

Rational Inattention via Ignorance Equivalence

Michèle Müller-Itten*, Roc Armenter†, Zachary R. Stangebye‡

July 9, 2025

Abstract

We introduce the concept of the *ignorance equivalent* to effectively summarize the payoff possibilities of a rationally inattentive (RI) agent. The ignorance equivalent is a unique fictitious action that does not distort learning incentives when added to the agent’s menu, but also makes ignorance optimal. By restating the RI problem as a choice over a richer menu without learning, the approach provides new insights for menu expansion, the formation of consideration sets, and belief elicitation. When RI agents with different learning costs interact, their ignorance equivalents emerge in equilibrium, facilitating trade that allows agents to emulate the first-best learning strategy.

Keywords: Rational inattention, information acquisition, learning.

JEL: D81, D83, C63

*University of St.Gallen: michele.mueller-itten@unisg.ch

†Federal Reserve Bank of Philadelphia: roc.armenter@phil.frb.org

‡University of Notre Dame: zstangeb@nd.edu

The views expressed in this paper are solely those of the authors and do not necessarily reflect the views of the Federal Reserve Bank of Philadelphia or the Federal Reserve System. The authors would like to thank Nageeb Ali, Isaac Baley, Alex Bloedel, Tommaso Denti, Thomas Gresik, Maciej Kotowski, John Leahy, Bart Lipman, Filip Matějka, Jun Nie, Paolo Piacquadio, several anonymous referees and participants at various conferences and seminars for constructive feedback and comments. We are especially grateful to Ina Taneva, who served as a conference discussant, for her insightful comments and suggestions. Financial support from the Notre Dame Institute for Scholarship in the Liberal Arts is gratefully acknowledged. Philadelphia Fed working papers are free to download at <https://philadelphiafed.org/research-and-data/publications/working-papers>.

1 Introduction

Strategic attention and flexible learning are an integral part of human decision making. Yet, many economic models fail to incorporate agents’ extensive ability to shape the uncertainty that they take on. In part, this may be due to the inherent complexity of describing – and tracking – the optimal choice with learning. In this paper, we demonstrate that the agent’s ability to learn is *as if* she had access to a larger menu of options. Re-expressing the choice with learning as one without learning over a larger menu has tractability benefits for comparisons across menus, beliefs, and learning strategies. This connection between learning ability and fictional payoffs is made explicit in games that require surplus splitting, and enables agents to implement the first best joint strategy.

Our analysis rests on the Rational Inattention (RI) paradigm. This decision model posits that prior to choosing from a menu of actions with state-dependent payoffs, an agent can acquire any state-dependent signal but faces a utility cost from generating, accessing or processing (“learning”) that information [Sims, 2003]. The result is an endogenous information structure that responds to incentives and changes in the environment.

A new concept. We introduce the concept of an *ignorance equivalent* to summarize the most pertinent features of a finite RI problem for a broad class of information costs. The ignorance equivalent is a fictitious action with state-dependent payoffs that makes the agent no better or worse off when made available in addition to, or instead of, the original menu. That is, the ignorance equivalent is just attractive enough for the agent to forgo all existing learning opportunities, and yet generates no new profitable learning opportunities when added to the menu. We show that the ignorance equivalent exists and is unique (Theorem 1), and argue that it summarizes the most pertinent features of the menu.

For concreteness, we here illustrate the main ideas with a simple example.¹ A risk-neutral investor can place her wealth in one of three assets whose payouts depend on market conditions. Before investing, she can learn whether she is in a bull or a bear market, albeit at a cost. Figure 1(a) depicts her choice problem graphically, by plotting the state-dependent payouts on each axis. Asset $\mathbf{a}^1$ is pro-cyclical with a high payout in a bull market, asset $\mathbf{a}^2$ is a safe asset with a constant payout, and asset $\mathbf{a}^3$ is counter-cyclical with a relatively higher payout in a bear market. The vector $\boldsymbol{\pi}$ indicates the agent’s prior belief; in this case, she is leaning slightly towards a bull market. If learning were impossible, she would place her wealth in the pro-cyclical asset $\mathbf{a}^1$ to maximize her expected utility (as indicated by the ‘no info’ line). If learning were free, she would select $\mathbf{a}^1$ in a bull market and $\mathbf{a}^3$ in a bear market

¹The paper itself uses generic language that befits any finite RI problem with prior-concave, uniformly posterior separable costs.

(‘full info’). With costly learning, welfare is between these two benchmarks, for example at W_S .² The associated optimal strategy $\mathcal{S}$ has the investor partially learn the state and place her wealth either in $\mathbf{a}^1$ or the safe asset $\mathbf{a}^2$ (proof omitted³).

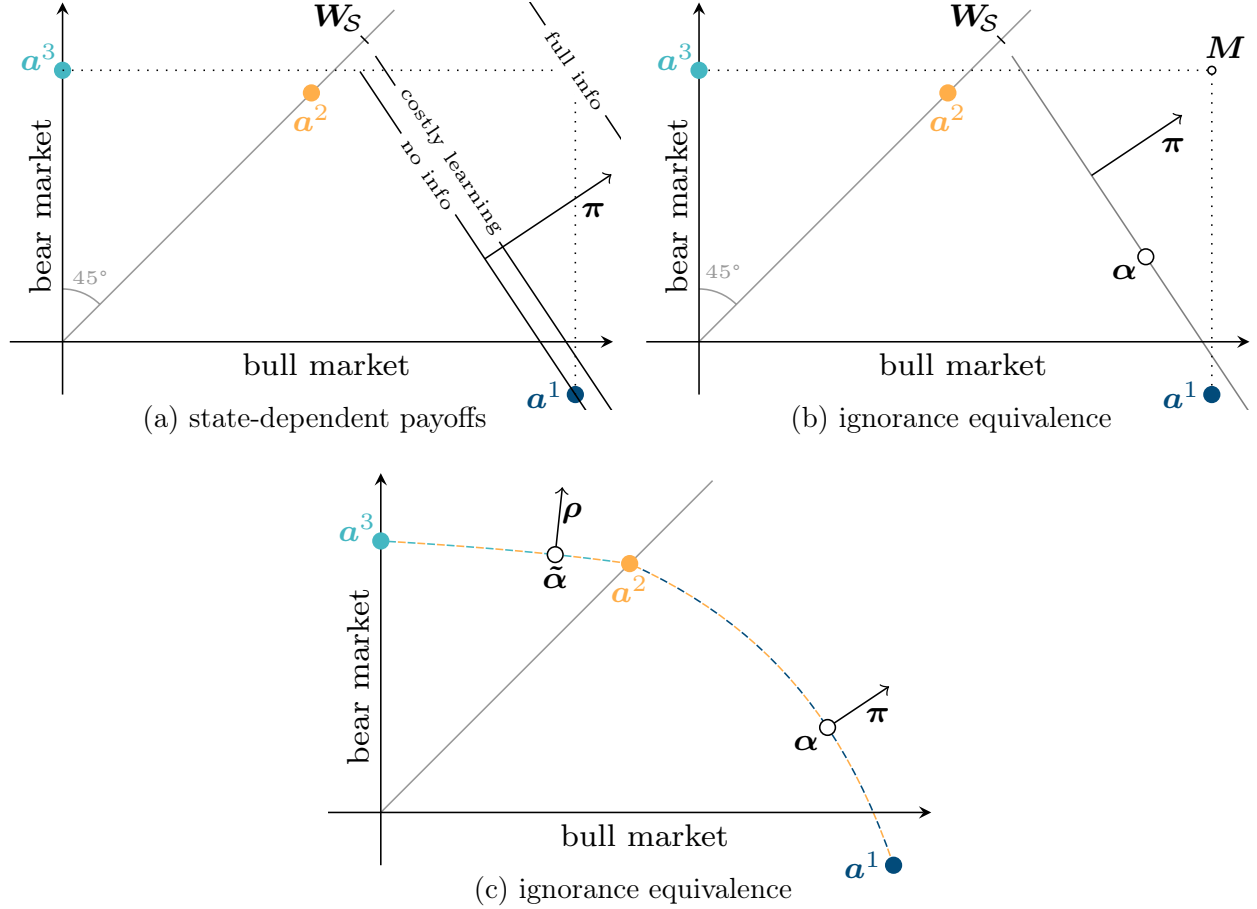


Figure 1: Investor Example

The crux of our paper is to argue that the investor’s ability to learn is ‘as if’ she also had access to the fictional asset α indicated on Figure 1(b). An unconditional purchase of this *ignorance equivalent* asset also yields welfare W_S , yet the availability of α enables no better investment strategies even with learning. Indeed, we show something stronger: Asset α is no more attractive than investment strategy $\mathcal{S}$ under *any* belief about the state. The availability of α thus brings no welfare gains for the investor – nor for another investor with a different prior.

²In the graph, the point $\mathbf{W}_S = (W_S, W_S)$ depicts the certain payoff equal to the investor’s welfare.

³Numerically, the example assumes that $\mathbf{a}^1 = (6, -1.5)$, $\mathbf{a}^2 = (2.5, 2.5)$, $\mathbf{a}^3 = (-0.8, 2.8)$, $\pi = (0.6, 0.4)$, and the investor can learn about the state according to (UPS+) costs with $\phi(\rho) = 4 \sum_{i=1}^2 \rho_i \ln(\rho_i)$. The optimal strategy can then be obtained by solving Equation (RI), which here means investing in asset $\mathbf{a}^1$ with probability 82.5% or 42.0% conditional on a bull or bear market respectively, and otherwise investing in asset $\mathbf{a}^2$. Net of costs, welfare is equal to $W_S = 3.21$.

Granted, an investor who initially suspects a bear market might not be very impressed by α . It is true that it would not bring any welfare gains for her, but she would certainly not be willing to purchase it unconditionally. If she is pessimistic enough, her own ignorance equivalent may well be counter-cyclical, such as $\tilde{\alpha}$ in [Figure 1\(c\)](#). We trace the collection of ignorance equivalents across all prior beliefs in the dashed line, and call it the *learning-proof menu* under the particular learning technology.⁴ This prior-free concept allows us to recast the original RI problem as a choice over a richer menu without learning. The key observation is that the investor is indifferent between the original and the learning-proof menu at all beliefs — and since unconditional choice is always optimal in the latter, learning becomes mute and the agent’s problem effectively turns into a ‘fixed-information’ expected utility maximization over the learning-proof menu.

A useful summary statistic. The learning-proof representation succinctly summarizes how an RI agent’s welfare and behavior respond to changes in the decision problem. In that sense, the ignorance equivalent is reminiscent of the certainty equivalent for choice problems under risk.⁵ Both concepts preserve key properties of the original problem but reduce its complexity. Yet contrary to the latter, the ignorance equivalence retains state-dependency in the payouts. This is helpful in situations with learning where the tailoring of actions to states takes center stage.

For instance, a researcher may want to know whether the investor will adjust her placement strategy if a new asset becomes available. A priori, the researcher would have to analyze the welfare impact of complex learning strategies that involve all available assets. As it turns out ([Theorem 3](#)), it is sufficient to consider only those that rely on the original ignorance equivalent and the new asset — resulting in a much more tractable condition. Conversely, if the researcher would like to determine whether the new asset will be chosen always (rather than sometimes), it is sufficient to check that the investor chooses it unconditionally against every single alternative ([Theorem 5](#)). If so, then the same is true when all assets are simultaneously available. This ability to extrapolate from binary menus to complex learning strategies significantly enhances tractability, particularly for rich menus of options. It also simplifies the computation of the ignorance equivalent itself ([Section 5.3](#)).

If the researcher would instead like to know which of the *existing* assets might be chosen in an expanded menu, he should look at the learning-proof menu. Quite simply, assets

⁴Colors indicate the assets chosen by an agent with a particular ignorance equivalent.

⁵In [Figure 1](#), the certainty equivalent of a single action $\mathbf{a}$ is located at the intersection between the diagonal (certain payoffs) and the hyperplane through $\mathbf{a}$ that is orthogonal to the prior (same expected utility). (For a risk averse agent, one would further need to invert the utility value to translate payoffs into payouts.) The certainty equivalent of α , for instance, is equal to W_S .

belonging to the learning-proof menu will become part of the optimal stochastic strategy if the right new asset is made available; all other assets are ignored in any expanded menu (Theorem 4). In Figure 1(c), this implies that asset $\mathbf{a}^3$ can become optimal in an expanded menu, but an asset below the dashed line would never be chosen, even if it is more attractive than all other individual assets under some beliefs.⁶

The learning-proof menu is also helpful to study the behavioral consequences of changes in the prior belief. We show that if two investors share ignorance equivalents, they also share the same optimal state-dependent investment strategies over the original menu (Corollary 2). In two dimensions, this result may seem trivial; after all, the investment strategy (and the ignorance equivalent) are unresponsive to local changes in belief only when there is no learning to begin with. More generally though, the ignorance equivalent remains locally constant in some directions whenever the agent relies on fewer actions than states, a common feature of many RI problems [Caplin et al., 2018]. The comparative statics of the ignorance equivalent thus identify which local information shocks would alter the behavior of the RI agent.

In terms of welfare, the learning-proof menu helps determine the value that an agent assigns to an exogenous signal structure. This value depends not just on the precision of the signal but also on whether the new information complements or replaces the agent’s own learning, which is determined endogenously by the menu at her disposal. For instance, even an informative signal may not induce any changes in the ignorance equivalent and thus may have no effect on her optimal behavior. In that sense, the local geometry of the learning proof menu characterizes which information is ‘actionable’ for the RI agent.

In addition, ignorance equivalence has relevance for incentive design. For instance, it informs a principal on how to design a menu of state-dependent transfers that incentivize an RI agent to acquire and disclose an arbitrary signal. The principal’s expected cost for this ‘indirect’ learning differs from the agent’s cost by only a constant that reflects outside options. And if the principal does *not* want the agent to learn (for example, because he wants to elicit the agent’s unknown prior as in Tsakas [2020]), he should design state-dependent transfers such that the resulting menu is learning-proof.

An equilibrium outcome. The ignorance equivalent takes on an important strategic role in multiplayer settings with learning and surplus splitting. To illustrate, we widen the lens of our example to include a professional wealth manager who can freely learn the market

⁶A point like $\boldsymbol{\alpha} - \varepsilon \mathbf{1}$ with small enough $\varepsilon > 0$ yields higher expected profit than any other asset under the prior, yet it would never be selected when the agent has access to (at least) menu $\{\mathbf{a}^1, \mathbf{a}^2, \mathbf{a}^3\}$.

conditions.⁷ If the manager places the wealth on the investor’s behalf, he can achieve the maximal payout $\mathbf{M}$. In exchange for the management opportunity, he can commit to any state-dependent transfers $\mathbf{t}$ once the payouts realize. From the investor’s perspective, the transfer offer generates another investment opportunity, since she stands to gain t_i in market condition i by entrusting her wealth to the manager. The remainder $\mathbf{M} - \mathbf{t}$ denotes the state-dependent surplus that accrues to the manager, thus one can imagine the manager looking at a rotated Figure 1(b) with origin $\mathbf{M}$, much like in a standard Edgeworth box.

A naive manager may be tempted to offer a transfer that guarantees the investor her expected autarky payout, depicted by point $\mathbf{W}_S$. However, this will not lead to the desired outcome of unconditional acceptance. The addition of $\mathbf{W}_S$ enriches the set of available assets, and the investor will adapt to this new menu. She will realize that outright acceptance of the manager’s terms is not in her best interest. Indeed, she is better off by maintaining her previous learning strategy and accepting the management contract only when she would have otherwise invested in the safe asset.⁸ This is bad news for the manager, for two reasons: First, he frequently misses out on the business opportunity. Second, conditional on receiving the offer, he now predominantly finds himself in a bear market, where his profit is in fact negative. In short: The learning ability of the investor generates an adverse selection problem for the manager.

Even if the manager tailors the transfer promise to his private knowledge of the state or uses information design to pre-empt agent learning, he cannot escape adverse selection. After all, any such strategy would merely reveal free information in equilibrium, thus increasing the ‘autarky welfare’ of the investor and reducing the manager’s profit in this zero-sum game.

Still, there is one way of combining maximal profits and unconditional acceptance: If the manager offers the state-dependent transfer $\boldsymbol{\alpha}$, the investor is willing to accept his offer outright. Together, they can then achieve the socially optimal payout $\mathbf{M}$. Now, the ignorance equivalent is no longer pure fiction: It emerges as the unique Bayesian Equilibrium transfer in the two-player game between investor and manager.

Even if the manager does not know the investor’s prior belief (for instance, because she might have been subject to a private information shock), he can extract the full surplus by offering the entire learning-proof menu. The investor is no better off than under autarky, but is now always willing to forgo learning in favor of the appropriate ignorance equivalent transfer.

Things can get more intricate in a higher-dimensional setting where the investor has a

⁷Free learning simplifies exposition, but the intuition generalizes to settings with costly learning by both agents and ‘restricted’ assets that only the manager has access to (see Section 6).

⁸She could do even better by adapting her learning strategy. In the new optimum, she hires the manager with (marginal) probability 52.8% and invests in asset $\mathbf{a}^1$ otherwise.

comparative learning advantage relative to the manager for some specific sectors, or the investor has access to restricted assets (e.g. company stock options) that are unavailable to the manager. In these situations, the first best solution requires collaboration that goes beyond a simple trade and sometimes involves learning by both agents. As we shall show in [Section 6](#), ignorance equivalence is still the key to achieving maximal welfare. Repeated trade at (appropriately chosen) ignorance equivalent terms aligns agents' incentives towards the socially optimal type of learning.

Paper structure. [Section 2](#) clarifies notation and recaps the standard Rational Inattention problem. [Section 3](#) introduces the ignorance equivalent and its key properties, and [Section 4](#) defines the learning-proof menu. [Section 5](#) and [Section 6](#) show how these tools are useful to analyze behavior in situations with a single or multiple rationally inattentive agents, respectively. [Section 7](#) discusses the related literature and [Section 8](#) concludes. All proofs are in the Appendix.

2 Rational Inattention Problem

The rationally inattentive decision maker has to implement an action from the finite menu $\mathcal{A}$.⁹ Payoffs of each action depend on an unknown state of the world $i \in \mathcal{I} := \{1, \dots, I\}$. The ex-ante likelihood of all states are captured by a full support prior $\boldsymbol{\pi} \in \text{int}(\Delta\mathcal{I})$. No two actions are payoff equivalent, and we identify an action $\mathbf{a} \in \mathcal{A}$ by its state-dependent payoffs $(a_1, \dots, a_I) \in \mathbb{R}^I$.

The agent can condition her choice on the outcome of a costly signal $\mathcal{S} = \langle S, \mathbf{q} \rangle$, where S refers to a finite signal realization space and $\mathbf{q} \in (\Delta S)^{\mathcal{I}}$ denotes the conditional probabilities $q_i(s)$ of realization s in state i . Upon observing a signal realization s , the agent updates her belief to $\boldsymbol{\pi}^s$ according to Bayes' rule and selects a utility-maximizing action in $\arg \max_{\mathbf{a} \in \mathcal{A}} \boldsymbol{\pi}^s \cdot \mathbf{a}$. We denote the cost of signal $\mathcal{S}$ under prior belief $\boldsymbol{\rho} \in \Delta\mathcal{I}$ by $c(\mathcal{S}, \boldsymbol{\rho}) \in [0, \infty)$, and discuss admissible cost functions in Subsection [2.1](#).

The welfare of a choice problem $(\mathcal{A}, \boldsymbol{\pi}, c)$ is obtained by optimizing over expected consumption utility net of learning costs,

$$W(\mathcal{A}, \boldsymbol{\pi}, c) = \sup_{\mathcal{S}=\langle S, \mathbf{q} \rangle} \sum_{s \in S} (\boldsymbol{\pi} \cdot \mathbf{q}(s)) \left(\max_{\mathbf{a} \in \mathcal{A}} \boldsymbol{\pi}^s \cdot \mathbf{a} \right) - c(\mathcal{S}, \boldsymbol{\pi}). \quad (\text{RI})$$

Since additional actions weakly increase achievable utility $\max_{\mathbf{a}} \boldsymbol{\pi}^s \cdot \mathbf{a}$, welfare is menu-

⁹Our notation carries over to compact menus, as long as learning is restricted to finite signals. For the cost functions that we consider, this is without loss of generality ([Lemma A.1](#)).

monotone, $W(\mathcal{A}, \boldsymbol{\pi}, c) \leq W(\mathcal{A}', \boldsymbol{\pi}, c)$ whenever $\mathcal{A} \subseteq \mathcal{A}'$. And since consumption utility and learning costs interact additively, optimal learning is unaffected when all payoffs are shifted from $\mathcal{A}$ to $\mathcal{A} + \{\mathbf{v}\}$ by the same vector $\mathbf{v} \in \mathbb{R}^I$.

To simplify language, we say that an agent “follows learning strategy $\mathcal{S} = \langle \mathcal{A}, \mathbf{q} \rangle$ ” if each signal realization identifies a recommended action that the agent implements. We say that “unconditional implementation of $\mathbf{a}$ is optimal” if the degenerate learning strategy with $\mathbf{q}(\mathbf{a}) = \mathbf{1}$ is optimal, and that “it is optimal to implement $\mathbf{a}$ with positive probability” if at least one learning strategy with $\mathbf{q}(\mathbf{a}) > \mathbf{0}$ is optimal.¹⁰

2.1 Admissible learning costs

All of our results hold when the agent’s cost function is uniformly posterior separable [Caplin et al., 2022] and prior-concave,¹¹ with

$$c(\langle S, \mathbf{q} \rangle, \boldsymbol{\pi}) = \sum_{s \in S} (\boldsymbol{\pi} \cdot \mathbf{q}(s)) \phi(\boldsymbol{\pi}^s) - \phi(\boldsymbol{\pi}) \quad (\text{UPS+})$$

for a differentiable, convex potential function $\phi : \Delta\mathcal{I} \rightarrow \mathbb{R}$. Within that class, leading examples are Mutual Information [Sims, 2003], some variants of the Tsallis costs [Caplin et al., 2022] and Total Information [Bloedel and Zhong, 2020], which subsumes the Wald cost by Morris and Strack [2019] and the Fisher Information of Hébert and Woodford [2020].

Axiomatization. To provide meaningful intuition, we do not derive our results directly from the (UPS+) specification. Instead, we base our proofs and discussions on five economically relevant properties of the learning cost. As we show in Lemma A.2, they are jointly equivalent to (UPS+); but most results hold even without the last property.

(C1) Cost is continuous: For any signal space S and any cutoff $\hat{c} \in \mathbb{R}$, the pre-images $\{(\mathbf{q}, \boldsymbol{\pi}) \in (\Delta S)^{\mathcal{I}} \times \Delta\mathcal{I} \mid c(\langle S, \mathbf{q} \rangle, \boldsymbol{\pi}) \geq \hat{c}\}$ are open.

Continuity is instrumental in ensuring that the choice problem (RI) admits an optimal solution, and that welfare is continuous in the prior belief (Lemma A.6).

(C2) Free information disposal: Cost function $c(\cdot, \boldsymbol{\pi})$ is non-decreasing in the Blackwell order and $c(\mathcal{S}, \cdot)$ is weakly concave in the prior for all $\mathcal{S}$.

¹⁰Throughout, we use the convention that $\mathbf{v} \geq \mathbf{w}$ if and only if $v_i \geq w_i \forall i$, that $\mathbf{v} > \mathbf{w}$ if and only if $\mathbf{v} \geq \mathbf{w}$ and $\mathbf{v} \neq \mathbf{w}$, and that $\mathbf{v} \gg \mathbf{w}$ if and only if $v_i > w_i \forall i$.

¹¹For twice-differentiable potentials, Bloedel and Zhong [2020] provide a test for prior-concavity that is based on the Hessian of ϕ .

Condition (C2) is best explained by introducing an ‘assistant’ to the decision maker. For Blackwell monotonicity,¹² suppose the assistant draws a signal and then, depending on its realization, communicates a garbled message to the agent. Since the garbling is uncorrelated with the state, the agent is weakly less informed than the assistant, and thus should incur a weakly lower cost. Blackwell monotonicity gives rise to the well known ‘obedience principle,’ which allows us to restrict attention to learning strategies (Lemma A.3). For prior-concavity, suppose the assistant privately observes a free signal about the state. The concavity of $c(\mathcal{S}, \cdot)$ then simply states that it should be no more expensive (in expectation) for the agent to implement a particular learning strategy $\mathcal{S}$ with access to the assistant’s information than it is without.¹³ Prior-concavity implies in particular that welfare W is convex in the prior belief (Lemma A.4).

The next property considers an RI agent who faces a binary choice between payoff vector $\mathbf{a}$ and an outside option $\mathbf{0}$. Both options are equally attractive ex ante, but the choice is payoff-relevant, $\boldsymbol{\pi} \cdot |\mathbf{a}| > 0$.

(C3) Ties are broken through learning: For any belief $\boldsymbol{\pi} \in \Delta\mathcal{I}$ and any payoff vector $\mathbf{a} \in \mathbb{R}^I$ that has zero expected utility $\boldsymbol{\pi} \cdot \mathbf{a} = 0$ but is nonzero with positive probability, $\boldsymbol{\pi} \cdot |\mathbf{a}| > 0$, there exists a binary signal $\mathcal{S} = \langle \{0, 1\}, \mathbf{q} \rangle$ whose cost is below the expected benefit, $c(\mathcal{S}, \boldsymbol{\pi}) < \sum_{i \in \mathcal{I}} \pi_i q_i(1) a_i$.

The property asserts that some learning is optimal, no matter how small the stakes $\mathbf{a}$ or how extreme the belief $\boldsymbol{\pi}$. This rules out fixed costs for learning and ensures that the agent can randomize across actions at no cost (Lemma A.5). Since the choice of $\mathbf{a}$ is arbitrary, the condition also guarantees that there exists a noisy enough signal whose expected benefit outweighs the cost by any arbitrary factor (Lemma A.8).

The final two properties compare costs across sequential learning strategies, where the agent first draws a signal $\mathcal{S} = \langle S, \mathbf{q} \rangle$ and, upon observing $s \in S$, draws a second signal $\mathcal{S}^s = \langle S^s, \mathbf{q}^s \rangle$. We denote this contingency plan as $(\mathcal{S}, \{\mathcal{S}^s\})$, and write $\mathcal{S} \times \{\mathcal{S}^s\}$ for the one-shot implementation of the same process.¹⁴ It is appealing to assume that $c(\mathcal{S}, \boldsymbol{\pi})$ is already the cost-minimizing envelope over all sequential information strategies [Bloedel and Zhong, 2020], thus avoiding more cumbersome notation.

(C4) Sequential information acquisition brings no cost savings: In expectation, any contingency plan $(\mathcal{S}, \{\mathcal{S}^s\})$ costs no less than the one-shot signal $\mathcal{S} \times \{\mathcal{S}^s\}$, $c(\mathcal{S}, \boldsymbol{\pi}) + \sum_{s \in S} (\boldsymbol{\pi} \cdot$

¹²Signal $\mathcal{S} = \langle S, \mathbf{q} \rangle$ is Blackwell more informative than signal $\tilde{\mathcal{S}} = \langle \tilde{S}, \tilde{\mathbf{q}} \rangle$ if for each $s \in S$, there exists a lottery $m^s \in \Delta \tilde{S}$ such that $\tilde{\mathbf{q}} = \sum_{s \in S} m^s \mathbf{q}(s)$ [Blackwell, 1953].

¹³The link between prior-concavity and the agent’s value of information has previously been discussed by Denti et al. [2022] and Bloedel and Zhong [2020].

¹⁴Formally, $\mathcal{S} \times \{\mathcal{S}^s\}$ is the signal $\langle S \times \bigcup_{s \in S} S^s, \tilde{\mathbf{q}} \rangle$ with $\tilde{q}_i((s, \tilde{s})) = q_i(s) q_i^s(\tilde{s})$.

$$q(s)) c(\mathcal{S}^s, \pi^s) \geq c(\mathcal{S} \times \{\mathcal{S}^s\}, \pi).$$

The property yields an optimal stopping result [Hébert and Woodford, 2019, Zhong, 2019]: If the agent implements $\mathbf{a} \in \mathcal{A}$ at some posterior $\boldsymbol{\rho}$, then she would do so unconditionally if her prior were $\boldsymbol{\rho}$ (Lemma A.7).

The opposite inequality feels more restrictive.

(C5) Sequential information acquisition incurs no extra costs: In expectation, any contingency plan $(\mathcal{S}, \{\mathcal{S}^s\})$ costs no more than the one-shot signal $\mathcal{S} \times \{\mathcal{S}^s\}$ $c(\mathcal{S}, \pi) + \sum_{s \in \mathcal{S}} (\pi \cdot q(s)) c(\mathcal{S}^s, \pi^s) \leq c(\mathcal{S} \times \{\mathcal{S}^s\}, \pi)$.

We do not mention this property anywhere in the proofs for Sections 3 and 4, making these results potentially more generic.

3 Ignorance Equivalent

The central concept of our paper is the notion of the ignorance equivalent. The ignorance equivalent of an RI problem $(\mathcal{A}, \pi, c)$ is a payoff vector $\boldsymbol{\alpha} \in \mathbb{R}^I$ that, when treated as a fictitious action, leaves the agent no worse off as a replacement of menu $\mathcal{A}$ and yet delivers no welfare gains as an addition to the menu $\mathcal{A}$.

Definition 1. The payoff vector $\boldsymbol{\alpha} \in \mathbb{R}^I$ is an **ignorance equivalent** of the RI problem $(\mathcal{A}, \pi, c)$ if and only if

$$W(\{\boldsymbol{\alpha}\}, \pi, c) \geq W(\mathcal{A}, \pi, c) \quad \text{and} \quad W(\mathcal{A}, \pi, c) \geq W(\mathcal{A} \cup \{\boldsymbol{\alpha}\}, \pi, c).$$

Intuitively, the first condition means the agent would be willing to commit to *always* implement $\boldsymbol{\alpha}$, forgoing any learning opportunities that are present in $\mathcal{A}$. The second condition means that she would also commit to *never* implement $\boldsymbol{\alpha}$, forgoing any learning opportunities that arise when $\boldsymbol{\alpha}$ is added to the original menu. Together, the two conditions imply that $W(\{\boldsymbol{\alpha}\}, \pi, c) \geq W(\mathcal{A} \cup \{\boldsymbol{\alpha}\}, \pi, c)$, but since larger menus weakly raise welfare, all inequalities are binding. In this sense, the payoff vector $\boldsymbol{\alpha}$ is such that the agent can forgo all learning opportunities, old or new, without loss or gain. It is thus appropriately called an ‘ignorance equivalent.’

The ignorance equivalent is reminiscent of the certainty equivalent for lotteries. Neither is typically available to the agent, unless we are considering degenerate lotteries or a decision problem where no learning is optimal. Yet, both concepts allow us to abstract from the crux of the underlying economic problem (risk, learning) to reduce its complexity, all while

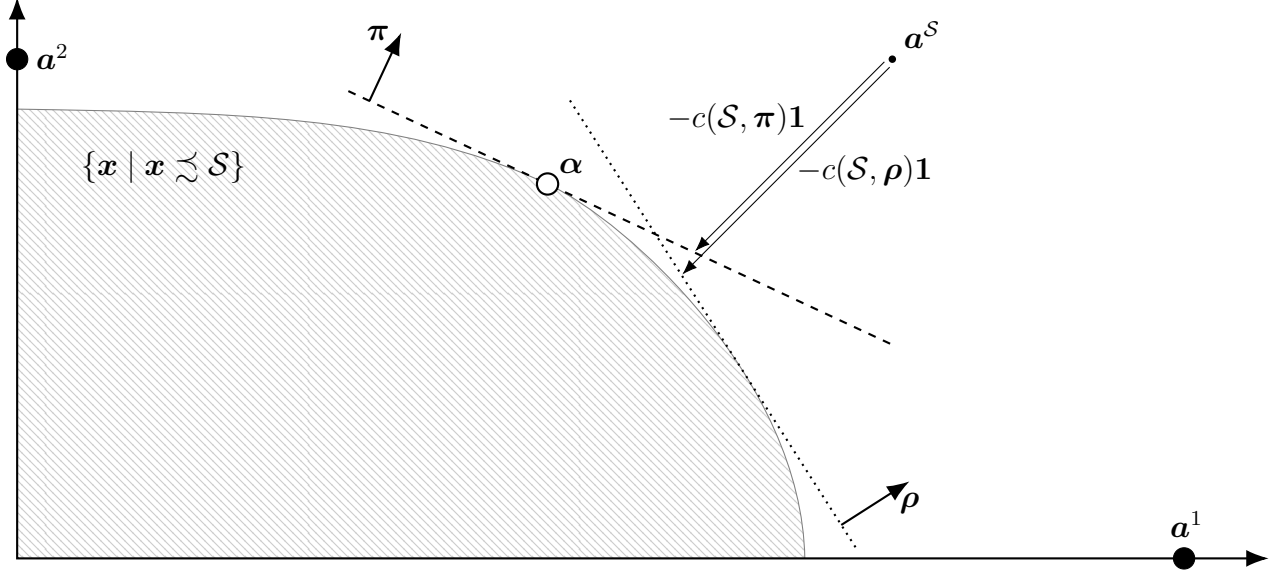


Figure 2: Dominated payoff vectors and construction of the ignorance equivalent.

preserving its key properties to enable comparative statics. The similarity between the two concepts extends to their construction: Just like the certainty equivalent is equal to the highest unconditional payment that is dominated by the lottery, we now show that the ignorance equivalent is equal to the payoff vector with the highest expected utility that is dominated by an optimal learning strategy.

Formally, this is what we mean with dominance by a learning strategy: An agent with belief $\rho \in \Delta\mathcal{I}$ obtains expected utility $\rho \cdot x$ from implementing action x unconditionally and incurs no learning costs. If instead the agent follows the learning strategy $\mathcal{S} = \langle \mathcal{A}, q \rangle$, she achieves expected consumption utility $a_i^S = \sum_{a \in \mathcal{A}} q_i(a) a_i$ in each state i . Welfare is obtained by weighing these state-wise expectations by the agent's belief and subtracting the signal cost. Payoff vector x is dominated by learning strategy $\mathcal{S}$ if the latter yields weakly larger welfare for any belief ρ . As we shall see later, the reference to other beliefs is what rules out profitable new learning opportunities that may arise from the addition of x to the menu.

Definition 2. Payoff vector x is **dominated** by a learning strategy $\mathcal{S} = \langle \mathcal{A}, q \rangle$, denoted $x \preceq \mathcal{S}$, if and only if

$$\rho \cdot x \leq \rho \cdot a^S - c(\mathcal{S}, \rho) \quad \forall \rho \in \Delta\mathcal{I}.$$

Figure 2 sketches a sample RI problem $(\mathcal{A}, \pi, c)$ with two states (on either axis) and two actions (a^1 and a^2) to illustrate the concept of dominance. The learning strategy $\mathcal{S}$ implements action a^2 with probability $1/4$ in state 1 and with probability 1 in state 2, leading to expected consumption utility $a_1^S = \frac{1}{4}a_1^2 + \frac{3}{4}a_1^1$ and $a_2^S = a_2^2$. For each belief, the

signal cost determines the maximal expected utility for any dominated payoff vector. We indicate this upper bound for prior π as a dashed line and for belief ρ as a dotted line. The intersection of all such lower half-spaces forms the set of dominated payoff vectors, which is thus naturally closed, convex and bounded above. The concavity of $c(\mathcal{S}, \cdot)$ determines the curvature of the boundary, and the cost $c(\mathcal{S}, \cdot)$ determines its distance to $\mathbf{a}^{\mathcal{S}}$. As a special case of interest, Total Information [Bloedel and Zhong, 2020] describes (UPS+) costs that are linear in the prior, $c(\mathcal{S}, \pi) = \pi \cdot \Delta^{\mathcal{S}}$ for an appropriately defined $\Delta^{\mathcal{S}} \in [0, \infty)^I$. In this case, dominance $\mathbf{x} \preceq \mathcal{S}$ is simply equivalent to the vector inequality $\mathbf{x} \leq \mathbf{a}^{\mathcal{S}} - \Delta^{\mathcal{S}}$ and the boundary is ‘maximally curved’ around a single point, so much so that it is actually nondifferentiable.

Our first result shows how to obtain the ignorance equivalent of $(\mathcal{A}, \pi, c)$ from any optimal signal using the dominance relationship, and establishes the existence and uniqueness of the ignorance equivalent.

Theorem 1. *Each RI problem $(\mathcal{A}, \pi, c)$ admits a unique ignorance equivalent $\alpha \in \mathbb{R}^I$. It is obtained from any optimal learning strategy $\mathcal{S}$ as the payoff vector that maximizes expected utility over all $\mathcal{S}$ -dominated payoff vectors,*

$$\alpha = \arg \max_{\mathbf{x} \preceq \mathcal{S}} \pi \cdot \mathbf{x}. \quad (1)$$

Three arguments are key to the result:

First, unconditional implementation of the best $\mathcal{S}$ -dominated payoff vector α is as attractive as following strategy $\mathcal{S}$. We use a duality argument to establish that the set of dominated payoff vectors $\{\mathbf{x} \mid \mathbf{x} \preceq \mathcal{S}\}$ in Figure 2 touches the dashed line that describes the net utility from implementing $\mathcal{S}$ under prior π .

Second, dominance $\alpha \preceq \mathcal{S}$ is sufficient to rule out new learning opportunities that arise from adding α to the original menu $\mathcal{A}$: Whenever a potential strategy over menu $\mathcal{A} \cup \{\alpha\}$ recommends implementation of α at some posterior ρ , the agent is just as well off by instead continuing with strategy $\mathcal{S}$ and relying only on actions in menu $\mathcal{A}$. By sequential optimality (C4), the agent can thus achieve at least as much welfare by restricting attention to menu $\mathcal{A}$, and $W(\mathcal{A} \cup \{\alpha\}, \pi, c) = W(\mathcal{A}, \pi, c)$. Together, these two steps construct a payoff vector that satisfies both conditions of Definition 1 and thus establish existence of the ignorance equivalent.

The final step establishes that the ignorance equivalent is unique, even if the (RI) problem admits multiple optimal signals. The proof relies heavily on the flexibility of learning, which drives the agent to break all ties through learning (C3). This property allows us to construct a welfare-enhancing sequential strategy from any two distinct candidate ignorance equivalents.

And by a similar sequential learning argument as above, these welfare gains could be realized even when only one of the ignorance equivalents is added to the menu.

Necessity of dominance. [Theorem 1](#) implies that imposing dominance is not just sufficient to rule out learning opportunities – it is also necessary.¹⁵ It may surprise that there is no loss in requiring the optimal signal to be preferable to the unconditional implementation of α for *all* beliefs, including those occurring with zero probability. Intuition can be provided in two steps: First, prior-concavity [\(C2\)](#) relates welfare gains at ‘far away’ and ‘nearby’ posteriors, so that local dominance becomes synonymous with global dominance. Indeed, if $\mathbf{x}$ achieves the same welfare as $\mathcal{S}$ under prior π and is preferable to $\mathcal{S}$ at *some* posterior ρ , then it also delivers welfare gains at the local perturbation $\rho^\varepsilon := (1 - \varepsilon)\pi + \varepsilon\rho$ for arbitrarily small $\varepsilon > 0$.¹⁶ Second, the flexibility of learning [\(C3\)](#) ensures that the agent could build a profitable learning strategy from *any* local welfare gains, even if they occur in a direction in which the agent does not, currently, update her beliefs.

Connection with optimal strategies. While the ignorance equivalent is always unique for any RI problem $(\mathcal{A}, \pi, c)$, there may exist multiple optimal learning strategies. Nevertheless, the construction in [Theorem 1](#) establishes a one-to-one relationship between the agent’s ignorance equivalent α and the set of all optimal strategies, which are exactly those that dominate α .

Continuity of the ignorance equivalent is a direct consequence of this one-to-one correspondence. Since the optimal RI signals are upper hemicontinuous by Berge’s Theorem ([Lemma A.6](#)), the ignorance equivalent is continuous in the prior.

Corollary 1. *The mapping $\pi \mapsto \alpha^{(\mathcal{A}, \pi, c)}$ is continuous at any prior $\pi \in \text{int}(\Delta\mathcal{I})$.*

The link between the ignorance equivalent and the set of optimal learning strategies also has implications for comparisons across RI agents with different priors. Since the dominance relationship does not depend on the prior, any two RI problems $(\mathcal{A}, \pi, c)$ and $(\mathcal{A}, \pi', c)$ that share the same ignorance equivalent also share all optimal learning strategies.

Corollary 2. *If the ignorance equivalent of RI problems $(\mathcal{A}, \pi, c)$ and $(\mathcal{A}, \pi', c)$ is the same, then so is the set of optimal learning strategies.*

¹⁵In situations with multiple optimal signals, [Theorem 1](#) requires that *each* of them dominates α .

¹⁶Formally, $\rho^\varepsilon \cdot \mathbf{x} = \varepsilon\rho \cdot \mathbf{x} + (1 - \varepsilon)\pi \cdot \mathbf{x}$. If $\mathbf{x}$ is strictly preferable to $\mathcal{S}$ under prior ρ and as good as $\mathcal{S}$ under π , the expected payoff under ρ^ε is strictly larger than $\varepsilon[\rho \cdot \mathbf{a}^\mathcal{S} - c(\mathcal{S}, \rho)] + (1 - \varepsilon)[\pi \cdot \mathbf{a}^\mathcal{S} - c(\mathcal{S}, \pi)]$, which by prior-concavity [\(C2\)](#) is at least $\rho^\varepsilon \cdot \mathbf{a}^\mathcal{S} - c(\mathcal{S}, \rho^\varepsilon)$.

Even though their ignorance equivalents coincide, the posterior beliefs of the agents will still diverge in general – but they seek out the same information sources and display the same state-dependent stochastic behavior.

Relation to the posterior-based approach. The RI literature typically views decision problems through a belief-based lens, similar to Bayesian Persuasion [Kamenica and Gentzkow, 2011]. This approach expresses learning as Bayes-plausible distributions over posterior beliefs, and states Equation (RI) equivalently as

$$W(\mathcal{A}, \pi, c) - \phi(\pi) = \sup_{\substack{\mu \in \Delta \Delta \mathcal{I} \\ \mathbb{E}[\mu] = \pi}} \int_{\Delta \mathcal{I}} (\max_{\mathbf{a} \in \mathcal{A}} \boldsymbol{\rho} \cdot \mathbf{a} - \phi(\boldsymbol{\rho})) d\mu(\boldsymbol{\rho}).$$

Visually, the optimal behavior is characterized implicitly through the concave upper envelope V over the ‘net payoff’ functions $\boldsymbol{\rho} \mapsto \boldsymbol{\rho} \cdot \mathbf{a} - \phi(\boldsymbol{\rho})$ associated with each action $\mathbf{a}$ [Caplin and Dean, 2013]. The optimal posteriors are then identified via the tangency points of V and the individual net payoff curves, as done in Figure 3 for the investment choice from in the introduction.

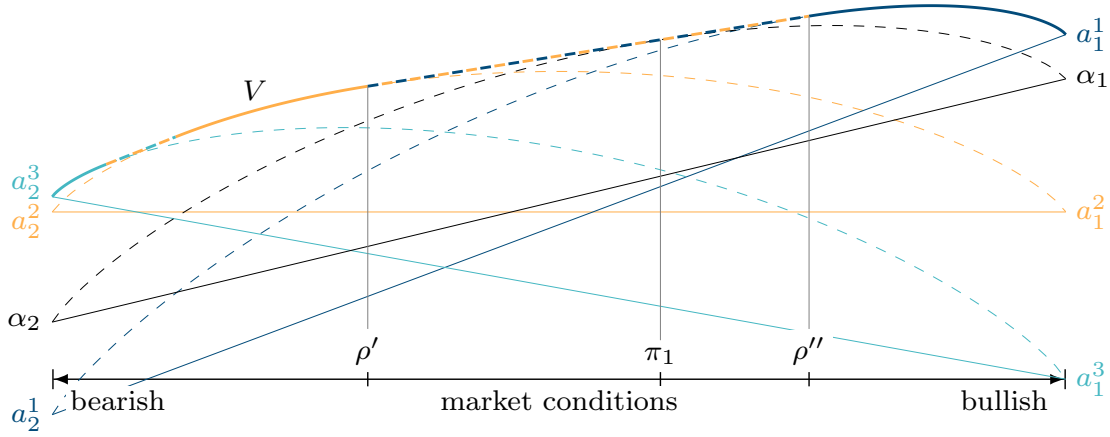


Figure 3: The introductory investor example through the posterior-based lens. Monochrome dashed lines represent the net payoff for each action $\mathbf{a}$.

In this view, the net payoff of the ignorance equivalent (depicted as a black dashed line) is tangent with V at π . It coincides with the upper envelope at the prior, implying that unconditional implementation is just as attractive as the original menu. Yet it also lies entirely below the upper envelope, so its availability does not raise welfare.¹⁷ Theorem 1 can be seen as the dual of the posterior-based approach. Concavification determines the least

¹⁷A similar construction (based on Caplin et al. [2022]) establishes existence and uniqueness of the ignorance equivalent for generic Posterior-Separable costs, as done in Appendix A.4. There, we also provide intuition for some of the other results from this vantage point. Our focus remains on (UPS+) costs primarily

upper bound to $W(\mathcal{A}, \boldsymbol{\pi}, c) - \phi(\boldsymbol{\pi})$ for each belief, while Equation (1) maximizes expected utility over the set of dominated payoff vectors.

Corollary 2 may, at first, remind of the well-known ‘locally invariant posterior’ property [Caplin et al., 2022]: Under moderate changes in the prior, the investor still picks the same assets $\mathbf{a}^2$ and $\mathbf{a}^1$ at the posterior beliefs ρ' and ρ'' , respectively. However, note that the optimal stochastic behavior changes continuously over that interval of beliefs: As the prior approaches ρ'' , asset $\mathbf{a}^1$ is picked more frequently in each state of the world. This change is reflected in the ignorance equivalent; it also approaches $\mathbf{a}^1$. Instead, Corollary 2 speaks to situations where the ignorance equivalent and the stochastic behavior remain constant, even though beliefs change. In Figure 3, this happens for priors to the right of ρ'' where $\mathbf{a}^1$ is implemented unconditionally. For a less trivial example, suppose that there is a third state, occurring with probability q , where all assets pay out zero. It can be shown that so long as the prior likelihood ratio between a bull and bear market is maintained, the investor’s optimal stochastic behavior is independent of q under Shannon entropy costs.¹⁸ The absence of behavioral consequences is reflected in the common ignorance equivalent across these priors. Nevertheless, q is not altogether inconsequential: It affects the investor’s posteriors, welfare, and the cost of the optimal signal.

Self-selection. By definition, the ignorance equivalent generates no additional learning under the agent’s prior. Theorem 1 further implies that the ignorance equivalent must be dominated, and thus it generates no additional learning opportunities under any prior, regardless of whether it is the one it was designed for. We refer to this result as the ‘self-selection’ property of the ignorance equivalent for the following reason: Suppose two RI agents with different priors face the same menu and cost function. Adding both ignorance equivalents to the menu would not be welfare-enhancing for either agent, yet both would now be willing to forgo learning by unconditionally implementing their respective ignorance equivalent.

Corollary 3. *Let $\boldsymbol{\alpha}$ denote the ignorance equivalent of RI problem $(\mathcal{A}, \boldsymbol{\pi}, c)$. For any belief $\boldsymbol{\rho} \in \Delta\mathcal{I}$, the ignorance equivalent of $(\mathcal{A}, \boldsymbol{\rho}, c)$ is equal to that of $(\mathcal{A} \cup \{\boldsymbol{\alpha}\}, \boldsymbol{\rho}, c)$.*

Self-selection is a direct consequence of dominance, $\boldsymbol{\alpha} \preceq \mathcal{S}$. Dominance implies that no matter the agent’s interim belief, continuing with strategy $\mathcal{S}$ is always at least as attractive as

because it is ultimately the comparisons across beliefs that makes the ignorance equivalent into a ‘summary’ of the learning opportunities embedded in the RI decision problem.

¹⁸For the specification in Footnote 3, she acquires just enough information to acquire asset $\mathbf{a}^1$ with probability 82.5% (42%) in a bull (bear) market, and 66.3% in the ‘third state’, irrespective of that state’s prior likelihood. Otherwise, she acquires asset $\mathbf{a}^2$.

implementing α . An agent who follows the optimal strategy for $(\mathcal{A} \cup \{\alpha\}, \rho, c)$ can achieve the same welfare in $(\mathcal{A}, \rho, c)$ by simply continuing with $\mathcal{S}$ whenever her strategy calls for α . For that reason, adding ignorance equivalents does not distort the choice. All strategies that were once optimal remain so (potentially alongside newly feasible strategies).

On a technical note, our uniqueness result relies on the full-support assumption of the prior. For beliefs on the boundary $\rho \in \Delta\mathcal{I} \setminus \text{int}(\Delta\mathcal{I})$, the ignorance equivalent is unique with probability one but can have arbitrary payoffs in the zero-probability states,¹⁹ and [Corollary 3](#) should be understood to say that the *set* of ignorance equivalents is unchanged when α is added to the menu.

4 Learning-Proof Menu

By design, the agent’s learning ability becomes obsolete once her ignorance equivalent is added to the menu. The agent is no better off than she was before the menu expansion — but she is now also no worse off if she were to lose access to her learning technology. The collection of ignorance equivalents across all priors thus forms a fictitious menu that captures all the payoff possibilities of the RI agent regardless of her belief.

Definition 3. Letting $\alpha^{(\mathcal{A}, \pi, c)}$ denote the ignorance equivalent of RI problem $(\mathcal{A}, \pi, c)$, the **learning-proof menu** for menu $\mathcal{A}$ under cost c is given by

$$\bar{\mathcal{A}} := \{\alpha^{(\mathcal{A}, \pi, c)} \mid \pi \in \text{int}(\Delta\mathcal{I})\}.$$

By studying the agent’s choice over this fictitious menu $\bar{\mathcal{A}}$ rather than the original menu $\mathcal{A}$, one can speak to situations with flexible learning without having to account for state-dependent stochastic choices. The learning-proof menu owes much of its appeal to properties that it inherits from its various equivalent characterizations.

Theorem 2. *For any menu $\mathcal{A}$ and cost function c , the learning-proof menu $\bar{\mathcal{A}} \supseteq \mathcal{A}$ is the smallest set such that for any prior $\pi \in \text{int}(\Delta\mathcal{I})$,*

- (a) *Ignorance is an optimal strategy in $(\bar{\mathcal{A}}, \pi, c)$, $W(\bar{\mathcal{A}}, \pi, c) = \max_{a \in \bar{\mathcal{A}}} \pi \cdot a$.*
- (b) *Any signal that is optimal in $(\mathcal{A}, \pi, c)$ is also optimal in $(\bar{\mathcal{A}}, \pi, c)$.*

[Theorem 2](#) formalizes why $\bar{\mathcal{A}}$ is ‘learning-proof’: Regardless of the agent’s prior, learning brings no welfare gains over the best unconditional action choice. At the same time, $\bar{\mathcal{A}}$ does

¹⁹Formally, one can recast the problem by dropping the dimensions with zero probability and redefining costs as the infimum over all signals that are Blackwell-equivalent over the remaining states.

not distort the agent’s choice and maintains all optimal learning strategies over the original menu $\mathcal{A}$.²⁰

The learning-proof menu accurately captures the agent’s payoff possibilities because it can be constructed from the agent’s welfare function. Indeed, each ignorance equivalent is uniquely maximal in some direction $\boldsymbol{\pi} \in \Delta\mathcal{I}$. The learning-proof menu thus represents the upper boundary²¹ of the strictly convex set

$$\bigcap_{\boldsymbol{\pi} \in \Delta\mathcal{I}} \{ \mathbf{x} \in \mathbb{R}^I \mid \boldsymbol{\pi} \cdot \mathbf{x} \leq W(\mathcal{A}, \boldsymbol{\pi}, c) \}. \quad (2)$$

This characterization has direct implications for the comparative statics of $\bar{\mathcal{A}}$: Any change in the problem parameters that raises agent welfare across all priors will move this boundary outward. Examples of such changes include the addition of new actions to the menu $\mathcal{A}$ or reductions in the learning cost c .

Characterization (2) also yields an analogy to standard producer theory: Hotelling’s lemma states that if a firm maximizes expected profit over a strictly convex production set, its optimal production vector is equal to the price-gradient of its profit function. In our case, the agent maximizes expected utility over (the upper boundary of) a strictly convex set of ‘feasible’ payoffs. The optimum identifies the ignorance equivalent, which is thus equal to the gradient of welfare with respect to the prior.

Similarly, the learning-proof menu accounts for all learning opportunities of the agent because it can be constructed from the set of payoff vectors that are dominated by at least one learning strategy,

$$\bigcup_{\mathbf{q} \in (\Delta\mathcal{A})^{\mathcal{I}}} \{ \mathbf{x} \in \mathbb{R}^I \mid \mathbf{x} \preceq \langle \mathcal{A}, \mathbf{q} \rangle \}. \quad (3)$$

By [Theorem 1](#), the upper boundary of this set is also equal to $\bar{\mathcal{A}}$.

Reduced-Form approach. If the learning-proof menu $\bar{\mathcal{A}}$ is known, one can identify the solution to the original (RI) problem in two steps by first locating the ignorance equivalent $\boldsymbol{\alpha}$ through the expected utility maximization problem $\max_{\mathbf{a} \in \bar{\mathcal{A}}} \boldsymbol{\pi} \cdot \mathbf{a}$. [Theorem 1](#) then identifies the full set of optimal learning strategies as exactly those that dominate $\boldsymbol{\alpha}$. The convenience of this approach is that the first step does not have to account for learning at all, and the second step is independent of the prior belief.

This logic may at first sound circular, as we define the learning-proof menu from the igno-

²⁰Since the learning strategies of the original menu remain optimal, the problems $(\mathcal{A}, \boldsymbol{\pi}, c)$ and $(\bar{\mathcal{A}}, \boldsymbol{\pi}, c)$ also share the same ignorance equivalent for any prior $\boldsymbol{\pi}$.

²¹Point $\mathbf{x}$ is part of the upper boundary of $X \subseteq \mathbb{R}^I$ if and only if $\mathbf{x} \in Y$ for any closed superset $Y \supseteq X$ and $\mathbf{x} + \mathbf{y} \notin X$ for any $\mathbf{y} > \mathbf{0}$.

rance equivalent, which in turn is obtained from the optimal learning strategy. In [Sections 5](#) and [6](#), we however show that ignorance equivalent can be constructed from binary strategies only, and that the learning-proof menu emerges endogenously in economic applications. And so to the extent that $\bar{\mathcal{A}}$ can be observed through the equilibrium contract terms, it allows inferences on the agent’s optimal learning or problem fundamentals such as π .

Anchor Actions. The intersection $A = \mathcal{A} \cap \bar{\mathcal{A}}$ contains all actions that are implemented unconditionally under at least one prior. We call these actions *anchors*, because they connect the ‘fictitious’ learning-proof menu $\bar{\mathcal{A}}$ to the ‘physical’ menu $\mathcal{A}$. There are several equivalent characterizations for these actions.

Corollary 4. *Fix a menu $\mathcal{A}$ and a cost function c . For any available action $\mathbf{a} \in \mathcal{A}$, the following are equivalent:*

- (a) *Action $\mathbf{a}$ is an anchor of the learning-proof menu, $\mathbf{a} \in \bar{\mathcal{A}}$.*
- (b) *It is optimal to implement $\mathbf{a}$ unconditionally for some prior $\pi \in \Delta\mathcal{I}$.*
- (c) *It is optimal to implement $\mathbf{a}$ with positive probability for some prior $\pi \in \Delta\mathcal{I}$.*
- (d) *There exists no prior $\rho \in \Delta\mathcal{I}$ such that the ignorance equivalent of $(\mathcal{A}, \rho, c)$ dominates $\mathbf{a}$ statewise, $\alpha^{(\mathcal{A}, \rho, c)} > \mathbf{a}$.*

Characterization (c) points out that the RI agent always restricts her attention to anchor actions, even at priors where ignorance is not an optimal strategy. The literature uses the term *consideration set* [[Caplin et al., 2018](#)] to refer to the (typically small) submenu of actions that are implemented with positive probability. [Corollary 4](#) implies that the union of consideration sets across priors yields exactly the set of anchors.²²

Characterization (d) describes in what sense non-anchor actions $\mathbf{a} \in \mathcal{A} \setminus \bar{\mathcal{A}}$ are suboptimal: It is not just that each RI agent, depending on her prior, finds some other learning strategy more attractive. It is also true that the learning-proof menu contains a payoff vector $\alpha^{(\mathcal{A}, \rho, c)}$ that dominates $\mathbf{a}$ statewise. By [Theorem 1](#), this also implies that there exists a specific learning strategy $\mathcal{S}$ (any one that is optimal under prior ρ) which all agents, irrespective of their prior, strictly prefer to $\mathbf{a}$.

²²The equivalence between (b) and (c) follows directly from the optimal stopping rule ([Lemma A.7](#)) that has previously been derived by [Hébert and Woodford \[2019\]](#), [Zhong \[2019\]](#). We include both characterizations for clarity.

5 Applications with a Single RI Agent

Recasting access to learning as access to a better menu yields new theoretical insights on the comparative statics of choice. In this section, we focus on situations in which a researcher or principal either seeks to learn from or about a single rationally inattentive agent.

5.1 Menu Expansion

When new actions are added to the menu, a rationally inattentive agent re-calibrates her entire learning strategy. As a result, the comparative statics of the consideration set depend in complex ways on the full menu $\mathcal{A}$ and the cost function. Fortunately, the ignorance equivalent and the learning-proof menu bring structure to menu expansion. To determine whether a new action is implemented with positive probability, it is without loss of generality to replace the original menu with its ignorance equivalent.²³

Theorem 3. *Let α denote the ignorance equivalent of RI problem $(\mathcal{A}, \pi, c)$. The following hold for any payoff vector $\mathbf{a}^+ \in \mathbb{R}^I$:*

- (a) $W(\mathcal{A} \cup \{\mathbf{a}^+\}, \pi, c) > W(\mathcal{A}, \pi, c) \iff W(\{\alpha, \mathbf{a}^+\}, \pi, c) > W(\{\alpha\}, \pi, c).$
- (b) $W(\mathcal{A} \cup \{\mathbf{a}^+\}, \pi, c) \geq W(\{\alpha, \mathbf{a}^+\}, \pi, c).$

Since the ignorance equivalent can be derived from the optimal learning strategy by [Equation \(1\)](#), the result also means that unchosen actions do not affect *whether* a new action is attractive. They may however affect *how often* and in which contingencies the new action is implemented. After all, the absolute welfare gains from the new addition can be larger in the full menu (claim b), since the diversity of action payoffs in the menu $\mathcal{A}$ presents more opportunities for profitable learning.

Moreover, anchor actions describe a ‘latent’ consideration set not only because appropriate changes in the prior can induce the agent to implement any anchor action with positive probability ([Corollary 4\(c\)](#)). The same is true if we keep the prior fixed and instead introduce a new action to the menu.

Theorem 4. *Given an RI problem $(\mathcal{A}, \pi, c)$, the following two are equivalent for any action $\mathbf{a} \in \mathcal{A}$:*

- (a) *Action $\mathbf{a}$ is an anchor action $\mathbf{a} \in \mathcal{A} \cap \bar{\mathcal{A}}$.*

²³In the case of Shannon entropy costs, this result is mathematically related to the ‘market entry condition’ in [Caplin et al. \[2018\]](#).

- (b) *There exists a payoff vector $\mathbf{a}^+ \in \mathbb{R}^I$ such that it is optimal to implement $\mathbf{a}$ with positive probability in RI problem $(\mathcal{A} \cup \{\mathbf{a}^+\}, \boldsymbol{\pi}, c)$.*

Matějka and McKay [2015] show by example that adding a new action to the RI agent’s menu may ‘activate’ a previously unchosen action, which now is implemented with positive probability. This activation is a distinguishing feature of Rational Inattention and is absent in fixed-information or random-utility models. Theorem 4 shows that it is exactly the anchor actions that can be activated in this way.

It is worth highlighting that it is too vague to ask whether a novel action $\mathbf{a}^+$ would be ‘attractive’ to an agent facing RI problem $(\mathcal{A}, \boldsymbol{\pi}, c)$: The answer may be negative if only $\mathbf{a}^+$ is added to the menu but positive if $\mathbf{a}^+$ is added alongside other actions. Learning introduces complementarities between actions, and thus the attractiveness of any single action is always menu dependent. For example, suppose payoff a_j^+ is large and positive, but the action has negative payoff consequences in all other states, $a_i^+ < 0$ for all $i \neq j$. When the agent is close to certain that she is in state j , $\mathbf{a}^+$ can bring large consumption gains relative to the outside option $\mathcal{A} = \{\mathbf{0}\}$. Yet, the cost of a sufficiently informative signal may remain prohibitive unless the agent also has access to an action that does well in the opposite contingency where the likelihood of j is revised downwards.

By combining Theorems 3 and 4, we nevertheless obtain a comprehensive answer for both interpretations of the question: To identify actions that are attractive if added in isolation, one has to look no further than the ignorance equivalent. To identify actions that are attractive in some supermenu $\mathcal{A}' \supseteq \mathcal{A}$, the learning-proof menu is the answer: If $\mathbf{a}^+$ is located on or above $\bar{\mathcal{A}}$,²⁴ the action becomes an anchor in $\mathcal{A} \cup \{\mathbf{a}^+\}$ and can be activated by the simultaneous addition of some complementary action. Conversely, if $\mathbf{a}^+$ is located below $\bar{\mathcal{A}}$, activation is impossible.

5.2 Buyouts and Menu Choice

In our introductory example, the wealth manager wants to offer the investor terms that are attractive enough that she unconditionally accepts and forgoes all of her other options. A similar situation arises when an incumbent wants to buy out a startup that has other avenues for growth, or when a sports team wants to hire a star athlete with multiple employment offers.

In each situation, the objective is to find a payoff vector $\mathbf{a}^+$ that will be chosen unconditionally from menu $\mathcal{A} \cup \{\mathbf{a}^+\}$, something that we will call a ‘buyout’. The menu expansion

²⁴Formally, $\mathbf{a}^+$ is on or above $\bar{\mathcal{A}}$ if there exists $\mathbf{v} \gg 0$ with $\mathbf{v} \cdot \mathbf{a}^+ \geq \mathbf{v} \cdot \mathbf{a}$ for all $\mathbf{a} \in \bar{\mathcal{A}}$.

result can be used in reverse to characterize the set of feasible buyouts through binary learning strategies only. The core of the proof is an inductive application of [Theorem 3\(a\)](#).

Theorem 5. *For any finite menu $\mathcal{A}$ and $\mathbf{a}^+ \in \mathbb{R}^I$, the following are equivalent,*

(a) *Action $\mathbf{a}^+$ is chosen unconditionally from menu $\mathcal{A} \cup \{\mathbf{a}^+\}$,*

$$W(\{\mathbf{a}^+\}, \boldsymbol{\pi}, c) \geq W(\mathcal{A} \cup \{\mathbf{a}^+\}, \boldsymbol{\pi}, c).$$

(b) *Action $\mathbf{a}^+$ is chosen unconditionally against any single action in $\mathcal{A}$,*

$$W(\{\mathbf{a}^+\}, \boldsymbol{\pi}, c) \geq W(\{\mathbf{a}^+, \mathbf{a}\}, \boldsymbol{\pi}, c) \quad \forall \mathbf{a} \in \mathcal{A}.$$

Not having to worry about more complicated learning deviations brings tractability benefits. It delineates the space of feasible buyouts using only $|\mathcal{A}|$ inequalities. In [Section 6](#), we make use of this technique in a multiplayer setting that generalizes the introductory example to situations where the investor has some bargaining power.

It also allows us to figure out whether the RI agent would choose to learn early when facing two candidate menus $\mathcal{A}^1$ and $\mathcal{A}^2$. A priori, one would have to determine the optimal learning strategy in the joint menu $\mathcal{A}^1 \cup \mathcal{A}^2$ in order to determine whether the same welfare can be reached by restricting access to actions in only one of the menus. Luckily, it is sufficient to consider binary learning strategies in conjunction with the ignorance equivalent.

Corollary 5. *The agent unconditionally prefers menu $\mathcal{A}^1$ to menu $\mathcal{A}^2$,*

$$W(\mathcal{A}^1, \boldsymbol{\pi}, c) \geq W(\mathcal{A}^1 \cup \mathcal{A}^2, \boldsymbol{\pi}, c),$$

if and only if she unconditionally prefers the ignorance equivalent $\boldsymbol{\alpha}^1$ of menu $\mathcal{A}^1$ to any single action in $\mathcal{A}^2$, $W(\{\boldsymbol{\alpha}^1\}, \boldsymbol{\pi}, c) \geq W(\{\boldsymbol{\alpha}^1, \mathbf{a}\}, \boldsymbol{\pi}, c)$ for all $\mathbf{a} \in \mathcal{A}^2$.

The result is a simple consequence of [Theorem 5](#), which can be restated as

$$W(\{\boldsymbol{\alpha}^1\}, \boldsymbol{\pi}, c) \geq W(\mathcal{A}^1 \cup \mathcal{A}^2 \cup \{\boldsymbol{\alpha}^1\}, \boldsymbol{\pi}, c) \tag{4}$$

if and only if

$$W(\{\boldsymbol{\alpha}^1\}, \boldsymbol{\pi}, c) \geq W(\{\boldsymbol{\alpha}^1, \mathbf{a}\}, \boldsymbol{\pi}, c) \quad \forall \mathbf{a} \in \mathcal{A}^1 \cup \mathcal{A}^2. \tag{5}$$

By definition of the ignorance equivalent, the left side in [Equation \(4\)](#) is equal to $W(\mathcal{A}^1, \boldsymbol{\pi}, c)$, and [Equation \(5\)](#) automatically holds whenever $\mathbf{a} \in \mathcal{A}^1$. [Corollary 5](#) then follows from the fact that $\boldsymbol{\alpha}^1$ lies (weakly) below the learning-proof menu of $\mathcal{A}^1 \cup \mathcal{A}^2$, and hence removing $\boldsymbol{\alpha}^1$ from the left side of [Equation \(4\)](#) carries no welfare consequences.

5.3 Computation

The binary evaluation of learning potential also allows for an alternative characterization of the ignorance equivalent that, contrary to Equation (1), does not rely on the optimal learning strategy. Indeed, by definition, adding a buyout to the menu raises agent welfare – except if said buyout is the unique ignorance equivalent, in which case welfare remains constant. Hence, the ignorance equivalent is the payoff vector with minimal expected utility that satisfies all inequalities in Theorem 5(b),

$$\begin{aligned} \boldsymbol{\alpha} = \arg \min_{\substack{\mathbf{a}^+ \in \mathbb{R}^I \\ W(\{\mathbf{a}^+\}, \boldsymbol{\pi}, c) \geq W(\{\mathbf{a}^+, \mathbf{a}\}, \boldsymbol{\pi}, c) \quad \forall \mathbf{a} \in \mathcal{A}}} \boldsymbol{\pi} \cdot \mathbf{a}^+ \end{aligned} \quad (6)$$

For (UPS+) costs, we can further simplify condition (b) by spelling out the welfare from a generic binary strategy. (Details in Appendix A.5.) The argument relies on finding, for each action, the ‘most tempting’ posterior

$$\boldsymbol{\rho}^{\mathbf{a}} := \arg \max_{\boldsymbol{\rho} \in \Delta \mathcal{I}} \boldsymbol{\rho} \cdot (\mathbf{a} - \boldsymbol{\alpha}) - \phi(\boldsymbol{\rho}) + \nabla \phi(\boldsymbol{\pi}) \cdot \boldsymbol{\rho} \quad (7)$$

and then ensuring that no welfare can be gained by putting even the smallest marginal weight on this posterior. Ultimately, we are left with a characterization that does not require finding any optimal learning strategies – not even in a binary submenu. For Shannon entropy costs with $\phi(\boldsymbol{\pi}) = \lambda \boldsymbol{\pi} \cdot \ln(\boldsymbol{\pi})$ for instance, expression (6) simplifies to

$$\begin{aligned} \boldsymbol{\alpha} = \arg \min_{\substack{\mathbf{a}^+ \in \mathbb{R}^I \\ \sum_{i \in \mathcal{I}} \pi_i e^{(a_i - a_i^+)/\lambda} \leq 1 \quad \forall \mathbf{a} \in \mathcal{A}}} \boldsymbol{\pi} \cdot \mathbf{a}^+ \end{aligned} \quad (8)$$

In Armenter et al. [2024], we exploit this characterization to derive a more efficient and robust algorithm to numerically compute optimal RI behavior. It is reasonable to conjecture, though beyond the scope of the current paper, that Equation (6) can yield similar efficiency gains for other (UPS+) cost functions.

5.4 Trading Information

Since the learning-proof menu allows us to abstract away from state-dependent choices, it is particularly well suited to study situations in which a principal wants to trade information with a rationally inattentive agent. We consider three scenarios:

Selling information. A principal wants to sell an exogenous signal at the highest possible price to a rationally inattentive agent. The agent's willingness to pay depends on whether this information would replace or complement her own information acquisition, and whether it helps her discriminate between the choices at her disposal.²⁵

Corollary 6. *An agent facing RI problem $(\mathcal{A}, \boldsymbol{\pi}, c)$ is willing to pay at most*

$$\sum_{s \in S} (\boldsymbol{\pi} \cdot \mathbf{q}(s)) \max_{\mathbf{a} \in \bar{\mathcal{A}}} (\boldsymbol{\pi}^s \cdot \mathbf{a}) - \max_{\mathbf{a} \in \bar{\mathcal{A}}} \boldsymbol{\pi} \cdot \mathbf{a}$$

for access to a signal $\mathcal{S} = \langle S, \mathbf{q} \rangle$.

In particular, the local geometry of the learning-proof menu around the ignorance equivalent determines whether the agent benefits from access to noisy information, and thus whether her willingness to pay for a diffusion signal is positive. The tangent space of the surface $\bar{\mathcal{A}}$ at the ignorance equivalent, V , determines what local information is *actionable* for the agent. If a signal realization updates her belief orthogonally, $\boldsymbol{\pi}^s - \boldsymbol{\pi} \perp V$, then small enough belief updates merely cause the supporting hyperplane to rotate, but it remains tangent at $\boldsymbol{\alpha}$. Such a signal realization changes the agent's belief about the world, but not in a way that causes her to change her ignorance equivalent or, by [Corollary 2](#), her optimal learning strategy. If this is true for all signal realizations, then the information has no value to the agent.

Buying information. A principal wants to purchase signal $\mathcal{S} = \langle S, \mathbf{q} \rangle$ at minimal cost from a risk-neutral, rationally inattentive agent. Both share the prior $\boldsymbol{\pi}$. The principal can offer a set of state-dependent transfers to the agent but cannot independently verify the agent's learning. Thus, he has to design the payoffs in a way that incentivizes the agent to pick payoff $\mathbf{a}^s$ with probability $q_i(s)$ conditional on state i . To further complicate his task, any menu offer from the principal has to contain the non-empty set $\mathcal{A}^0$. In the simplest case, $\mathcal{A}^0 = \{\mathbf{0}\}$ simply captures that the agent's participation is voluntary.

Corollary 7. *The cost of incentivizing an RI agent to reveal signal $\mathcal{S}$ is given by*

$$c(\mathcal{S}, \boldsymbol{\pi}) + W(\mathcal{A}^0, \boldsymbol{\pi}, c),$$

where $\mathcal{A}^0$ contains all mandatory offers, $\boldsymbol{\pi}$ is their common prior and c is the agent's learning cost.

²⁵In [Online Appendix B.2.4](#), we illustrate the interaction between exogenous information and endogenous learning with a concrete example.

Since the cost function of the principal and the agent differ only by a constant, this means that the principal will either not engage with the agent or learn exactly the same amount as if he had direct access to the agent’s learning technology. Moreover, the proof in [Appendix A.3](#) is constructive in that it identifies exactly how the principal can structure the optimal transfer menu $\mathcal{A}$.

In mechanism design and delegation problems, ‘mandatory offers’ may appear endogenously, out of the principal’s desire to contract with multiple distinct agent types. For instance, [Yoder \[2022\]](#) considers a two-state setup where the principal is unaware of the agent’s cost. When learning is unobservable,²⁶ information acquisition and truth-telling both have to be incentivized through the design of the transfer menu $\mathcal{A}$. Letting $\mathcal{A}^0$ denote the transfers intended for other agent types and the outside option $\mathbf{0}$, the formulation in [Corollary 7](#) implies that the resulting information rent for an agent with cost c is exactly equal to $W(\mathcal{A}^0, \boldsymbol{\pi}, c)$.

Eliciting beliefs. A principal wants to identify an RI agent’s unknown prior belief without learning distortions as in [Tsakas \[2020\]](#). For example, a marketing company may want to gauge the subjective beliefs about a new product within a representative sample of consumers. Unincentivized belief elicitation is subject to all the standard pitfalls of stated preference, but too strong incentives might lead the agent to invest in learning, making her less representative of the population as whole. The solution is to tie each announcement to a state-dependent payment, and the analyst’s job is to design this menu of transfers $\mathcal{A}$ in a way that discourages learning but incentivizes truth-telling. The analyst is successful if $\mathcal{A}$ itself is learning-proof, $\mathcal{A} = \bar{\mathcal{A}}$. Truth-telling ensues as long as each belief announcement $\boldsymbol{\pi}$ maps to the utility-maximizing transfer $\mathbf{a}^\pi = \arg \max_{\mathbf{a} \in \mathcal{A}} \boldsymbol{\pi} \cdot \mathbf{a}$.

If $\mathcal{A}$ has points of non-differentiability, several announcements map to the same transfer. The extreme is that of unincentivized belief elicitation, which collapses all announcements to the single payoff vector $\mathcal{A} = \{0\}$. Here, truth-telling is optimal, but not (strictly) incentivized. By introducing more spread in $\mathcal{A}$, the analyst can widen the difference in payoff when the agent announces various beliefs, and thus decrease the appeal of lying. At the same time, menus with a larger spread also provide incentives for learning — lest the curvature of the menu be sufficiently high. The method in this paper tells the analyst how to ‘learning-proof’ any candidate set of transfers $\mathcal{A}^0$ into a menu $\bar{\mathcal{A}}^0$ that maintains the same spread but avoids learning.

²⁶This is in contrast to [Yoder \[2022\]](#), who focuses on situations with hard information and allows the principal to condition the transfer either directly on the acquired signal or on the realized posterior belief, thereby removing the need to incentivize truth-telling.

6 Optimal Allocation with Multiple RI Agents

Ignorance equivalence also plays an important role in multiplayer settings. We describe a generic allocation problem and analyze whether the first-best solution is implementable when agents can learn. We study two scenarios with one common theme: State-dependent terms informed by the ignorance equivalent align agents' learning incentives with those of the social planner.

Generic Setup. A single opportunity has to be allocated to one of finitely many RI agents $k \in \{1, \dots, K\}$. The agent who eventually executes the opportunity gets access to a menu of actions, $\mathcal{A}^k$, and enjoys the resulting state-dependent payoffs. All other agents receive a payoff of zero. All agents share a common prior π^0 about the state of the world, and cost functions c^k capture their respective learning costs. The opportunity initially rests with agent 1, but it is transferable as long as it has not been executed. Agents may learn at any time, and such learning does not execute the opportunity.

To determine the social optimum, we consider a planner with direct access to the joint menu $\mathcal{A}^P = \bigcup_{k=1}^K \mathcal{A}^k$ and all learning technologies with costs $\{c^k\}_{k=1}^K$. The planner may save on learning costs through sequential optimization, in which the realization of one signal determines the design of the next. [Bloedel and Zhong \[2020\]](#) define the planner's *indirect* cost c^P as the one-shot equivalent of the optimization over sequential experiments with *direct* cost function $\min_k c^k$. By optimally assigning learning and opportunity ownership, the planner can generate social surplus

$$\Delta = W(\mathcal{A}^P, \pi^0, c^P) - W(\mathcal{A}^1, \pi^0, c^1) \quad (9)$$

relative to the autarky allocation, in which agent 1 executes the opportunity.

The goal of this section is to show when and how agents can emulate the planner's choices through trade. To establish a trade between agents k and ℓ , both agents need to agree on *terms* $\mathbf{t} \in \mathbb{R}^I$. If they do, agent k releases the opportunity to agent ℓ , who in turn pays the former t_i once the state i realizes. If at least one agent objects, the opportunity remains with agent k . Both agents can learn and selectively participate only in those contingencies that are most favorable for them.

Examples. Pre-trade learning opportunities abound in a range of economic applications. We give four concrete examples:

- *Finance.* Agent 1 is looking to invest his wealth, and the menu $\mathcal{A}^1$ captures the state-dependent returns of each feasible portfolio. Agents $k > 1$ are investment managers

who are raising funds. Each manager could invest the wealth into portfolios $\mathcal{A}^k$, yielding a potentially larger return due to their improved fund access and lower research costs.

- *Procurement.* Agent 1's business needs to produce some parts. In-house, he has access to technologies with state-dependent costs in $\mathcal{A}^1$. Agents $k > 1$ are suppliers with access to technologies $\mathcal{A}^k$, and their expertise gives them easier access to relevant information on the availability of inputs.
- *Real Estate.* Home owner 1 is contemplating a range of last-minute fixes $\mathcal{A}^1$ that could affect the sale price of his house. Agents $k > 1$ specialize in property redevelopment, and each has grandiose plans $\mathcal{A}^k$ for the same lot. Although Agent 1's fixes do not technically preclude the realization of another agent's vision, it is socially wasteful to renovate a house before demolition. Our model applies as long as the planner would select a single project in $\mathcal{A}^P$: The choice of the project can depend on learning, but we do not allow the planner to assign a sequence of projects.
- *Teams.* A firm is considering the purchase of a new technology that will benefit workers across many specialized units. Some workers are uniquely qualified to learn about specific characteristics of the technology, and the optimal teams solution may involve a sequence of cost-benefit investigations by multiple workers before settling on a final choice.

The 'opportunity' grants its owner the right to purchase any available technology and full control over access. Once the purchase is complete, the state realizes, and the owner can sell non-exclusive access rights to all his collaborators. In that way, the opportunity owner's payoff from a technology ends up being equal to the state-dependent total surplus from the purchase.

Absolute Advantage. We start with a situation in which agent K has weakly lower costs than everyone else and can access a weakly larger menu, so that $c^P = c^K$ and $\mathcal{A}^P = \mathcal{A}^K$. Under this ranking, the first best surplus is achieved whenever trade is unconditional. (The details of any transfer agreements are irrelevant for social welfare.) Since the first best is now easy to characterize, we can focus entirely on implementability constraints. Our goal is to characterize the entire set of terms T that incentivize agents 1 and K to participate unconditionally in a trade.

Agent 1 can either participate for a payoff of $\mathbf{t}$ or decline trade in favor of any other option $\mathbf{a} \in \mathcal{A}^1$. Unconditional participation is optimal if and only if $\mathbf{t}$ is the ignorance equivalent

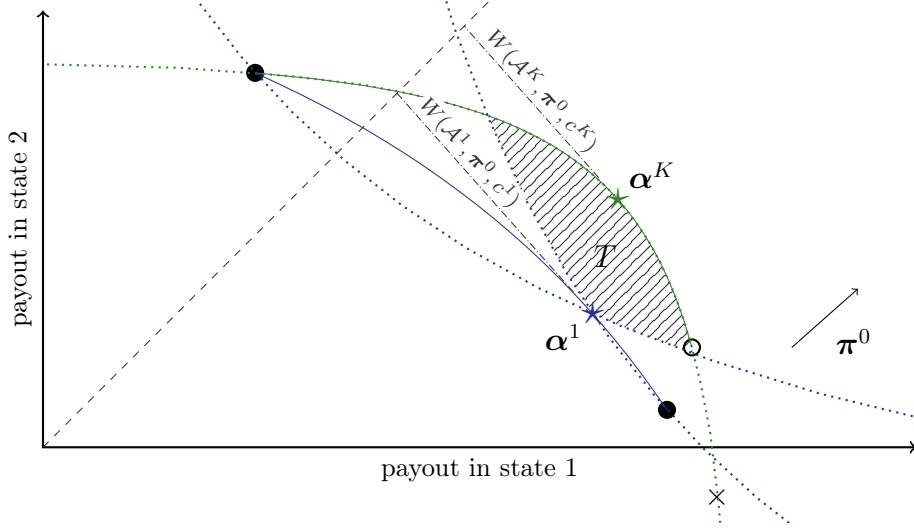


Figure 4: Transfers T that implement unconditional trade.

of the RI problem $W(\mathcal{A}^1, \pi^0, c^1)$. By [Theorem 5](#), this is equivalent to

$$W(\{t\}, \pi^0, c^1) \geq W(\{a, t\}, \pi^0, c^1) \quad \forall a \in \mathcal{A}^1, \quad (10)$$

which imposes lower bounds on the terms t .

Agent K can either participate and then implement any $a \in \mathcal{A}^K$ for a payoff of $a - t$, or decline trade for a zero payoff. Unconditional participation is optimal if and only if removing the option to reject does not lower welfare,

$$W((\mathcal{A}^K - \{t\}) \cup \{0\}, \pi^0, c^K) \leq W(\mathcal{A}^K - \{t\}, \pi^0, c^K).$$

Since the optimal strategy is unaffected by an unconditional (additive) shift of all payoff vectors, this can equivalently be stated as $W(\mathcal{A}^K \cup \{t\}, \pi^0, c^K) \leq W(\mathcal{A}^K, \pi^0, c^K)$, or by [Theorem 3](#) with reference to the ignorance equivalent α^K of $(\mathcal{A}^K, \pi^0, c^K)$,

$$W(\{\alpha^K, t\}, \pi^0, c^K) \leq W(\{\alpha^K\}, \pi^0, c^K), \quad (11)$$

imposing an upper bound on the terms t .

Together, [Equations \(10\) and \(11\)](#) fully characterize the set T of terms that implement unconditional trade in line with the social optimum. They do so by relying on binary learning strategies only, even for complex problems with many states and rich menus. [Figure 4](#) illustrates the construction in a simple two-state example. Agent 1's menu comprises all payoffs marked $\bullet$; agent K 's menu also contains those marked $\circ$. Agent K 's learning advantage manifests itself by the higher curvature of his learning-proof menu ($\backslash$ for agent 1, $\curvearrowright$ for

agent K). Equation (10) imposes a lower bound on the transfers for each of agent 1's actions () , and Equation (11) imposes an upper bound through agent K 's ignorance equivalent () . The set is bounded because learning incentives depend on the full vector $\mathbf{t}$ rather than just its expected payoff. In particular, trade can fail even at terms that yield positive surplus to both agents, with $W(\mathcal{A}^1, \boldsymbol{\pi}^0, c^1) < \boldsymbol{\pi}^0 \cdot \mathbf{t} < W(\mathcal{A}^K, \boldsymbol{\pi}^0, c^K)$.

The selection of terms within T determines the split of the optimal surplus Δ across the two agents and thus captures their relative bargaining power. In particular, if one agent can extend a take-it-or-leave-it (TIOLI) offer, the equilibrium features unconditional trade at terms equal to his opponent's ignorance equivalent. In this way, the agent in power secures the entire surplus Δ . He can do no better by proposing a plan with partial trade either: Incentive compatibility requires nonnegative surplus for the opponent, lest outright rejection be a profitable deviation.

The surplus can be split between more than two agents. An example of this arises when agent $2 < K$ plays the role of a platform with $\mathcal{A}^2 = \emptyset$, and her ability to form or restrict matches grants her some bargaining power. In that setting, the full surplus is realized no matter at what terms $\mathbf{t}^1 \in T$ platform 2 buys the opportunity from agent 1 and at what terms $\mathbf{t}^K \in T$ she sells it to agent K . The difference $\boldsymbol{\pi} \cdot (\mathbf{t}^K - \mathbf{t}^1)$ constitutes the expected platform profits. In particular, if the platform has TIOLI power, she can secure the entire surplus Δ by buying at terms $\boldsymbol{\alpha}^1$ and selling at terms $\boldsymbol{\alpha}^K$. She could also claim a smaller flat fee τ by choosing terms with $\mathbf{t}^K - \mathbf{t}^1 = \tau \mathbf{1}$, thereby shielding platform profits entirely from uncertainty.

The construction also speaks to whether 'rich' opportunities are harder to trade. Although a menu with diverse options is no more attractive than one with just the ignorance equivalent $\boldsymbol{\alpha}$, it offers more incentives for learning. This can shrink the set of terms that implement the first best: Indeed, each additional action in Agent 1's menu $\mathcal{A}^1$ imposes another lower bound through Equation (10). This is true even when his ignorance equivalent remains unchanged, as is the case for action $\times$ in Figure 4. In contrast, diversity of options for Agent K does not affect implementability of the first best, since his menu only enters Equation (11) through the ignorance equivalent $\boldsymbol{\alpha}^K$. Economically speaking, this is because the richness in Agent 1's menu offers him potential *alternatives* to trade, whereas the richness in Agent K 's menu matters only *conditional* on trade.

Finally, the construction of T also answers whether state-dependent payment terms are *necessary* to achieve the first-best surplus. In Figure 4, the menu payoffs are heavily skewed in favor of state 1. This asymmetry generates an incentive to learn unless the suggested transfer is also higher in that state. In general, unconditional trade can be achieved with constant payment terms only when T contains a point on the diagonal $\{\tau \mathbf{1} \mid \tau \in \mathbb{R}\}$. A

natural follow-up question is what happens when states are unverifiable and payment terms need to be constant across states. Would Agent 1 sometimes make an offer so generous that Agent K always accepts, or always make a less generous offer that Agent K sometimes declines? — In [Online Appendix B.2.5](#), we answer this question for the specific two-state example of [Figure 4](#) by relying on the equivalence between learning ability and access to fictitious payoffs. We show that when Agent 1 gets to make a take-it-or-leave-it offer and states are unverifiable, the Perfect Bayesian Equilibrium features partial trade and pre-trade learning by both agents.

Comparative Advantages. To showcase the versatility of the ignorance equivalent, we now allow for non-nested menus and non-ordered cost functions, thereby incorporating agents’ comparative advantages in terms of feasible payoffs or learning. To focus on learning incentives rather than truth-telling constraints, we assume that learning produces hard information in line with [Yoder \[2022\]](#): Learning itself is non-contractible, but the realization of any signal draw is publicly observable, so that agents share the same beliefs at all times.²⁷ In [Appendix A.6](#), we formally define the dynamic game that results from allowing agents to acquire costly information and trade the opportunity among themselves until someone executes it. We show that the first-best allocation remains implementable: Agents can realize surplus Δ in a subgame-perfect equilibrium that features repeated trades with terms chosen from the learning proof menu $\bar{\mathcal{A}}^P$ resulting from menu $\mathcal{A}^P$ under cost c^P . The key intuition is that trading at terms $\arg \max_{\mathbf{a} \in \bar{\mathcal{A}}^P} \boldsymbol{\rho} \cdot \mathbf{a}$ at all (joint) beliefs $\boldsymbol{\rho}$ aligns agents’ learning incentives with those of the planner.

Theorem 6. *Agents emulate the planner’s strategy in a Perfect Bayesian Equilibrium by repeatedly trading at terms chosen from the planner’s learning-proof menu.*

This result speaks to the ‘teams’ literature and its focus on optimal compensation schemes that induce cooperation among workers with a common task [[Bergemann and Välimäki, 2002](#), [Holmström and Milgrom, 1990](#), [Itoh, 1991](#), among others]. If the task is one of pure information acquisition, the result implies that heterogeneous workers can effectively organize themselves as long as they are able to write binding, state-contingent contracts. In equilibrium, the ‘decision power’ (viz. ‘opportunity’) changes hands repeatedly, and information is acquired sequentially by various agents working together. This simplifies things for the manager: As long as she ensures that transfer promises are honored ex-post, she does not have to worry about tailoring incentives to the relative learning skills of each individual worker. Instead, she can address the team *as if* it were one worker with menu $\mathcal{A}^P$ and cost

²⁷[Yoder \[2022\]](#) calls this results-based contracting.

function c^P and let workers organize themselves. So long as each worker is familiar with their own learning cost and the ‘joint’ learning proof menu $\bar{\mathcal{A}}^P$, they can implement the first best through trade.

7 Related Literature

Rational inattention models are widely used to model flexible learning in both the applied and theoretical literature, in fields ranging from finance to political economy (see [Maćkowiak et al. \[2023\]](#) for a survey). The most common approaches to RI models fall broadly into two categories: For tractability reasons, most applied research relies on special cases that admit an analytic solution (e.g., a Linear-Quadratic Gaussian setup and Shannon-entropy costs, or costs that are infinite except for Gaussian signals). Theory research typically relies on the ‘posterior-based approach’ [[Caplin et al., 2022](#)]. In this section, we discuss how ignorance equivalence stands to benefit both applied and theoretic research.

First, ignorance equivalence does not impose any functional form or domain restrictions on the agent’s learning technology. As such, it identifies robust features of rationally inattentive behavior, thereby generalizing and extending results from prior research. For instance, [Caplin et al. \[2018\]](#) derive a ‘market entry condition’ for the specific case of Shannon entropy costs to judge whether an agent stands to benefit from a newly available action. [Theorem 3](#) derives an analogous expression for any (UPS+) cost function, including those that do not satisfy the Invariance under Compression axiom of the Shannon cost function [[Caplin et al., 2022](#)]. [Matějka and McKay \[2015\]](#) show in a numeric example (also with Shannon costs) that menu expansion can raise the appeal of a previously unchosen action, a distinguishing feature that is absent in random utility models of stochastic choice. [Theorem 4](#) proves that it is exactly the set of anchor actions that can be ‘activated’ through menu expansion.

By tying our proofs to economically meaningful properties of the cost function, our work follows in the vein of [[Bloedel and Zhong, 2020](#), [Denti et al., 2022](#), [Hébert and Woodford, 2020](#), [Mensch, 2018](#), [Pomatto et al., 2018](#)], who axiomatize various cost functions. Our cost assumptions are compatible with some, but not all, of the studied cost specifications, and we rely heavily on [Bloedel and Zhong \[2020\]](#)’s characterization of costs that are sequential learning proof. Yet, our primary focus is not on characterizing which cost functions are compatible with (C1) to (C5) — rather, we argue that these intuitively appealing properties enable summary statistics that abstract away from learning. None of these other papers condense the ‘payoff possibilities’ into a single vector in the spirit of our ignorance-equivalent approach.

Second, ignorance equivalence is related to the posterior-based approach, but lends itself

to different questions. We believe that both views are useful. From a technical perspective, the posterior-based approach is very general. It does not require prior-concavity and can even accommodate potentials that depend on the prior belief [Caplin et al., 2022]. The ignorance-equivalent approach, on the other hand, is more parsimonious because it transforms the high-dimensional search for an optimal signal (a distribution over posteriors, an element of $\Delta\Delta\mathcal{I}$) into a lower-dimensional choice (an action from the learning-proof menu, an element of $\mathbb{R}^{\mathcal{I}}$). This reduction in complexity and the geometry of the learning-proof menu can be exploited to improve numerical solution methods that stand to benefit the applied literature, as we outline in Section 5.3 and do in Armenter et al. [2024] for Shannon entropy costs.

From a conceptual perspective, the posterior-based approach puts the agent’s beliefs front and center, and its obvious parallels to Bayesian Persuasion [Gentzkow and Kamenica, 2014, Kamenica and Gentzkow, 2011] are helpful in information-design applications. In contrast, the ignorance-equivalence approach is particularly promising in scenarios where payoffs are determined endogenously, as in the applications that we discuss in Sections 5 and 6. By describing signals via the induced state-dependent choices $\mathbf{q}(\mathbf{a})$ rather than the induced posteriors $\boldsymbol{\pi}^{\mathbf{a}}$, we obtain a clear distinction between prior beliefs and behavior. This separation is helpful in particular when comparing across agents with different beliefs, as in Corollary 2.

There is also a large literature that still models learning much more rigidly, where the agent simply chooses whether to ‘uncover’ the state at a cost. Sometimes, this binary choice reflects the actual mechanics of the information acquisition process. Other times, the modeling decision is motivated more by tractability rather than realism and could fail to capture important features, like the intensive margin of learning or the agent’s endogenous focus on the most salient aspects of a multifaceted decision. By recasting the agent’s ability to learn *as if* she had access to a richer menu, ignorance equivalence opens new doors to incorporate flexible learning in multiplayer settings with transfers. Importantly, the applications in Sections 5 and 6 highlight that ignorance equivalence is useful not only in settings where learning is, in fact, to be avoided. Indeed, a principal can use it specifically to elicit the ‘right kind’ of learning from an agent, and a team can use it to coordinate learning across members of differing expertise. The fundamental reason is that the learning-proof menu does not just make *ignorance* optimal; it also maintains optimality of the original learning strategy — as well as any ‘intermediate’ strategy that involves less informative signals. In this way, the learning-proof menu can induce a continuum of information acquisition strategies, including the one that maximizes the principal’s (or social) welfare.

8 Conclusion

The ignorance-equivalent approach simplifies the description of optimal agent behavior under costly learning. In essence, it points out that the agent's ability to learn acts *as if* she instead had access to a fictitious action whose payoffs are given by the ignorance equivalent. In allocation problems in which multiple agents interact, this ignorance-equivalent action often emerges as part of the surplus-maximizing trading scheme. The same is true in experimental settings in which the analyst is free to design payoffs in a way that boosts his ability to, for instance, identify the agent's prior. Yet, even in menus where the ignorance equivalent is only a conceptual shortcut, it characterizes the full set of optimal signals and allows for parsimonious comparisons across learning strategies, menus, and beliefs. Much like the certainty equivalent reduces the complexity of economic problems with exogenous uncertainty, the ignorance equivalent allows us to abstract away from learning by simply maximizing expected utility over a larger menu.

A Proofs

A.1 Properties of the cost function

We first establish that as long as the state space is finite, it is without loss of generality to assume that the optimal signal is finite, even when the menu contains infinitely many actions.²⁸ Since we are deriving this result algebraically from the (UPS+) specification, it is more convenient to represent a signal (equivalently) as a distribution over posteriors $\sigma \in \Delta\Delta\mathcal{I}$.

Lemma A.1. *For any compact menu $\mathcal{A}$, welfare under (UPS+) cost c ,*

$$W(\mathcal{A}, \boldsymbol{\pi}, c) = \sup_{\substack{\sigma \in \Delta\Delta\mathcal{I} \\ \mathbb{E}[\sigma] = \boldsymbol{\pi}}} \int_{\Delta\mathcal{I}} (\max_{\mathbf{a} \in \mathcal{A}} \boldsymbol{\rho} \cdot \mathbf{a} - \phi(\boldsymbol{\rho})) d\sigma(\boldsymbol{\rho}) + \phi(\boldsymbol{\pi}),$$

is maximized at a signal with finite support.

Proof. The integral is affine in σ , and thus maximized at one of the extreme points of the space $\mathcal{M} = \{\sigma \in \Delta\Delta\mathcal{I} \mid \mathbb{E}[\sigma] = \boldsymbol{\pi}\}$. By Winkler [1988], extreme points of $\mathcal{M}$ are exactly those probability distributions whose support contains at most $|\mathcal{I}|$ posteriors. $\square$

Next, we show that the properties (C1) to (C5) are jointly equivalent to the (UPS+) family. The proof draws heavily on Bloedel and Zhong [2020]’s characterization of sequentially optimal cost functions, and often refers to them verbatim.

Lemma A.2. *Conditions (C1) to (C5) are jointly equivalent to the family of (UPS+) costs.*

Proof. Suppose c belongs to the (UPS+) family, which explicitly requires prior-concavity and implies Blackwell-monotonicity (C2) by convexity of ϕ . Continuity (C1) follows directly from that of the potential ϕ . Bloedel and Zhong [2020] establish that any uniformly posterior separable cost satisfies indifference to sequential learning, implying that the inequalities in (C4) and (C5) both hold. For (C3), fix any $\boldsymbol{\pi} \in \Delta\mathcal{I}$ and $\mathbf{a} \in \mathbb{R}^I$ with $\boldsymbol{\pi} \cdot \mathbf{a} = 0$ and $\boldsymbol{\pi} \cdot |\mathbf{a}| > 0$. For any $\varepsilon > 0$ small enough, consider the binary signal $\mathcal{S}^\varepsilon = \langle \{0, 1\}, \mathbf{q}^\varepsilon \rangle$ with $q_i^\varepsilon(1) = \frac{1}{2}(1 + \varepsilon a_i) \in (0, 1)$, and note that $\boldsymbol{\pi} \cdot \mathbf{q}^\varepsilon(0) = \boldsymbol{\pi} \cdot \mathbf{q}^\varepsilon(1) = 1/2$ since $\boldsymbol{\pi} \cdot \mathbf{a} = 0$. By Taylor’s theorem, cost grows less than linearly in ε ,

$$\lim_{\varepsilon \rightarrow 0} \frac{c(\mathcal{S}^\varepsilon, \boldsymbol{\pi})}{\varepsilon} = \lim_{\varepsilon \rightarrow 0} \frac{c(\mathcal{S}^0, \boldsymbol{\pi})}{\varepsilon} + \left. \frac{\partial c(\mathcal{S}^\varepsilon, \boldsymbol{\pi})}{\partial \varepsilon} \right|_{\varepsilon=0} = \lim_{\varepsilon \rightarrow 0} \frac{1}{\varepsilon} 0 + \frac{1}{2} \sum_{i \in \mathcal{I}} \nabla_i \phi(\boldsymbol{\pi}) (a_i \pi_i - a_i \pi_i) = 0,$$

²⁸We thank Tommaso Denti for pointing us in the right direction.

while the benefit grows linearly in ε , $\sum_{i \in \mathcal{I}} \pi_i q_i^\varepsilon(1) a_i = \frac{\varepsilon}{2} \sum_{i \in \mathcal{I}} \pi_i a_i^2$. In particular, benefits strictly outweigh costs for ε small enough.

Conversely, suppose c satisfies (C1) to (C5). By Bloedel and Zhong [2020], (C4) and (C5) jointly restrict costs to the uniformly posterior separable family with a convex potential ϕ , which means it only remains to show that ϕ is differentiable. For convenience, we extend ϕ to all of $\mathbb{R}_{\geq 0}^{\mathcal{I}}$ by imposing homogeneity of degree one $\phi(k\boldsymbol{\rho}) = k\phi(\boldsymbol{\rho})$ [as in Hébert and Woodford, 2019]. This extension is differentiable in $\mathbb{R}_{\geq 0}^{\mathcal{I}}$ if and only if ϕ itself is differentiable in $\Delta\mathcal{I}$. Fix any belief $\boldsymbol{\pi} \in \Delta\mathcal{I}$. Continuity of ϕ is implied by (C1), or an arbitrarily noisy belief perturbation would have costs bounded below by a positive constant, strictly above the free uninformative signal. To show differentiability, let G denote the set of subgradients of ϕ at $\boldsymbol{\pi}$, defined as the set of all vectors $\mathbf{g} \in \mathbb{R}^I$ such that

$$\phi(\boldsymbol{\rho}) - \phi(\boldsymbol{\pi}) \geq \mathbf{g} \cdot (\boldsymbol{\rho} - \boldsymbol{\pi}) \quad \forall \boldsymbol{\rho} \in \Delta\mathcal{I}.$$

Continuity and convexity of ϕ imply that G is nonempty, and ϕ is differentiable at $\boldsymbol{\pi}$ if and only if G is a singleton. (If $\boldsymbol{\pi}$ is on the boundary, G needs to be unique only on the positive probability states.) Note that homogeneity of degree one implies in particular that

$$\phi((1+k)\boldsymbol{\pi}) - \phi(\boldsymbol{\pi}) = k\phi(\boldsymbol{\pi}) \geq k\mathbf{g} \cdot \boldsymbol{\pi} \quad \forall k \in (-1, \infty),$$

which is binding at $k = 0$ and thus implies that $\boldsymbol{\pi} \cdot \mathbf{g} = \phi(\boldsymbol{\pi})$ for all subgradients $\mathbf{g} \in G$. Assume now by contradiction that $\mathbf{g}$ and $\mathbf{h}$ are both elements of G , and they differ in at least one state to which $\boldsymbol{\pi}$ assigns positive probability. Define the payoff vector $\mathbf{a} = \mathbf{g} - \mathbf{h}$, which by the above yields expected utility zero, but is nonzero with positive probability. By (C3), the agent would break her indifference between $\mathbf{a}$ and $\mathbf{0}$ with a binary signal. For convenience, we refer to the signal using the marginal likelihood q of implementing $\mathbf{a}$, and the shift in posterior $\boldsymbol{\Delta} = \boldsymbol{\pi}^{\mathbf{a}} - \boldsymbol{\pi}$, so we can write the inequality in (C3) as

$$0 > q [\phi(\boldsymbol{\pi} + \boldsymbol{\Delta}) - \phi(\boldsymbol{\pi})] + (1 - q) \left[\phi \left(\boldsymbol{\pi} - \frac{q}{1 - q} \boldsymbol{\Delta} \right) - \phi(\boldsymbol{\pi}) \right] - q(\boldsymbol{\pi} + \boldsymbol{\Delta}) \cdot \mathbf{a}.$$

However, substituting both terms in square brackets with a lower bound from one of the subgradients, the right side is bounded below by

$$q\boldsymbol{\Delta} \cdot (\mathbf{g} - \mathbf{h} - \mathbf{a}) - q\boldsymbol{\pi} \cdot \mathbf{a} = 0.$$

As a result, ϕ is differentiable, and the c belongs to the (UPS+) family. $\square$

Lemma A.3. *Under (C2), it is without loss of generality in RI problem (RI) to restrict*

attention to learning strategies $\langle \mathcal{A}, \mathbf{q} \rangle$.

Proof. For any signal $\mathcal{S} = \langle S, \mathbf{q} \rangle$ and conditional selection $\mathbf{a}^s \in \mathcal{A}$ for each $s \in S$, we can define the learning strategy $\hat{\mathcal{S}} = \langle \mathcal{A}, \hat{\mathbf{q}} \rangle$ with $\hat{q}_i(\mathbf{a}) = \sum_{s \in S: \mathbf{a}^s = \mathbf{a}} q_i(s)$ for each $\mathbf{a} \in \mathcal{A}$. This learning strategy achieves the same expected consumption utility and is Blackwell less informative than the original signal $\mathcal{S}$. By (C2), it thus achieves a weakly higher welfare. It is therefore without loss of optimality to restrict attention to learning strategies only. $\square$

Lemma A.4. Under (C2), welfare W is convex in the prior belief.

Proof. Convexity follows readily from the linearity of the consumption utility and the fact that signal costs are prior-concave (C2). Formally, let $\mathcal{S}$ be the optimal learning strategy for RI problem $(\mathcal{A}, t\boldsymbol{\pi} + (1-t)\boldsymbol{\pi}')$. The welfare is bounded above by the linear interpolation of the welfare achieved in $(\mathcal{A}, \boldsymbol{\pi})$ and $(\mathcal{A}, \boldsymbol{\pi}')$ when the same strategy is used.

$$\begin{aligned} W(\mathcal{A}, t\boldsymbol{\pi} + (1-t)\boldsymbol{\pi}', c) &= \sum_{i=1}^I (t\pi_i + (1-t)\pi'_i) \sum_{\mathbf{a} \in \mathcal{A}} q(\mathbf{a}|i) a_i - c(\mathcal{S}, t\boldsymbol{\pi} + (1-t)\boldsymbol{\pi}') \\ &\leq t \left[\sum_{i=1}^I \pi_i \sum_{\mathbf{a} \in \mathcal{A}} q(\mathbf{a}|i) a_i - c(\mathcal{S}, \boldsymbol{\pi}) \right] + (1-t) \left[\sum_{i=1}^I \pi'_i \sum_{\mathbf{a} \in \mathcal{A}} q(\mathbf{a}|i) a_i - c(\mathcal{S}, \boldsymbol{\pi}') \right] \\ &\leq tW(\mathcal{A}, \boldsymbol{\pi}, c) + (1-t)W(\mathcal{A}, \boldsymbol{\pi}', c). \end{aligned} \quad \square$$

Lemma A.5. Properties (C2) and (C3) imply that any pure-noise signal $\mathcal{S} = \langle S, \mathbf{q} \rangle$ with $q_i(s) \equiv q_j(s) \forall i, j \in \mathcal{I}$ and $s \in S$ is free, and the agent can freely randomize across actions.

Proof. All pure-noise signals are Blackwell equivalent, and thus have the same cost c^0 by (C2). Suppose by contradiction that $c^0 > 0$ and consider any two states $j, k \in \text{support}(\boldsymbol{\pi})$. Let $\mathbf{a} \in \mathbb{R}^I$ be such that

$$a_j = \frac{c^0}{2\pi_j}, a_k = -\frac{c^0}{2\pi_k} \quad \text{and} \quad a_i \equiv 0 \quad \forall i \in \mathcal{I} \setminus \{j, k\}.$$

By (C3), there exists a binary signal $\mathcal{S}' = \langle \{0, 1\}, \mathbf{q} \rangle$ with

$$c(\mathcal{S}', \boldsymbol{\pi}) < \sum_{i \in \mathcal{I}} \pi_i q_i(1) a_i = (q_j(1) - q_k(1)) \frac{c^0}{2} \leq c^0.$$

However, $\mathcal{S}'$ is Blackwell more informative than any pure-noise signal, and so monotonicity (C2) implies that the left side is weakly bounded below by c^0 , creating a contradiction. $\square$

Lemma A.6. Under (C1) and (C2), welfare W is finite and continuous in the prior belief $\boldsymbol{\pi}$. Moreover, there exists an upper hemicontinuous correspondence $Q^* : \Delta\mathcal{I} \rightarrow (\Delta\mathcal{A})^{\mathcal{I}}$ with

nonempty and compact values, such that a learning strategy $\mathcal{S} = \langle \mathcal{A}, \mathbf{q} \rangle$ is optimal in RI problem $(\mathcal{A}, \boldsymbol{\pi}, c)$ if and only if $\mathbf{q} \in Q^*(\boldsymbol{\pi})$.

Proof. Since it is without loss of generality to restrict attention to learning strategies (Lemma A.3), we restate (RI) as a choice over conditionals

$$W(\mathcal{A}, \boldsymbol{\pi}, c) = \sup_{\mathbf{q} \in (\Delta \mathcal{A})^{\mathcal{I}}} \sum_{\mathbf{a} \in \mathcal{A}} \sum_{i \in \mathcal{I}} \pi_i q_i(\mathbf{a}) a_i - c(\langle \mathcal{A}, \mathbf{q} \rangle, \boldsymbol{\pi})$$

and define $Q^*(\boldsymbol{\pi})$ to be the set of conditionals that achieve this supremum. Since the objective is continuous over the compact domain $(\Delta \mathcal{A})^{\mathcal{I}} \times \Delta \mathcal{I}$, the claim then follows by Berge's Theorem of the Maximum. $\square$

Lemma A.7. Suppose (C4) holds. Consider an optimal learning strategy $\mathcal{S} = \langle \mathcal{A}, \mathbf{q} \rangle$ to RI problem $(\mathcal{A}, \boldsymbol{\pi}, c)$. Let $\mathbf{a} \in \text{support}(\mathbf{q})$ be an action that is implemented with positive probability, and $\boldsymbol{\pi}^{\mathbf{a}}$ the associated posterior belief, $\pi_i^{\mathbf{a}} = \frac{\pi_i q_i(\mathbf{a})}{\boldsymbol{\pi} \cdot \mathbf{q}(\mathbf{a})}$. In RI problem $(\mathcal{A}, \boldsymbol{\pi}^{\mathbf{a}}, c)$, unconditional implementation of $\mathbf{a}$ is optimal, $W(\mathcal{A}, \boldsymbol{\pi}^{\mathbf{a}}, c) = \boldsymbol{\pi}^{\mathbf{a}} \cdot \mathbf{a}$.

Proof. Since unconditional implementation is feasible, clearly $W(\mathcal{A}, \boldsymbol{\pi}^{\mathbf{a}}, c) \geq \boldsymbol{\pi}^{\mathbf{a}} \cdot \mathbf{a}$. By contradiction, suppose $W(\mathcal{A}, \boldsymbol{\pi}^{\mathbf{a}}, c) > \boldsymbol{\pi}^{\mathbf{a}} \cdot \mathbf{a}$ and consider the sequential strategy where the agent first draws $\mathcal{S}$ and follows its recommendation except for when it evaluates to $\mathbf{a}$, when she instead continues with an optimal strategy for $(\mathcal{A}, \boldsymbol{\pi}^{\mathbf{a}}, c)$. This sequential strategy achieves utility

$$\sum_{i=1}^I \pi_i \sum_{\mathbf{a}' \in \mathcal{A} \setminus \{\mathbf{a}\}} q_i(\mathbf{a}') a'_i + (\boldsymbol{\pi} \cdot \mathbf{q}(\mathbf{a})) W(\mathcal{A}, \boldsymbol{\pi}^{\mathbf{a}}, c) - c(\mathcal{S}, \boldsymbol{\pi}).$$

Since $W(\mathcal{A}, \boldsymbol{\pi}^{\mathbf{a}}, c) > \boldsymbol{\pi}^{\mathbf{a}} \cdot \mathbf{a}$, this is strictly larger than

$$\sum_{i=1}^I \pi_i \sum_{\mathbf{a} \in \mathcal{A}} q_i(\mathbf{a}) a_i - c(\mathcal{S}, \boldsymbol{\pi}),$$

which implies that this sequential strategy achieves strictly higher utility than signal $\mathcal{S}$ alone. By (C4), the same is true for its one-shot equivalent, contradicting the optimality of $\mathcal{S}$ in $(\mathcal{A}, \boldsymbol{\pi}, c)$. $\square$

Lemma A.8. Consider any belief $\boldsymbol{\pi} \in \Delta \mathcal{I}$ and any two actions with the same expected utility, $u := \mathbf{a} \cdot \boldsymbol{\pi} = \mathbf{a}' \cdot \boldsymbol{\pi}$, that differ in at least one positive probability state. Under (C2) and (C3), and any $L > 0$, there exists a learning strategy $\langle \{\mathbf{a}, \mathbf{a}'\}, \mathbf{q} \rangle$ such that the change

in consumption utility outweighs the signal cost by more than a factor L ,

$$\sum_{i=1}^I \pi_i (q_i(\mathbf{a})a_i + q_i(\mathbf{a}')a'_i) - u > L c(\langle \{\mathbf{a}, \mathbf{a}'\}, \mathbf{q} \rangle, \boldsymbol{\pi}).$$

Proof. The payoff vector $\tilde{\mathbf{a}} = \frac{1}{L}(\mathbf{a} - \mathbf{a}')$ has zero expected utility $\boldsymbol{\pi} \cdot \tilde{\mathbf{a}} = \frac{1}{L}(u - u) = 0$ and is non-zero with positive probability since $a_i \neq a'_i$ for some i with $\pi_i > 0$. By condition (C3), there exists a signal $\mathcal{S} = \langle \{0, 1\}, \mathbf{q} \rangle$ with $\sum_{i \in \mathcal{I}} \pi_i q_i(1) \tilde{a}_i > c(\mathcal{S}, \boldsymbol{\pi})$. By (C2), the agent incurs the same cost for the Blackwell-equivalent learning strategy where she implements $\mathbf{a}'$ after signal 0 and $\mathbf{a}$ after signal 1. Multiplying both sides with L , we obtain

$$\sum_{i \in \mathcal{I}} \pi_i q_i(1)(a_i - a'_i) = \sum_{i \in \mathcal{I}} \pi_i (q_i(1)a_i + q_i(0)a'_i) - u > L c(\mathcal{S}, \boldsymbol{\pi}). \quad \square$$

A.2 Ignorance equivalent

Duality. Fenchel's duality theorem allows us to apply strong duality results to optimization problems with infinitely many constraints. For our purposes, the following variant is strong enough – we formally derive it from Borwein and Zhu [2005, Theorem 4.4.3] in Online Appendix B.1.1.

Lemma A.9 (Fenchel Duality for Linear Constraints). *For any Lebesgue measurable function $\Phi : \Delta\mathcal{I} \rightarrow \mathbb{R}$ and any belief $\boldsymbol{\pi} \in \text{int}(\Delta\mathcal{I})$, the two optimization problems*

$$d^* = \sup_{\mathbf{a} \in \mathbb{R}^I} \{ \boldsymbol{\pi} \cdot \mathbf{a} \mid \boldsymbol{\rho} \cdot \mathbf{a} \leq \Phi(\boldsymbol{\rho}) \ \forall \boldsymbol{\rho} \in \Delta\mathcal{I} \} \quad \text{and} \quad p^* = \inf_{\substack{\boldsymbol{\nu} \in \Delta\Delta\mathcal{I} \\ \mathbb{E}[\boldsymbol{\nu}] = \boldsymbol{\pi}}} \mathbb{E}[\Phi(\boldsymbol{\rho}) \mid \boldsymbol{\rho} \sim \boldsymbol{\nu}]$$

achieve the same bound, $p^* = d^*$. Moreover, the supremum is attained whenever that bound is finite.

When Φ is convex, the primal is trivially optimized at the degenerate distribution that puts full weight on $\boldsymbol{\pi}$, and hence the optimum in the dual attains $\Phi(\boldsymbol{\pi})$.

Lemma A.10. *Let $\Phi : \Delta\mathcal{I} \rightarrow \mathbb{R}$ denote a convex function over $\Delta\mathcal{I}$. Then for any $\boldsymbol{\pi} \in \text{int}(\Delta\mathcal{I})$, the optimization problem*

$$\sup_{\mathbf{a} \in \mathbb{R}^I} \{ \boldsymbol{\pi} \cdot \mathbf{a} \mid \boldsymbol{\rho} \cdot \mathbf{a} \leq \Phi(\boldsymbol{\rho}) \ \forall \boldsymbol{\rho} \in \Delta\mathcal{I} \} \tag{12}$$

admits a maximum at some $\mathbf{a}^* \in \mathbb{R}^I$ with objective value $\boldsymbol{\pi} \cdot \mathbf{a}^* = \Phi(\boldsymbol{\pi})$.

Proof. Since (12) corresponds to the dual d^* in Lemma A.9, its objective value is equal to that of

$$\inf_{\substack{\nu \in \Delta \Delta \mathcal{I} \\ \mathbb{E}[\nu] = \pi}} \mathbb{E}[\Phi(\rho) \mid \rho \sim \nu].$$

In turn, convexity of Φ and Jensen's inequality imply that the infimum is attained when full weight is placed on π , implying $d^* = \Phi(\pi) \in \mathbb{R}$. By duality, the supremum is also attained and achieves objective value $\Phi(\pi)$. $\square$

Existence and Uniqueness. Put in economically relevant terms, Lemma A.9 establishes a mathematical duality between optimization problems

$$\bar{w}^{\mathcal{S}} = \sup_{\mathbf{x} \in \mathbb{R}^I} \{\pi \cdot \mathbf{a} \mid \mathbf{x} \preceq \mathcal{S}\} \quad \text{and} \quad \underline{w}^{\mathcal{S}} = \inf_{\substack{\nu \in \Delta \Delta \mathcal{I} \\ \mathbb{E}[\nu] = \pi}} \mathbb{E}[\rho \cdot \mathbf{a}^{\mathcal{S}} - c(\mathcal{S}, \rho) \mid \rho \sim \nu]. \quad (13)$$

The former ($\bar{w}^{\mathcal{S}}$) maximizes the agent's welfare, assuming she implements a dominated payoff vector $\mathbf{x} \preceq \mathcal{S}$ unconditionally. The latter ($\underline{w}^{\mathcal{S}}$) computes the agent's worst-case welfare from following strategy $\mathcal{S}$, when an adversarial nature can reveal free information before the signal is drawn – or, equivalently, nature triggers a belief update to priors ρ drawn from any Bayes-plausible distribution ν . This belief update may change expected signal costs, but only to the agent's benefit because of prior-concavity (C2). The worst case occurs under no free information, and hence $\underline{w}^{\mathcal{S}} = \pi \cdot \mathbf{a}^{\mathcal{S}} - c(\mathcal{S}, \pi)$. This is the learning interpretation of Lemma A.10.²⁹

Strong duality asserts that both problems attain the same welfare, $\bar{w}^{\mathcal{S}} = \underline{w}^{\mathcal{S}}$. In the case of an optimal signal — which exists by continuity of the cost (C1) — we obtain $\underline{w}^{\mathcal{S}} = W(\mathcal{A}, \pi, c)$. Since this value is finite, duality also guarantees that the supremum is attained, ensuring that $\alpha = \arg \max_{\mathbf{x} \preceq \mathcal{S}} \pi \cdot \mathbf{x}$ is well-defined with

$$W(\{\alpha\}, \pi, c) = \bar{w}^{\mathcal{S}} = \underline{w}^{\mathcal{S}} = W(\mathcal{A}, \pi, c). \quad (14)$$

We now use this to establish existence of the ignorance equivalent.

Proof of Theorem 1: We start with existence of the ignorance equivalent and then focus on uniqueness.

Existence. Continuity of the cost function (C1) ensures that the RI problem $(\mathcal{A}, \pi, c)$ admits an optimal learning strategy $\mathcal{S}$ (Lemma A.6). By definition, dominance $\mathbf{x} \preceq \mathcal{S}$ is equivalent to $\rho \cdot \mathbf{x} \leq \Phi(\rho) \forall \rho \in \Delta \mathcal{I}$, with an upper bound function $\Phi(\rho) = \rho \cdot \mathbf{a}^{\mathcal{S}} - c(\mathcal{S}, \rho)$ that is

²⁹We will refer to the result again later in a different context, which is why we state it as a separate mathematical lemma.

convex by prior-concavity of the cost (C2). By Lemma A.10, there exists a point $\alpha \precsim \mathcal{S}$ such that the first inequality in Definition 1 binds, $W(\{\alpha\}, \pi, c) = \pi \cdot \alpha = W(\mathcal{A}, \pi, c)$.

To show that the second inequality in Definition 1 also holds, let $\tilde{\mathcal{S}} = \langle \mathcal{A} \cup \{\alpha\}, q \rangle$ denote an optimal learning strategy for RI problem $(\mathcal{A} \cup \{\alpha\}, \pi, c)$. If $q(\alpha) = \mathbf{0}$, then the strategy does not rely on the presence of α , and is thus also feasible in $(\mathcal{A}, \pi, c)$. If $q(\alpha) > \mathbf{0}$, let ρ denote the agent's posterior belief upon observing realization α . Consider the sequential strategy where the agent draws $\tilde{\mathcal{S}}$ and follows its recommendation except for realization α , when she instead draws and follows $\mathcal{S}$. The only change in the agent's payoff occurs conditional on realization α , when she achieves expected utility $\rho \cdot \alpha^{\mathcal{S}} - c(\mathcal{S}, \rho)$ rather than $\rho \cdot \alpha$. Since $\alpha \precsim \mathcal{S}$ and hence $\rho \cdot \alpha \leq \rho \cdot \alpha^{\mathcal{S}} - c(\mathcal{S}, \rho)$, the sequential strategy weakly increases welfare. A one-shot implementation of this same strategy is weakly cheaper by (C4) and thus forms a lower bound for $W(\mathcal{A}, \pi, c)$ that weakly exceeds $W(\mathcal{A} \cup \{\alpha\}, \pi, c)$.

Uniqueness. By contradiction, suppose that there exist $\alpha^1 \neq \alpha^2$ that both satisfy Definition 1. Since

$$u := W(\{\alpha^1\}, \pi, c) = W(\mathcal{A}, \pi, c) = W(\{\alpha^2\}, \pi, c),$$

the two payoff vectors achieve the same expected utility u .

By Lemma A.8, there exists a signal $\mathcal{S}^0 = \langle \{\alpha^1, \alpha^2\}, q^0 \rangle$ such that

$$\sum_{i=1}^I \pi_i (q_i^0(\alpha^1) \alpha_i^1 + q_i^0(\alpha^2) \alpha_i^2) - u > 2c(\mathcal{S}^0, \pi). \quad (15)$$

We use this signal to construct an improved strategy in menu $\mathcal{A} \cup \{\alpha^k\}$ for either $k = 1$ or $k = 2$. Specifically, suppose that the agent first draws signal $\mathcal{S}^0$. If its realization α^ℓ is available, $\ell = k$, the agent implements that action and otherwise proceeds with the optimal strategy for $(\mathcal{A}, \pi^\ell, c)$, where π^ℓ is the posterior belief after observing α^ℓ . The welfare of this strategy in menu $\mathcal{A} \cup \{\alpha^k\}$ is

$$V^k := \left[\sum_{i=1}^I \pi_i q_i^0(\alpha^k) \alpha_i^k \right] + (q^0(\alpha^{-k}) \cdot \pi) W(\mathcal{A}, \pi^{-k}, c) - c(\mathcal{S}^0, \pi).$$

It is comprised of the agent's expected continuation utility after either of the two outcomes of the binary signal $\mathcal{S}^0$, net its information costs.

The sum can be written as

$$\begin{aligned} V^1 + V^2 &= \sum_{i=1}^I \pi_i [q_i^0(\boldsymbol{\alpha}^1) \alpha_i^1 + q_i^0(\boldsymbol{\alpha}^2) \alpha_i^2] - 2c(\mathcal{S}^0, \boldsymbol{\pi}) \\ &\quad + (\mathbf{q}^0(\boldsymbol{\alpha}^1) \cdot \boldsymbol{\pi}) W(\mathcal{A}, \boldsymbol{\pi}^{\alpha^1}, c) + (\mathbf{q}^0(\boldsymbol{\alpha}^2) \cdot \boldsymbol{\pi}) W(\mathcal{A}, \boldsymbol{\pi}^{\alpha^2}, c). \end{aligned}$$

The first line is strictly larger than $W(\mathcal{A}, \boldsymbol{\pi}, c)$ by Equation (15), and the second is weakly larger than that by prior-convexity of W (Lemma A.4). As a consequence, $V^k > W(\mathcal{A}, \boldsymbol{\pi}, c)$ for at least one k . Since the strategy is feasible, it also follows that $W(\mathcal{A} \cup \{\boldsymbol{\alpha}^k\}, \boldsymbol{\pi}, c) \geq V^k$. Because the addition of $\boldsymbol{\alpha}^k$ to the menu $\mathcal{A}$ generates additional learning opportunities, it is not an ignorance equivalent. $\square$

Corollaries.

Proof of Corollary 2: Let $\boldsymbol{\alpha}$ denote the ignorance equivalent, and consider any learning strategy $\mathcal{S}$. Uniqueness of the ignorance equivalent implies that $\boldsymbol{\alpha}$ is dominated by any optimal signal according to Definition 2. So when this inequality does not hold, $\mathcal{S}$ cannot be not optimal in either RI problem.

Conversely, unconditional implementation of the ignorance equivalent must, by Definition 1, achieve at least as much utility as following signal $\mathcal{S}$ under both priors $\boldsymbol{\pi}$ and $\boldsymbol{\pi}'$. Dominance $\boldsymbol{\alpha} \succsim \mathcal{S}$ implies that the opposite inequality also holds at both $\boldsymbol{\pi}$ and $\boldsymbol{\pi}'$, and so $\mathcal{S}$ achieves maximal welfare in both RI problems. $\square$

Lemma A.11. *The ignorance equivalent of RI problem $(\mathcal{A}, \boldsymbol{\pi}, c)$ can be computed from an optimal learning strategy $\mathcal{S}$ as*

$$\boldsymbol{\alpha} = \mathbf{a}^{\mathcal{S}} - c(\mathcal{S}, \boldsymbol{\pi}) \mathbf{1} + (\mathbf{1} \boldsymbol{\pi}^T - \text{diag}(\mathbf{1})) \nabla_{\boldsymbol{\pi}} c(\mathcal{S}, \boldsymbol{\pi})$$

where $\mathbf{a}^{\mathcal{S}}$ captures the expected consumption utility that the agent achieves by following $\mathcal{S}$ conditional on state i , $a_i^{\mathcal{S}} = \sum_{\mathbf{a} \in \mathcal{A}} q_i(\mathbf{a}) a_i$, $\mathbf{1}$ is an I -dimensional vector of ones and $\text{diag}(\mathbf{1})$ is the I -dimensional identity matrix.

Proof. Unconditional implementation of the payoff vector $\boldsymbol{\alpha}$ achieves the same welfare as following signal $\mathcal{S}$, since

$$\boldsymbol{\pi} \cdot \boldsymbol{\alpha} = \boldsymbol{\pi} \cdot \mathbf{a}^{\mathcal{S}} - c(\mathcal{S}, \boldsymbol{\pi}) + \underbrace{(\boldsymbol{\pi}^T - \boldsymbol{\pi}^T) \nabla_{\boldsymbol{\pi}} c(\mathcal{S}, \boldsymbol{\pi})}_{=0}.$$

Nevertheless, the payoff vector α is dominated by signal $\mathcal{S}$, since for all $\rho \in \Delta\mathcal{I}$,

$$\rho \cdot \alpha = \rho \cdot \mathbf{a}^{\mathcal{S}} - c(\mathcal{S}, \pi) + (\pi^\top - \rho^\top) \nabla_\pi c(\mathcal{S}, \pi) \leq \rho \cdot \mathbf{a}^{\mathcal{S}} - c(\mathcal{S}, \rho),$$

where the inequality follows by prior-concavity (C2) of the cost function. Indeed, since the inequality

$$(1-t)c(\mathcal{S}, \pi) + tc(\mathcal{S}, \rho) \leq c(\mathcal{S}, (1-t)\pi + t\rho) \quad \forall t \in [0, 1]$$

is binding at $t = 0$, the local derivative with respect to t is steeper on the right,

$$c(\mathcal{S}, \rho) - c(\mathcal{S}, \pi) \leq (\rho^\top - \pi^\top) \nabla_\pi c(\mathcal{S}, \pi). \quad \square$$

Proof of Corollary 1: Consider a sequence of beliefs $\{\pi^n\}_{n=0}^\infty$ that converges to a prior $\pi^0 \in \text{int}(\Delta\mathcal{I})$, and let α^n denote the ignorance equivalent of RI problem $(\mathcal{A}, \pi^n, c)$, and $\mathcal{S}^n$ an optimal learning strategy. By Lemma A.6, the correspondence of optimal learning strategies is upper hemicontinuous with nonempty and compact values. In other words, there exists a convergent subsequence $\mathcal{S}^{n_k}$ such that $\mathcal{S}^0 = \lim_{k \rightarrow \infty} \mathcal{S}^{n_k}$ is optimal in RI problem $(\mathcal{A}, \pi^0, c)$. The $\mathcal{S}^0$ -dominated payoff vector that maximizes expected payoff under π is thus, by Theorem 1, equal to the ignorance equivalent α^0 .

Moreover, by uniqueness of the ignorance equivalent, any convergent subsequence of signals generates the exact same limit vector α^0 . It is a well-known result from real analysis that uniqueness of the limit implies that any bounded sequence,³⁰ and hence $\{\alpha^n\}_{n=0}^\infty$ itself, converges to α^0 . $\square$

Proof of Corollary 3: Let $\mathcal{S}$ and $\mathcal{S}^+$ denote optimal learning strategies for RI problems $(\mathcal{A}, \pi, c)$ and $(\mathcal{A} \cup \{\alpha\}, \rho, c)$ respectively. By Theorem 1, α is weakly dominated by $\mathcal{S}$ under any belief. In particular, this implies that whenever $\mathcal{S}^+$ recommends implementation of α at some posterior σ , the agent achieves weakly higher welfare by relying on $\mathcal{S}$ instead. The one-shot implementation $\mathcal{S}^\rho$ of this strategy is admissible in $(\mathcal{A}, \rho, c)$, yet achieves weakly higher welfare than $\mathcal{S}^+$ by (C4). This implies

$$W(\mathcal{A}, \rho, c) \geq \rho \cdot \mathbf{a}^{\mathcal{S}^\rho} - c(\mathcal{S}^\rho, \rho) \geq \rho \cdot \mathbf{a}^{\mathcal{S}^+} - c(\mathcal{S}^+, \rho) = W(\mathcal{A} \cup \{\alpha\}, \rho, c),$$

with the opposite inequality binding by menu-monotonicity. In particular, $\mathcal{S}^\rho$ is optimal in

³⁰By contradiction, suppose the sequence $\{\alpha^n\}_{n=0}^\infty$ does not converge. By definition, this implies that there exists $\varepsilon > 0$ and a subsequence $\{\alpha^{n_k}\}$ with $\|\alpha^{n_k} - \alpha^0\| > \varepsilon$ for all $k \in \mathbb{N}$. Still, the associated learning strategies have bounded conditionals, and thus admit a convergent subsequence. The Bolzano-Weierstrass Theorem asserts that this bounded subsequence admits a convergent subsubsequence, but its limit payoff vector must be different from α^0 .

both RI problems $(\mathcal{A}, \rho, c)$ and $(\mathcal{A} \cup \{\alpha\}, \rho, c)$. By [Theorem 1](#), the ignorance equivalent for both problems is thus equal to the unique $\arg \max_{x \preceq_{S\rho}} \pi \cdot x$. $\square$

A.3 Learning-Proof Menu

Proof of [Theorem 2](#): Let $\mathcal{A}'$ denote the smallest set that satisfies the two properties. Fix any prior $\pi \in \text{int}(\Delta\mathcal{I})$ and let $\mathcal{S}$ denote a learning strategy that is optimal in $(\mathcal{A}, \pi, c)$. By property (a), $\mathcal{A}'$ needs to contain the ignorance equivalent of RI problem $(\mathcal{A}', \pi, c)$. By property (b), strategy $\mathcal{S}$ is also optimal in $(\mathcal{A}', \pi, c)$. The most attractive $\mathcal{S}$ -dominated payoff vector thus identifies the ignorance equivalent under both RI problems $(\mathcal{A}', \pi, c)$ and $(\mathcal{A}, \pi, c)$ by [Theorem 1](#). In particular, the two ignorance equivalents are equal to each other, and so $\mathcal{A}'$ contains all the elements of $\bar{\mathcal{A}}$ according to [Definition 3](#).

Conversely, fix any belief π and let $\mathcal{S}'$ denote an optimal strategy in RI problem $(\mathcal{A} \cup \bar{\mathcal{A}}, \pi, c)$. By [Lemma A.1](#), it is without loss of generality to assume that $\mathcal{S}'$ only relies on finitely many actions $A \subseteq \bar{\mathcal{A}}$ in addition to those in $\mathcal{A}$. By iterative application of [Corollary 3](#), the ignorance equivalent of $(\mathcal{A} \cup A, \pi, c)$ is the same as that of $(\mathcal{A}, \pi, c)$, and hence the same is true of welfare $W(\mathcal{A} \cup \bar{\mathcal{A}}, \pi, c) = W(\mathcal{A}, \pi, c)$. By menu-monotonicity, the former is weakly larger than $W(\bar{\mathcal{A}}, \pi, c)$. The agent can achieve this upper bound through unconditional implementation of $\alpha^{(\mathcal{A}, \pi, c)}$, making ignorance optimal in $(\bar{\mathcal{A}}, \pi, c)$ and establishing property (a). For property (b), note that the equality of welfare across the two menus implies that any optimal strategy $\mathcal{S}$ in $(\mathcal{A}, \pi, c)$ remains optimal in $(\bar{\mathcal{A}}, \pi, c)$ – provided that it is still feasible. And indeed, whenever $\mathcal{S}$ recommends an action $a \in \mathcal{A}$ at some posterior ρ , [Lemma A.7](#) guarantees that a itself is the ignorance equivalent of RI problem $(\mathcal{A}, \rho, c)$, and hence is included in $\bar{\mathcal{A}}$. $\square$

Proof of [Corollary 4](#): We proceed in steps.

(a) $\Leftrightarrow$ (b): [Definition 3](#) states that $a \in \bar{\mathcal{A}}$ if and only if it is the ignorance equivalent of $(\mathcal{A}, \pi, c)$ for some prior $\pi \in \Delta\mathcal{I}$. Since $a \in \mathcal{A}$, [Definition 1](#) collapses to just requiring that unconditional implementation of a is optimal, $W(\mathcal{A}, \pi, c) = W(\{a\}, \pi, c)$.

(b) $\Leftrightarrow$ (c): Since unconditional implementation implies that a is chosen with positive probability, (b) trivially implies (c). [Lemma A.7](#) formalizes the converse claim.

(b) $\Rightarrow$ (d): Proving the contrapositive claim, assume also that there exists $\rho \in \Delta\mathcal{I}$ such that $\alpha^{(\mathcal{A}, \rho, c)} > a$. In particular, for any prior π , we have $\pi \cdot a < \pi \cdot \alpha^{(\mathcal{A}, \rho, c)}$. By self-selection of the ignorance equivalent ([Corollary 3](#)), adding $\alpha^{(\mathcal{A}, \rho, c)}$ to the menu is not welfare-enhancing under π , implying in particular that its unconditional implementation can achieve at most

utility $W(\mathcal{A}, \boldsymbol{\pi}, c)$. Taken together, at any prior $\boldsymbol{\pi}$,

$$\boldsymbol{\pi} \cdot \mathbf{a} < \boldsymbol{\pi} \cdot \boldsymbol{\alpha}^{(\mathcal{A}, \boldsymbol{\rho}, c)} \leq W(\mathcal{A}, \boldsymbol{\pi}, c),$$

proving that unconditional implementation of $\mathbf{a}$ is suboptimal.

(d) $\Rightarrow$ (a): Proving the contrapositive claim, assume that $\mathbf{a}$ is not part of the learning-proof menu $\bar{\mathcal{A}}$. Since we can write the learning-proof menu as the upper boundary of an intersection of half-spaces with positive orthogonality vectors by Equation (2), $\mathbf{a}$ lies strictly below each individual half-space. As a consequence, the ray $\{\mathbf{a} + t\mathbf{1} \mid t \geq 0\}$ crosses the learning-proof menu at some point $\boldsymbol{\alpha} \in \bar{\mathcal{A}}$, and by Definition 3, this point represents the ignorance equivalent under some prior. $\square$

Menu Expansion. In this subsection, we rely on all cost properties (C1) to (C5). For any RI problem $(\mathcal{A}, \boldsymbol{\pi}, c)$, we define the set of $\boldsymbol{\pi}$ -dominated payoff vectors

$$\bar{\mathcal{A}}^\pi = \{\mathbf{x} \in \mathbb{R}^I \mid W(\mathcal{A} \cup \{\mathbf{x}\}, \boldsymbol{\pi}, c) \leq W(\mathcal{A}, \boldsymbol{\pi}, c)\} \quad (16)$$

as those that do not increase the welfare of an agent with prior $\boldsymbol{\pi}$ when added to the menu. We first establish Theorem 3, which ensures that $\bar{\mathcal{A}}^\pi$ can equivalently be stated by replacing all references to $\mathcal{A}$ with the ignorance equivalent $\{\boldsymbol{\alpha}\}$.

Proof of Theorem 3: We start by proving part (b). Let $\mathcal{S}^0$ denote an optimal learning strategy for RI problem $(\{\boldsymbol{\alpha}, \mathbf{a}^+\}, \boldsymbol{\pi}, c)$ and $\mathcal{S}^1$ an optimal learning strategy for $(\mathcal{A}, \boldsymbol{\pi}, c)$. Consider a sequential strategy where the agent first draws $\mathcal{S}^0$ and follows its recommendation except for when it evaluates to $\boldsymbol{\alpha}$, when she draws and follows $\mathcal{S}^1$. The sequential strategy is weakly more attractive because $\mathcal{S}^1$ dominates $\boldsymbol{\alpha}$, implying $W(\mathcal{A} \cup \{\mathbf{a}^+\}, \boldsymbol{\pi}, c) \geq W(\{\boldsymbol{\alpha}, \mathbf{a}^+\}, \boldsymbol{\pi}, c)$.

We now establish part (a). The backwards implication is a direct consequence of the argument we just made, since $W(\{\boldsymbol{\alpha}, \mathbf{a}^+\}, \boldsymbol{\pi}, c) > W(\{\boldsymbol{\alpha}\}, \boldsymbol{\pi}, c)$ implies that

$$W(\mathcal{A} \cup \{\mathbf{a}^+\}, \boldsymbol{\pi}, c) \geq W(\{\boldsymbol{\alpha}, \mathbf{a}^+\}, \boldsymbol{\pi}, c) > W(\{\boldsymbol{\alpha}\}, \boldsymbol{\pi}, c) = W(\mathcal{A}, \boldsymbol{\pi}, c),$$

where the first inequality restates claim (b) and the last equality follows from the definition of the ignorance equivalent.

Next, we provide a direct proof of the forward implication,

$$W(\mathcal{A} \cup \{\mathbf{a}^+\}, \boldsymbol{\pi}, c) > W(\mathcal{A}, \boldsymbol{\pi}, c) \implies W(\{\boldsymbol{\alpha}, \mathbf{a}^+\}, \boldsymbol{\pi}, c) > W(\{\boldsymbol{\alpha}\}, \boldsymbol{\pi}, c).$$

To do so, let $\mathcal{S}^+ = \langle \mathcal{A} \cup \{\mathbf{a}^+\}, \mathbf{q}^+ \rangle$ denote an optimal learning strategy for RI problem $(\mathcal{A} \cup \{\mathbf{a}^+\}, \boldsymbol{\pi}, c)$, and $\mathcal{S} = \langle \mathcal{A}, \mathbf{q} \rangle$ one for $(\mathcal{A}, \boldsymbol{\pi}, c)$. Let $\Delta = W(\mathcal{A} \cup \{\mathbf{a}^+\}, \boldsymbol{\pi}, c) - W(\mathcal{A}, \boldsymbol{\pi}, c) > 0$ denote the difference in the welfare between the two strategies.

Consider an agent who relies on $\mathcal{S}^+$ with probability ε and on $\mathcal{S}$ otherwise. Since the gains from $\mathcal{S}^+$ are realized only with probability ε , this strategy yields welfare $W(\mathcal{A}, \boldsymbol{\pi}, c) + \varepsilon\Delta$. Note also that mixing changes the marginal likelihood that $\mathbf{a}^+$ is implemented, but not its associated posterior $\boldsymbol{\pi}^+$.

We now suggest a Blackwell-equivalent implementation strategy that proceeds in two steps: First, the agent draws a binary signal $\mathcal{S}_\varepsilon^0 = \langle \{-, +\}, \mathbf{q}_\varepsilon^0 \rangle$ which pools all realizations other than $\mathbf{a}^+$ by returning $-$ with probability $\mathbf{q}_\varepsilon^0(-) = 1 - \varepsilon \mathbf{q}^+(\mathbf{a}^+)$. If $\mathcal{S}_\varepsilon^0$ evaluates to $+$, the agent implements $\mathbf{a}^+$ at the same posterior $\boldsymbol{\pi}^+$ as above. If $\mathcal{S}_\varepsilon^0$ evaluates to $-$, the agent updates her belief to $\boldsymbol{\pi}_\varepsilon^-$ and draws $\mathcal{S}_\varepsilon^1$, which conditions on ‘not $\mathbf{a}^+$ ’ by returning $\mathbf{a} \in \mathcal{A}$ with probability $\frac{\varepsilon q_i^+(\mathbf{a}) + (1-\varepsilon)q_i(\mathbf{a})}{1-\varepsilon q_i^+(\mathbf{a})}$ in state i . Since the agent is indifferent across all Blackwell-equivalent sequential information strategies under (C4) and (C5), this strategy too yields welfare $W(\mathcal{A}, \boldsymbol{\pi}, c) + \varepsilon\Delta$. Welfare can only increase if we replace signal $\mathcal{S}_\varepsilon^1$ by the ignorance equivalent $\boldsymbol{\alpha}^\varepsilon$ of the corresponding RI problem $(\mathcal{A}, \boldsymbol{\pi}_\varepsilon^-, c)$, hence

$$(\boldsymbol{\pi} \cdot \mathbf{q}_\varepsilon^0(-))(\boldsymbol{\pi}_\varepsilon^- \cdot \boldsymbol{\alpha}^\varepsilon) + (\boldsymbol{\pi} \cdot \mathbf{q}_\varepsilon^0(+))(\boldsymbol{\pi}^+ \cdot \mathbf{a}^+) - c(\mathcal{S}_\varepsilon^0, \boldsymbol{\pi}) \geq W(\mathcal{A}, \boldsymbol{\pi}, c) + \varepsilon\Delta. \quad (17)$$

By the law of total probability, the average posterior is equal to the prior, allowing us to express the posterior $\boldsymbol{\pi}_\varepsilon^-$ as a function of the prior $\boldsymbol{\pi}$ and the posterior $\boldsymbol{\pi}^+$,

$$(\boldsymbol{\pi} \cdot \mathbf{q}_\varepsilon^0(-))\boldsymbol{\pi}_\varepsilon^- = \boldsymbol{\pi} - (\boldsymbol{\pi} \cdot \mathbf{q}_\varepsilon^0(+))\boldsymbol{\pi}^+ = \boldsymbol{\pi} - \varepsilon(\boldsymbol{\pi} \cdot \mathbf{q}^+(\mathbf{a}^+))\boldsymbol{\pi}^+,$$

with both $\boldsymbol{\pi}$ and $\boldsymbol{\pi}^+$ independent of ε .

We now show that replacing $\boldsymbol{\alpha}^\varepsilon$ with $\boldsymbol{\alpha}$ still achieves a positive welfare gain for small enough ε . First, prior-continuity of the ignorance equivalent (Corollary 1) implies that there exists $\varepsilon > 0$ such that

$$(\boldsymbol{\pi} \cdot \mathbf{q}^+(\mathbf{a}^+))\boldsymbol{\pi}^+ \cdot (\boldsymbol{\alpha} - \boldsymbol{\alpha}^\varepsilon) < \frac{\Delta}{2}.$$

Second, the self-selection property of the ignorance equivalent (Corollary 3) further implies that $\boldsymbol{\pi} \cdot (\boldsymbol{\alpha} - \boldsymbol{\alpha}^\varepsilon) \geq 0$, and so the welfare loss of replacing $\boldsymbol{\alpha}^\varepsilon$ with $\boldsymbol{\alpha}$ is bounded below by

$$(\boldsymbol{\pi} \cdot \mathbf{q}_\varepsilon^0(-))\boldsymbol{\pi}_\varepsilon^- \cdot (\boldsymbol{\alpha} - \boldsymbol{\alpha}^\varepsilon) = \boldsymbol{\pi} \cdot (\boldsymbol{\alpha} - \boldsymbol{\alpha}^\varepsilon) - \varepsilon(\boldsymbol{\pi} \cdot \mathbf{q}^+(\mathbf{a}^+))\boldsymbol{\pi}^+ \cdot (\boldsymbol{\alpha} - \boldsymbol{\alpha}^\varepsilon) > -\varepsilon \frac{\Delta}{2}. \quad (18)$$

Taken together, [Equations \(17\) and \(18\)](#) imply that the agent can achieve welfare

$$(\boldsymbol{\pi} \cdot \mathbf{q}_\varepsilon^0(-))\boldsymbol{\pi}_\varepsilon^- \cdot \boldsymbol{\alpha} + (\boldsymbol{\pi} \cdot \mathbf{q}_\varepsilon^0(+))(\boldsymbol{\pi}^+ \cdot \mathbf{a}^+) - c(\mathcal{S}_\varepsilon^0, \boldsymbol{\pi}) \geq W(\mathcal{A}, \boldsymbol{\pi}, c) + \varepsilon \frac{\Delta}{2}$$

by drawing $\mathcal{S}_\varepsilon^0$ and implementing $\boldsymbol{\alpha}$ upon realization $-$ and $\mathbf{a}^+$ otherwise. Since this strategy is feasible in RI problem $(\{\boldsymbol{\alpha}, \mathbf{a}^+\}, \boldsymbol{\pi}, c)$, it implies in particular that $W(\{\boldsymbol{\alpha}, \mathbf{a}^+\}, \boldsymbol{\pi}, c) > W(\{\boldsymbol{\alpha}\}, \boldsymbol{\pi}, c)$. $\square$

We now turn back to our definition of $\boldsymbol{\pi}$ -dominated payoff vectors in [Equation \(16\)](#) and rewrite it in a way that fits [Lemma A.10](#). To do so, we let $\mathcal{S}_\pi^{(k, \sigma)}$ denote the binary signal that, starting from prior $\boldsymbol{\pi}$, reaches posterior $\boldsymbol{\sigma}$ with marginal likelihood $k \in (0, 1)$. By the law of total probability, the same signal reaches the opposing posterior $\boldsymbol{\sigma}^{-k} := \frac{1}{1-k}\boldsymbol{\pi} - \frac{k}{1-k}\boldsymbol{\sigma}$ with marginal likelihood $1 - k$ for a cost of

$$c(\mathcal{S}_\pi^{(k, \sigma)}, \boldsymbol{\pi}) = k\phi(\boldsymbol{\sigma}) + (1 - k)\phi(\boldsymbol{\sigma}^{-k}) - \phi(\boldsymbol{\pi}).$$

For $k > 0$ small enough, all probabilities are non-degenerate. We also define function $\Phi_\pi : \Delta\mathcal{I} \rightarrow \mathbb{R}$ as

$$\begin{aligned} \Phi_\pi(\boldsymbol{\sigma}) &:= \boldsymbol{\sigma} \cdot \boldsymbol{\alpha}^{(\mathcal{A}, \boldsymbol{\pi}, c)} + \lim_{k \downarrow 0} \frac{1}{k} c(\mathcal{S}_\pi^{(k, \sigma)}, \boldsymbol{\pi}) \\ &= \boldsymbol{\sigma} \cdot \boldsymbol{\alpha}^{(\mathcal{A}, \boldsymbol{\pi}, c)} + \phi(\boldsymbol{\sigma}) - \phi(\boldsymbol{\pi}) + \nabla \phi(\boldsymbol{\pi}) \cdot (\boldsymbol{\pi} - \boldsymbol{\sigma}), \end{aligned} \tag{19}$$

and use it to restate $\bar{\mathcal{A}}^\pi$.

Lemma A.12. *The function Φ_π given in [Equation \(19\)](#) is convex and finite-valued, and $\bar{\mathcal{A}}^\pi = \{\mathbf{x} \in \mathbb{R}^I \mid \boldsymbol{\sigma} \cdot \mathbf{x} \leq \Phi_\pi(\boldsymbol{\sigma}) \ \forall \boldsymbol{\sigma} \in \Delta\mathcal{I}\}$.*

Proof. By [Theorem 3\(a\)](#), $\bar{\mathcal{A}}^\pi$ can be stated with reference to the ignorance equivalent $\boldsymbol{\alpha} = \boldsymbol{\alpha}^{(\mathcal{A}, \boldsymbol{\pi}, c)}$ alone, $\bar{\mathcal{A}}^\pi = \{\mathbf{x} \in \mathbb{R}^I \mid W(\{\boldsymbol{\alpha}, \mathbf{x}\}, \boldsymbol{\pi}, c) \leq W(\{\boldsymbol{\alpha}\}, \boldsymbol{\pi}, c)\}$. And since the menu $\{\boldsymbol{\alpha}, \mathbf{x}\}$ is binary, any feasible strategy can be described as drawing $\mathcal{S}_\pi^{(k, \sigma)}$ for some k , and implementing $\mathbf{x}$ at posterior $\boldsymbol{\rho}$ and $\boldsymbol{\alpha}$ otherwise. Relative to unconditional implementation of $\boldsymbol{\alpha}$, relying on this costly signal improves welfare by

$$k\boldsymbol{\sigma} \cdot \mathbf{x} + (\boldsymbol{\pi} - k\boldsymbol{\sigma}) \cdot \boldsymbol{\alpha} - c(\mathcal{S}_\pi^{(k, \sigma)}, \boldsymbol{\pi}) - \boldsymbol{\pi} \cdot \boldsymbol{\alpha} = k\boldsymbol{\sigma} \cdot (\mathbf{x} - \boldsymbol{\alpha}) - c(\mathcal{S}_\pi^{(k, \sigma)}, \boldsymbol{\pi}),$$

which has the same sign as $\boldsymbol{\sigma} \cdot (\mathbf{x} - \boldsymbol{\alpha}) - \frac{1}{k}c(\mathcal{S}_\pi^{(k, \sigma)}, \boldsymbol{\pi})$. As k converges to zero from above, $k \downarrow 0$, the weighted cost term shrinks. This is because for any $t \in [0, 1]$, $\mathcal{S}_\pi^{(tk, \sigma)}$ can be implemented in two steps by an uninformative coin flip that triggers a draw of $\mathcal{S}_\pi^{(k, \sigma)}$ with probability

t , and otherwise recommends action α . By Blackwell monotonicity (C2) and sequential optimality (C4), this implies the monotonic ranking $\frac{1}{tk}c(\mathcal{S}_\pi^{(tk, \sigma)}, \pi) \leq \frac{1}{k}c(\mathcal{S}_\pi^{(k, \sigma)}, \pi)$. In other words, the signal is welfare-enhancing for some $k > 0$ if and only if

$$\sigma \cdot x > \inf_{k>0} \left\{ \sigma \cdot \alpha + \frac{1}{k}c(\mathcal{S}_\pi^{(k, \sigma)}, \pi) \right\} = \Phi_\pi(\sigma) \quad \forall \sigma \in \Delta\mathcal{I},$$

warranting the suggested formula for $\bar{\mathcal{A}}^\pi$. And since costs are bounded below by zero, the infimum $\Phi_\pi(\sigma)$ is finite-valued.

It remains to show that the function $\Phi_\pi : \Delta\mathcal{I} \rightarrow \mathbb{R}$ is convex. Consider a two-stage process that first flips an uninformative coin that triggers a draw of $\mathcal{S}_\pi^{(k, \sigma)}$ with probability t , and otherwise triggers a draw of $\mathcal{S}_\pi^{(k, \sigma')}$. Overall, this process recommends action x at posterior $t\sigma + (1-t)\sigma'$ with marginal likelihood k . Again, the one-shot implementation is weakly cheaper by (C2) and (C4), $c(\mathcal{S}^{(k, t\sigma + (1-t)\sigma')}, \pi) \leq tc(\mathcal{S}_\pi^{(k, \sigma)}, \pi) + (1-t)c(\mathcal{S}_\pi^{(k, \sigma')}, \pi)$. Convexity follows since

$$\begin{aligned} \Phi_\pi(t\sigma + (1-t)\sigma') &= (t\sigma + (1-t)\sigma') \cdot \alpha + \lim_{k \downarrow 0} \frac{1}{k}c(\mathcal{S}_\pi^{(k, t\sigma + (1-t)\sigma')}, \pi) \\ &\leq t \left[\sigma \cdot \alpha + \lim_{k \downarrow 0} \frac{1}{k}c(\mathcal{S}_\pi^{(k, \sigma)}, \pi) \right] + (1-t) \left[\sigma' \cdot \alpha + \lim_{k \downarrow 0} \frac{1}{k}c(\mathcal{S}_\pi^{(k, \sigma')}, \pi) \right] \\ &= t\Phi_\pi(\sigma) + (1-t)\Phi_\pi(\sigma'). \quad \square \end{aligned}$$

We now show that under any other prior ρ , the set $\bar{\mathcal{A}}^\pi$ contains a point that is implemented unconditionally when present.

Lemma A.13. *For any prior $\rho \in \text{int}(\Delta\mathcal{I})$, unconditional implementation of the payoff vector $\mathbf{a}^\rho = \arg \max \{ \rho \cdot x \mid x \in \bar{\mathcal{A}}^\pi \}$ is optimal in any RI problem $(\mathcal{A}', \rho, c)$ where the finite menu $\mathcal{A}' \subseteq \mathcal{A} \cup \bar{\mathcal{A}}^\pi$ contains $\mathbf{a}^\rho$.*

Proof. By Lemmas A.10 and A.12, the point $\mathbf{a}^\rho$ satisfies $\rho \cdot \mathbf{a}^\rho = \Phi_\pi(\rho)$. By contradiction, assume that there exists a learning strategy $\mathcal{S} = \langle \mathcal{A}', \mathbf{q} \rangle$ that achieves higher welfare than unconditional implementation of $\mathbf{a}^\rho$ under prior ρ ,

$$\sum_{i \in \mathcal{I}} \rho_i \sum_{\mathbf{a} \in \mathcal{A}'} q_i(\mathbf{a}) a_i - c(\mathcal{S}, \rho) > \rho \cdot \mathbf{a}^\rho = \Phi_\pi(\rho).$$

By definition of Φ_π , this implies that there exists $k > 0$ small enough such that $\sum_{i \in \mathcal{I}} \rho_i \sum_{\mathbf{a} \in \mathcal{A}'} q_i(\mathbf{a}) a_i - c(\mathcal{S}, \rho) > \rho \cdot \alpha + \frac{1}{k}c(\mathcal{S}_\pi^{(k, \rho)}, \pi)$. Rearranging terms, we obtain

$$k \sum_{i \in \mathcal{I}} \rho_i \sum_{\mathbf{a} \in \mathcal{A}'} q_i(\mathbf{a}) a_i + (\pi - k\rho) \cdot \alpha - c(\mathcal{S}_\pi^{(k, \rho)}, \pi) - kc(\mathcal{S}, \rho) > \pi \cdot \alpha.$$

This strict inequality implies that unconditional implementation of α is not optimal in RI problem $(\mathcal{A}' \cup \{\alpha\}, \pi, c)$; it is strictly dominated by a sequential strategy where $\mathcal{S}_\pi^{(k, \rho)}$ is drawn first, and upon realization $\mathbf{x}$, the agent draws and follows signal $\mathcal{S}$, otherwise she implements α . This contradicts the binary characterization for optimality of ignorance (Theorem 5) since by definition of $\bar{\mathcal{A}}^\pi$, none of the actions $\mathbf{a} \in \mathcal{A}' \subseteq \mathcal{A} \cup \bar{\mathcal{A}}^\pi$ yields a welfare improvement by itself. $\square$

Building on this, we prove Theorem 4, which states that any anchor can be ‘activated’ by adding the right action to the menu.

Proof of Theorem 4: Suppose first that $\mathbf{a}$ is an anchor action. Since anchor $\mathbf{a}$ is part of the learning-proof menu, there exists a prior $\rho \in \text{int}(\Delta\mathcal{I})$ such that $\mathbf{a}$ is implemented unconditionally in $(\mathcal{A}, \rho, c)$, and thus denotes the ignorance equivalent of that learning problem. Fix any $t \in (0, 1)$ small enough such that the belief $\rho^+ = \rho + \frac{1}{1-t}(\pi - \rho)$ has full support. By Lemmas A.10 and A.12, there exist payoff vectors $\alpha, \mathbf{a}^+ \in \bar{\mathcal{A}}^\rho$ such that

$$\pi \cdot \alpha = \Phi_\rho(\pi) \quad \text{and} \quad \rho^+ \cdot \mathbf{a}^+ = \Phi_\rho(\rho^+). \quad (20)$$

Moreover, by Lemma A.13, unconditional implementation of α and $\mathbf{a}^+$ is optimal in RI problems $(\mathcal{A} \cup \{\mathbf{a}^+, \alpha\}, \pi, c)$ and $(\mathcal{A} \cup \{\mathbf{a}^+, \alpha\}, \rho^+, c)$, respectively. In particular,

$$W(\mathcal{A} \cup \{\mathbf{a}^+\}, \pi, c) \leq W(\mathcal{A} \cup \{\mathbf{a}^+, \alpha\}, \pi, c) = \pi \cdot \alpha.$$

We now show that the agent can achieve this upper bound for welfare by following a learning strategy $\mathcal{S}_\pi^{(t, \rho)}$ that implements action $\mathbf{a}$ at posterior ρ with marginal likelihood t , and otherwise implements action $\mathbf{a}^+$ at posterior ρ^+ .

To do so, note that for any $k > 0$ small enough, the following two sequential learning strategies are Blackwell equivalent for an agent with prior ρ :

- Draw $\mathcal{S}_\rho^{(k, \pi)}$ and then, conditional on reaching posterior π , draw $\mathcal{S}_\pi^{(t, \rho)}$.
- Flip an uninformative coin that triggers a draw of $\mathcal{S}_\rho^{(k(1-t)/(1-tk), \rho^+)}$ with probability $1 - tk$ and reveals no information otherwise.

Since sequential implementation yields neither gains nor losses by (C4) and (C5), both strategies have the same expected cost,

$$c(\mathcal{S}_\rho^{(k, \pi)}, \rho) + kc(\mathcal{S}_\pi^{(t, \rho)}, \pi) = (1 - tk)c(\mathcal{S}_\rho^{(k(1-t)/(1-tk), \rho^+)}, \rho).$$

Rearranging terms and taking the limit $k \downarrow 0$, we find

$$\begin{aligned}
c(\mathcal{S}_{\pi}^{(t,\rho)}, \pi) &= (1-t) \lim_{k \downarrow 0} \frac{1-tk}{k(1-t)} c(\mathcal{S}_{\rho}^{(k(1-t)/(1-tk), \rho^+)}, \rho) - \lim_{k \downarrow 0} \frac{1}{k} c(\mathcal{S}_{\rho}^{(k,\pi)}, \rho) \\
&\stackrel{(19)}{=} (1-t)[\Phi_{\rho}(\rho^+) - \rho^+ \cdot \mathbf{a}] - [\Phi_{\rho}(\pi) - \pi \cdot \mathbf{a}] \\
&\stackrel{(20)}{=} (1-t)\rho^+ \cdot (\mathbf{a}^+ - \mathbf{a}) - \pi \cdot (\alpha - \mathbf{a}).
\end{aligned}$$

Subtracting this cost from the consumption utility $t\rho \cdot \mathbf{a} + (1-t)\rho^+ \cdot \mathbf{a}^+$, we obtain the welfare under signal $\mathcal{S}_{\pi}^{(t,\rho)}$,

$$\begin{aligned}
&t\rho \cdot \mathbf{a} + (1-t)\rho^+ \cdot \mathbf{a}^+ - [(1-t)\rho^+ \cdot (\mathbf{a}^+ - \mathbf{a}) - \pi \cdot (\alpha - \mathbf{a})] \\
&= \underbrace{[t\rho + (1-t)\rho^+] \cdot \mathbf{a}}_{=\pi} + \pi \cdot (\alpha - \mathbf{a}) = \pi \cdot \alpha.
\end{aligned}$$

As a result, it is optimal for the agent to implement $\mathbf{a}$ with positive probability $t > 0$ in RI problem $(\mathcal{A} \cup \{\mathbf{a}^+\}, \pi, c)$.

Conversely, suppose $\mathbf{a} \in \mathcal{A} \setminus \bar{\mathcal{A}}$ is not an anchor. By [Corollary 4\(d\)](#), there exists a prior ρ such that $\mathbf{a} < \alpha^{(\mathcal{A}, \rho, c)}$. By construction of the ignorance equivalent in [Theorem 1](#), $\mathbf{a}$ is thus strictly dominated by any optimal learning strategy $\mathcal{S}$ of RI problem $(\mathcal{A}, \rho, c)$. In particular, as long as the menu contains $\mathcal{A}$, the agent is always strictly better off following strategy $\mathcal{S}$ rather than implementing $\mathbf{a}$, no matter her belief. By sequential optimality [\(C4\)](#), implementing $\mathbf{a}$ is therefore also strictly suboptimal among one-shot strategies. $\square$

Proof of [Corollary 6](#): The agent's willingness to pay is equal to the expected change in welfare, $\sum_{s \in S} (\pi \cdot \mathbf{q}(s))W(\mathcal{A}, \pi^s, c) - W(\mathcal{A}, \pi, c)$. The claim then follows because [Theorem 2](#) implies that for all priors $\rho \in \text{int}(\Delta \mathcal{I})$,

$$W(\mathcal{A}, \rho, c) = \rho \cdot \alpha^{(\mathcal{A}, \rho, c)} \stackrel{2(b)}{=} \rho \cdot \alpha^{(\bar{\mathcal{A}}, \rho, c)} = W(\bar{\mathcal{A}}, \rho, c) \stackrel{2(a)}{=} \max_{\mathbf{a} \in \bar{\mathcal{A}}} \rho \cdot \mathbf{a}.$$

The result generalizes to the boundary of the simplex by continuity of the welfare function [Lemma A.6](#). $\square$

Proof of [Corollary 7](#): Suppose the principal offers the menu $\mathcal{A}^0 \cup \{\mathbf{a}^s \mid s \in S\}$ with $\mathbf{a}^s = \arg \max \{\pi^s \cdot \mathbf{a} \mid \mathbf{a} \in \bar{\mathcal{A}}^{\pi}\}$. Since none of the added options are welfare-enhancing, the agent's welfare is bounded above by $\pi \cdot \alpha$, and we now show that the agent achieves this upper bound by drawing $\mathcal{S}$ and selecting action $\mathbf{a}^s$ after realization s . Doing so yields consumption utility $\pi^s \cdot \mathbf{a}^s$ with marginal likelihood $q(s) = \pi \cdot \mathbf{q}(s)$ for a cost of $c(\mathcal{S}, \pi)$. By [Lemmas A.10](#)

and A.12, we can write

$$\begin{aligned}
\sum_{s \in S} q(s)(\pi^s \cdot \mathbf{a}^s) - c(\mathcal{S}, \pi) &= \sum_{s \in S} q(s) \Phi_\pi(\pi^s) - c(\mathcal{S}, \pi) \\
&= \sum_{s \in S} q(s) \left(\pi^s \cdot \boldsymbol{\alpha} + \lim_{k \downarrow 0} \frac{1}{k} c(\mathcal{S}_\pi^{(k, \pi^s)}, \pi) \right) - c(\mathcal{S}, \pi) \\
&= \pi \cdot \boldsymbol{\alpha} + \lim_{k \downarrow 0} \frac{1}{k} \left(\sum_{s \in S} q(s) c(\mathcal{S}_\pi^{(k, \pi^s)}, \pi) \right) - c(\mathcal{S}, \pi)
\end{aligned} \tag{21}$$

The term in parentheses captures the cost of a sequential strategy where an assistant first draws an uninformative signal that returns each $s \in S$ with unconditional probability $q(s)$, and then flips an informative coin that yields heads at posterior π^s with probability k . Suppose the agent announces s if the coin returns heads and says nothing otherwise. The assistant's announcement is then Blackwell equivalent to a strategy where the agent draws $\mathcal{S}$ with probability k and maintains prior π otherwise. Since the agent is indifferent across all sequential strategies (C4) and (C5) and costs are Blackwell-monotone (C2), the term $\sum_{s \in S} q(s) c(\mathcal{S}_\pi^{(k, \pi^s)}, \pi)$ is equal to $kc(\mathcal{S}, \pi)$. From Equation (21), it follows that $\mathcal{S}$ achieves $\pi \cdot \boldsymbol{\alpha}$ and is thus indeed optimal for the agent. And since the agent receives the payoff $\pi \cdot \boldsymbol{\alpha} = W(\mathcal{A}^0, \pi, c)$ that he is guaranteed in any offer $A \supseteq \mathcal{A}^0$, the principal is making minimal transfers. $\square$

Buyouts. Inductive application of Theorem 3 yields a binary characterization of situations where ignorance is optimal.

Proof of Theorem 5: Suppose first that unconditional implementation of $\mathbf{a}^+ \in \mathcal{A}$ is optimal in RI problem $(\mathcal{A} \cup \{\mathbf{a}^+\}, \pi, c)$. By menu-monotonicity, $W(\mathcal{A} \cup \{\mathbf{a}^+\}, \pi, c) \geq W(\{\mathbf{a}^+, \mathbf{a}\}, \pi, c)$ for each $\mathbf{a} \in \mathcal{A}$, and so (b) follows trivially from (a).

Conversely, suppose condition $W(\{\mathbf{a}^+\}, \pi, c) \geq W(\{\mathbf{a}^+, \mathbf{a}\}, \pi, c)$ holds for each $\mathbf{a} \in \mathcal{A}$. We prove by induction that unconditional implementation of $\mathbf{a}^+$ is optimal in RI problem (A, π, c) , starting with the trivial case $A = \{\mathbf{a}^+\}$ and adding actions one-by-one. For the inductive step, assume $W(\{\mathbf{a}^+\}, \pi, c) = W(A, \pi, c)$ for some subset $A \subseteq \mathcal{A} \cup \{\mathbf{a}^+\}$, and consider what happens when $\mathbf{a} \in \mathcal{A} \setminus A$ is added to the menu. Since $\mathbf{a}^+ \in A$, the assumption satisfies all the conditions of Definition 1, and thus $\mathbf{a}^+$ denotes the ignorance equivalent of (A, π, c) . By Theorem 3(a), the condition $W(\{\mathbf{a}^+, \mathbf{a}\}, \pi, c) \leq W(\{\mathbf{a}^+\}, \pi, c)$ then implies that $W(A \cup \{\mathbf{a}\}, \pi, c) \leq W(A, \pi, c)$. Menu-monotonicity implies the opposite inequality and thus

$$W(A \cup \{\mathbf{a}\}, \pi, c) = W(A, \pi, c) = W(\{\mathbf{a}^+\}, \pi, c). \quad \square$$

A.4 Posterior-separable Costs

Although our proof of [Theorem 1](#) relies on prior-concavity (C2) and sequential learning (C4), several results also hold for the larger class of posterior-separable cost,³¹ with

$$c(\langle S, \mathbf{q} \rangle, \boldsymbol{\pi}) = \sum_{s \in S} (\boldsymbol{\pi} \cdot \mathbf{q}(s)) \phi^\pi(\boldsymbol{\pi}^s) - \phi^\pi(\boldsymbol{\pi}). \quad (\text{PS})$$

This specification differs from [Equation \(UPS+\)](#) in that the convex and differentiable potential ϕ^π can depend on the prior, and does not need to induce prior-concavity.

Let $\mathcal{S}$ denote a learning strategy that is optimal in RI problem $(\mathcal{A}, \boldsymbol{\pi}, c)$. By [Caplin et al. \[2022, Lemma 1\]](#), there exists a vector $\boldsymbol{\theta} \in \mathbb{R}^I$ such that, for any $\mathbf{a}$ that is chosen with positive probability and for any optimal posterior $\boldsymbol{\pi}^s$,

$$\boldsymbol{\rho} \cdot \mathbf{a}' - \phi^\pi(\boldsymbol{\rho}) - \boldsymbol{\rho} \cdot \boldsymbol{\theta} \leq \boldsymbol{\pi}^s \cdot \mathbf{a} - \phi^\pi(\boldsymbol{\pi}^s) - \boldsymbol{\pi}^s \cdot \boldsymbol{\theta} \quad \forall \boldsymbol{\rho} \in \Delta \mathcal{I}, \forall \mathbf{a}' \in \mathcal{A}. \quad (22)$$

Moreover, the existence of such a $\boldsymbol{\theta}$ guarantees optimality of the signal $\mathcal{S}$.

By symmetry, the value on right side of inequality (22) needs to be the same, independently of the choice of the optimal action $\mathbf{a}$ or the optimal posterior $\boldsymbol{\pi}^s$. Moreover, since the inequality also holds for any $\boldsymbol{\theta} + k\mathbf{1}$ with $k \in \mathbb{R}$, it is without loss of generality to normalize the right side to zero. Under this normalization, $\boldsymbol{\theta}$ corresponds, in [Figure 3](#), to the end-points of the hyperplane that is tangent to V at the prior. We can now also express welfare succinctly as

$$W(\mathcal{A}, \boldsymbol{\pi}, c) = \sum_{\mathbf{a} \in \mathcal{A}} (\boldsymbol{\pi} \cdot \mathbf{q}(\mathbf{a})) (\boldsymbol{\pi}^{\mathbf{a}} \cdot \mathbf{a} - \phi^\pi(\boldsymbol{\pi}^{\mathbf{a}})) + \phi^\pi(\boldsymbol{\pi}) = \boldsymbol{\pi} \cdot \boldsymbol{\theta} + \phi^\pi(\boldsymbol{\pi}).$$

The ignorance equivalent can be obtained from this optimal Lagrangian vector as

$$\boldsymbol{\alpha} := \nabla \phi^\pi(\boldsymbol{\pi}) + \boldsymbol{\theta} + (\phi^\pi(\boldsymbol{\pi}) - \boldsymbol{\pi} \cdot \nabla \phi^\pi(\boldsymbol{\pi})) \mathbf{1}. \quad (23)$$

Indeed, this payoff vector $\boldsymbol{\alpha}$ achieves the same welfare as the optimal signal since

$$\boldsymbol{\pi} \cdot \boldsymbol{\alpha} = \boldsymbol{\pi} \cdot \boldsymbol{\theta} + \phi^\pi(\boldsymbol{\pi}) = W(\mathcal{A}, \boldsymbol{\pi}, c). \quad (24)$$

Moreover, the addition of $\boldsymbol{\alpha}$ to the menu does not introduce any new learning opportunities since

$$\boldsymbol{\rho} \cdot \boldsymbol{\alpha} - \phi^\pi(\boldsymbol{\rho}) - \boldsymbol{\rho} \cdot \boldsymbol{\theta} = \phi^\pi(\boldsymbol{\pi}) - \phi^\pi(\boldsymbol{\rho}) - (\boldsymbol{\pi} - \boldsymbol{\rho}) \cdot \nabla \phi^\pi(\boldsymbol{\pi}) \leq 0 \quad (25)$$

³¹We thank an anonymous referee for pointing this out.

by concavity of ϕ . In other words, Equation (22) holds with the same θ for the degenerate strategy that puts full weight on α , making this optimal even in the large menu $\mathcal{A} \cup \{\alpha\}$. Moreover, both constraints (24) and (25) have to be satisfied, making the payoff vector in (23) the unique candidate.

Theorems 3 and 4 can also (partly) be derived from the posterior-based representation in Figure 3. For Theorem 3, note that a new action raises welfare only if it pushes the upper envelope V upwards at the prior π . This happens if and only if its net utility lies above the tangent hyperplane $\rho \cdot \theta$, which is common to V and the net utility curve of α . For Theorem 4, note that the net payoff of actions chosen with positive probability each touch the tangent hyperplane at some posterior ρ^a , and the prior can be written as a convex combination that puts positive weight on each ρ^a . Thus, to activate an anchor a through menu expansion, one can draw a tangent hyperplane H^ρ at a posterior ρ where the net payoff of a coincides with V . Adding an action whose net payoff touches H^ρ on the ‘other side’ of π , e.g. at $\pi - \varepsilon(\rho - \pi)$ for $\varepsilon > 0$, thus ‘activates’ action a . These results, too, hold therefore for all posterior-separable functions, provided ‘anchor actions’ are well-defined. (More on that below.)

The results that rely on comparisons across priors do not readily generalize to the larger family of posterior-separable costs. Still, for prior-concave and *uniformly* posterior separable costs, it is intuitively correct to summarize Theorem 2 as saying that the maximum over net utility curves of all ignorance equivalents generate the same upper envelope as V (without the need for concavification, i.e. learning regions).

Despite this important intuition, some caution is warranted: Not all (uniformly) posterior separable costs satisfy self-selection. Online Appendix B.2.2 provides a counterexample where prior-concavity is the only missing property, yet one agent’s ignorance equivalent ends up being welfare-enhancing for an agent with a different prior. Visually, the rounded cone that represents the net payoff of the ignorance equivalent is not everywhere below the upper envelope V in this three state example, and thus agents strictly prefer the union of ignorance equivalents to the original menu.

The failure of self-selection undercuts the very construction of the learning-proof menu and its related results (including the characterization of anchor actions). Our results thus emphasize the importance of prior-concavity in a way other papers do not — and since prior-concavity cannot readily be assessed from posterior-based illustrations such as Figure 3, some care must be taken to separately verify this property before relying on the graphs for intuition.

A.5 Optimality conditions

Here, we restate the ignorance equivalent characterization (6) first for generic (UPS+) costs, and then for the most-studied case of Shannon entropy.

To do so, we spell out the welfare from a generic binary strategy that implements $\mathbf{a}$ at some posterior $\boldsymbol{\rho} \in \Delta\mathcal{I}$ with marginal likelihood t and $\boldsymbol{\alpha}$ otherwise. (By Bayes plausibility, $\boldsymbol{\alpha}$ is implemented at posterior $\boldsymbol{\rho}_t^- := \boldsymbol{\rho} - \frac{1}{1-t}(\boldsymbol{\pi} - \boldsymbol{\rho})$, which is a valid probability for small enough t). We obtain

$$\begin{aligned} W(\{\mathbf{a}^+\}, \boldsymbol{\pi}, c) &\geq W(\{\mathbf{a}^+, \mathbf{a}\}, \boldsymbol{\pi}, c) \\ \Leftrightarrow \quad \boldsymbol{\pi} \cdot \mathbf{a}^+ &\geq t(\boldsymbol{\rho} \cdot \mathbf{a}) + (1-t)(\boldsymbol{\rho}_t^- \cdot \mathbf{a}^+) - t\phi(\boldsymbol{\rho}) - (1-t)\phi(\boldsymbol{\rho}_t^-) + \phi(\boldsymbol{\pi}) \\ \Leftrightarrow \quad 0 &\geq \chi(t, \boldsymbol{\rho}) := t\boldsymbol{\rho} \cdot (\mathbf{a} - \mathbf{a}^+) - t\phi(\boldsymbol{\rho}) + \phi(\boldsymbol{\pi}) - (1-t)\phi\left(\boldsymbol{\rho} + \frac{\boldsymbol{\pi} - \boldsymbol{\rho}}{1-t}\right). \end{aligned}$$

From the convexity of ϕ , it follows that χ is concave in t ,³² and $\chi(0, \boldsymbol{\rho}) \equiv 0$.

Thus, it is necessary and sufficient for the derivative $\frac{d}{dt}\chi$ to be negative at $t = 0$,

$$\left. \frac{d}{dt}\chi(t, \boldsymbol{\rho}) \right|_{t=0} = \boldsymbol{\rho} \cdot (\mathbf{a} - \mathbf{a}^+) - \phi(\boldsymbol{\rho}) + \phi(\boldsymbol{\pi}) - \nabla\phi(\boldsymbol{\pi}) \cdot (\boldsymbol{\pi} - \boldsymbol{\rho}) \leq 0.$$

In turn, this expression is concave in $\boldsymbol{\rho}$, and admits a unique maximum. By referring to this ‘most tempting posterior’ $\boldsymbol{\rho}^{\mathbf{a}}$ for each action $\mathbf{a} \in \mathcal{A}$, as given in Equation (7), we can thus also characterize the ignorance equivalent as

$$\boldsymbol{\alpha} = \arg \min_{\substack{\mathbf{a}^+ \in \mathbb{R}^I \\ \left. \frac{d}{dt}\chi(t, \boldsymbol{\rho}^{\mathbf{a}}) \right|_{t=0} \leq 0}} \boldsymbol{\pi} \cdot \mathbf{a}^+ \quad \forall \mathbf{a} \in \mathcal{A}. \quad (26)$$

To illustrate how this characterization can be helpful for computational methods, we now simplify it for Shannon entropy costs with $\phi(\boldsymbol{\pi}) = \lambda \boldsymbol{\pi} \cdot \ln(\boldsymbol{\pi})$. Substituting this potential in Equation (7), we find the most tempting posterior by solving $\boldsymbol{\rho}_i^{\mathbf{a}} = \arg \max_{\boldsymbol{\rho}} \boldsymbol{\rho} \cdot (\mathbf{a} - \boldsymbol{\alpha} + \lambda(\ln(\boldsymbol{\pi}) + \mathbf{1})) - \lambda \sum_{i \in \mathcal{I}} \rho_i \ln(\rho_i)$. This belief is interior since the slope of the last term approaches $+\infty$ as $\rho_i \rightarrow 0$ for any i . Moreover, optimality implies that the impact of a

³²It is a simple exercise in algebra to check that $\frac{1}{1-\lambda t_1 - (1-\lambda)t_2} = g(\lambda)\frac{1}{1-t_1} + (1-g(\lambda))\frac{1}{1-t_2}$ for $g(\lambda) = \frac{\lambda(1-t_1)}{1-\lambda t_1 - (1-\lambda)t_2}$. From there, it follows that for $t = \lambda t_1 + (1-\lambda)t_2$ and $\boldsymbol{\rho}^-(t) := \boldsymbol{\rho} - \frac{1}{1-t}(\boldsymbol{\pi} - \boldsymbol{\rho})$,

$$\begin{aligned} -(1-t)\phi(\boldsymbol{\rho}^-(t)) &= -(1-t)\phi(g(\lambda)\boldsymbol{\rho}^-(t_1) + (1-g(\lambda))\boldsymbol{\rho}^-(t_2)) \\ &\geq -\lambda(1-t_1)\phi(\boldsymbol{\rho}^-(t_1)) - (1-\lambda)(1-t_2)\phi(\boldsymbol{\rho}^-(t_2)) \end{aligned}$$

by convexity of ϕ . Since all other terms in χ are linear, concavity follows.

marginal deviation $dt\Delta\boldsymbol{\rho}$ with $\Delta\boldsymbol{\rho} \cdot \mathbf{1} = 0$ must be zero,

$$\begin{aligned} 0 &= \frac{d}{dt} \left[(\boldsymbol{\rho}^a + t\Delta\boldsymbol{\rho}) \cdot (\mathbf{a} - \boldsymbol{\alpha} + \lambda[1 + \ln(\boldsymbol{\pi})]) - \lambda \sum_{i \in \mathcal{I}} (\rho_i + t\Delta\rho_i) \ln(\rho_i + t\Delta\rho_i) \right] \Big|_{t=0} \\ &= \Delta\boldsymbol{\rho} \cdot (\mathbf{a} - \boldsymbol{\alpha} + \lambda(\ln(\boldsymbol{\pi}) - \ln(\boldsymbol{\rho}^a))). \end{aligned}$$

In other words, the vector $(\mathbf{a} - \boldsymbol{\alpha} + \lambda(\ln(\boldsymbol{\pi}) - \ln(\boldsymbol{\rho}^a)))$ is colinear with $\mathbf{1}$, and hence $\boldsymbol{\rho}^a$ is colinear with $\boldsymbol{\pi}e^{(\mathbf{a}-\boldsymbol{\alpha})/\lambda}$. The scaling factor ensures that $\boldsymbol{\rho}^a$ is a proper probability,

$$\rho_i^a \equiv \frac{\pi_i e^{(a_i - \alpha_i)/\lambda}}{\sum_{i \in \mathcal{I}} \pi_i e^{(a_i - \alpha_i)/\lambda}}.$$

It is then a straightforward exercise in algebra to confirm that

$$\begin{aligned} \phi(\boldsymbol{\rho}^a) &= \boldsymbol{\rho}^a \cdot \nabla \phi(\boldsymbol{\pi}) - \lambda + \boldsymbol{\rho}^a \cdot (\mathbf{a} - \boldsymbol{\alpha}) - \lambda \ln \left(\sum_{j \in \mathcal{I}} \pi_j e^{(a_j - \alpha_j)/\lambda} \right), \\ \phi(\boldsymbol{\pi}) &= \boldsymbol{\pi} \cdot \nabla \phi(\boldsymbol{\pi}) - \lambda \end{aligned}$$

and hence

$$\frac{d}{dt} \chi(t, \boldsymbol{\rho}^a) \Big|_{t=0} = \lambda \ln \left(\sum_{j \in \mathcal{I}} \pi_j e^{(a_j - \alpha_j)/\lambda} \right).$$

Substituting into [Equation \(26\)](#), we finally obtain the characterization given in [Equation \(8\)](#), and used in [Armenter et al. \[2024\]](#).

A.6 Comparative Advantage and Repeated Trade

The goal of this section is to establish [Theorem 6](#) by formally defining the multi-period trading game and deriving a Perfect Bayesian Equilibrium that achieves the first-best surplus.

Game structure. We formally describe the trading protocol as an infinite-period dynamic game. Each period starts with of a simultaneous move game where all but one player can acquire any signal of their choice. (Drawing an uninformative signal models inaction.) Then, one player can either learn or act: If there is no outstanding offer, the owner of the opportunity κ can extend a TIOLI offer to some other agent ℓ at terms $\mathbf{t} \in \mathbb{R}^I$, or implement an action $\mathbf{a} \in \mathcal{A}^\kappa$. If there is an outstanding offer at the beginning of the period, agent ℓ can accept or reject the offer. A rejection deactivates the offer, and the game proceeds to a period without outstanding offers. The game starts without any outstanding offer and with agent 1 as the owner. The game ends when an action is implemented, at which point the

state i becomes public knowledge. Learning costs are incurred immediately and privately by the agent who draws the signal, but all agents are able to observe the resulting information, leading to a common belief update at the end of each period. Trades represent transfer commitments that are due at the end of the game: Each agent pays t_i for any offers $\mathbf{t}$ that he has accepted and receives t'_i from anyone who has accepted one of his offers $\mathbf{t}'$. In addition, the last owner reaps the payoffs from the implemented action, a_i .

Preliminaries. In order to guarantee that the learning-proof menu $\bar{\mathcal{A}}^P$ is well-defined, we need to verify that c^P satisfies properties (C1) to (C4). Continuity, Blackwell monotonicity and prior-concavity are maintained under minimization. Theorem 1 and Lemma 9 in [Bloedel and Zhong \[2020\]](#) ensure that these properties are also preserved under sequential optimization and that sequential learning cannot further improve upon c^P . And since in particular $c^P \leq c^1$, the principal also breaks all payoff-relevant ties through learning. Therefore, $\bar{\mathcal{A}}^P$ is well-defined and [Theorem 2](#) holds.³³

Notation. For each belief $\pi \in \Delta\mathcal{I}$, we let $\mathcal{S}^\pi$ denote the first signal acquired by the planner with prior π in an optimal learning sequence that minimizes the expected number of draws. We denote the associated distribution over posterior beliefs as $\nu^\pi \in \Delta\Delta\mathcal{I}$. By optimality,

$$W(\mathcal{A}^P, \pi, c^P) = \int_{\rho \in \Delta\mathcal{I}} W(\mathcal{A}^P, \rho, c^P) d\nu^\pi - \min_{k \in \{1, \dots, K\}} c^k(\mathcal{S}^\pi, \pi), \quad (27)$$

and since the number of draws is minimized across all optimal learning sequences, a degenerate distribution indicates that the principal implements the best action unconditionally at this belief,

$$\nu^\pi = \delta_\pi \implies W(\mathcal{A}^P, \pi, c^P) = \max_{\mathbf{a} \in \mathcal{A}^P} \pi \cdot \mathbf{a}. \quad (28)$$

We also let $\mathcal{K}^\pi$ denote the set of agents who can emulate the next step of the planner's strategy. Formally, if $\nu^\pi = \delta_\pi$, then agents in $\mathcal{K}^\pi = \arg \max_k \max_{\mathbf{a} \in \mathcal{A}^k} \pi \cdot \mathbf{a}$ are those with access to the optimal action; otherwise agents in $\mathcal{K}^\pi = \arg \min_k c^k(\mathcal{S}^\pi, \pi)$ are those who can implement $\mathcal{S}^\pi$ at minimal cost. The minimality of draws also implies that $\mathcal{K}^\pi \cap \mathcal{K}^\rho = \emptyset$ for any $\rho \in \text{support}(\nu^\pi)$, for otherwise the principal could collapse signals $\mathcal{S}^\pi$ and $\mathcal{S}^\rho$ into a single draw.

³³Although not needed, property (C5) also holds: The sequential implementation incurs no extra costs for the particular agent with the lowest cost for the one-shot implementation $\mathcal{S} \times \{\mathcal{S}^s\}$, and so is even cheaper for the principal who can optimize over learning technologies.

Candidate equilibrium. We define a strategy profile by indexing each history by the agents' common belief π and the current owner κ , as well as potentially an outstanding offer from κ to some other agent ℓ with terms $\mathbf{t}$. Our candidate equilibrium is as follows:

- (i) At a history (π, κ) without an outstanding offer, but with $\kappa \in \mathcal{K}^\pi$ and $\nu^\pi = \delta_\pi$, owner κ implements the optimal action $\arg \max_{\mathbf{a} \in \mathcal{A}^\kappa} \pi \cdot \mathbf{a}$. No other agent does any learning.
- (ii) At a history (π, κ) without an outstanding offer, but with $\kappa \in \mathcal{K}^\pi$ and $\nu^\pi \neq \delta_\pi$, owner κ draws signal $\mathcal{S}^\pi$. No other agent does any learning.
- (iii) At a history with an outstanding offer from owner κ to agent ℓ at terms $\mathbf{t}$ under belief π , agent ℓ analyzes RI problem $(\bar{\mathcal{A}}^P \cup \{\mathbf{t}\}, \pi, c^\ell)$. He accepts the offer whenever there exists an optimal strategy that does not rely on $\mathbf{t}$. In particular, this includes all instances where the terms are at or below $\bar{\mathcal{A}}^P$. Agent ℓ rejects the offer if unconditional implementation of $\mathbf{t}$ is optimal. (If indifferent, he defaults to acceptance.) If neither is optimal, he draws an optimal signal that maximizes $\pi \cdot \mathbf{q}(\mathbf{t})$. No other agent does any learning.
- (iv) At a history (π, κ) without an outstanding offer and with $\kappa \notin \mathcal{K}^\pi$, owner κ extends a take-it-or-leave-it offer with terms $\arg \max_{\mathbf{t} \in \bar{\mathcal{A}}^P} \pi \cdot \mathbf{t}$ to one of the agents in $\mathcal{K}^\pi$. No other agent does any learning.

Subgame payoffs. We claim that the expected continuation payoff after any history (π, κ) without an outstanding offer is equal to $W(\mathcal{A}^P, \pi, c^P)$ for owner κ and zero for everyone else. (If there is an outstanding offer, the owner's payoff is bounded above by $W(\mathcal{A}^P, \pi, c^P)$ but is possibly smaller.) As a consequence, everyone values opportunity ownership as granting access to the learning-proof menu $\bar{\mathcal{A}}^P$. We establish this claim by backwards induction:

- (i) At a terminal history (π, κ) without an outstanding offer, but with $\kappa \in \mathcal{K}^\pi$ and $\nu^\pi = \delta_\pi$, Equation (28) guarantees that owner κ earns an expected payoff of $W(\mathcal{A}^P, \pi, c^P)$ by implementing his best action. Everyone else is inactive and earns a payoff of zero.
- (ii) At a history (π, κ) without an outstanding offer, but with $\kappa \in \mathcal{K}^\pi$ and $\nu^\pi \neq \delta_\pi$, owner κ incurs learning costs of $c^\kappa(\mathcal{S}^\pi, \pi)$. The game then proceeds with histories (ρ, κ) , with ρ distributed according to ν^π . By backwards induction and Equation (27), the continuation payoff for owner κ is thus equal to

$$\int_{\rho \in \Delta \mathcal{I}} W(\mathcal{A}^P, \rho, c^P) d\nu^\pi - c^\kappa(\mathcal{S}^\pi, \pi) = W(\mathcal{A}^P, \pi, c^P),$$

and zero for everyone else. No other agent does any learning.

- (iii) At a history with an outstanding offer from owner κ to agent ℓ at terms $\mathbf{t}$ under belief $\boldsymbol{\pi}$, trade happens whenever $W(\bar{\mathcal{A}}^P \cup \{\mathbf{t}\}, \boldsymbol{\pi}, c^\ell) \leq W(\bar{\mathcal{A}}^P, \boldsymbol{\pi}, c^\ell)$. By backwards induction, this results in a continuation payoff of $\boldsymbol{\pi} \cdot \mathbf{t}$ for κ , and $W(\mathcal{A}^P, \boldsymbol{\pi}, c^P) - \boldsymbol{\pi} \cdot \mathbf{t}$ for ℓ . Since this implies in particular that $\boldsymbol{\pi} \cdot \mathbf{t} \leq W(\bar{\mathcal{A}}^P, \boldsymbol{\pi}, c^\ell) \leq W(\bar{\mathcal{A}}^P, \boldsymbol{\pi}, c^P)$, the inductive hypothesis holds.

If the terms are so high that $\mathbf{t}$ is the ignorance equivalent of RI problem $(\bar{\mathcal{A}}^P \cup \{\mathbf{t}\}, \boldsymbol{\pi}, c^\ell)$, no learning occurs and the game simply proceeds to history $(\boldsymbol{\pi}, \kappa)$ without an outstanding offer. Again, the inductive hypothesis holds.

At all other terms, agent ℓ draws a costly signal $\mathcal{S}^\ell$ that updates beliefs to the Bayes-plausible distribution ν^ℓ . Regardless of whether trade occurs or not, the sum of the continuation payoffs between agents ℓ and κ is bounded above by

$$\int_{\boldsymbol{\rho} \in \Delta \mathcal{I}} W(\mathcal{A}^P, \boldsymbol{\rho}, c^P) d\nu^\ell - c^\ell(\mathcal{S}^\ell, \boldsymbol{\pi}) \quad (29)$$

by backwards induction. Agent ℓ can guarantee himself a continuation payoff of zero by rejecting, so his optimal payoff is nonnegative. In turn, this bounds the continuation payoff of owner κ above by (29), which is at most $W(\mathcal{A}^P, \boldsymbol{\pi}, c^P)$, thus confirming the inductive hypothesis.

- (iv) At a history $(\boldsymbol{\pi}, \kappa)$ without an outstanding offer and with $\kappa \notin \mathcal{K}^\pi$, the suggested terms are on $\bar{\mathcal{A}}^P$ and represent an expected transfer of $\boldsymbol{\pi} \cdot \mathbf{t} = W(\mathcal{A}^P, \boldsymbol{\pi}, c^P)$ by Theorem 2. By backwards induction, this results in payoffs $\boldsymbol{\pi} \cdot \mathbf{t} = W(\mathcal{A}^P, \boldsymbol{\pi}, c^P)$ for agent κ and $W(\mathcal{A}^P, \boldsymbol{\pi}, c^P) - \boldsymbol{\pi} \cdot \mathbf{t} = 0$ for agent ℓ . All other agents earn zero.

Optimality. From the subgame payoffs, it also follows that no agent has a profitable deviation. The owner earns $\int_{\boldsymbol{\rho} \in \Delta \mathcal{I}} W(\mathcal{A}^P, \boldsymbol{\rho}, c^P) d\nu - c^\kappa(\mathcal{S}, \boldsymbol{\pi})$ by deviating to a learning strategy $\mathcal{S}$ with associated belief distribution ν , $\boldsymbol{\pi} \cdot \mathbf{a}$ by deviating to a game-ending move of implementing action $\mathbf{a}$, and at most $W(\mathcal{A}^P, \boldsymbol{\pi}, c^P)$ by extending any other offer. All are bounded above by his subgame payoff of $W(\mathcal{A}^P, \boldsymbol{\pi}, c^P)$. The acceptance strategy in (iii) is already chosen to be optimal across agent ℓ 's feasible subgame payoffs. And any other agent who would engage in learning will face only costs but no benefits.

Moreover, agents emulate the first-best strategy on the equilibrium path, and thus reap social surplus Δ . This establishes optimality and proves Theorem 6.

References

- Roc Armenter, Michèle Müller-Itten, and Zachary R Stangebye. Geometric methods for finite rational inattention. *Quantitative Economics*, 15(1):115–144, 2024.
- Dirk Bergemann and Juuso Välimäki. Information acquisition and efficient mechanism design. *Econometrica*, 70(3):1007–1033, 2002. doi: <https://doi.org/10.1111/1468-0262.00317>. URL <https://onlinelibrary.wiley.com/doi/abs/10.1111/1468-0262.00317>.
- David Blackwell. Equivalent comparisons of experiments. *The annals of mathematical statistics*, pages 265–272, 1953.
- Alexander W. Bloedel and Weijie Zhong. The cost of optimally acquired information. *Unpublished working paper*, 2020. Available at <https://drive.google.com/file/d/1HC5WxR6p5sB-6FvzXwf73rrFzx39UXYw/view?usp=sharing>.
- Jonathan M. Borwein and Qiji J. Zhu. *Techniques of Variational Analysis*. CMS Books in Mathematics. Springer New York, New York, NY, 2005. ISBN 0387242988.
- Andrew Caplin and Mark Dean. Behavioral Implications of Rational Inattention with Shannon Entropy. *NBER Working Paper 19318*, 2013.
- Andrew Caplin, Mark Dean, and John Leahy. Rational inattention, optimal consideration sets, and stochastic choice. *Review of Economic Studies*, 86(3):1061–1094, 07 2018. ISSN 0034-6527. doi: 10.1093/restud/rdy037. URL <https://doi.org/10.1093/restud/rdy037>.
- Andrew Caplin, Mark Dean, and John Leahy. Rationally inattentive behavior: Characterizing and generalizing shannon entropy. *Journal of Political Economy*, 130(6):1676–1715, 2022. doi: 10.1086/719276. URL <https://doi.org/10.1086/719276>.
- D.L. Cohn. *Measure Theory*. Birkhäuser Boston, 1997. ISBN 9780817630034. URL <https://books.google.com/books?id=vRxV2FwJvoAC>.
- Tommaso Denti, Massimo Marinacci, and Aldo Rustichini. Experimental cost of information. *American Economic Review*, 112(9):3106–23, September 2022. doi: 10.1257/aer.20210879. URL <https://www.aeaweb.org/articles?id=10.1257/aer.20210879>.
- Matthew Gentzkow and Emir Kamenica. Costly persuasion. *American Economic Review*, 104(5):457–462, 2014.

- Benjamin M. Hébert and Michael Woodford. Rational inattention when decisions take time. Technical report, National Bureau of Economic Research, 2019.
- Benjamin M. Hébert and Michael Woodford. Neighborhood-based information costs. Technical report, National Bureau of Economic Research, 2020.
- Bengt Holmström and Paul Milgrom. Regulating trade among agents. *Journal of Institutional and Theoretical Economics (JITE) / Zeitschrift für die gesamte Staatswissenschaft*, 146(1):85–105, 1990. ISSN 09324569. URL <http://www.jstor.org/stable/40751306>.
- Hideshi Itoh. Incentives to help in multi-agent situations. *Econometrica*, 59(3):611–636, 1991. ISSN 00129682, 14680262. URL <http://www.jstor.org/stable/2938221>.
- Emir Kamenica and Matthew Gentzkow. Bayesian persuasion. *American Economic Review*, 101(6):2590–2615, 2011.
- Bartosz Maćkowiak, Filip Matějka, and Mirko Wiederholt. Rational inattention: A review. *Journal of Economic Literature*, 61(1):226–273, 2023.
- Filip Matějka and Alisdair McKay. Rational inattention to discrete choices: A new foundation for the multinomial logit model. *American Economic Review*, 105(1):272–98, January 2015. doi: 10.1257/aer.20130047. URL <http://www.aeaweb.org/articles?id=10.1257/aer.20130047>.
- Jeffrey Mensch. Cardinal representations of information. *Unpublished working paper*, 2018. Available at <https://ssrn.com/abstract=3148954>.
- Stephen Morris and Philipp Strack. The Wald problem and the relation of sequential sampling and ex-ante information costs. *Unpublished working paper*, 2019. Available at <https://ssrn.com/abstract=2991567>.
- Luciano Pomatto, Philipp Strack, and Omer Tamuz. The cost of information. *arXiv preprint arXiv:1812.04211*, 2018.
- Christopher A. Sims. Implications of rational inattention. *Journal of Monetary Economics*, 50(3):665–690, 2003.
- Elias Tsakas. Robust scoring rules. *Theoretical Economics*, 15(3):955–987, 2020. URL <https://econtheory.org/ojs/index.php/te/article/view/20200955/0>.

- Gerhard Winkler. Extreme points of moment sets. *Mathematics of Operations Research*, 13(4):581–587, 1988. ISSN 0364765X, 15265471. URL <http://www.jstor.org/stable/3689944>.
- Nathan Yoder. Designing incentives for heterogeneous researchers. *Journal of Political Economy*, 130(8):2018–2054, 2022. doi: 10.1086/720072. URL <https://doi.org/10.1086/720072>.
- Weijie Zhong. Optimal dynamic information acquisition. *Unpublished working paper*, 2019. Available at https://wjzhong.com/workingpapers/info_acquisition/Dynamic_info_acquisition_main.pdf.

B Online Appendix

B.1 Additional Explanations

B.1.1 Duality

Proof of Lemma A.9: Let V denote the Banach space of all finite signed measures over $\Delta\mathcal{I}$, and $V_+ = \{v \in V \mid v \geq 0\}$ the subset of nonnegative measures. Along with the total variation norm $\|\cdot\|_V$, this forms a Banach space [Cohn, 1997, p.219]. We endow $\mathbb{R}^I$ with the ℓ_1 norm $\|\cdot\|_1$ for convenience, and let $M : V \rightarrow \mathbb{R}^I$ denote the linear map that assigns to each measure ν the Lebesgue integral $M\nu = \int_{\Delta\mathcal{I}} \boldsymbol{\rho} d\nu$. The map M is bounded since $\int_{\Delta\mathcal{I}} \rho_i d\nu \in [-\nu^-(\Delta\mathcal{I}), \nu^+(\Delta\mathcal{I})]$ for each i and hence $\|M\nu\|_1 \leq |\mathcal{I}|(\nu^+(\Delta\mathcal{I}) + \nu^-(\Delta\mathcal{I})) = |\mathcal{I}|\|\nu\|_V$, where ν^+ and $-\nu^-$ denote the upper and lower variation respectively. For any set B , we further define the indicator function $\chi_B(b)$ that is 0 whenever $b \in B$ and ∞ otherwise. This function is lower semi-continuous and convex whenever B is closed and convex.

We now define two functions $f : V \rightarrow \mathbb{R} \cup \{+\infty\}$ and $g : \mathbb{R}^I \times \mathbb{R} \rightarrow \mathbb{R} \cup \{+\infty\}$ as

$$f(\nu) = \int_{\Delta\mathcal{I}} \Phi(\boldsymbol{\rho}) d\nu + \chi_{V_+}(\nu) \quad \text{and} \quad g(\mathbf{x}) = \chi_{\{\boldsymbol{\pi}\}}(\mathbf{x}).$$

Since convexity and lower semi-continuity are preserved under addition, f and g are both convex and lower-semicontinuous. Moreover, in the notation of Borwein and Zhu [2005, Theorem 4.4.3],

$$\mathbf{0} \in \text{core}(\text{dom}g - M\text{dom}f) = \text{core}(\{\boldsymbol{\pi}\} - [0, \infty)^{I+1}),$$

and therefore strong duality holds:

$$\inf_{\nu \in V} f(\nu) + g(M\nu) = \sup_{\langle \mathbf{x}^*, \cdot \rangle \in (\mathbb{R}^I)^*} -f^*(M^*\langle \mathbf{x}^*, \cdot \rangle) - g^*(-\langle \mathbf{x}^*, \cdot \rangle), \quad (30)$$

where $(\mathbb{R}^I)^*$ denotes the space of all linear forms over $\mathbb{R}^I$, M^* is the adjoint operator and the convex conjugates assign to each linear form a supremum

$$\begin{aligned} f^*(\langle \nu^*, \cdot \rangle) &= \sup_{v \in V} \langle \nu^*, v \rangle - f(v) = \sup_{v \in V_+} \langle \nu^*, v \rangle - \int_{\Delta\mathcal{I}} \Phi(\boldsymbol{\rho}) dv \\ &= \begin{cases} 0 & \text{if } \langle \nu^*, v \rangle \leq \int_{\Delta\mathcal{I}} \Phi(\boldsymbol{\rho}) dv \quad \forall v \in V_+ \\ \infty & \text{otherwise,} \end{cases} \\ g^*(\langle \mathbf{x}^*, \cdot \rangle) &= \sup_{\mathbf{x} \in \mathbb{R}^I} \langle \mathbf{x}^*, \mathbf{x} \rangle - g(\mathbf{x}) = \langle \mathbf{x}^*, \boldsymbol{\pi} \rangle. \end{aligned}$$

Using these expressions, we can equivalently write the duality (30) as

$$\begin{cases} \inf_{\nu \in V_+} & \int_{\Delta\mathcal{I}} \Phi(\boldsymbol{\rho}) d\nu \\ \text{s.t.} & \int_{\Delta\mathcal{I}} \boldsymbol{\rho} d\nu = \boldsymbol{\pi} \end{cases} = \begin{cases} \sup_{\langle \mathbf{x}^*, \cdot \rangle \in (\mathbb{R}^{\mathcal{I}})^*} & \langle \mathbf{x}^*, \boldsymbol{\pi} \rangle \\ \text{s.t.} & \langle M^* \langle \mathbf{x}^*, \cdot \rangle, v \rangle \leq \int_{\Delta\mathcal{I}} \Phi(\boldsymbol{\rho}) d\nu \quad \forall v \in V_+. \end{cases}$$

The constraint on the left side implies in particular that

$$\nu(\Delta\mathcal{I}) = \int_{\Delta\mathcal{I}} 1 d\nu = \mathbf{1} \cdot \int_{\Delta\mathcal{I}} \boldsymbol{\rho} d\nu = \mathbf{1} \cdot \boldsymbol{\pi} = 1,$$

and hence it is without loss of generality to restrict attention to probability measures $\nu \in \Delta\Delta\mathcal{I}$, making this infimum equal to p^* from the claim.

As for the right side, the definition of the adjoint states that $\langle M^* \langle \mathbf{x}^*, \cdot \rangle, v \rangle = \langle \mathbf{x}^*, Mv \rangle$, and since any linear form over $\mathbb{R}^{\mathcal{I}}$ can be expressed as a scalar product with a vector in $\mathbb{R}^{\mathcal{I}}$, we can also write the duality (30) as

$$p^* = \begin{cases} \sup_{\mathbf{a} \in \mathbb{R}^{\mathcal{I}}} & \mathbf{a} \cdot \boldsymbol{\pi} \\ \text{s.t.} & \mathbf{a} \cdot Mv = \mathbf{a} \cdot \int_{\Delta\mathcal{I}} \boldsymbol{\rho} d\nu \leq \int_{\Delta\mathcal{I}} \Phi(\boldsymbol{\rho}) d\nu \quad \forall v \in V_+. \end{cases}$$

We rewrite the constraint as

$$0 \leq \int_{\Delta\mathcal{I}} (\Phi(\boldsymbol{\rho}) - \mathbf{a} \cdot \boldsymbol{\rho}) d\nu \quad \forall v \in V_+,$$

and make a few simplifying observations: First, the inequality is unaffected by positive scaling of v , so it is without loss of generality to restrict attention to probability measures $v \in \Delta\Delta\mathcal{I}$. Moreover, the right side is most binding when v is a degenerate distribution, so it is also without loss of generality to simply write (30) as

$$p^* = \begin{cases} \sup_{\mathbf{a} \in \mathbb{R}^{\mathcal{I}}} & \mathbf{a} \cdot \boldsymbol{\pi} \\ \text{s.t.} & 0 \leq \Phi(\boldsymbol{\rho}) - \mathbf{a} \cdot \boldsymbol{\rho} \quad \forall \boldsymbol{\rho} \in \Delta\mathcal{I}, \end{cases}$$

which is now equal to d^* from the claim.

Moreover, the Fenchel duality theorem also asserts that whenever $p^* = d^*$ is finite, then the supremum is attained. $\square$

B.2 Examples

B.2.1 Non-Existence of an Ignorance Equivalence

In this section, we show by example that an ignorance equivalent may fail to exist when the cost function does not satisfy (C4).

In particular, we consider prior-independent costs

$$c(\langle S, \mathbf{q} \rangle, \boldsymbol{\pi}) = \left[\frac{1}{I} \sum_{i \in \mathcal{I}} H(q_i) - H \left(\frac{1}{I} \sum_{i \in \mathcal{I}} q_i \right) \right]^2,$$

where $H(p) = -\sum_{s \in S} p(s) \ln(p(s))$ denotes the Shannon entropy of probability mass function p . This cost function measures the square of the mutual information between the conditional probabilities q_i and the mean probability $\frac{1}{I} \sum_{i \in \mathcal{I}} q_i$. It is well known that mutual information is continuous and strictly increasing in the Blackwell order. By taking the unconditional mean rather than (as is standard) weighting each term by π_i , the cost function becomes prior-independent and therefore trivially prior-concave. In addition, let $\mathbf{a}^1 = (1, 0)$, $\mathbf{a}^2 = (0, 1)$, $\mathcal{A} = \{\mathbf{a}^1, \mathbf{a}^2\}$ and $\boldsymbol{\pi} = (\frac{1}{2}, \frac{1}{2})$. The exact choice of the menu and the prior matter only in so far as we want to guarantee that some learning is optimal in RI problem $(\mathcal{A}, \boldsymbol{\pi}, c)$. We claim that this RI problem does not admit an ignorance equivalent.

Indeed, let $\mathcal{S} = \langle \mathcal{A}, \mathbf{q} \rangle$ denote an optimal learning strategy, with expected statewise consumption of $\mathbf{a}^{\mathcal{S}}$. For convenience, we determine $\mathcal{S}$ using Mathematica (see code at the end of this section) and verify that it involves a strictly positive amount of learning, $c(\mathcal{S}, \boldsymbol{\pi}) > 0$. By definition, the expected utility of any candidate ignorance equivalent $\boldsymbol{\alpha}$ has to be equal to the net utility as $\mathcal{S}$,

$$\boldsymbol{\pi} \cdot \boldsymbol{\alpha} = W(\{\boldsymbol{\alpha}\}, \boldsymbol{\pi}, c) = W(\mathcal{A}, \boldsymbol{\pi}, c) = \boldsymbol{\pi} \cdot \mathbf{a}^{\mathcal{S}} - c(\mathcal{S}, \boldsymbol{\pi}).$$

At the same time, an agent facing RI problem $(\mathcal{A} \cup \{\boldsymbol{\alpha}\}, \boldsymbol{\pi}, c)$ could follow signal $\tilde{\mathcal{S}} = \langle \mathcal{A} \cup \{\boldsymbol{\alpha}\}, \tilde{\mathbf{q}} \rangle$ with $q_i(\boldsymbol{\alpha}) \equiv t \in (0, 1)$ and $\mathbf{q}(\mathbf{a}^k) \equiv (1 - t)\mathbf{q}(\mathbf{a}^k)$. This strategy is Blackwell equivalent to a two-stage strategy where the agent first tosses an uninformative coin. The coin returns heads with probability t , prompting the agent to implement $\boldsymbol{\alpha}$, and tails with probability $1 - t$, prompting her to follow $\mathcal{S}$. Beliefs are unaffected by the coin toss, and so the agent's expected consumption utility is equal to $t\boldsymbol{\pi} \cdot \boldsymbol{\alpha} + (1 - t)\boldsymbol{\pi} \cdot \mathbf{a}^{\mathcal{S}}$. However, the cost of strategy $\tilde{\mathcal{S}}$ is strictly below that of its two-stage implementation, with $c(\tilde{\mathcal{S}}, \boldsymbol{\pi}) = (1 - t)^2 c(\mathcal{S}, \boldsymbol{\pi})$. Consequently, the net utility of $\tilde{\mathcal{S}}$ exceeds $W(\mathcal{A}, \boldsymbol{\pi}, c)$, implying that adding

α to the menu is strictly welfare enhancing,

$$\begin{aligned} W(\mathcal{A} \cup \{\alpha\}, \pi, c) &\geq t\pi \cdot \alpha + (1-t)\pi \cdot \alpha^S - (1-t)^2 c(\mathcal{S}, \pi) \\ &= tW(\{\alpha\}, \pi, c) + (1-t)W(\mathcal{A}, \pi, c) + t(1-t)c(\mathcal{S}, \pi) \\ &> W(\mathcal{A}, \pi, c). \end{aligned}$$

The incompatibility between the two properties of Definition 1 is linked to the fact that c violates sequential-learning proofness (C4). This means that even though $\mathcal{S}$ dominates α , the agent is not able to realize the full welfare potential of $(\mathcal{A} \cup \{\alpha\}, \pi, c)$ without relying on the ‘moderate’ payoffs α .³⁴

Code: Non-existence under violation of (C4).

```
1 Clear["Global`*"];
```

We consider a setup with two states and two actions:

```
2 a1={1,0};
```

```
3 a2={0,1};
```

The prior is equal to

```
4 pi1=0.5;
```

```
5 pi={pi1,1-pi1};
```

A learning strategy is written as a matrix where entry (i,a) denotes the likelihood of implementing action a conditional on state i. Costs are given by the prior-independent function

```
6 p={1/2,1/2};
```

```
7 c[q_?(MatrixQ[#,NumericQ]&),pi_]:= Sum[ If[ q[[i,a]]>0, p[[i]]q[[i,a]]
  Log[q[[i,a]] / p.q[[;;,a]] , 0] , {i,1,Dimensions[q][[1]]} ,
  {a,1,Dimensions[q][[2]]} ]^2
```

This function denotes the squared reduction in Shannon entropy, but relative to the fixed state distribution p rather than the prior pi. This has the advantage of ensuring that the function trivially satisfies (C1) and (C2).

We now compute the optimal signal:

```
8 q={{q1a1,1-q1a1},{1-q2a2,q2a2}};
```

```
9 W[q_?(MatrixQ[#,NumericQ]&),pi_] := pi.Total[{a1,a2}*Transpose[q]]-c[q,pi]
```

```
10 NMaximize[{W[q,pi],0<=q1a1<=1,0<=q2a2<=1},{q1a1,q2a2}]
```

³⁴To avoid confusion: The issue is not that $\tilde{\mathcal{S}}$ is strictly cheaper than the two-step implementation described above. (C4) allows for that. The issue is that $\mathcal{S}$ is strictly more expensive than a two-stage implementation where the agent first follows $\tilde{\mathcal{S}}$ and replaces action α with another instance of $\mathcal{S}$. The former costs $c(\mathcal{S}, \pi)$ while the latter merely costs $(1-t)^2 c(\mathcal{S}, \pi) + t c(\mathcal{S}, \pi)$.

```

» {0.777046, {q1a1 → 0.855844, q2a2 → 0.855844}}
11 qopt={{q1a1,1-q1a1},{1-q2a2,q2a2}}/.%[[2]];
    This involves a strictly positive amount of learning,
12 c[qopt,c]>0
» True

```

B.2.2 Failure of Self-Selection

To illustrate the role that (C2) plays in our results, we here show that self-selection (Corollary 3) fails for a concrete example with a cost function that satisfies all properties save for prior-concavity. Self-selection is crucial for our subsequent results, as the construction of the learning-proof menu rests upon the absence of welfare gains.

Specifically, we consider a uniformly posterior-separable cost function with the convex and differentiable potential $\phi(\boldsymbol{\pi}) = -\sum_{i \in \mathcal{I}} \pi_i(1 - \pi_i)$. We study a three-state problem with two actions $\mathcal{A} = \{(1, 0, 0), (0, 1, 0)\}$ and a uniform prior $\boldsymbol{\pi} = (\frac{1}{3}, \frac{1}{3}, \frac{1}{3})$. To streamline exposition, we verify all claims algebraically with Mathematica and provide the code at the end of this section.

Using the tools from Appendix A.4, we locate the unique ignorance equivalent equivalent for this posterior separable function at $\boldsymbol{\alpha} = (\frac{31}{24}, \frac{31}{24}, \frac{19}{24})$. We then show that adding this action $\boldsymbol{\alpha}$ to the menu is strictly welfare-enhancing for an agent who puts a higher likelihood on state 3, with prior $\boldsymbol{\rho} = (\frac{1}{10}, \frac{1}{10}, \frac{8}{10})$. In particular,

$$W(\mathcal{A}, \boldsymbol{\rho}, c) = 0.18 < \boldsymbol{\rho} \cdot \boldsymbol{\alpha} = 0.225 \leq W(\mathcal{A} \cup \{\boldsymbol{\alpha}\}, \boldsymbol{\rho}, c),$$

thus contradicting Corollary 3. Finally, we show that that this cost function violates (only) prior-concavity (C2). Intuitively, prior-concavity allows us to connect global welfare gains to local welfare gains in the ‘necessity of dominance’ argument (Section 3). Without it, there may be strict welfare gains at prior $\boldsymbol{\rho}$ even though there are none for local priors around $\boldsymbol{\pi}$.

Code: Violation of Corollary 3 without (C2).

```

1 Clear["Global`*"];
    There are three states and two actions with payoffs
2 a1={1,0,0};
3 a2={0,1,0};

```

The UPS potential is given by

4 $\phi[\{r1_ , r2_ , r3_ \}] := \text{Sum}[-r (1-r), \{r, \{r1, r2, r3\}\}]$

For notational convenience, we explicitly define its gradient as

5 $D\phi[\{p1_ , p2_ , p3_ \}] = D[\phi[\{p1, p2, p3\}], \{\{p1, p2, p3\}, 1\}];$

and verify that the potential is convex by checking for positive Eigenvalues of the Hessian,

6 $\text{Eigenvalues}[D[\phi[\{p1, p2, p3\}], \{\{p1, p2, p3\}, 2\}]]$

» $\{2, 2, 2\}$

We claim that the ignorance equivalent at

7 $\text{pi} = \{1/3, 1/3, 1/3\};$

is equal to

8 $\alpha = \{5/8, 5/8, 1/8\};$

and hence using formula (23),

9 $\theta = \alpha - D\phi[\text{pi}] - \{1, 1, 1\}(\phi[\text{pi}] - \text{pi} \cdot D\phi[\text{pi}])$

» $\{31/24, 31/24, 19/24\}$

This satisfies the constraints (22) since for any beliefs $\mathbf{r}$ and any available action $\mathbf{a}$, the expression $\mathbf{r} \cdot \mathbf{a} - \phi[\mathbf{r}] - \mathbf{r} \cdot \theta$ is maximized at 0.

10 $\mathbf{r} = \{r1, r2, r3\};$

11 $\text{Maximize}[\{\mathbf{r} \cdot \mathbf{a1} - \phi[\mathbf{r}] - \mathbf{r} \cdot \theta, \{0 \leq r1, 0 \leq r2, 0 \leq r3, r1 + r2 + r3 == 1\}\}, \{\mathbf{r1}, \mathbf{r2}, \mathbf{r3}\}]$

12 $\text{Maximize}[\{\mathbf{r} \cdot \mathbf{a2} - \phi[\mathbf{r}] - \mathbf{r} \cdot \theta, \{0 \leq r1, 0 \leq r2, 0 \leq r3, r1 + r2 + r3 == 1\}\}, \{\mathbf{r1}, \mathbf{r2}, \mathbf{r3}\}]$

» $\{0, \{r1 \rightarrow 7/12, r2 \rightarrow 1/12, r3 \rightarrow 1/3\}\}$

» $\{0, \{r1 \rightarrow 1/12, r2 \rightarrow 7/12, r3 \rightarrow 1/3\}\}$

For this reason, we know that we have found the correct ignorance equivalent.

We now show that adding this α to the menu is strictly welfare enhancing for an agent with prior

13 $\text{rho} = \{1/10, 1/10, 8/10\};$

Sticking to the posterior-based approach, we know we can find $W(\{\mathbf{a1}, \mathbf{a2}\}, \text{rho}, c)$ by maximizing over Bayes-plausible posteriors. In particular, we maximize $q(\mathbf{r} \cdot \mathbf{a1} - \phi[\mathbf{r}]) + (1-q)(\mathbf{s} \cdot \mathbf{a2} - \phi[\mathbf{s}])$, which is the convex combination over the agent's net utility at $\mathbf{r}$ and $\mathbf{s}$, weighed according to q . The constraints ensure that both posteriors are valid, and that the average posterior equals the prior rho . Doing so determines the maximal net utility of the agent, and we obtain actual welfare by adding $\phi[\text{rho}]$ back in.

14 $\mathbf{r} = \{r1, r2, r3\};$

15 $\mathbf{s} = \{s1, s2, s3\};$

16 $\text{Maximize}[\{q(\mathbf{r} \cdot \mathbf{a1} - \phi[\mathbf{r}]) + (1-q)(\mathbf{s} \cdot \mathbf{a2} - \phi[\mathbf{s}]), \{0 \leq r1, 0 \leq r2, 0 \leq r3, r1 + r2 + r3 == 1, 0 \leq s1, 0 \leq s2, 0 \leq s3, s1 + s2 + s3 == 1, 0 \leq q \leq 1, q \mathbf{r} + (1-q)\mathbf{s} == \text{rho}\}\}, \{\mathbf{r1}, \mathbf{r2}, \mathbf{r3}, \mathbf{s1}, \mathbf{s2}, \mathbf{s3}, q\}]$

```

    rho}}, {r1, r2, r3, q, s1, s2, s3}]
17 W=%[[1]]+phi[rho] (* In mathematica, % refers to the last computed output. *)
>> {13/25, {r1->1/5, r2->0, r3->4/5, q->1/2, s1->0, s2->1/5, s3->4/5}}
>> 9/50

```

In particular, the agent's welfare without access to α is equal to

```

18 W //N
>> 0.18

```

However, if α is available to the agent, she can in particular choose to implement it unconditionally, which yields a lower bound for $W(\{a1, a2, \alpha\}, \text{rho}, c)$,

```

19 rho.alpha//N
>> 0.225

```

This shows that adding α affects the welfare of an agent who differs only in prior belief, contrary to the result we obtain in [Corollary 3](#).

Lastly, we verify that prior-concavity is the only property not satisfied by the UPS cost function considered here. Without loss of generality, we denote a signal by a list of state-dependent probabilities, where $qs=\{qs1, qs2, qs3\}$ denotes the likelihood for realization $s \in \{x, y, z\}$ conditional on each of the states $i \in \{1, 2, 3\}$,

```

20 c[q_, pi_] := Sum[(pi.qs) phi[(pi*qs)/pi.qs], {qs, q}] - phi[pi]

```

Continuity (C1) holds by continuity of the expectation and that of ϕ . Blackwell monotonicity holds because ϕ is convex (as verified above). [Bloedel and Zhong \[2020\]](#) show that indifference to sequential information acquisition (C4) and (C5) hold for any UPS functions, including this one. The differentiability of the potential further implies that ties are broken through learning (C3), as detailed in [Lemma A.2](#).

However, prior-concavity does not hold. We illustrate this by focusing on the optimal signal under prior pi ,

```

21 q={{r1, r2, r3}/(2pi) /.Maximize[{r.a1-phi[r]-r.theta, {0<=r1, 0<=r2, 0<=r3, r1+r2+r3==1}},
    {r1, r2, r3}][[2]], {r1, r2, r3}/(2pi)}
    /.Maximize[{r.a2-phi[r]-r.theta, {0<=r1, 0<=r2, 0<=r3, r1+r2+r3==1}}, {r1, r2, r3}][[2]]
>> {{7/8, 1/8, 1/2}, {1/8, 7/8, 1/2}}

```

This is optimal because it achieves the same welfare at pi as unconditional implementation of α ,

```

22 (pi*q[[1, ;;]]) .a1 + (pi*q[[2, ;;]]) .a2 - c[q, pi] == pi.alpha
>> True

```

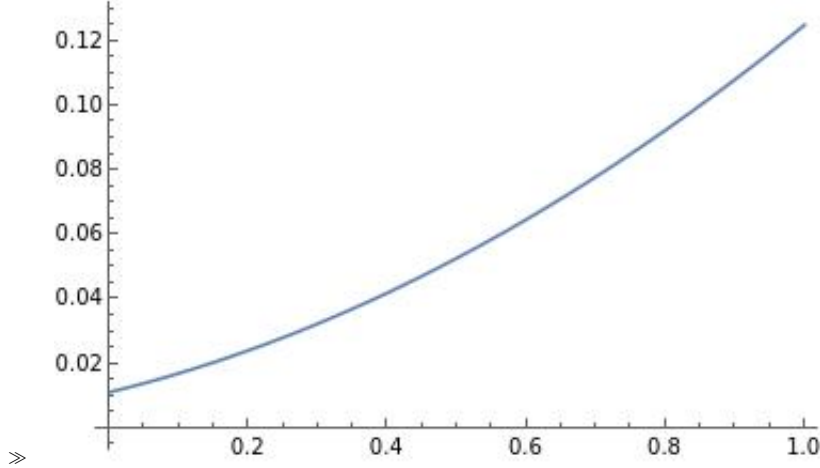
However, α is no longer dominated by this signal q because

```
23 (rho*q[[1,;;]]) .a1+(rho*q[[2,;;]]) .a2-c[q,rho] < rho.alpha
```

```
>> True
```

And indeed, we can verify that prior-concavity fails here,

```
24 Plot[c[q,t pi +(1-t)rho],{t,0,1}]
```



This breaks the step in the ‘dominance of necessity’ argument, where we go from an arbitrary prior where α is welfare-enhancing to a local perturbation with the same property.

B.2.3 Learning Cost inferences

Suppose that two RI agents differ only in their cost of learning, and these the costs can be globally ranked, with either $c_1 \leq c_2$ or $c_1 \geq c_2$. Environments with globally ranked information costs include in particular those where the cost follows some parametric family with an unknown scaling factor. In that case, it is possible to detect which agent has access to the superior learning technology by considering only their ignorance equivalents.

Lemma 1. *Suppose that agents 1 and 2 have the same prior π and globally ranked information costs. If α is such that $\{\alpha\} \sim_1 \mathcal{A} \sim_1 \mathcal{A} \cup \{\alpha\}$ and $\mathcal{A} \cup \{\alpha\} \succ_2 \mathcal{A}$, then $c^1 < c^2$.*

Proof. When $c^i \leq c^j$, every learning strategy offers weakly higher welfare to agent i , hence $W(\mathcal{A}, \rho, c^i) \geq W(\mathcal{A}, \rho, c^j)$ for all beliefs. By Equation (2), this is equivalent to stating that the learning-proof menu of agent i , $\bar{\mathcal{A}}^i$, lies weakly above that of agent j , $\bar{\mathcal{A}}^j$. Since α is agent 1’s ignorance equivalent, it is part of $\bar{\mathcal{A}}^1$ by Theorem 2. Since α is welfare-enhancing for agent 2, it achieves higher expected utility at some posterior than any strategy involving actions in $\mathcal{A}^2$. Consequently, it lies strictly above $\bar{\mathcal{A}}^2$. With globally ranked costs, the only possibility where $\bar{\mathcal{A}}^1$ lies strictly above $\bar{\mathcal{A}}^2$ at α is when $c^1 < c^2$. $\square$

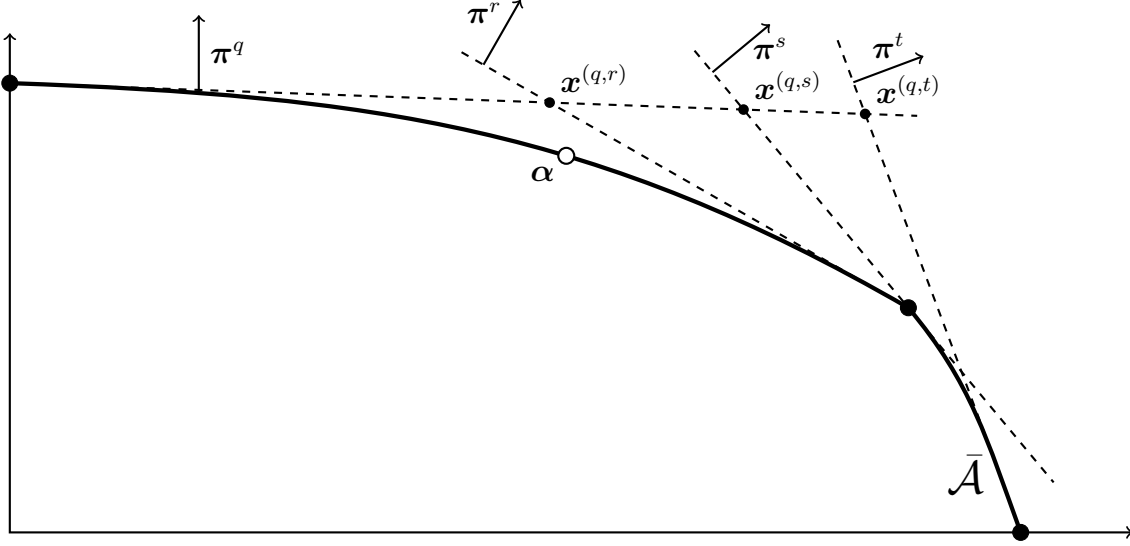


Figure 5: Willingness to pay for exogenous information.

B.2.4 Exogenous Information and Endogenous Learning

Figure 5 illustrates the interaction between exogenous and endogenous information for an RI agent of prior π with access to three actions and a learning technology that results in ignorance equivalent α and learning-proof menu $\bar{\mathcal{A}}$. Left to her own devices, the agent employs learning strategy $\mathcal{S}^{(q,r)}$ that yields posteriors π^q and π^r . If she instead obtains this signal from the principal, her welfare increases from $\pi \cdot \alpha$ to $\pi \cdot x^{(q,r)}$, and she would be willing to pay at most her endogenous information acquisition cost $c(\mathcal{S}^{(q,r)}, \pi)$. By sequential optimality (C4), the same is true for any Blackwell less informative signal: The potential welfare gain is exactly equal to the agent's own learning cost as the signal purchase *replaces* endogenous learning. Suppose the principal offers a signal that yields more conclusive evidence in favor of state 1, yielding posteriors π^q and π^s . The agent still relies on the same two actions but implements the high-payoff action more often, so that welfare increases further to $\pi \cdot x^{(q,s)}$. This signal is therefore more valuable to her than the first one, albeit only if she can acquire it at a cost strictly below her endogenous information acquisition cost. If the possible evidence in favor of state 1 grows even further, to posteriors π^q and π^t , the agent may eventually implement different actions, which further raises the welfare gains. Still, her willingness to pay remains strictly below her endogenous information acquisition cost. It is also possible to design signals that have no value to her. Here, this is the case if we add a third state that occurs with likelihood $\pi_3 = \frac{1}{2}$ in all previously discussed beliefs, and where all actions achieve a zero payoff. If, in addition, costs are linear in regards to that state, $c(\mathcal{S}, t\pi + (1-t)\pi') = tc(\mathcal{S}, \pi) + (1-t)c(\mathcal{S}, \pi')$ whenever $\pi_1/\pi_2 = \pi'_1/\pi'_2$, then the agent

derives no benefit from learning about the relative likelihood of state 3: It does not help her identify the most profitable action, nor does it make her learning any easier.

B.2.5 Constrained Optimum

In [Figure 6\(a\)](#), we illustrate how to determine the optimal offer and acceptance strategies for one particular $\mathbf{t} \notin T$. This state-independent transfer is above the threshold for unconditional acceptance $\underline{t}$ but below the threshold for unconditional rejection $\bar{t}$.³⁵ As discussed above, Agent K can either accept the trade and then implement one of the actions $\mathbf{a} \in \mathcal{A}^K$ for a payoff of $\mathbf{a} - \mathbf{t}$, or he can settle for his outside option $\mathbf{0}$. Shifting all payoffs by $\mathbf{t}$, we can thus identify his optimal acceptance strategy from the learning-proof menu over $\mathcal{A}^K \cup \{\mathbf{t}\}$ under cost c^K (indicated in orange). Agent K rejects the offer at a posterior $\boldsymbol{\rho}^R$ orthogonal to the learning-proof menu at $\mathbf{t}$, and accepts the offer at a posterior $\boldsymbol{\rho}^A$ orthogonal at the other anchor action.

Using backwards induction, Agent 1 realizes that offering terms $\mathbf{t}$ will yield expected payoff $\mathbf{t} \cdot \boldsymbol{\rho}^A$ conditional on trade happening, and $W(\mathcal{A}^1, \boldsymbol{\rho}^R, c^1)$ otherwise. By drawing equiprobability lines through $\mathbf{t}$ and tangent to his learning-proof menu, he determines the ‘full-trade equivalent’ transfers $\mathbf{t}^W$. As far as Agent 1 is concerned, offering $\mathbf{t}$ is equivalent to certain trade at $\mathbf{t}^W$ — and this is true for all beliefs between $\boldsymbol{\rho}^A$ and $\boldsymbol{\rho}^R$. However, $\mathbf{t}^W$ is not attractive enough to warrant an unconditional offer, as can be seen by the fact that it lies below one of the lower bounds from [Equation \(10\)](#). The learning-proof menu over $\mathcal{A}^1 \cup \{\mathbf{t}^W\}$ (drawn in dark red) reveals the equilibrium offer strategy: Agent 1 draws a noisy signal and either offers the trade $\mathbf{t}$ at posterior $\boldsymbol{\rho}^O$ or implements the action that favors state one at posterior $\boldsymbol{\rho}^\emptyset$. If the trade is offered, Agent K draws a more accurate signal and either rejects at posterior $\boldsymbol{\rho}^R$ or accepts at posterior $\boldsymbol{\rho}^A$, in which case he goes on to implement action $\circ$.

[Figure 6\(b\)](#) plots the resulting payoff to Agent 1 across a range of state-independent transfers. For low prices $t < \underline{t}$, trade happens whenever Agent 1 wants it to, and thus his payoff increases monotonically with the price t . For high prices $t > \bar{t}$, trade never happens and Agent 1 receives his autarky payoff $W(\mathcal{A}^1, \boldsymbol{\pi}, c^1)$. Agent 1’s optimal price includes partial trade and corresponds to the equilibrium constructed in [Figure 6\(a\)](#).

³⁵These thresholds can be computed analogous to [Equation \(11\)](#) and [Equation \(10\)](#), respectively, using Agent K ’s menu and cost.

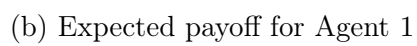
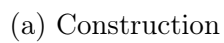


Figure 6: Partial trade when states are unverifiable