model Validation.” *The Review of Economic Studies* 82 (1): 45–82.
- Cohen, Shani, and Shengwu Li.** 2023. “Sequential Cursed Equilibrium.”
- Eliaz, Kfir, and Ran Spiegler.** 2020. “A Model of Competing Narratives.” *American Economic Review* 110 (12): 3786–3816. [10.1257/aer.20180624](https://doi.org/10.1257/aer.20180624).
- Eliaz, Kfir, Ran Spiegler, and Heidi C. Thyssen.** 2021a. “Persuasion with endogenous misspecified beliefs.” *European Economic Review* 134 103712. [10.1016/j.eurocorev.2021.103712](https://doi.org/10.1016/j.eurocorev.2021.103712).

- Eliaz, Kfir, Ran Spiegler, and Heidi C. Thyssen.** 2021b. “Strategic interpretations.” *Journal of Economic Theory* 192 105192. [10.1016/j.jet.2021.105192](https://doi.org/10.1016/j.jet.2021.105192).
- Eliaz, Kfir, Ran Spiegler, and Yair Weiss.** 2021c. “Cheating with Models.” *American Economic Review: Insights* 3 (4): 417–434. [10.1257/aeri.20200635](https://doi.org/10.1257/aeri.20200635).
- Esponda, Ignacio.** 2008. “Behavioral Equilibrium in Economies with Adverse Selection.” *American Economic Review* 98 (4): 1269–1291. [10.1257/aer.98.4.1269](https://doi.org/10.1257/aer.98.4.1269).
- Esponda, Ignacio, and Demian Pouzo.** 2016. “Berk–Nash Equilibrium: A Framework for Modeling Agents With Misspecified Models.” *Econometrica* 84 (3): 1093–1130. [10.3982/ECTA12609](https://doi.org/10.3982/ECTA12609).
- Eyster, Erik, and Matthew Rabin.** 2005. “Cursed Equilibrium.” *Econometrica* 73 (5): 1623–1672. [10.1111/j.1468-0262.2005.00631.x](https://doi.org/10.1111/j.1468-0262.2005.00631.x).
- Gonçalves, Duarte.** 2023. “Sequential Sampling Equilibrium.” Working paper, arXiv, 2212.07725.
- He, Kevin, and Jonathan Libgober.** 2023. “Evolutionarily Stable (Mis)specifications: Theory and Applications.” Working paper.
- Heller, Yuval, and Eyal Winter.** 2020. “Biased-Belief Equilibrium.” *American Economic Journal: Microeconomics* 12 (2): 1–40. [10.1257/mic.20180319](https://doi.org/10.1257/mic.20180319).
- Jehiel, Philippe.** 2005. “Analogy-Based Expectation Equilibrium.” *Journal of Economic Theory* 123 (2): 81–104. doi.org/10.1016/j.jet.2003.12.003.
- Jehiel, Philippe.** 2022. “Analogy-based expectation equilibrium and related concepts: Theory, applications, and beyond.” Working paper.
- Jehiel, Philippe, and Frédéric Koessler.** 2008. “Revisiting games of incomplete information with analogy-based expectations.” *Games and Economic Behavior* 62 (2): 533–557. [10.1016/j.geb.2007.06.006](https://doi.org/10.1016/j.geb.2007.06.006).
- Jehiel, Philippe, and Giacomo Weber.** 2024. “Endogenous Clustering and Analogy-Based Expectation Equilibrium.” Working paper.
- Kops, Christopher, and Illia Pasichnichenko.** 2023. “Testing negative value of information and ambiguity aversion.” *Journal of Economic Theory* 213 105730.
- Laurence, Emily, and Nicholas Temple.** 2023. “The Psychology Behind FOMO.” *Forbes Health*, <https://www.forbes.com/health/mind/the-psychology-behind-fomo/>.
- Li, Jian.** 2020. “Preferences for partial information and ambiguity.” *Theoretical Economics* 15 (3): 1059–1094.
- Maćkowiak, Bartosz, Filip Matějka, and Mirko Wiederholt.** 2023. “Rational

- Inattention: A Review.” *Journal of Economic Literature* 61 (1): 226–273. [10.1257/jel.20211284](https://doi.org/10.1257/jel.20211284).
- Milyavskaya, Marina, Mark Saffran, Nadia Hope, and Richard Koestner.** 2018. “Fear of missing out: prevalence, dynamics, and consequences of experiencing FOMO.” *Motivation and Emotion* 42 (5): 725–737. [10.1007/s11031-018-9683-5](https://doi.org/10.1007/s11031-018-9683-5).
- Osborne, Martin, and Ariel Rubinstein.** 1998. “Games with Procedurally Rational Players.” *American Economic Review* 88 (4): 834–847.
- Salant, Yuval, and Josh Cherry.** 2020. “Statistical Inference in Games.” *Econometrica* 88 (4): 1725–1752. <https://doi.org/10.3982/ECTA17105>.
- Schwartzstein, Joshua, and Adi Sunderam.** 2021. “Using Models to Persuade.” *American Economic Review* 111 (1): 276–323. [10.1257/aer.20191074](https://doi.org/10.1257/aer.20191074).
- Shishkin, Denis, and Pietro Ortoleva.** 2023. “Ambiguous information and dilation: An experiment.” *Journal of Economic Theory* 208 105610.
- Sims, Christopher A.** 2003. “Implications of rational inattention.” *Journal of Monetary Economics* 50 (3): 665–690. [10.1016/S0304-3932\(03\)00029-1](https://doi.org/10.1016/S0304-3932(03)00029-1).
- Spiegler, Ran.** 2011. *Bounded Rationality and Industrial Organization*. Oxford: Oxford University Press, . [10.1093/acprof:oso/9780195398717.001.0001](https://doi.org/10.1093/acprof:oso/9780195398717.001.0001).
- Spiegler, Ran.** 2016. “Bayesian networks and boundedly rational expectations.” *The Quarterly Journal of Economics* 131 (3): 1243–1290. [10.1093/qje/qjw017](https://doi.org/10.1093/qje/qjw017).
- Spiegler, Ran.** 2020. “Behavioral Implications of Causal Misperceptions.” *Annual Review of Economics* 12 (Volume 12, 2020): 81–106. <https://doi.org/10.1146/annurev-economics-072219-111921>.