

Complementary Information and Learning Traps^{*}

Annie Liang[†] Xiaosheng Mu[‡]

September 30, 2019

Abstract

We develop a model of social learning from complementary information: Short-lived agents sequentially choose from a large set of flexibly correlated information sources for prediction of an unknown state, and information is passed down across periods. Will the community collectively acquire the best kinds of information? Long-run outcomes fall into one of two cases: (1) efficient information aggregation, where the community eventually learns as fast as possible; (2) “learning traps,” where the community gets stuck observing suboptimal sources and information aggregation is inefficient. Our main results identify a simple property of the underlying informational complementarities that determines which occurs. In both regimes, we characterize which sources are observed in the long run and how often.

JEL codes: D81, D83, D62, O32.

^{*}We are grateful to Nageeb Ali, Aislinn Bohren, Tilman Börgers, Drew Fudenberg, Ben Golub, David Hirshleifer, Emir Kamenica, George Mailath, Paul Milgrom, Andrew Postlewaite, Ilya Segal, Carlos Segura, Rajiv Sethi, Andrzej Skrzypacz, Vasilis Syrgkanis, and Yuichi Yamamoto for comments and suggestions that improved this paper. We also thank four anonymous referees for very helpful suggestions. Xiaosheng Mu acknowledges the hospitality of the Cowles Foundation at Yale University, which hosted him during parts of this research.

[†]University of Pennsylvania. **Address:** Perelman Center for Political Science and Economics, 133 S. 36th Street, Philadelphia, PA 19104. **Email:** anliang@upenn.edu.

[‡]Columbia University. **Address:** International Affairs Building Room 1209, 420 West 118th Street, New York, NY 10027. **Email:** xm2230@columbia.edu.

1 Introduction

Societies accumulate knowledge over time through the efforts of many individuals. The quality of current knowledge—e.g., about science and engineering—translates into the quality of decisions that are made, and so the rates at which societies accumulate knowledge matter for their welfare. These rates vary substantially across contexts. What helps a society efficiently refine its knowledge over time? What, in contrast, can cause the social process of knowledge acquisition to become trapped in slower, less productive paths?

As a leading example, we focus on research as a channel of knowledge accumulation. The process of research is inherently “path-dependent,” since past research affects what current researchers know, and thus which investigations they consider most valuable to undertake. If research were centrally planned, then a planner’s directives could account for the externalities of researchers’ activities. But researchers are often motivated by more immediate and individual goals, such as making the biggest possible discoveries now (rather than enabling others to make better discoveries later). How does the socially optimal information acquisition process compare to a decentralized one that results from the choices of short-lived agents, who do not internalize the externalities of their information acquisitions?

In this paper, we consider a sequence of researchers working on a shared scientific problem, which we model as acquiring information over time about the value of an unknown parameter. We study the rate at which the society learns about this parameter. Our main results demonstrate that the structure of the interdependencies across the available signals—more precisely, *informational complementarities*—is crucial in determining whether the decentralized process is essentially efficient, or whether it gets trapped at a learning rate inferior to the optimal one.^{1,2}

To fix ideas, consider as a stylized example a particular research question: whether dopamine levels are predictive of the onset and severity of Parkinson’s disease (PD). Each researcher can conduct a study to learn more about the typical level of dopamine in PD patients. A study is modeled as a Gaussian signal, whose realization is informative about the value of the parameter. For example, researchers have available to them various technolo-

¹The basic feature that information can be complementary appears in many settings besides research—for example, informational complementarities are relevant for auctions (Milgrom and Weber, 1982a,b), for firms (Athey and Schmutzler, 1995), for team composition (Chade and Eeckhout, 2018), and within markets (Goldstein and Yang, 2015; Chen and Waggoner, 2016). We build on this literature—which has primarily focused on *one-time* information acquisitions—by asking how informational complementarities affect learning in a *dynamic* setting.

²Prior work studying the rate of learning include Vives (1992), DeMarzo, Vayanos and Zwiebel (2003), Golub and Jackson (2012), Hann-Caruthers, Martynov and Tamuz (2017), and Harel et al. (2018) among others.

gies for measuring dopamine levels, such as different imaging methods. We model limits on researchers’ resources by restricting each researcher to obtain one signal about the unknown parameter from a finite but potentially large set of signals: In the example, this corresponds to a choice of which measurement to use in the study.³ The researcher’s goal is to maximize immediate reduction in uncertainty about the unknown parameter. This framework is a social learning model, but our paper departs from the classic model (Banerjee, 1992; Bikhchandani, Hirshleifer and Welch, 1992; Smith and Sørensen, 2000) by assuming that all information is *public*, thus turning off the inference problem essential to informational cascades in the previous literature.⁴ Additionally, we suppose that information is endogenously acquired—as in Burguet and Vives (2000), Mueller-Frank and Pai (2016), and Ali (2018)—with the new feature that agents choose from a finite set of complementary signals.

In our model, the choices of successive researchers are linked because of informational complementarities, which we now describe: Each signal observation is modeled as a linear combination of the parameter of interest, various confounding variables, and idiosyncratic Gaussian noise.⁵ How informative the observation is depends on what is known about the confounding variables. For example, sensors used to measure dopamine may be sensitive to other chemicals in the brain in addition to dopamine. The informativeness of the sensor’s measurement is improved by a better understanding of which confounds are picked up by sensors, and to what extent they bias the reading. In this sense, information about the confounding variables, and observations of measurements confounded by those variables, are complementary.

Formally, we propose a definition for *complementary sets of signals*, which builds on prior work by Börgers, Hernando-Veciana and Krahmer (2013). A complementary set is a set of signals that, (i) if observed infinitely often, reveals the value of the payoff-relevant parameter, and (ii) has no proper subset that also reveals the parameter. Removing any source from such a set therefore makes it impossible to learn the parameter of interest.

Our main result is that if the *smallest complementary set of signals is of size K* —where

³This feature relates to a number of recent papers also studying dynamic information acquisition from different *kinds* of information—see for example Che and Mierendorff (2019) and Mayskaya (2019), who study learning from two Poisson signals, and Sethi and Yildiz (2016) and Fudenberg, Strack and Strzalecki (2018), who study learning from multiple Gaussian signals. Although the framework bears some resemblance to classic multi-armed bandit models (Gittins, 1979; Easley and Kiefer, 1988), the signals in our setting do not directly produce payoffs.

⁴A recent paper that also uses a social learning model to study the take-up of new technologies is Wolitzky (2018). This model departs from the classic herding model by assuming that agents observe the *outcomes* of previous agents, and not their actions, which leads to quite different learning dynamics relative to our assumption that agents observe the information of past agents.

⁵Online Appendix G contains a discussion of how our main results might be extended beyond normal signals.

K is the number of unknown variables (including both the parameter of interest and all of the confounding variables)—then a decentralized research process achieves efficient learning in the long run. In contrast, if the smallest complementary set contains fewer than K signals, then early suboptimal information acquisitions can propagate across time, creating persistent inefficiencies in information gathering. We call such outcomes “learning traps.”

The main intuition is as follows. The tension between the interests of short-lived agents and a patient social planner arises because optimal learning may require investments in information about confounding variables that will be useful only later in the learning process. If the smallest complementary set has cardinality equal to the number of unknowns, then the payoff-relevant parameter can only be learned if all confounds are also learned. Thus even short-lived agents will choose to acquire information that reveals the confounding variables, leading them to eventually discover the best set of signals (that leads to efficient learning). By the same logic, if there is a complementary set of size smaller than the number of unknowns, then it is possible to learn the payoff-relevant parameter even if some confounding variables are never learned. This allows the wedge between short- and long-run incentives to persist.

The main technical innovation behind our results is to establish a formal connection between society’s learning dynamics and a limiting dynamical system (see Section 7 for details). We demonstrate that a set of signals is observed in the long-run (from some prior beliefs) if and only if it corresponds to a stable point in this system. This connection enables us to completely characterize the possible long-run observation sets.

Next we consider the size of welfare losses associated with learning traps. We show that the rate of information aggregation can be arbitrarily slower than the efficient benchmark. Measured in terms of absolute discounted payoffs, inefficiency can also be arbitrarily large in a learning trap. However, if measured on a per-period basis (i.e., considering *average discounted payoffs*), the payoff loss caused by a learning trap vanishes in the patient limit, because society eventually learns (albeit slowly) the value of the payoff-relevant parameter. We present a generalization of our model, where the unknown variables are slowly changing over time.⁶ In this model, we show that even the *average* discounted loss can be large. Thus, except in the special case of perfect persistence, learning traps can lead not only to inefficiencies in the rate of information aggregation but also to inefficiencies in payoff terms.

Having shown that learning traps can exist, and that the welfare loss under these traps can be large, we turn next to considering different interventions for avoiding learning traps. We show that policymakers can restore efficient information aggregation by providing in-

⁶This is a technically challenging setting to analyze, and correspondingly prior work is limited: Jovanovic and Nyarko (1996), Moscarini, Ottaviani and Smith (1998), Frongillo, Schoenebeck and Tamuz (2011), Vivi Alatas and Olken (2016), and Dasaratha, Golub and Hak (2018) are the only social learning settings with a dynamic state that we are aware of.

formation about the relevant confounding variables (for example, a forward-looking funding agency can support research about confounding variables that are not of direct societal interest). Another effective intervention is to reshape the payoff structure so that agents can be rewarded for information acquired over many periods. These observations are consistent with practices that have arisen in academic research, including the establishment of third-party funding agencies (e.g., the NSF) to support basic science and methodological research, and the evaluation of researchers based on advancements developed across several papers (e.g., tenure and various awards).

The learning traps we demonstrate connect to a body of work regarding dynamic investment in human capital ([Jovanovic and Nyarko, 1996](#); [Cunha and Heckman, 2007](#); [Lizzeri and Siniscalchi, 2008](#)), which studies complementarities in production technologies (rather than complementarities in information). There are certain high-level connections: For example, this literature shows that early investment choices impact the incentives to invest in the future, and that early misallocations potentially lead to worse long-run outcomes (similar to our learning traps). And related to our [Section 9](#), interventions can push agents onto better skill-acquisition paths. However, the specific form of long-run inefficiencies we obtain follows from the structure of informational complementarities in our setting.

Our framework, and the long-run learning patterns we identify, are particularly related to [Sethi and Yildiz \(2016, 2019\)](#). These papers study a setting in which individuals can learn from one another over time, choosing which of many individuals to listen to in each period. A key force in [Sethi and Yildiz \(2016\)](#) is that repeatedly listening to a given individual is informative about that individual’s “perspective,” which allows for improved (future) de-biasing of that individual’s information. This pushes individuals to return to those they have listened to before. The same force plays an important role in our paper: Repeated acquisition of a given signal helps agents to learn not just the payoff-relevant parameter, but also all of the variables confounding that signal. Unlike in [Sethi and Yildiz \(2016\)](#), we allow for different signals to *share* confounding variables. Thus, observations of a given signal not only help agents interpret future observations of signals of the same kind, but also potentially aid in the interpretation of *other* signals. It is exactly the structure of these learning spillovers that determines whether inefficient learning can obtain in the long run. Our finding of learning traps in the present paper is related to the observation of long-run homophily in [Sethi and Yildiz \(2019\)](#), where individuals eventually listen only to others whose perspectives are correlated with their own.

Finally, in related work ([Liang, Mu and Syrgkanis, 2017, 2019](#)), we study dynamic learning from correlated Gaussian signals, focusing on informational environments with a *single* complementary set, where agents must attend to all signals to learn the payoff-relevant parameter. We find there that myopic information acquisition is in some cases efficient from

period 1 (Liang, Mu and Syrgkanis, 2017), and that under conditions on the prior belief, the optimal path of information acquisitions admits a simple and exact characterization (Liang, Mu and Syrgkanis, 2019). The present paper considers a substantially more general class of informational environments—in particular, allowing for society to have multiple ways for eventually learning the state—but focuses on *asymptotic* efficiency. We find here a necessary and sufficient condition for long-run efficiency, which nests the environment considered in Liang, Mu and Syrgkanis (2017, 2019). In environments that do not satisfy this condition, we demonstrate that in contrast, learning traps may obtain.

2 Examples

Our main results demonstrate that in some informational environments, society maximizes the long-run speed of learning, while in others, agents can persistently acquire information from inefficient sources. The examples below illustrate properties of the informational environment that differentiate these two cases.

2.1 Existence of Learning Traps

Suppose agents sequentially acquire information to learn about an unknown parameter $\omega \sim \mathcal{N}(\mu_\omega, \sigma_\omega^2)$, and have access to three kinds of signals. The first is

$$X_1 = 3\omega + b + \varepsilon_1,$$

where $b \sim \mathcal{N}(\mu_b, \sigma_b^2)$ is a persistent and unknown bias that is independent of ω . The noise term $\varepsilon_1 \sim \mathcal{N}(0, 1)$ is independent of ω and b , and is redrawn each time an agent acquires signal X_1 . The second signal provides information about the bias term b :

$$X_2 = b + \varepsilon_2.$$

Finally, there is an unbiased signal about the parameter of interest:

$$X_3 = \omega + \varepsilon_3.$$

Like ε_1 , both ε_2 and ε_3 are standard Gaussian noise terms independent of ω and b (and of one another).

Agents are indexed by discrete time and act in order. Each agent chooses to acquire one independent observation of either X_1 , X_2 , or X_3 , where his choice is determined by which observation would be most informative about ω (i.e., maximally reduces the uncertainty about ω). This signal realization is then made public.

Repeated acquisitions of X_3 suffice for agents to eventually learn ω . But acquisitions of X_1 , if de-biased via acquisitions of X_2 , can lead to much faster learning. This follows from our subsequent Claim 1, but one can also observe that $X_1 - X_2$ is Blackwell more-informative than two realizations of X_3 ; thus, agents learn faster by alternating between X_1 and X_2 than by acquiring X_3 exclusively.⁷

But suppose that agents have large initial uncertainty about the bias term b —specifically, the prior variance $\sigma_b^2 > 8$. In this case, the signal X_3 is initially most informative about ω , as the biased signal X_1 is noisier due to the uncertainty about b , and the signal X_2 about b is completely uninformative (recall that b is independent of ω in the prior). The first agent’s acquisition of X_3 does not provide any information about b , so the above arguments show that X_3 remains most informative for the second agent. Iterating this logic, it follows that *every* agent chooses to acquire X_3 .

This example demonstrates that although it is socially optimal for agents to invest in understanding the bias b , short-sighted agents will choose instead to repeatedly acquire X_3 , maximizing immediate gains but leading to inefficiently slow long-run learning about ω . We refer to the set $\{X_3\}$ as a “learning trap.”⁸

2.2 Efficient Information Aggregation

In contrast, consider the following informational environment with signals

$$\begin{aligned} X_1 &= \omega + b_1 + \varepsilon_1, \\ X_2 &= b_1 + b_2 + \varepsilon_2, \\ X_3 &= b_2 + \varepsilon_3, \\ X_4 &= 10\omega + b_1 + 2b_2 + \varepsilon_4. \end{aligned}$$

The payoff-relevant parameter ω and confounding variables b_1, b_2 are persistent and jointly normally distributed, while the (independent) noise terms $\varepsilon_1, \varepsilon_2, \varepsilon_3, \varepsilon_4$ are standard Gaussian and i.i.d. across realizations.

As in the previous environment, there are multiple sets of signals that permit long-run learning of ω . Specifically, repeated observations of signals in any of the four sets $\{X_1, X_2, X_3\}$, $\{X_1, X_2, X_4\}$, $\{X_1, X_3, X_4\}$, and $\{X_2, X_3, X_4\}$ will lead agents to eventually learn ω . In Section 4, we will formalize the sense in which these are “complementary” sets of signals. However, the rates of learning permitted by each of these sets are not the same, and society’s long-run speed of learning is strictly maximized when agents repeatedly acquire X_2 ,

⁷Note that $X_1 - X_2 = 3\omega + \varepsilon$, where $\varepsilon \sim \mathcal{N}(0, 2)$, whereas two independent realizations of X_3 provide the same information as $2\omega + \varepsilon$, with $\varepsilon \sim \mathcal{N}(0, 2)$. The former is clearly more informative.

⁸Online Appendix G contains a similar example of a learning trap with binary states and signals.

X_3 , and X_4 .⁹ It turns out that starting from *any* prior belief over the unknown variables, agents will eventually discover this best set and exclusively acquire these signals.

The key difference between the two environments is that the inefficient sets $\{X_1, X_2, X_3\}$, $\{X_1, X_2, X_4\}$ and $\{X_1, X_3, X_4\}$ in this example share their confounding variables (b_1 and b_2) with the efficient set $\{X_2, X_3, X_4\}$. Repeated observation of signals from any inefficient set will thus lead agents to learn b_1 and b_2 , given which they can de-bias signals in the efficient set. This informational spillover from learning about ω to learning about the confounding variables means that inefficient sets fail to be self-reinforcing—the information they generate helps agents to realize that there are more informative signals. Put another way, the complementarities between signals in an inefficient set are eventually outweighed by stronger complementarities with signals outside of the set. Only the efficient set $\{X_2, X_3, X_4\}$ has the property that complementarities within the set remain strongest (see the discussion after Theorem 1 for more detail).

In contrast, in our example in Section 2.1, repeated acquisition of the signal X_3 provides no information about the variable b_1 that confounds the signals $\{X_1, X_2\}$, and so repeated observation of X_3 is self-reinforcing. Our subsequent main results make this contrast precise, and explain in general how the structure of complementarities across signals determines the efficiency of long-run information aggregation.

3 Setup

Informational Environment. There are K persistent unknowns: a real-valued payoff-relevant state ω and $K - 1$ real-valued confounding variables $b_1, \dots, b_{K-1}$.¹⁰ We assume that the state vector $\theta := (\omega, b_1, \dots, b_{K-1})'$ follows a multivariate normal distribution $\mathcal{N}(\mu^0, \Sigma^0)$ where μ^0 is a $K \times 1$ real-valued vector, and the prior covariance matrix Σ^0 has full rank.¹¹

There are N (fixed) *kinds* or *sources* of information available at each discrete period $t \in \mathbb{Z}_+$. Observation of source i in period t produces a realization of the random variable

$$X_i^t = \langle c_i, \theta \rangle + \varepsilon_i^t = \omega c_{i1} + b_1 c_{i2} + \dots + b_{K-1} c_{iK} + \varepsilon_i^t, \quad \varepsilon_i^t \sim \mathcal{N}(0, 1)$$

where $c_i = (c_{i1}, \dots, c_{iK})'$ is a vector of constants, and the noise terms ε_i^t are independent from each other and across periods. Normalizing these noise terms to have unit variance is without loss of generality, since the coefficients c_i are unrestricted. We will often drop the time indices on the random variables, associating $X_i = \langle c_i, \theta \rangle + \varepsilon_i$ with source i and

⁹This follows from the later Proposition 2: It can be checked that $\text{val}(\{X_2, X_3, X_4\}) = 100/9$ is largest across the four sets, so the “best set” is $\mathcal{S}^* = \{X_2, X_3, X_4\}$.

¹⁰Online Appendix D.3 discusses the case of multiple payoff-relevant states.

¹¹The full rank assumption is without loss of generality: If there is linear dependence across the states, the model can be reduced to a lower dimensional state space that satisfies full rank.

understanding that the noise term is independently realized with each new observation. For notational ease, we use $[N]$ to denote the set $\{1, \dots, N\}$ of all sources.

The payoff-irrelevant unknowns $b_1, \dots, b_{K-1}$ produce correlations across the sources, and can be interpreted for example as confounding variables. The difference between the terms b_i and the terms ε_i is that the former are persistent over time while the latter are i.i.d.—so the variances of b_i can be reduced over time, but the variances of ε_i are fixed. Separating these terms allows us to distinguish between reducible and irreducible noise (with respect to learning about ω) in the signals.

Decision Environment. Agents are indexed by discrete time t and move sequentially. Each agent first chooses one of the N sources, observing an independent realization of the corresponding signal. He then predicts ω , selecting an action $a \in \mathbb{R}$ and receiving the payoff $-\mathbb{E}[(a - \omega)^2]$. The agent’s optimal prediction is the posterior mean of ω , and the payoff is the negative of the posterior variance of ω . We note that the action a is not necessary for specifying the model: The analysis is unchanged if we directly assume that each agent chooses to acquire the signal that maximally reduces posterior variance of ω . In some applications, the latter formulation (where the agent acquires information but does not take an action) may be more natural—e.g., if the agent’s goal is simply to progress scientific understanding—while in others, it may be natural to suppose that the agent does take an action on the basis of his information (e.g., recommends a treatment) and receives a higher payoff when his belief is more precise. We note additionally that the specific payoff function of quadratic loss is not crucial: Our subsequent analysis, with the exception of the result in part (b) of Proposition 2, goes through for arbitrary payoff functions $u(a, \omega)$. See Online Appendix D.1 for more detail.

We assume throughout that all signal realizations are public. Thus, each agent t faces a history $h^{t-1} \in ([N] \times \mathbb{R})^{t-1} = H^{t-1}$ consisting of all past signal choices and their realizations, and his signal acquisition strategy is a mapping from histories to sources. At every history h^{t-1} , the agent’s payoffs are maximized by choosing the signal that minimizes his posterior variance of ω .¹²

Society’s Information Acquisitions. Since the environment is Gaussian, the posterior variance of ω is a deterministic function $V(q_1, \dots, q_N)$ of the number of times q_i that each signal i has been observed so far.¹³ Thus, each agent’s signal acquisition is a function of the prior and past signal acquisitions only (and does not depend on the signal realizations).

¹²Online Appendix D.2 shows that our results generalize to agents who are slightly forward-looking.

¹³See Appendix A.1 for the complete closed-form expression for V , which depends on the prior Σ^0 and signal coefficient vectors $\{c_i\}_{i=1}^N$.

This allows us to track society’s acquisitions as deterministic *count vectors*

$$m(t) = (m_1(t), \dots, m_N(t))' \in \mathbb{Z}_+^N$$

where $m_i(t)$ is the number of times that signal i has been observed up to and including period t . The count vector $m(t)$ evolves according to the following rule: $m(0)$ is the zero vector, and for each time $t \geq 0$ there exists $i^* \in \operatorname{argmin}_i V(m_i(t) + 1, m_{-i}(t))$ such that

$$m_i(t+1) = \begin{cases} m_i(t) + 1 & \text{if } i = i^* \\ m_i(t) & \text{otherwise} \end{cases}$$

That is, the count vector increases by 1 in the coordinate corresponding to the signal that yields the greatest immediate reduction in posterior variance. We allow ties to be broken arbitrarily, and there may be multiple possible paths $\{m(t)\}_{t=0}^\infty$.

We are interested in the *long-run frequencies* of observation $\lim_{t \rightarrow \infty} \frac{m_i(t)}{t}$ for each source i —that is, the fraction of periods eventually devoted to each source. As we show later in Section 6, these limits exist under a weak technical assumption. But note that the limit may depend on the prior belief, as already illustrated by the example in Section 2.1. Which long-run outcomes are possible is one of the central questions we seek to understand in this paper.

4 Complementary Set of Sources

We first introduce a definition for a *complementary set of sources*. These sets will play an important role in the subsequent results. As a preliminary step, we assign to each set of sources $\mathcal{S} \subseteq [N] := \{1, \dots, N\}$ an *informational value*.

Informational Value. Write $\tau(q_1, \dots, q_N) = 1/V(q_1, \dots, q_N)$ for the posterior precision about the payoff-relevant state ω given q_i observations of each source i , where the prior precision is $\tau_0 := \tau(0, \dots, 0)$. The informational value of $\mathcal{S}$, denoted $\operatorname{val}(\mathcal{S})$, is defined to be the largest feasible improvement on precision (averaged across many periods), when signals are acquired from $\mathcal{S}$ alone.¹⁴

¹⁴This definition of informational value closely resembles the definition of the value of a team in Chade and Eeckhout (2018), although we consider belief precision instead of negative posterior variance. Using posterior variances in Definition 1 would yield a value given by $\operatorname{val}(\mathcal{S}) = \limsup_{t \rightarrow \infty} \left[\max_{q^t \in Q_{\mathcal{S}}^t} (-tV(q^t)) \right]$. This together with Definition 2 would return a similar notion of complementarity, but would present the technical issue of evaluating $\infty - \infty$ since value as defined this way (is always negative and) can be $-\infty$.

Definition 1. The (asymptotic) informational value of the set $\mathcal{S}$ is the maximal per-period increase in the precision about ω over a long horizon:

$$\text{val}(\mathcal{S}) = \limsup_{t \rightarrow \infty} \left[\max_{q^t \in Q_{\mathcal{S}}^t} \left(\frac{\tau(q^t) - \tau_0}{t} \right) \right]$$

where

$$Q_{\mathcal{S}}^t = \left\{ q \in \mathbb{Z}_+^N : \sum_{i=1}^N q_i = t \text{ and } \text{supp}(q) \subset \mathcal{S} \right\}$$

is the set of all count vectors that allocate t observations across (only) the sources in $\mathcal{S}$.

The informational value is defined with respect to learning about ω , but we omit this dependence since the payoff-relevant state is fixed throughout this paper. The informational value turns out to be *prior-independent*, as we show in Proposition 1 below. Finally, note that $\text{val}(\mathcal{S})$ exceeds zero only if it is possible to completely learn ω given infinite observations from $\mathcal{S}$, as otherwise the reduction in variance is finite, and hence the *average* improvement of each signal observation (taking the total number of observations to infinity) must be zero.

Complementary Set. Our definition for a *complementary set* is based on [Börger, Hernando-Veciana and Krahmer \(2013\)](#); see Online Appendix H for an extended comparison.

Definition 2. The set $\mathcal{S}$ is complementary if $\text{val}(\mathcal{S}) > 0$ and

$$\text{val}(\mathcal{S}) > \text{val}(\mathcal{S}') + \text{val}(\mathcal{S} \setminus \mathcal{S}') \quad (1)$$

for all nonempty proper subsets $\mathcal{S}'$ of $\mathcal{S}$.

The condition in (1) requires that the marginal value of having access to the sources in any $\mathcal{S}' \subset \mathcal{S}$ is increased by also having access to sources $\mathcal{S} \setminus \mathcal{S}'$.¹⁵ We note that the first condition that $\text{val}(\mathcal{S}) > 0$ is implied by (1) whenever $\mathcal{S}$ is not a singleton.¹⁶

Characterization of Complementary Sets. The proposition below shows that a set $\mathcal{S}$ is complementary if and only if its signals uniquely combine to produce an unbiased signal about ω .

Proposition 1. $\mathcal{S}$ is a complementary set if and only if the first coordinate vector in $\mathbb{R}^K$ admits a unique decomposition $(1, 0, \dots, 0)' = \sum_{i \in \mathcal{S}} \beta_i^{\mathcal{S}} \cdot c_i$, where all coefficients $\beta_i^{\mathcal{S}}$ are nonzero.

¹⁵Since $\text{val}(\emptyset) = 0$, the inequality can be rewritten as $\text{val}(\mathcal{S}) - \text{val}(\mathcal{S} \setminus \mathcal{S}') > \text{val}(\mathcal{S}') - \text{val}(\emptyset)$.

¹⁶The condition $\text{val}(\mathcal{S}) > 0$ has bite when $|\mathcal{S}| = 1$. Specifically, it rules out all singleton sets $\mathcal{S}$, except those consisting of a single unbiased signal $X = c\omega + \varepsilon$.

This characterization makes clear a second, equivalent, way of understanding complementary sets: They are sets with the property that observing all signals within that set infinitely often reveals the value of the payoff-relevant state, and moreover, all signals are crucial for this recovery—that is, removing any source(s) from a complementary set makes recovery of the payoff-relevant state impossible.

Proposition 1 allows us to identify complementary sets based on their signal coefficient vectors:

Example 1. The set $\{X_1, X_2, X_3\}$ consisting of signals $X_1 = \omega + b_1 + \varepsilon_1$, $X_2 = b_1 + b_2 + \varepsilon_2$, and $X_3 = b_2 + \varepsilon_3$ is complementary. To see this, observe that $(1, 0, 0)' = c_1 - c_2 + c_3$ (where $c_1 = (1, 1, 0)'$ is the coefficient vector associated with X_1 , $c_2 = (0, 1, 1)'$ is the coefficient vector associated with X_2 , and $c_3 = (0, 0, 1)'$ is the coefficient vector associated with X_3). In contrast, the set of signals $\{X_4, X_5\}$ with $X_4 = \omega + \varepsilon_1$ and $X_5 = 2\omega + \varepsilon_2$ is not complementary, since many different linear combinations of c_4 and c_5 produce $(1, 0)'$. The set $\{X_1, X_2, X_3, X_4\}$ is also not complementary, although it contains two complementary subsets.

Best Complementary Set $\mathcal{S}^*$. The informational value for any complementary set can be computed using the following claim:

Claim 1. *Let $\mathcal{S}$ be a complementary set. Then, $\text{val}(\mathcal{S}) = 1 / (\sum_{i \in \mathcal{S}} |\beta_i^{\mathcal{S}}|)^2$, where $\beta_i^{\mathcal{S}}$ are the ones given in Proposition 1.¹⁷*

More generally, $\text{val}(\mathcal{S})$ can be determined for an arbitrary set $\mathcal{S}$ as follows: If $\mathcal{S}$ contains at least one complementary subset, then its value is equal to the *highest value among its complementary subsets*; otherwise the value of $\mathcal{S}$ is zero. This result will follow from Proposition 2 part (a) in the next section.

Throughout the paper, we assume that there is at least one complementary set, and also that complementary sets can be completely ordered based on their informational values.

Assumption 1. *There is at least one complementary set $\mathcal{S} \subseteq [N]$.*

Assumption 2. *Each complementary set has a distinct informational value; that is, $\text{val}(\mathcal{S}) \neq \text{val}(\mathcal{S}')$ for any pair of complementary sets $\mathcal{S} \neq \mathcal{S}'$.*

Assumption 1 guarantees that the payoff-relevant parameter ω is identifiable given the available signals, and hence it is possible to learn it eventually. Our main results extend even when this assumption fails, and we refer the reader to Online Appendix C for details and discussion of some subtleties. Assumption 2 is *generically* satisfied.

¹⁷This claim and the definition of informational value together imply that the minimum posterior variance at time t (when restricting to signals in $\mathcal{S}$) vanishes like $\frac{(\sum_{i \in \mathcal{S}} |\beta_i^{\mathcal{S}}|)^2}{t}$ asymptotically.

Together, these assumptions imply the existence of a “best” complementary set, whose informational value is largest among complementary sets. This set plays an important role, and we denote it by $\mathcal{S}^*$ in the remainder of this paper.

5 Optimal Long-Run Observations

We show next that optimal information acquisitions eventually concentrate on the best complementary set $\mathcal{S}^*$. Specifically, consider the distribution

$$\lambda_i^* = \begin{cases} \frac{|\beta_i^{\mathcal{S}^*}|}{\sum_{j \in \mathcal{S}^*} |\beta_j^{\mathcal{S}^*}|} & \forall i \in \mathcal{S}^* \\ 0 & \text{otherwise} \end{cases}$$

which assigns zero frequency to sources outside of the best set $\mathcal{S}^*$, and samples sources within $\mathcal{S}^*$ proportionally to the magnitude of $\beta_i^{\mathcal{S}^*}$. That is, each signal in $\mathcal{S}^*$ receives frequency proportional to its contribution to an unbiased signal about ω , as defined in Proposition 1. The result below shows two senses in which λ^* is the optimal long-run frequency over signals.

Proposition 2. (a) Optimal Information Aggregation: $\text{val}([N]) = \text{val}(\mathcal{S}^*)$. Moreover, for any sequence $q(t)$ such that $\lim_{t \rightarrow \infty} \frac{\tau(q^t) - \tau_0}{t} = \text{val}([N])$, it holds that $\lim_{t \rightarrow \infty} \frac{q(t)}{t} = \lambda^*$.

(b) Social Planner Problem: For any δ , let $d_\delta(t)$ be the vector of signal counts (up to period t) associated with any signal path that maximizes the δ -discounted average payoff

$$U_\delta := \mathbb{E} \left[- \sum_{t=1}^{\infty} (1 - \delta) \delta^{t-1} \cdot (a_t - \omega)^2 \right]$$

Then there exists $\underline{\delta} < 1$ such that $\lim_{t \rightarrow \infty} \frac{d_\delta(t)}{t} = \lambda^*$ for every $\delta \geq \underline{\delta}$.

Part (a) says that the informational value of $\mathcal{S}^*$ is the same as the informational value of the entire set of available sources. In this sense, having access to all available sources does not improve upon the speed of learning achievable from the best complementary set $\mathcal{S}^*$ alone. Moreover, this speed of learning is attainable *only if* the long-run frequency over sources is the distribution λ^* .¹⁸ Part (b) of Proposition 2 says that a (patient) social planner—who maximizes a discounted average of agent payoffs—will eventually observe sources in the proportions described by λ^* . Based on these results, we subsequently use λ^* as the optimal benchmark against which to compare society’s long-run information acquisitions.

¹⁸This result builds on Chaloner (1984), who shows that a “ c -optimal simultaneous experiment design” exists on at most K points. Part (a) additionally supplies a characterization of the optimal design itself and demonstrates uniqueness, with a minor technical difference that we impose an integer constraint on signal counts. We are not aware of prior work on the discounted payoff criterion studied in Part (b).

6 Main Results

We now ask whether society’s acquisitions converge to the optimal long-run frequencies λ^* characterized above. We show that informational environments can be classified into two kinds—those for which efficient information aggregation is guaranteed (long-run frequencies are λ^* from all prior beliefs), and those for which “learning traps” are possible (there are prior beliefs from which agents end up exclusively observing some set of sources different from the efficient set $\mathcal{S}^*$). Separation of these two classes depends critically on the size of the smallest complementary set.

6.1 Learning Traps vs. Efficiency

Our first result uses an assumption on the signal structure, which requires that every set of $k \leq K$ signals are linearly independent:

Assumption 3 (Strong Linear Independence). *Every $k \leq K$ signal coefficient vectors $c_{i_1}, c_{i_2}, \dots, c_{i_k}$ are linearly independent.*

If there are at least K signals (i.e., $N \geq K$), then Assumption 3 requires every K signal coefficient vectors to be linearly independent. If instead $N < K$, then all of the signal coefficient vectors should be linearly independent.

Strong Linear Independence will be assumed in part (a) of the following result, although not in part (b), nor in any of our subsequent results.

Theorem 1. (a) *Assume Strong Linear Independence. Then for every complementary set $\mathcal{S}$ with $|\mathcal{S}| < K$, there exists an open set of prior beliefs given which agents exclusively observe signals from $\mathcal{S}$.*

(b) *If there are no complementary sets with fewer than K sources, then starting from any prior belief, $\lim_{t \rightarrow \infty} m_i(t)/t = \lambda_i^*$ holds for every signal i .*

Part (a) of the theorem generalizes our example in Section 2.1. It says that every small complementary set (fewer than K signals) is exclusively observed in the long-run from some set of priors.¹⁹ When that complementary set is not the best one, then it is a “learning trap.”²⁰

¹⁹In the example in Section 2.1, the set $\{X_3\}$ is a complementary set of size 1, while $K = 2$.

²⁰Note that the dynamics here are deterministic: Instead of “bad signal realizations” causing a failure of learning, efficient learning either fails or succeeds depending on the prior and signal structure. This is a key difference between our result, and the more classic learning frictions in Banerjee (1992), Bikhchandani, Hirshleifer and Welch (1992), and Smith and Sørensen (2000). We thank an anonymous referee for pointing this out.

In contrast, if no complementary sets are smaller than size K ,²¹ then a very different long-run outcome obtains: Starting from *any* prior, society’s information acquisitions eventually approximate the optimal frequency. Thus, even though agents are short-lived (“myopic”), they end up acquiring information efficiently. We mention that the conclusion of part (b) can be strengthened to $m_i(t) - \lambda_i^* \cdot t$ being bounded as $t \rightarrow \infty$ (see Online Appendix B). Thus, in particular, signals outside of the efficient set $\mathcal{S}^*$ are observed only finitely often. This result provides a stronger sense in which society’s information acquisitions will be asymptotically efficient under the stated assumptions in part (b) of Theorem 1.

We now provide a brief intuition for Theorem 1, and in particular for the importance of the number K : Since each agent chooses the signal whose marginal value (reduction of posterior variance of ω) is highest, any set $\mathcal{S}$ on which signal acquisitions concentrate must satisfy two properties. First, all signals in the set must repeatedly have their turn as “most valuable,” so that agents do not focus on a strict subset of $\mathcal{S}$. This condition requires that the long-run outcome is a complementary set, where (by definition) all sources are critical to the value of the set.

Second, the marginal values of signals in that set must be persistently higher than marginal values of other signals. Not all complementary sets satisfy this criterion, since the sources within a given complementary set $\mathcal{S}$ could have even stronger complementarities with sources outside of the set. If that were the case, observation of sources within $\mathcal{S}$ would eventually push agents to acquire information outside of $\mathcal{S}$.

No complementary set $\mathcal{S}$ consisting of K sources can satisfy this second property unless it is the best set: As observations accumulate from such a set, agents learn about all of the confounding variables and come to evaluate all sources according to “objective” (i.e., prior-independent) asymptotic values. Repeated acquisitions of signals from $\mathcal{S}$ improve the value of signals in $\mathcal{S}^*$ over the value of signals in $\mathcal{S}$. Thus, agents eventually turn to the sources in $\mathcal{S}^*$, achieving efficient information aggregation as predicted by part (b) of Theorem 1.

In contrast, if agents observe only $k < K$ sources, then they can have persistent uncertainty about some confounding variables. This may cause society to persistently undervalue those sources confounded by these variables and to continually observe signals from a small complementary set. We saw this already in the example in Section 2.1, where agents failed to obtain any information about the confounding variable b_1 , and thus persistently undervalued the sources X_1 and X_2 . The same intuition applies to part (a) of Theorem 1.

One may argue that the condition that no complementary set has fewer than K sources is generically satisfied. However, if we expect that sources are endogenous to design or strategic motivations, the relevant informational environments may not fall under this condition. For

²¹It follows from Proposition 1 that there are no complementary sets with more than K sources, so this is equivalent to assuming that all complementary sets are of size K .

example, signals that partition into different groups with group-specific confounding variables (as studied in [Sethi and Yildiz \(2019\)](#)) are economically interesting but non-generic. Part (a) of Theorem 1 shows that inefficiency is a possible outcome in these cases.

Finally, we use a few examples below to illustrate some implications of Theorem 1. First, it is possible that agents end up concentrating on an inefficient set of *higher* cardinality than the optimal set.

Example 2. Suppose the available signals are $X_1 = \omega + b_1 + \varepsilon_1$, $X_2 = b_1 + b_2 + \varepsilon_2$, $X_3 = b_2 + \varepsilon_3$, $X_4 = \omega + b_3 + \varepsilon_4$, and $X_5 = b_3 + \varepsilon_5$; Strong Linear Independence is satisfied. There are two complementary sets, $\{X_1, X_2, X_3\}$ and $\{X_4, X_5\}$, and the set $\{X_4, X_5\}$ is efficient. But part (a) of Theorem 1 tells us that both complementary sets are potential long-run outcomes (since $K = 4$). Thus from some set of priors, agents will end up exclusively observing the inefficient set $\{X_1, X_2, X_3\}$, which is of larger size than the optimal set.

In Section 2.1, we already saw that agents may concentrate on a set of lower cardinality than the efficient complementary set. Thus, we cannot in general compare sizes of learning traps versus efficient sets.

It is also straightforward to see that adding sources can worsen overall learning:

Example 3. Suppose the available signals are $X_1 = 3\omega + b_1 + \varepsilon_1$ and $X_2 = b_1 + \varepsilon_2$. Then, $\{X_1, X_2\}$ is the only complementary set. It follows from part (b) of Theorem 1 that agents will achieve the efficient benchmark with a value of $\text{val}(\{X_1, X_2\}) = 9/4$. Now suppose we add $X_3 = \omega + \varepsilon_3$, returning the example in Section 2.1. The efficient benchmark does not change, but now there are priors that lead to exclusive observation of X_3 , achieving $\text{val}(\{X_3\}) = 1$.

Relatedly, worsening the information content of a signal can *improve* the speed of long-run learning:²²

Example 4. Consider the environment of Section 2.1 with signals $X_1 = 3\omega + b_1 + \varepsilon_1$ and $X_2 = b_1 + \varepsilon_2$, and $X_3 = \omega + \varepsilon_3$. Now suppose we degrade X_3 by replacing it with $X'_3 = \omega + b_1 + \varepsilon_3$. In contrast to Section 2.1, once signal X_3 has been degraded in this way, all priors lead to long-run information acquisition in the efficient frequency, which concentrates on X_1 and X_2 .

We note however that in general, adjustments to the signal structure can result in changes to the efficient speed of learning, so the welfare comparison is not straightforward.

²²We thank an anonymous referee for this example.

6.2 General Characterization of Long-run Outcomes

We now generalize Theorem 1, providing a complete characterization of the possible long-run observations (as the prior belief varies) for an arbitrary signal structure. We need a new definition, which strengthens the notion of a complementary set:

Definition 3. $\mathcal{S}$ is a *strongly complementary set* if it is complementary, and $\text{val}(\mathcal{S}) > \text{val}(\mathcal{S}')$ for all sets $\mathcal{S}'$ such that $|\mathcal{S} - \mathcal{S}'| = |\mathcal{S}' - \mathcal{S}| = 1$.²³

The property of *strongly complementary* can be understood as requiring that the set is complementary and also something more: These complementarities are “locally best” in the sense that it is not possible to obtain stronger complementarities by swapping out just one source. We point out that while the definition of complementary sets does not depend on the ambient set (i.e., $[N]$) of available sources, the notion of strong complementarity does.

Example 5. Suppose the available signals are $X_1 = \omega + b_1 + \varepsilon_1$, $X_2 = b_1 + \varepsilon_2$, and $X_3 = 2b_1 + \varepsilon_3$. Then the set $\{X_1, X_2\}$ is complementary but not strongly complementary, as $\text{val}(\{X_1, X_3\}) > \text{val}(\{X_1, X_2\})$.

Theorem 2 below says that long-run information acquisitions concentrate on a set $\mathcal{S}$ (starting from some prior belief) *if and only if* $\mathcal{S}$ is strongly complementary. This generalizes Theorem 1 to signal structures that need not satisfy Strong Linear Independence.

Theorem 2. *The set $\mathcal{S}$ is strongly complementary if and only if there exists an open set of prior beliefs given which agents eventually exclusively observe signals from $\mathcal{S}$ (that is, long-run frequencies exist and have support in $\mathcal{S}$).*²⁴

When there is a unique strongly complementary set, then all priors must eventually lead to this set. Part (b) of Theorem 1 provides a sufficient condition that implies uniqueness, and moreover gives that the single strongly complementary set is the *best* complementary set. When there are multiple strongly complementary sets, then different priors lead to different long-run outcomes, some of which are inefficient. Part (a) of Theorem 1 describes a sufficient condition for such multiplicity.

Theorem 2 implies that learning must always end in a complementary set, so that (myopic) learners will eventually recover the payoff-relevant state, albeit potentially slowly. This is so even in settings such as the following: Agents have access to $X_1 = \omega + b_1 + \varepsilon_1$ and

²³In fact, the requirement that $\mathcal{S}$ is complementary is extraneous: One can show using Proposition 1 that if $\text{val}(\mathcal{S}) > \text{val}(\mathcal{S}')$ for all sets $\mathcal{S}'$ with $|\mathcal{S} - \mathcal{S}'| = |\mathcal{S}' - \mathcal{S}| = 1$, then $\mathcal{S}$ must be complementary.

²⁴The “if” part of this statement can be strengthened as follows: The set $\mathcal{S}$ is strongly complementary if there exists *any* prior belief given which agents eventually choose from $\mathcal{S}$ (see Appendix A.6). Thus, the regions of prior beliefs that would lead to different strongly complementary sets cover the whole space.

$X_2 = \omega/100 + \varepsilon_2$, where the prior belief over ω and b_1 is standard normal.²⁵ Here, the initial agents will acquire X_1 , viewing X_2 as less informative about ω . But what Theorem 2 tells us is that agents must eventually switch over to acquiring (the efficient signal) X_2 —this is because once agents have learned the sum $\omega + b_1$ very well, the marginal value to learning more about this biased sum will be smaller than the marginal value to learning directly about ω . In Section 8.2, we revisit this observation that agents eventually learn the state starting from all prior beliefs, and show that this can fail when we allow for arbitrarily small amounts of evolution in the state.

Obtaining a complete characterization of the sets of prior beliefs associated with different long-run outcomes is challenging, since society’s signal path may in general exhibit complex dynamics—for example, switching multiple times between different complementary sets (including the efficient set). This makes it difficult to relate the initial prior to long-run learning behavior. In the next section we discuss the technical details that go into the proof of Theorem 2. In particular, we explain how we are able to determine the *range* of possible long-run outcomes despite incomplete knowledge about which occurs under a given prior.²⁶

7 Proof Outline for Theorem 2

7.1 Limiting Dynamical System

We first introduce the following *normalized asymptotic* posterior variance function V^* , which takes frequency vectors $\lambda \in \Delta^{N-1}$ as input, where the i -th coordinate of λ is the proportion of total acquisitions devoted to source i .²⁷

$$V^*(\lambda) = \lim_{t \rightarrow \infty} t \cdot V(\lambda t).$$

The RHS is well-defined because we can extend the domain of the posterior variance function V so that it takes positive real numbers (and not just natural numbers) as arguments (see Appendix A.1). The asymptotic variance function $V^*(\lambda)$ is convex in λ and its unique minimizer is the optimal frequency vector λ^* (see Lemma 5 in Appendix A.2).

For simplicity of explanation, we will assume throughout this section that at large t , the

²⁵We thank an anonymous referee for this example.

²⁶For some signal structures, such as the example in Section 2.1, a partial characterization is feasible. We showed previously that if ω and b_1 are independent under the prior, and the prior variance of b_1 exceeds 8, then every agent observes X_3 and lead to a learning trap. If instead the prior variance of b_1 is smaller than 8, then we can show that every agent chooses from the efficient set $\{X_1, X_2\}$. See Online Appendix E for the analysis as well as a related example.

²⁷Note that the coordinates of λ must sum to 1.

signal choice that minimizes V also minimizes V^* .²⁸ Then, the frequency vector $\lambda(t) := \frac{m(t)}{t}$ evolves in the coordinate direction that minimizes V^* . We will refer to this as *coordinate descent*. Unlike the usual gradient descent, coordinate descent is restricted to move in coordinate directions. This restriction reflects our assumption that each agent can only acquire a discrete signal (rather than a mixture of signals).

One case where the rest point of coordinate descent coincides with that of gradient descent is when V^* is everywhere differentiable, since differentiability ensures that directional derivatives can be written as convex combinations of partial derivatives along coordinate directions. In this case, evolution of $\lambda(t)$ necessarily ends at the global minimizer λ^* , implying efficient information aggregation.

7.2 Differentiability of V^*

The function V^* , however, is *not* guaranteed to be differentiable everywhere. Consider our example from Section 2.1 with signals $X_1 = 3\omega + b_1 + \varepsilon_1$, $X_2 = b_1 + \varepsilon_2$, $X_3 = \omega + \varepsilon_3$, and fix the frequency vector to be $\lambda = (0, 0, 1)$. It is easy to verify that the asymptotic posterior variance $V^*(\lambda)$ is increased if we perturb λ by re-assigning weight from X_3 to X_1 , or from X_3 to X_2 . But V^* is reduced if we re-assign weight from X_3 to *both* X_1 and X_2 in an even manner.²⁹ So V^* is not differentiable at λ . Such points of non-differentiability are exactly why learning traps are possible: Coordinate descent can become stuck at these vectors λ , so that agents repeatedly sustain an inefficient frequency of information acquisitions.³⁰

A sufficient condition for V^* to be differentiable at a frequency vector turns out to be that the signals receiving positive frequencies at that vector span all of $\mathbb{R}^K$. This explains the result in part (b) of Theorem 1: When each complementary set consists of K signals, society *has to* observe K signals in order to learn the payoff-relevant state ω . Thus, driven by learning about ω , agents end up observing signals that span $\mathbb{R}^K$, which leads to efficient information aggregation.

²⁸This is not in fact generally correct, and the potential gap is one of the technical challenges in the proof. Nevertheless, we do show that at large t , the signal choice that minimizes V *approximately* minimizes V^* .

²⁹This follows from the formula $V^*(\lambda_1, \lambda_2, \lambda_3) = \lambda_3 + \frac{9}{1/\lambda_1 + 1/\lambda_2}$. The derivative of V^* in either direction $(1, 0, -1)$ or $(0, 1, -1)$ is positive, while its derivative in the direction $(\frac{1}{2}, \frac{1}{2}, -1)$ is in fact negative.

³⁰The above intuition connects to a literature on learning convergence in potential games (Monderer and Shapley, 1996; Sandholm, 2010). Define an N -player game where each player i chooses a number $\lambda_i \in \mathbb{R}_+$ and receives payoff $-\left(\sum_{j=1}^N \lambda_j\right) \cdot V^*(\lambda) = -V^*\left(\lambda / \sum_{j=1}^N \lambda_j\right)$. Then, we have a potential game with (exact) potential function $-V^*$, and our long-run observation sets correspond to equilibria of this game. This is an infinite potential game with a non-differentiable potential function. It is known that Nash equilibria in such games need not occur at extreme points, and this is consistent with our observation of learning traps. Nonetheless, we note that the connection to potential games is not sufficient to derive our main results, since our agents receive payoff $-V(\lambda t)$ rather than its asymptotic variant V^* .

7.3 Generalization to Arbitrary Subspaces

Now observe that our arguments above were not special to considering the whole space $\mathbb{R}^K$. If we restrict the available sources to some subset of $[N]$, and look at the subspace of $\mathbb{R}^K$ spanned by these sources, our previous analysis applies to this restricted space.

Specifically, given any prior belief, define $\mathcal{S}$ to be the set of sources that agents eventually observe. Let $\bar{\mathcal{S}}$ be the available signals that can be reproduced as linear combinations of signals from $\mathcal{S}$ —these sources belong to the “subspace spanned by $\mathcal{S}$.” We can consider the restriction of the function V^* to all frequency vectors with support in $\bar{\mathcal{S}}$. Parallel to the discussion above, the restricted version of V^* is both convex and differentiable in this subspace (at frequency vectors that assign positive weights to signals in $\mathcal{S}$). Thus, coordinate descent must lead to the minimizer of V^* in this subspace.

Just as the overall optimal frequency vector λ^* is supported on the best complementary set $\mathcal{S}^*$, the frequency vector that minimizes V^* in the restricted subspace is also supported on the best complementary set within $\bar{\mathcal{S}}$. So agents can eventually concentrate signal acquisitions on the set $\mathcal{S}$ only if $\mathcal{S}$ is *best in its subspace*; that is, $\text{val}(\mathcal{S}) = \text{val}(\bar{\mathcal{S}})$.

7.4 An Equivalence Result

Next we demonstrate that a set is “best in its subspace” if and only if it is strongly complementary.

Lemma 1. *The following conditions are equivalent for a complementary set $\mathcal{S}$:*

- (a) $\text{val}(\mathcal{S}) = \text{val}(\bar{\mathcal{S}})$.
- (b) $\mathcal{S}$ is strongly complementary.
- (c) For any $i \in \mathcal{S}$ and $j \notin \mathcal{S}$, $\partial_i V^*(\lambda^{\mathcal{S}}) < \partial_j V^*(\lambda^{\mathcal{S}})$, where $\lambda^{\mathcal{S}}$ (proportional to $|\beta^{\mathcal{S}}|$) is the optimal frequency vector supported on $\mathcal{S}$.

This lemma states that a strongly complementary set $\mathcal{S}$ is “locally best” in three different senses. Condition (a) says such a set has the highest informational value in its subspace. Condition (b) says its informational value is higher than any set obtained by swapping out one source. Condition (c) says that starting from the optimal sampling rule over $\mathcal{S}$, re-allocating frequencies from signals in $\mathcal{S}$ to any other signal increases posterior variance and reduces speed of learning. The rest of this subsection is devoted to the proof of Lemma 1.

The implication from (a) to (b) is straightforward: Suppose $\mathcal{S}$ is best in its subspace, and $\mathcal{S}'$ is obtained from $\mathcal{S}$ by removing signal i and adding signal j . Then the informational value of $\mathcal{S}'$ is either zero, or equal to the highest value among its complementary subsets. In

the latter case, such a complementary subset necessarily includes signal j , and Proposition 1 implies that j belongs to the subspace spanned by $\mathcal{S}$. Thus $\mathcal{S}' \subset \overline{\mathcal{S}}$, and $\text{val}(\mathcal{S}') \leq \text{val}(\overline{\mathcal{S}}) = \text{val}(\mathcal{S})$. The inequality is in fact strict, because complementary sets have different values by Assumption 2.

We next show that (b) implies (c). Suppose Condition (c) fails, so some perturbation that shifts weight from source $i \in \mathcal{S}$ to source $j \notin \mathcal{S}$ decreases V^* . Then, by definition of informational value, we would have $\text{val}(\mathcal{S} \cup \{j\}) > \text{val}(\mathcal{S})$. But as Proposition 2 part (a) suggests, the value of $\mathcal{S} \cup \{j\}$ is equal to the highest value among its complementary subsets. Strong complementarity of $\mathcal{S}$ ensures that $\mathcal{S}$ is the best complementary subset of $\mathcal{S} \cup \{j\}$. Thus we obtain $\text{val}(\mathcal{S} \cup \{j\}) = \text{val}(\mathcal{S})$, leading to a contradiction.

Finally, Condition (c) implies that $\lambda^{\mathcal{S}}$ is a *local* minimizer of V^* in the subspace spanned by $\mathcal{S}$ (where the restriction of V^* is differentiable). Since V^* is convex, the frequency vector $\lambda^{\mathcal{S}}$ must in fact be a *global* minimizer of V^* in this subspace. Hence $\mathcal{S}$ is best in its subspace and (a) holds.

7.5 Completing the Argument

The arguments above tell us that information acquisitions eventually concentrate on a strongly complementary set, delivering one direction of Theorem 2: $\mathcal{S}$ is a long-run outcome only if $\mathcal{S}$ is strongly complementary.

To prove the “if” direction, we directly construct priors such that a given strongly complementary set $\mathcal{S}$ is the long-run outcome. The construction generalizes the idea in the example in Section 2.1, where we assign high uncertainty to those confounding variables that do not afflict signals in $\mathcal{S}$, and low uncertainty to those that do. This asymmetry guarantees that signals in $\overline{\mathcal{S}}$ have persistently higher marginal values than the remaining signals. Lastly, we use part (c) of the above Lemma 1 to show that agents focus on observing signals from $\mathcal{S}$, rather than the potentially larger set $\overline{\mathcal{S}}$. Indeed, if the historical frequency of acquisitions is close to $\lambda^{\mathcal{S}}$, then signals in $\mathcal{S}$ have higher marginal values than the remaining signals in their subspace; and as these signals in $\mathcal{S}$ continue to be chosen, society’s frequency vector remains close to $\lambda^{\mathcal{S}}$. This completes the proof of Theorem 2.

8 Welfare Loss Under Learning Traps

The previous sections demonstrate that long-run learning is sometimes inefficient; how large can this inefficiency be? In this section, we study the welfare loss under learning traps, and in the process, develop a generalization of our model in which the unknown states evolve over time.

8.1 Welfare Criteria

Two classic welfare criteria are the *speed of information aggregation* and the *discounted average payoff* achieved by agents within the community.

According to the first criterion, the welfare loss under learning traps can be arbitrarily large. Specifically, as the following example shows, the informational value of the best complementary set can be arbitrarily large compared to the set that agents eventually observe (and thus, the achieved speed of learning can be arbitrarily slow compared to what is feasible).

Example 6. There are three available sources: $X_1 = \omega + b_1 + \varepsilon_1$, $X_2 = b_1 + \varepsilon_2$, and $X_3 = \frac{1}{L}\omega + \varepsilon_3$, where $L > 0$ is a constant. In this example, the ratio $\text{val}(\{X_1, X_2\}) / \text{val}(\{X_3\}) = \frac{L^2}{4}$ increases without bound as $L \rightarrow \infty$. But for every choice of L , there is a set of priors given which X_3 is exclusively observed.³¹

For the second criterion, define

$$U_\delta^M = \mathbb{E}_M \left[- \sum_{t=1}^{\infty} (1 - \delta) \delta^{t-1} \cdot (a_t - \omega)^2 \right]$$

to be the δ -discounted average payoff across agents who follow a “myopic” signal acquisition strategy with optimal predictions a_t . Also define U_δ^{SP} to be the maximum δ -discounted average payoff, where the social planner can use any signal acquisition strategy. Note that both payoff sums are negative, since flow payoffs are quadratic loss at every period.

Again from Example 6, we see that for every constant $c > 0$, there is a signal structure and prior belief such that the limiting payoff ratio satisfies³²

$$\liminf_{\delta \rightarrow 1} U_\delta^M / U_\delta^{SP} > c.$$

Thus, the *payoff ratio* can be arbitrarily large. Note that because payoffs are negative, larger values of the ratio $U_\delta^M / U_\delta^{SP}$ correspond to greater payoff inefficiencies.

On the other hand, the *payoff difference* vanishes in the patient limit; that is,

$$\lim_{\delta \rightarrow 1} (U_\delta^{SP} - U_\delta^M) = 0$$

³¹The region of inefficient priors (that result in suboptimal learning) does decrease in size as the level of inefficiency increases. Specifically, as L increases, the prior variance of b_1 has to increase correspondingly in order for the first agent to choose X_3 .

³²Example 6 implies the ratio of flow payoffs at late periods can be arbitrarily large. As $\delta \rightarrow 1$, these later payoffs dominate the total payoffs from the initial periods (since the harmonic series diverges). So the ratio of aggregate discounted payoffs is also large.

in all environments. To see this, note that agents eventually learn ω even while in a learning trap, albeit slowly. Thus flow payoffs converge to zero at large periods, implying $\lim_{\delta \rightarrow 1} U_\delta^{SP} = \lim_{\delta \rightarrow 1} U_\delta^M = 0$.

In what follows, we show that this conclusion critically depends on the assumption that unknown states are perfectly persistent. We outline a sequence of autocorrelated models that converge to our main model (with perfect state persistence). At near perfect persistence, welfare losses under learning traps can be large according to all of the above measures.

8.2 Extension: Autocorrelated Model

In our main model, the state vector $\theta = (\omega, b_1, \dots, b_{K-1})'$ is persistent across time. Consider now a state vector θ^t that evolves according to the following law:

$$\theta^1 \sim \mathcal{N}(0, \Sigma^0); \quad \theta^{t+1} = \sqrt{\alpha} \cdot \theta^t + \sqrt{1 - \alpha} \cdot \eta^t, \quad \text{where } \eta^t \sim \mathcal{N}(0, M).$$

Above, means are normalized to zero, and the prior covariance matrix of the state vector at time $t = 1$ is Σ^0 . We restrict the autocorrelation coefficient $\sqrt{\alpha}$ to belong to $(0, 1)$. Choice of $\alpha = 1$ returns our main model, and we will be interested in approximations where α is close to but strictly less than 1. The *innovation* $\eta^t \sim \mathcal{N}(0, M)$ captures the additional noise terms that emerge under state evolution, which we assume to be i.i.d. across time.³³ Fixing signal coefficients $\{c_i\}$, every autocorrelated model is indexed by the triple (M, Σ^0, α) .

In each period, the available signals are

$$X_i^t = \langle c_i, \theta^t \rangle + \varepsilon_i^t, \quad \varepsilon_i^t \sim \mathcal{N}(0, 1).$$

The signal noises ε_i^t are i.i.d. and further independent from the innovations in state evolution. The agent in period t chooses the signal that minimizes the posterior variance of ω^t , while the social planner seeks to minimize a discounted sum of such posterior variances.

We have the following result:

Theorem 3. *Suppose $\mathcal{S}$ is strongly complementary. Then there exist M and Σ^0 such that for every $\varepsilon > 0$, there is an $\underline{\alpha}(\varepsilon) < 1$ with the following property: In each autocorrelated model (M, Σ^0, α) with $\alpha > \underline{\alpha}(\varepsilon)$,*

1. *all agents only observe signals in $\mathcal{S}$;*
2. *the resulting discounted average payoff satisfies*

$$\limsup_{\delta \rightarrow 1} U_\delta^M \leq -(1 - \varepsilon) \cdot \sqrt{(1 - \alpha) \left(\frac{M_{11}}{\text{val}(\mathcal{S})} \right)},$$

³³The coefficient $\sqrt{1 - \alpha}$ in front of η^t is chosen so that when no signals are observed, society's posterior covariance matrix about θ^t will converge to M . This allows us to meaningfully consider the limit as $\alpha \rightarrow 1$ while keeping M fixed.

while it is feasible to achieve a patient payoff of

$$\liminf_{\delta \rightarrow 1} U_{\delta}^{SP} \geq -(1 + \varepsilon) \cdot \sqrt{(1 - \alpha) \left(\frac{M_{11}}{\text{val}(\mathcal{S}^*)} \right)}$$

by sampling from $\mathcal{S}^*$.

Part (1) generalizes Theorem 2, showing that every strongly complementary set is a potential long-run observation set given imperfect persistence. This suggests that the notion of strong complementarity and its importance extend beyond our main model with unchanging states.

Part (2) shows that whenever $\mathcal{S}$ is different from the best complementary set $\mathcal{S}^*$, then social acquisitions result in significant payoff inefficiency as measured by the payoff ratio. Indeed, for α close to 1 the ratio $\lim_{\delta \rightarrow 1} U_{\delta}^M / U_{\delta}^{SP}$ is at least $\sqrt{\text{val}(\mathcal{S}^*) / \text{val}(\mathcal{S})}$, which can be arbitrarily large depending on the signal structure.

The following proposition strengthens this statement, using Example 6 to show that the payoff difference between optimal and social acquisitions can also be arbitrarily large:

Proposition 3. *For every $\varepsilon > 0$, there exists a signal structure as in Example 6 and a corresponding autocorrelated model (M, Σ^0, α) such that*

$$\liminf_{\delta \rightarrow 1} U_{\delta}^{SP} \geq -\varepsilon;$$

$$\limsup_{\delta \rightarrow 1} U_{\delta}^M \leq -\frac{1}{\varepsilon}.$$

From this analysis, we take away that learning traps can result in average payoff losses (and potentially large losses) so long as unknown states are not perfectly persistent over time.

9 Interventions

We have now shown that learning traps are possible, and can lead to large welfare loss. This naturally suggests a question of what kinds of policies could preclude learning traps, or free agents from an inefficient path of learning. We compare several possible policy interventions in this section: increasing the *quality* of information acquisition (so that each signal realization is more informative); restructuring incentives so that agents' payoffs are based on information obtained over several periods (equivalent to acquisition of *multiple signals* each period); and providing a one-shot release of *free information*, which can guide subsequent acquisitions.

9.1 More Precise Information

Consider first an intervention in which the precision of each signal draw is uniformly increased. We model this intervention by supposing that each signal acquisition now produces B independent observations from that source (where the main model is nested as $B = 1$). The result below shows that providing more informative signals is of limited effectiveness: All potential learning traps for $B = 1$ remain potential learning traps under arbitrary improvements to signal precision.

Corollary 1. *Suppose that for $B = 1$, there is a set of priors given which signals in $\mathcal{S}$ are exclusively viewed in the long run. Then, for every $B \in \mathbb{Z}_+$, there is a set of priors given which these signals are exclusively viewed in the long run.*³⁴

This corollary follows directly from Theorem 2.³⁵

9.2 Batches of Signals

Another possibility is to restructure the incentive scheme so that agents' payoffs are based on information acquired from multiple signals. In practice, this might mean that payoffs are determined after a given time interval: For example, researchers may be evaluated based on a set of papers, so that they maximize the impact of the entire set. Alternatively, agents might be given the means to acquire multiple signals each period: For example, researchers may be arranged in labs, with a principal investigator directing the work of multiple individuals at once.

Formally, suppose that each agent can allocate B observations across the sources (where $B = 1$ returns the main model). Note the key difference from the previous intervention: It is now possible for the B observations to be allocated across *different* signals. This distinction enables agents to take advantage of the presence of complementarities, and we show that efficient information aggregation can be guaranteed in this case:

Proposition 4. *For sufficiently large B , if each agent acquires B signals every period, then society's long-run frequency vector is λ^* starting from every prior belief.*

³⁴However, the *set* of prior beliefs that yield $\mathcal{S}$ as a long-run outcome need not be the same as B varies. For a fixed prior belief, subsidizing higher quality acquisitions may or may not move society out of a learning trap. To see this, consider first the signal structure and prior belief from the example in Section 2.1. Increasing the precision of signals is ineffective there: As long as the prior variance on b is larger than 8, each agent still chooses signal X_3 regardless of signal precision. In Online Appendix F, we provide a contrasting example in which increasing the precision of signals indeed breaks agents out of a learning trap from a specified prior.

³⁵To see this, observe that B independent observations reduce the noise variance of each signal to $\frac{1}{B}$. Thus the model with signal coefficient vectors $\{c_i\}$ and $B > 1$ observations is equivalent to our main model ($B = 1$) with scaled coefficient vectors $\{\sqrt{B} \cdot c_i\}$. Since scaling does not change the family of strongly complementary sets, this model produces the same set of learning traps as in our main model.

Thus, given sufficiently many observations each period, agents will allocate observations in a way that approximates the optimal frequency.

The number of observations needed for long-run efficiency, however, depends on details of the informational environment. In particular, the required B cannot be bounded as a function of the number of states K and number of signals N .³⁶ See Appendix A.7.1 for further details.

9.3 Free Information

Finally, we consider provision of free information to the agents. We can interpret this either as release of information that a policymaker knows, or as a reduced form for funding specific kinds of research, the results of which are then made public.

Formally, the policymaker chooses several signals $Y_j = \langle p_j, \theta \rangle + \mathcal{N}(0, 1)$, where each $\|p_j\|_2 \leq \gamma$ so that signal precisions are bounded by γ^2 . At time $t = 0$, independent realizations of these signals are made public. All subsequent agents update their prior beliefs based on this free information in addition to the history of signal acquisitions thus far.

We show that given a sufficient number of (different kinds of) signals, efficient learning can be guaranteed. Specifically, if $k \leq K$ is the size of the best set $\mathcal{S}^*$, then $k - 1$ precise signals are sufficient to guarantee efficient learning:

Proposition 5. *Let $k := |\mathcal{S}^*|$. There exists a $\gamma < \infty$, and $k - 1$ signals $Y_j = \langle p_j, \theta \rangle + \mathcal{N}(0, 1)$ with $\|p_j\|_2 \leq \gamma$, such that with these free signals provided at $t = 0$, society's long-run frequency vector is λ^* starting from every prior belief.*

The proof is by construction. We show that as long as agents understand those confounding variables that appear in the best set of signals (these variables have dimension $k - 1$), they will come to acquire information from this set.³⁷

We point out the following converse to Proposition 5: Whenever agents begin with sufficiently low prior uncertainty about the confounding variables that afflict the signals in $\mathcal{S}^*$, it is *impossible* for a malevolent third-party to provide free information and induce a learning trap as the long-run outcome.

³⁶The required B depends on two properties: First, it depends on how well the optimal frequency λ^* can be approximated by B (discrete) observations. Second, it depends on the difference in learning speed between the best set and the next best complementary set, which determines the slack that is permitted in the approximation of λ^* .

³⁷This intervention requires knowledge of the full correlation structure as well as which set $\mathcal{S}^*$ is best. An alternative intervention, with higher demands on information provision but lower demands on knowledge of the environment, is to provide $K - 1$ (sufficiently precise) signals about all of the confounding variables.

10 Conclusion

We conclude with brief mention of additional directions and interpretations of the model. First, although we have focused on a sequence of decision-makers with a *common prior*, we might alternatively consider multiple communities of decision-makers, each seeded with a different prior belief. For example, in the absence of a global research community, researchers in different countries may share different prior beliefs and pass down their information within their country. Under this setup, our results can be interpreted as answering the question: *Will individuals from different communities end up observing the same (best) set of sources, or will they persistently acquire information from different sources?* Our main results show that when there is a unique strongly complementary set of sources, then different priors wash out; otherwise, different priors can result in persistent differences across communities in what sources are listened to and consequently differences in beliefs.³⁸

Second, although we have focused on research as the leading interpretation of the framework, the model is relevant to other settings of knowledge acquisition where: (1) information is passed down across time/generations, and (2) information acquisition is myopic at each period. For example, we may consider a sequence of managers within a company (or politicians within a state) who seek only to maximize profits during their tenure, but acquire information that has externalities for future managers. Alternatively, we may consider knowledge acquisition by a *single decision-maker* over time, e.g. an aspiring computer programmer’s choices of what classes to take or blogs to read. Here, too, investment in certain skills (e.g., abstract math classes) may not be immediately useful, but may allow the individual to learn faster in future courses (e.g., an algorithms course). Our paper characterizes the cases in which a student who “only learns things that are useful right now” will nevertheless end up developing his abilities as fast as a student who recognized in advance the complementarities across courses.

Finally, while we consider choice between information sources, a more general model may consider choice between complementary actions. The concepts of *efficient information aggregation* and *learning traps* have natural generalizations (i.e., actions that maximize society’s long-term welfare, versus those that do not). Relative to the general setting, we study here a class of complementarities that are micro-founded in correlated signals. It is an interesting question of whether and how the forces we find here generalize to other kinds of complementarities.

³⁸In our main model with persistent states, beliefs about ω end up converging across the population to the truth. However, beliefs about other confounding variables need not converge. And when states are not fully persistent (as in Section 8.2), even beliefs about ω can diverge across communities.

A Proofs for the Main Model

The structure of the appendix follows that of the paper. In this appendix we provide proofs for the results in our main model, where states are perfectly persistent. The next appendix provides proofs for the autocorrelated model as discussed in Section 8.2. The only exception is that the proof of part (b) of Proposition 2 is more technical, so it is given in a separate Online Appendix, which also contains additional results and examples.

A.1 Preliminaries

A.1.1 Posterior Variance Function

Throughout, let C denote the $N \times K$ matrix of signal coefficients, whose i -th row is the vector c'_i associated with signal i . Here we review and extend a basic result from Liang, Mu and Syrgkanis (2017). Specifically, we show that the posterior variance of ω weakly decreases over time, and the marginal value of any signal decreases in its signal count.

Lemma 2. *Given prior covariance matrix Σ^0 and $q_i \in \mathbb{Z}_+$ observations of each signal i , society's posterior variance of ω is*

$$V(q_1, \dots, q_N) = [((\Sigma^0)^{-1} + C'QC)^{-1}]_{11} \quad (2)$$

where $Q = \text{diag}(q_1, \dots, q_N)$.

This function V admits an extension to the larger domain of non-negative real numbers q_i (beyond integers), and the extended function is decreasing and convex in each q_i .

Proof. Note that $(\Sigma^0)^{-1}$ is the prior precision matrix and $C'QC = \sum_{i=1}^N q_i \cdot [c_i c'_i]$ is the total precision from the observed signals. Thus (2) simply represents the fact that for Gaussian prior and signals, the posterior precision matrix is the sum of the prior and signal precision matrices. The RHS of (2) can be evaluated for any $q_i \in \mathbb{R}_+$, providing an extension of the function V to non-integral arguments.

To prove the monotonicity of V , consider the partial order $\succeq$ on positive semi-definite matrices where $A \succeq B$ if and only if $A - B$ is positive semi-definite. As q_i increases, the matrix Q and $C'QC$ increase in this order. Thus the posterior covariance matrix $((\Sigma^0)^{-1} + C'QC)^{-1}$ decreases in this order, which implies that the posterior variance of ω decreases.

To prove that V is convex, it suffices to prove that V is *midpoint-convex* since the function is clearly continuous.³⁹ Take $q_1, \dots, q_N, r_1, \dots, r_N \in \mathbb{R}_+$ and let $s_i = \frac{q_i + r_i}{2}$. Define the corresponding diagonal matrices to be Q, R, S . Note that $Q + R = 2S$. Thus by the AM-HM inequality for positive-definite matrices, we have

$$((\Sigma^0)^{-1} + C'QC)^{-1} + ((\Sigma^0)^{-1} + C'RC)^{-1} \succeq 2((\Sigma^0)^{-1} + C'SC)^{-1}.$$

³⁹ A function V is midpoint-convex if the inequality $V(a) + V(b) \geq 2V(\frac{a+b}{2})$ always holds. Every continuous function that is midpoint-convex is also convex.

Using (2), we conclude that $V(q_1, \dots, q_N) + V(r_1, \dots, r_N) \geq 2V(s_1, \dots, s_N)$. This proves the (midpoint) convexity of V . $\square$

A.1.2 Inverse of Positive Semi-definite Matrices

For future use, we provide a definition of $[X^{-1}]_{11}$ for positive *semi-definite* matrices X . When X is positive definite, its eigenvalues are strictly positive, and its inverse matrix is defined as usual. In general, we can apply the Spectral Theorem to write $X = UDU'$, where U is a $K \times K$ orthogonal matrix whose columns are eigenvectors of X , and $D = \text{diag}(d_1, \dots, d_K)$ is a diagonal matrix consisting of non-negative eigenvalues. When these eigenvalues are strictly positive, we have

$$X^{-1} = (UDU')^{-1} = UD^{-1}U' = \sum_{j=1}^K \frac{1}{d_j} \cdot [u_j u_j']$$

where u_j is the j -th column vector of U . In this case

$$[X^{-1}]_{11} = e_1' X^{-1} e_1 = \sum_{j=1}^K \frac{(\langle u_j, e_1 \rangle)^2}{d_j} \quad (3)$$

is well-defined. Even if some d_j are zero, we can still use the RHS above to define $[X^{-1}]_{11}$, applying the convention that $\frac{0}{0} = 0$ and $\frac{z}{0} = \infty$ for any $z > 0$. Note that by this definition,

$$[X^{-1}]_{11} = \lim_{\varepsilon \rightarrow 0_+} \left(\sum_{j=1}^K \frac{(\langle u_j, e_1 \rangle)^2}{d_j + \varepsilon} \right) = \lim_{\varepsilon \rightarrow 0_+} [(X + \varepsilon I_K)^{-1}]_{11},$$

since the matrix $X + \varepsilon I_K$ has the same set of eigenvectors as X (with eigenvalues increased by ε). Hence our definition of $[X^{-1}]_{11}$ is a continuous extension of the usual definition to positive semi-definite matrices.

A.1.3 Asymptotic Posterior Variance

We can approximate the posterior variance as a function of the frequencies with which each signal is observed. Specifically, as mentioned in Section 7, we can define

$$V^*(\lambda) := \lim_{t \rightarrow \infty} t \cdot V(\lambda t)$$

for any $\lambda \in \mathbb{R}_+^N$. The following result shows V^* to be well-defined and computes its value:

Lemma 3. *Let $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_N)$. Then⁴⁰*

$$V^*(\lambda) = [(C' \Lambda C)^{-1}]_{11} \quad (4)$$

The value of $[(C' \Lambda C)^{-1}]_{11}$ is well-defined, see (3).

⁴⁰Note that $C' \Lambda C$ is the Fisher Information Matrix when signals are observed according to frequencies λ . So this lemma can also be seen as an application of the Bayesian Central Limit Theorem.

Proof. Recall that $V(q_1, \dots, q_N) = [((\Sigma^0)^{-1} + C'QC)^{-1}]_{11}$ with $Q = \text{diag}(q_1, \dots, q_N)$. Thus

$$t \cdot V(\lambda_1 t, \dots, \lambda_N t) = \left[\left(\frac{1}{t} (\Sigma^0)^{-1} + C' \Lambda C \right)^{-1} \right]_{11}.$$

Hence the lemma follows from the continuity of $[X^{-1}]_{11}$ in the matrix X . $\square$

A.2 Key Object ϕ

We now define an object that will play a central role in the proofs. For each set of signals $\mathcal{S}$, consider writing the first coordinate vector $e_1 \in \mathbb{R}^K$ (corresponding to the payoff-relevant state ω) as a linear combination of signals in $\mathcal{S}$:

$$e_1 = \sum_{i \in \mathcal{S}} \beta_i^{\mathcal{S}} \cdot c_i.$$

Definition 4. $\phi(\mathcal{S}) := \min_{\beta} \sum_{i \in \mathcal{S}} |\beta_i^{\mathcal{S}}|$.

That is, $\phi(\mathcal{S})$ measures the size of the “smallest” (in the l_1 norm) linear combination of the signals in $\mathcal{S}$ to produce an unbiased estimate of the payoff-relevant state. In case ω is not spanned by $\mathcal{S}$, this definition sets $\phi(\mathcal{S}) = \infty$.

When $\mathcal{S}$ *minimally spans* ω (so that no subset spans), the coefficients $\beta_i^{\mathcal{S}}$ are unique and nonzero. In this case $\phi(\mathcal{S})$ is easy to compute. In general, we have the following characterization:

Lemma 4. *For any set $\mathcal{S}$ that spans ω , $\phi(\mathcal{S}) = \min_{\mathcal{T} \subset \mathcal{S}} \phi(\mathcal{T})$ where the minimum is over subsets $\mathcal{T}$ that “minimally span” ω .*

This lemma is a standard result in linear programming, so we omit the proof.

As a corollary, $\phi([N]) = \phi(\mathcal{S}^*)$ where $\mathcal{S}^*$ is the set of signals that minimize ϕ among all sets that minimally span ω . The following proposition, which generalizes Claim 1 in the main text, makes clear that this set $\mathcal{S}^*$ also has the greatest informational value.

Proposition 6. *For any set of signals $\mathcal{S}$, $\text{val}(\mathcal{S}) = \frac{1}{\phi(\mathcal{S})^2}$.*

In what follows (before Proposition 6 is proved), we will abuse definition and let $\mathcal{S}^*$ denote the minimal spanning set of signals that minimizes ϕ . Accordingly, λ^* denotes the frequency vector supported on $\mathcal{S}^*$ that is proportional to $|\beta^{\mathcal{S}^*}|$. Once we prove Proposition 6 and Proposition 1, it will follow that $\mathcal{S}^*$ is exactly the best complementary set defined in the main text, and there will be no confusion.

The proof of Proposition 6 uses the following three lemmata:

Lemma 5. *λ^* is the unique minimizer of $V^*(\lambda)$ as λ varies in Δ^{N-1} .*

Lemma 6. $V^*(\lambda^*) = \phi(\mathcal{S}^*)^2$.

Lemma 7. *Suppose Lemma 5 and Lemma 6 hold. Then $\text{val}([N]) = \frac{1}{\phi(\mathcal{S}^*)^2}$.*

To see why these lemmata imply Proposition 6, recall that Lemma 4 gives $\phi(\mathcal{S}^*) = \phi([N])$. So Lemma 7 implies

$$\text{val}([N]) = \frac{1}{\phi([N])^2}.$$

More generally, if we take any set of signals $\mathcal{S}$ that span ω as the set of *all* available signals “[N]”, then the same analysis yields

$$\text{val}(\mathcal{S}) = \frac{1}{\phi(\mathcal{S})^2}.$$

This proves Proposition 6 whenever $\mathcal{S}$ spans ω . But in case $\mathcal{S}$ does not span ω , the posterior variance of ω is bounded away from zero when agents are constrained to observe from $\mathcal{S}$. Thus the posterior precision $\tau(q^t)$ is bounded above and $\text{val}(\mathcal{S}) = \limsup_{t \rightarrow \infty} \frac{\tau(q^t) - \tau_0}{t} = 0$, which is also equal to $\frac{1}{\phi(\mathcal{S})^2}$ since in this case $\phi(\mathcal{S}) = \infty$ by definition.

A.2.1 Proof of Proposition 1

We can use Proposition 6 to show Proposition 1. Take any set of signals $\mathcal{S}$. If these signals do not span ω , then Proposition 6 implies $\text{val}(\mathcal{S}) = 0$ and $\mathcal{S}$ is not complementary by Definition 2. If a *proper* subset of $\mathcal{S}$ spans ω , then Proposition 6 together with Lemma 4 implies that the informational value of $\mathcal{S}$ is equal to the highest value among its subsets that minimally span ω . Let $\mathcal{S}'$ denote this subset that achieves this highest value. For this $\mathcal{S}'$ the inequality in Definition 2 is violated, and $\mathcal{S}$ is again not complementary.

Finally, suppose $\mathcal{S}$ itself minimally spans ω . In this case any nonempty proper subset of $\mathcal{S}$ does not span ω and have zero informational value, whereas $\mathcal{S}$ has positive value. So Definition 2 is satisfied and such sets $\mathcal{S}$ are complementary, as described in Proposition 1.

A.2.2 Proof of Lemma 6

It remains to establish Lemmata 5-7. Here we prove Lemma 6; that is, $V^*(\lambda^*) = \phi(\mathcal{S}^*)^2$. This proof will illustrate why the ϕ function (i.e., the ℓ_1 norm of β) plays an important role.

Without loss of generality we assume $\mathcal{S}^* = \{1, \dots, k\}$, which minimally span ω (so $k \leq K$). For $1 \leq i \leq k$, define a “transformed state” $\tilde{\theta}_i = \langle c_i, \theta \rangle$. Then the signals in $\mathcal{S}^*$ are individual $\tilde{\theta}_i$ plus standard normal noise. The payoff-relevant state ω can be written as

$$\omega = \sum_{1 \leq i \leq k} \beta_i \cdot \tilde{\theta}_i.$$

Since we are currently interested in the value of $V^*(\lambda^*)$ and λ^* is supported on $\mathcal{S}^*$, only the k signals in $\mathcal{S}^*$ matter. Thus we can work with this transformed model and focus on the beliefs about the transformed states. Note that the prior covariance matrix Σ^0 of the original state vector $\theta \in \mathbb{R}^K$ induces a prior covariance matrix $\tilde{\Sigma}^0$ of the transformed state vector $\tilde{\theta} \in \mathbb{R}^k$. $\tilde{\Sigma}^0$ has full rank because Σ^0 does and the signal coefficient vectors in $\mathcal{S}$ are linearly independent.

Working in this transformed model, we have that the posterior covariance matrix of $\tilde{\theta}$ is given by

$$\left[(\tilde{\Sigma}^0)^{-1} + \text{diag}(q_1, \dots, q_k) \right]^{-1},$$

where q_i is the number of observations of signal i . Thus the posterior variance of ω is

$$V(q_1, \dots, q_k, 0, \dots, 0) = \beta' \cdot \left[(\tilde{\Sigma}^0)^{-1} + \text{diag}(q_1, \dots, q_k) \right]^{-1} \cdot \beta.$$

It follows that for any frequency vector λ supported on $\mathcal{S}$,

$$\begin{aligned} V^*(\lambda) &= \lim_{t \rightarrow \infty} t \cdot V(\lambda t) = \lim_{t \rightarrow \infty} t \cdot \beta' \cdot \left[(\tilde{\Sigma}^0)^{-1} + \text{diag}(\lambda_1 t, \dots, \lambda_k t) \right]^{-1} \cdot \beta \\ &= \lim_{t \rightarrow \infty} \beta' \cdot \left[(\tilde{\Sigma}^0)^{-1}/t + \text{diag}(\lambda_1, \dots, \lambda_k) \right]^{-1} \cdot \beta \\ &= \beta' \cdot [\text{diag}(\lambda_1, \dots, \lambda_k)]^{-1} \cdot \beta \\ &= \sum_{i=1}^k \frac{\beta_i^2}{\lambda_i}. \end{aligned}$$

By the Cauchy-Schwartz inequality, whenever $\lambda_1, \dots, \lambda_k$ are non-negative and sum to 1, it holds that

$$\sum_{i=1}^k \frac{\beta_i^2}{\lambda_i} \geq \left(\sum_{i=1}^k |\beta_i| \right)^2 = \phi(\mathcal{S}^*)^2.$$

Moreover, equality holds if and only if each λ_i is proportional to $|\beta_i|$; that is, when $\lambda = \lambda^*$. We thus deduce that

$$V^*(\lambda^*) = \phi(\mathcal{S}^*)^2,$$

and that λ^* uniquely minimizes the value of V^* when the frequency vector is restricted to be supported on $\mathcal{S}^*$. Later we will prove Lemma 5, which shows that λ^* remains the unique minimizer without the restriction.

A.2.3 Proof of Lemma 7 and Proposition 2 Part (a)

We next prove Lemma 7. Intuitively, maximizing average increase in precision is equivalent to minimizing asymptotic posterior variance, leading to the relation $\text{val}[N] = \frac{1}{\min_{\lambda} V^*(\lambda)}$, which is in turn equal to $\frac{1}{\phi(\mathcal{S}^*)^2}$ by Lemma 5 and 6.

Toward Lemma 7, we first show $\text{val}([N]) \geq \frac{1}{\phi(\mathcal{S}^*)^2}$. By Definition 1, $\text{val}([N])$ is the maximal average increase in the precision about ω given all of the available signals. Choose a sequence of count vectors q^t such that $\lim_{t \rightarrow \infty} \frac{q^t}{t} = \lambda^*$, then by definition of the function V^* and by Lemma 6,

$$\lim_{t \rightarrow \infty} t \cdot V(q^t) = V^*(\lambda^*) = \phi(\mathcal{S}^*)^2.$$

Thus $\tau(q^t) = \frac{1}{V(q^t)} = \frac{(1+o(1))t}{\phi(\mathcal{S}^*)^2}$. It follows that

$$\text{val}([N]) \geq \limsup_{t \rightarrow \infty} \frac{\tau(q^t) - \tau_0}{t} = \frac{1}{\phi(\mathcal{S}^*)^2}.$$

In the opposite direction, take *any* sequence q^t with $\limsup_t \frac{\tau(q^t) - \tau_0}{t} = \text{val}([N])$. Since τ_0 is a constant, we equivalently have $\limsup_t \frac{\tau(q^t)}{t} = \text{val}([N])$, which gives

$$\liminf_{t \rightarrow \infty} t \cdot V(q^t) = \frac{1}{\text{val}([N])}.$$

Passing to a subsequence if necessary, we may assume the frequency vector $\lambda := \lim_{t \rightarrow \infty} \frac{q^t}{t}$ exists. Then by definition of V^* , the LHS of the above display is simply $V^*(\lambda)$. We therefore deduce $\text{val}([N]) = \frac{1}{V^*(\lambda)}$ for some $\lambda \in \Delta^{N-1}$. Since λ^* minimizes V^* , we conclude that

$$\text{val}([N]) = \frac{1}{V^*(\lambda)} \leq \frac{1}{V^*(\lambda^*)} = \frac{1}{\phi(\mathcal{S}^*)^2}.$$

This proves Lemma 7.

The second half of the above analysis additionally proves part (a) of Proposition 2. Indeed, the inequality in $\text{val}([N]) = \frac{1}{V^*(\lambda)} \leq \frac{1}{V^*(\lambda^*)}$ holds equal only if $\lambda = \lambda^*$, since λ^* is the unique minimizer of V^* by Lemma 5.

A.3 Proof of Lemma 5

A.3.1 Case 1: $|\mathcal{S}^*| = K$

In this subsection, we prove that λ^* is the unique minimizer of V^* whenever the set $\mathcal{S}^*$ contains exactly K signals. Later on we will prove the same result even when $|\mathcal{S}^*| < K$, but that proof will require additional techniques.

First, we assume $\mathcal{S}^* = \{1, \dots, K\}$ and let C^* be the $K \times K$ submatrix of C corresponding to the first K signals. Replacing c_i with $-c_i$ if necessary, we can assume $[(C^*)^{-1}]_{1i}$ is positive for $1 \leq i \leq K$. The following technical lemma is key to the argument:

Lemma 8. *Suppose $\mathcal{S}^* = \{1, \dots, K\}$ uniquely minimizes ϕ . Define C^* as above and further suppose $[(C^*)^{-1}]_{1i}$ is positive for $1 \leq i \leq K$. Then for any signal $j > K$, if we write $c_j = \sum_{i=1}^K \alpha_i \cdot c_i$ (which is a unique representation), then $|\sum_{i=1}^K \alpha_i| < 1$.*

Proof of Lemma 8. By assumption, we have the vector identity

$$e_1 = \sum_{i=1}^K \beta_i \cdot c_i \quad \text{with } \beta_i = [(C^*)^{-1}]_{1i} > 0.$$

Suppose for contradiction that $\sum_{i=1}^K \alpha_i \geq 1$ (the opposite case where the sum is ≤ -1 can be similarly treated). Then some α_i must be positive. Without loss of generality, we assume $\frac{\alpha_1}{\beta_1}$ is the largest among such ratios. Then $\alpha_1 > 0$ and

$$e_1 = \sum_{i=1}^K \beta_i \cdot c_i = \left(\sum_{i=2}^K \left(\beta_i - \frac{\beta_1}{\alpha_1} \cdot \alpha_i \right) \cdot c_i \right) + \frac{\beta_1}{\alpha_1} \cdot \left(\sum_{i=1}^K \alpha_i \cdot c_i \right)$$

This represents e_1 as a linear combination of the vectors $c_2, \dots, c_K$ and c_j , with coefficients $\beta_2 - \frac{\beta_1}{\alpha_1} \cdot \alpha_2, \dots, \beta_K - \frac{\beta_1}{\alpha_1} \cdot \alpha_K$ and $\frac{\beta_1}{\alpha_1}$. Note that these coefficients are non-negative: For each $2 \leq i \leq K$, $\beta_i - \frac{\beta_1}{\alpha_1} \cdot \alpha_i$ is clearly positive if $\alpha_i \leq 0$ (since $\beta_i > 0$). And if $\alpha_i > 0$, $\beta_i - \frac{\beta_1}{\alpha_1} \cdot \alpha_i$ is again non-negative by the assumption that $\frac{\alpha_i}{\beta_i} \leq \frac{\alpha_1}{\beta_1}$.

By definition, $\phi(\{2, \dots, K, j\})$ is the sum of the absolute value of these coefficients. This sum is

$$\sum_{i=2}^K \left(\beta_i - \frac{\beta_1}{\alpha_1} \cdot \alpha_i \right) + \frac{\beta_1}{\alpha_1} = \sum_{i=1}^K \beta_i + \frac{\beta_1}{\alpha_1} \cdot \left(1 - \sum_{i=1}^K \alpha_i \right) \leq \sum_{i=1}^K \beta_i.$$

But then $\phi(\{2, \dots, K, j\}) \leq \phi(\{1, 2, \dots, K\})$, contradicting the unique minimality of $\phi(\mathcal{S}^*)$. Hence the lemma must be true. $\square$

Proof of Lemma 5 using Lemma 8. Since $V(q_1, \dots, q_N)$ is convex in its arguments, $V^*(\lambda) = \lim_{t \rightarrow \infty} t \cdot V(\lambda_1 t, \dots, \lambda_N t)$ is also convex in λ . To show λ^* uniquely minimizes V^* , we only need to show λ^* is a *local minimum*. In other words, it suffices to show $V^*(\lambda^*) < V^*(\lambda)$ for any λ that belongs to an ε -neighborhood of λ^* . By definition, $\mathcal{S}^*$ minimally spans ω and so its signals are linearly independent. Under the additional assumption that $\mathcal{S}^*$ has size K , we deduce that its signals span the entire space $\mathbb{R}^K$. From this it follows that the $K \times K$ matrix $C' \Lambda^* C$ is positive definite, and by (4) the function V^* is differentiable near λ^* .

We claim that the partial derivatives of V^* satisfy the following inequality:

$$\partial_K V^*(\lambda^*) < \partial_j V^*(\lambda^*) \leq 0, \forall j > K. \quad (*)$$

Once this is proved, we will have, for λ close to λ^* ,

$$V^*(\lambda_1, \dots, \lambda_K, \lambda_{K+1}, \dots, \lambda_N) \geq V^* \left(\lambda_1, \dots, \lambda_{K-1}, \sum_{k=K}^N \lambda_k, 0, \dots, 0 \right) \geq V^*(\lambda^*). \quad (5)$$

The first inequality is based on (*) and differentiability of V^* , while the second inequality is because λ^* uniquely minimizes V^* when restricting to the first K signals.⁴¹ Moreover, when $\lambda \neq \lambda^*$, one of these inequalities is strict so that $V^*(\lambda) > V^*(\lambda^*)$ holds strictly.

To prove (*), we recall that

$$V^*(\lambda) = e'_1 (C' \Lambda C)^{-1} e_1.$$

Since $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_N)$, its derivative is $\partial_i \Lambda = \Delta_{ii}$, which is an $N \times N$ matrix whose (i, i) -th entry is 1 with all other entries equal to zero. Using properties of matrix derivatives, we obtain

$$\partial_i V^*(\lambda) = -e'_1 (C' \Lambda C)^{-1} C' \Delta_{ii} C (C' \Lambda C)^{-1} e_1.$$

As the i -th row vector of C is c'_i , $C' \Delta_{ii} C$ is the $K \times K$ matrix $c_i c'_i$. The above simplifies to

$$\partial_i V^*(\lambda) = -[e'_1 (C' \Lambda C)^{-1} c_i]^2.$$

⁴¹See the proof of Lemma 6 before.

At $\lambda = \lambda^*$, the matrix $C' \Lambda C$ further simplifies to $(C^*)' \cdot \text{diag}(\lambda_1^*, \dots, \lambda_K^*) \cdot (C^*)$, which is a product of $K \times K$ invertible matrices. We thus deduce that

$$\partial_i V^*(\lambda^*) = - \left[e'_1 \cdot (C^*)^{-1} \cdot \text{diag} \left(\frac{1}{\lambda_1^*}, \dots, \frac{1}{\lambda_K^*} \right) \cdot ((C^*)')^{-1} \cdot c_i \right]^2.$$

Crucially, note that the term in the brackets is a linear function of c_i . To ease notation, we write $v' = e'_1 \cdot (C^*)^{-1} \cdot \text{diag} \left(\frac{1}{\lambda_1^*}, \dots, \frac{1}{\lambda_K^*} \right) \cdot ((C^*)')^{-1}$ and $\gamma_i = \langle v, c_i \rangle$. Then

$$\partial_i V^*(\lambda^*) = -\gamma_i^2, \quad 1 \leq i \leq N. \quad (6)$$

For $1 \leq i \leq K$, $((C^*)')^{-1} \cdot c_i$ is just e_i . Thus, using the assumption $[(C^*)^{-1}]_{1i} > 0, \forall i$, we have

$$\gamma_i = e'_1 \cdot (C^*)^{-1} \cdot \text{diag} \left(\frac{1}{\lambda_1^*}, \dots, \frac{1}{\lambda_K^*} \right) \cdot e_i = \frac{[(C^*)^{-1}]_{1i}}{\lambda_i^*} = \frac{\beta_i}{\lambda_i^*} = \beta_1 + \dots + \beta_K = \phi(\mathcal{S}^*). \quad (7)$$

On the other hand, choosing any signal $j > K$, we can uniquely write the vector c_j as a linear combination of $c_1, \dots, c_K$. By Lemma 8,

$$\gamma_j = \langle v, c_j \rangle = \sum_{i=1}^K \alpha_i \cdot \langle v, c_i \rangle = \sum_{i=1}^K \alpha_i \cdot \gamma_i = \phi(\mathcal{S}^*) \cdot \sum_{i=1}^K \alpha_i, \quad (8)$$

where the last equality uses (7). Since $|\sum_{i=1}^K \alpha_i| < 1$, the absolute value of γ_j is strictly smaller than the absolute value of γ_K for any $j > K$. This together with (6) proves the desired inequality (*), and Lemma 5 follows. □

A.3.2 A Perturbation Argument

We have shown that when ϕ is uniquely minimized by a set $\mathcal{S}^*$ containing exactly K signals,

$$\min_{\lambda \in \Delta^{N-1}} V^*(\lambda) = V^*(\lambda^*) = \phi(\mathcal{S}^*)^2 = \phi([N])^2.$$

We now use a perturbation argument to show this equality holds more generally.

Lemma 9. *For any coefficient matrix C ,*

$$\min_{\lambda \in \Delta^{N-1}} V^*(\lambda) = \phi([N])^2. \quad (9)$$

Proof. In general, the set $\mathcal{S}^*$ that minimizes ϕ may not be unique or involve exactly K signals. However, we always have (by Lemma 4 and Lemma 6)

$$\min_{\lambda \in \Delta^{N-1}} V^*(\lambda) \leq V^*(\lambda^*) = \phi(\mathcal{S}^*)^2 = \phi([N])^2.$$

It remains to prove $V^*(\lambda) \geq \phi([N])^2$ for every $\lambda \in \Delta^{N-1}$. Below we fix λ . By Lemma 3, we need to show $[(C' \Lambda C)^{-1}]_{11} \geq \phi([N])^2$.

Note that we already proved this inequality for *generic* coefficient matrices C : specifically, when ϕ is uniquely minimized by a set of K signals, Lemma 5 holds and we have $V^*(\lambda) \geq V^*(\lambda^*) = \phi(\mathcal{S}^*)^2 = \phi([N])^2$. But even if C is “non-generic”, we can approximate it by a sequence of generic matrices C_m .⁴² Along this sequence, we have

$$[(C'_m \Lambda C_m)^{-1}]_{11} \geq \phi_m([N])^2$$

where ϕ_m is the analogue of ϕ for the coefficient matrix C_m .

As $m \rightarrow \infty$, the LHS above approaches $[(C' \Lambda C)^{-1}]_{11}$. We will show that on the RHS

$$\limsup_{m \rightarrow \infty} \phi_m([N]) \geq \phi([N]),$$

which then implies $[(C' \Lambda C)^{-1}]_{11} \geq \phi([N])^2$ and the lemma. Indeed, suppose $e_1 = \sum_i \beta_i^{(m)} \cdot c_i^{(m)}$ along the convergent sequence, then $e_1 = \sum_i \beta_i \cdot c_i$ for any limit point β of $\beta^{(m)}$. Using the definition of ϕ , this enables us to conclude $\liminf_{m \rightarrow \infty} \phi_m([N]) \geq \phi([N])$, which is more than sufficient. $\square$

A.3.3 Case 2: $|\mathcal{S}^*| < K$

We now consider the case where $\mathcal{S}^* = \{1, \dots, k\}$ with $k < K$. We will show that λ^* is still the unique minimizer of $V^*(\cdot)$. Since $V^*(\lambda^*) = \phi(\mathcal{S}^*)^2 = \phi([N])^2$, we know from Lemma 9 that λ^* does minimize V^* . It remains to show λ^* is the *unique* minimizer.

To do this, we will consider a perturbed informational environment in which signals $k+1, \dots, N$ are made slightly more precise. Specifically, let $\eta > 0$ be a small positive number. Consider an alternative signal coefficient matrix $\tilde{C}$ with $\tilde{c}_i = c_i$ for $i \leq k$ and $\tilde{c}_i = (1 + \eta)c_i$ for $i > k$. Let $\tilde{\phi}(\mathcal{S})$ be the analogue of ϕ for this alternative environment. It is clear that $\tilde{\phi}(\mathcal{S}^*) = \phi(\mathcal{S}^*)$, while $\tilde{\phi}(\mathcal{S})$ is slightly smaller than $\phi(\mathcal{S})$ for $\mathcal{S} \neq \mathcal{S}^*$. Thus with sufficiently small η , the set $\mathcal{S}^*$ remains the unique minimizer of ϕ (among sets that minimally span ω) in this perturbed environment, and the definition of λ^* is also maintained.

Let $\tilde{V}^*$ be the perturbed asymptotic posterior variance function, then our previous analysis shows that $\tilde{V}^*$ has minimum value $\phi(\mathcal{S}^*)^2$ on the simplex. Taking advantage of the connection between V^* and $\tilde{V}^*$, we thus have

$$\begin{aligned} V^*(\lambda_1, \dots, \lambda_N) &= \tilde{V}^* \left(\lambda_1, \dots, \lambda_k, \frac{\lambda_{k+1}}{(1+\eta)^2}, \dots, \frac{\lambda_N}{(1+\eta)^2} \right) \\ &\geq \frac{\phi(\mathcal{S}^*)^2}{\sum_{i \leq k} \lambda_i + \frac{1}{(1+\eta)^2} \sum_{i > k} \lambda_i}. \end{aligned}$$

The equality uses (4) and $C' \Lambda C = \sum_i \lambda_i c_i c_i' = \sum_{i \leq k} \lambda_i c_i c_i' + \sum_{i > k} \frac{\lambda_i}{(1+\eta)^2} \tilde{c}_i \tilde{c}_i'$. The inequality follows from the homogeneity of $\tilde{V}^*$.

⁴²First, we may add repetitive signals to ensure $N \geq K$. This does not affect the value of $\min_{\lambda} V^*(\lambda)$ or $\phi([N])$. Whenever $N \geq K$, it is generically true that every set that minimally spans ω contains exactly K signals. Moreover, the equality $\phi(\mathcal{S}) = \phi(\tilde{\mathcal{S}})$ for $\mathcal{S} \neq \tilde{\mathcal{S}}$ induces a non-trivial polynomial equation over the entries in C . This means we can always find C_m close to C such that for each coefficient matrix C_m , different subsets $\mathcal{S}$ of size K attain different values of ϕ , so that ϕ is uniquely minimized.

The above display implies that any frequency vector λ ,

$$V^*(\lambda) \geq \frac{\phi(\mathcal{S}^*)^2}{1 - \frac{2\eta + \eta^2}{(1+\eta)^2} \sum_{i>k} \lambda_i} \geq \frac{\phi(\mathcal{S}^*)^2}{1 - \eta \sum_{i>k} \lambda_i} \quad \text{for some } \eta > 0. \quad (10)$$

Hence $V^*(\lambda) > \phi(\mathcal{S}^*)^2 = V^*(\lambda^*)$ whenever λ puts positive weight outside of $\mathcal{S}^*$. But as shown before, $V^*(\lambda) > V^*(\lambda^*)$ also holds when λ is supported on $\mathcal{S}^*$ and different from λ^* .

We conclude that λ^* is the unique minimizer of V^* over the whole simplex. This proves Lemma 5, which also completes the proof of Proposition 6 and Proposition 1 as we showed before.

A.4 Proof of Theorem 1

We show here that Theorem 1 follows from Theorem 2, which we prove in the next appendix. Indeed, as we explained in Lemma 1, the set $\mathcal{S}$ is strongly complementary if and only if it has the greatest informational value among signals *in its subspace*. Under the assumption of Theorem 1 part (a), every complementary set $\mathcal{S}$ of size less than K is strongly complementary because no other signal is in that subspace (otherwise there would be linear dependence). Thus by the “only if” part of Theorem 2, there exists prior beliefs that lead to exclusive observation of signals from $\mathcal{S}$. As for Theorem 1 part (b), the assumption implies that the subspace spanned by every complementary set is the whole space, and thus the *only* strongly complementary set is $\mathcal{S}^*$. Hence the “if” part of Theorem 2 implies that long-run efficiency is guaranteed.

A.5 Proof of Theorem 2: “Only If” Part

Let signals $1, \dots, k$ (with $k \leq K$) be a strongly complementary set; by Lemma 1 in the main text, these signals are best in their subspace. We will demonstrate an open set of prior beliefs given which *all agents* observe these k signals. Since these signals are complementary, Proposition 1 implies they must be linearly independent. Thus we can consider linearly transformed states $\tilde{\theta}_1, \dots, \tilde{\theta}_K$ such that these k signals are simply $\tilde{\theta}_1, \dots, \tilde{\theta}_k$ plus standard Gaussian noise. This linear transformation is invertible, so any prior over the original states is bijectively mapped to a prior over the transformed states. Thus it is without loss to work with the transformed model and look for prior beliefs over the transformed states.

The payoff-relevant state ω becomes a linear combination $\lambda_1^* \tilde{\theta}_1 + \dots + \lambda_k^* \tilde{\theta}_k$ (up to a scalar multiple). Since the first k signals are best in their subspace, Lemma 8 before implies that any other signal belonging to this subspace can be written as $\sum_{i=1}^k \alpha_i \tilde{\theta}_i + \mathcal{N}(0, 1)$ with $|\sum_{i=1}^k \alpha_i| < 1$. On the other hand, if a signal does not belong to this subspace, it must take the form of $\sum_{i=1}^K \beta_i \tilde{\theta}_i + \mathcal{N}(0, 1)$ with $\beta_{k+1}, \dots, \beta_K$ *not all equal to zero*.

Now consider any prior belief with precision matrix P ; the inverse of P is the prior covariance matrix (in terms of the transformed states). Suppose ε is a very small positive number, and P satisfies the following conditions:

1. For $1 \leq i \leq k$, $P_{ii} \geq \frac{1}{\varepsilon^2}$;
2. For $1 \leq i \neq j \leq k$, $\frac{P_{ii}}{\lambda_i^*} \leq (1 + \varepsilon) \cdot \frac{P_{jj}}{\lambda_j^*}$;
3. For $k + 1 \leq i \leq K$, $P_{ii} \in [\varepsilon, 2\varepsilon]$;
4. For $1 \leq i \neq j \leq K$, $|P_{ij}| \leq \varepsilon^2$.

It is clear that any such P is positive definite, since on each row the diagonal entry has dominant size.⁴³ Moreover, P contains an open subset. Below we show that given any such prior, the myopic signal choice is among the first k signals, and that the posterior precision matrix also satisfies the same four conditions. As such, *all* agents would choose from the first k signals.

Let $V = P^{-1}$ be the prior covariance matrix. Applying Cramer's rule for the matrix inverse, the above conditions on P imply the following conditions on V :

1. For $1 \leq i \leq k$, $V_{ii} \leq 2\varepsilon^2$;
2. For $1 \leq i \neq j \leq k$, $V_{ii}\lambda_i^* \leq (1 + L\varepsilon) \cdot V_{jj}\lambda_j^*$;
3. For $k + 1 \leq i \leq K$, $V_{ii} \in [\frac{1}{4\varepsilon}, \frac{2}{\varepsilon}]$;
4. For $1 \leq i \neq j \leq K$, $|V_{ij}| \leq L\varepsilon \cdot V_{ii}$.

Here L is a constant depending only on K (but not on ε). For example, the last condition is equivalent to $\det(P_{-ij}) \leq L\varepsilon \cdot \det(P_{-ii})$. This is proved by expanding both determinants into multilinear sums, and using the fact that on each row of P the off-diagonal entries are at most ε -fraction of the diagonal entry.

Given this matrix V , the variance reduction of $\omega = \sum_{i=1}^k \lambda_i^* \tilde{\theta}_i$ by any signal $\sum_{i=1}^k \alpha_i \tilde{\theta}_i + \mathcal{N}(0, 1)$ can be computed as

$$\frac{(\sum_{i,j=1}^k \alpha_i \lambda_j^* V_{ij})^2}{1 + \sum_{i,j=1}^k \alpha_i \alpha_j V_{ij}},$$

where the denominator is the variance of the signal and the numerator is the covariance between the signal and ω . By the first and last conditions on V , the denominator here is $1 + O(\varepsilon^2)$. By the second and last condition, the numerator is

$$\left(\left(\sum_{i=1}^k \alpha_i + O(\varepsilon) \right) \cdot \lambda_1^* V_{11} \right)^2.$$

Since $|\sum_{i=1}^k \alpha_i| < 1$, we deduce that any other signal belonging to the subspace of the first k signals is myopically worse than signal 1, whose variance reduction is $\frac{(\lambda_1^* V_{11})^2}{1 + V_{11}}$.

⁴³Suppose P is a symmetric matrix s.t. $P_{ii} > \sum_{j \neq i} P_{ij}$, then for any vector $x \in \mathbb{R}^K$, it holds that

$$x'Px = \sum_{i=1}^K P_{ii}x_i^2 + \sum_{1 \leq i < j \leq K} 2P_{ij}x_i x_j \geq \sum_{i=1}^K P_{ii}x_i^2 - \sum_{1 \leq i < j \leq K} P_{ij}(x_i^2 + x_j^2) = \sum_{i=1}^K (P_{ii} - \sum_{j \neq i} P_{ij})x_i^2 \geq 0,$$

with equality only if x is the zero vector. This shows P is positive-definite.

Meanwhile, take any signal outside of the subspace. The variance reduction by such a signal $\sum_{i=1}^K \beta_i \tilde{\theta}_i + \mathcal{N}(0, 1)$ is

$$\frac{(\sum_{i=1}^K \sum_{j=1}^k \beta_i \lambda_j^* V_{ij})^2}{1 + \sum_{i,j=1}^K \beta_i \beta_j V_{ij}}$$

By the second and last condition on V , the numerator here is $O((\lambda_1^* V_{11})^2)$. If we can show that the denominator is very large, then such a signal would also be myopically worse than signal 1. Indeed, since $V_{ij} = O(\varepsilon^2)$ whenever $i \leq k$ or $j \leq k$, it is sufficient to show $\sum_{i,j>k} \beta_i \beta_j V_{ij}$ is large. This holds by the last two conditions on V and the assumption that $\beta_{k+1}, \dots, \beta_K$ are not all zero.⁴⁴

Hence, we have shown that given any prior precision matrix P satisfying the above conditions, the myopic signal choice is among the first k signals. It remains to check the resulting posterior precision matrix $\hat{P}$ also satisfies those four conditions. If the signal acquired is signal i ($1 \leq i \leq k$), then $\hat{P} = P + \Delta_{ii}$. Therefore we only need to show the second condition holds for $\hat{P}$; that is, $\frac{P_{ii}+1}{\lambda_i^*} \leq (1 + \varepsilon) \cdot \frac{P_{jj}}{\lambda_j^*}$ for each $1 \leq j \leq k$. To this end, we note that since signal i is myopically best given V , the following must hold:

$$\frac{(\lambda_i^* V_{ii})^2}{1 + V_{ii}} \geq \frac{(\lambda_j^* V_{jj})^2}{1 + V_{jj}}.$$

As $0 \leq V_{ii}, V_{jj} \leq 2\varepsilon^2$, this implies $\lambda_i^* V_{ii} \geq (1 - \varepsilon^2) \lambda_j^* V_{jj}$. Now applying Cramer's rule to $V = P^{-1}$ again, we can deduce $V_{ii} = \frac{1+O(\varepsilon^2)}{P_{ii}}$. So for ε small it holds that $\frac{P_{ii}}{\lambda_i^*} \leq (1 + \frac{\varepsilon}{2}) \cdot \frac{P_{jj}}{\lambda_j^*}$. As $P_{ii} \geq \frac{1}{\varepsilon^2}$, we also have $\frac{1}{\lambda_i^*} \leq \frac{\varepsilon}{2} \cdot \frac{P_{jj}}{\lambda_j^*}$. Adding up these two inequalities yields the second condition for $\hat{P}$ and completes the proof.

A.6 Proof of Theorem 2: “If” Part

A.6.1 Restated Version

Given any prior belief, let $\mathcal{A} \subset [N]$ be the set of all signals that are observed by infinitely many agents. Our goal is to show that $\mathcal{A}$ is strongly complementary. Toward that goal, we first show $\mathcal{A}$ spans ω .

Indeed, by definition we can find some period t after which agents *exclusively* observe signals from $\mathcal{A}$. Note that the variance reduction of any signal approaches zero as its signal count gets large. Thus, along society's signal path, the variance reduction is close to zero at sufficiently late periods. If $\mathcal{A}$ does not span ω , society's posterior variance remains bounded away from zero. Thus in the limit where each signal in $\mathcal{A}$ has infinite signal counts, there still exists some signal j outside of $\mathcal{A}$

⁴⁴Formally, we can without loss assume $\beta_K^2 V_{KK}$ is largest among $\beta_i^2 V_{ii}$ for $i > k$. Then for any $i \neq j$, the last condition implies

$$\beta_i \beta_j V_{ij} \geq -L\varepsilon \cdot \beta_i \beta_j \sqrt{V_{ii} V_{jj}} \geq -L\varepsilon \cdot \beta_K^2 V_{KK}.$$

This trivially also holds for $i = j \neq K$. Summing across all pairs $(i, j) \neq (K, K)$ yields $\sum_{i,j>k} \beta_i \beta_j V_{ij} > (1 - K^2 L\varepsilon) \beta_K^2 V_{KK}$, which must be large by the third condition on V .

whose variance reduction is strictly positive.⁴⁵ By continuity, we deduce that at any sufficiently late period, observing signal j is better than observing any signal in $\mathcal{A}$. This contradicts our assumption that later agents only observe signals in $\mathcal{A}$.

Now that $\mathcal{A}$ spans ω , we can take $\mathcal{S}$ to be the best complementary set in $\overline{\mathcal{A}}$, which is the subspace spanned by $\mathcal{A}$. By Lemma 1, $\mathcal{S}$ is strongly complementary. To prove Theorem 2 “if” part, we will show that long-run frequencies are positive precisely for the signals in $\mathcal{S}$. By ignoring the initial periods, we can assume without loss that *only signals in $\overline{\mathcal{A}}$ are available*. It thus suffices to show that whenever the signals observed infinitely often span a subspace, agents eventually focus on the best complementary set $\mathcal{S}$ in that subspace. To ease notation, we assume this subspace is the entire $\mathbb{R}^K$, and prove the following result:

Theorem 2 “If” Part Restated. *Suppose that the signals observed infinitely often span $\mathbb{R}^K$. Then society’s long-run frequency vector is λ^* .*

The next sections are devoted to the proof of this restatement.

A.6.2 Estimates of Derivatives

We introduce a few technical lemmata:

Lemma 10. *For any $q_1, \dots, q_N$, we have*

$$\left| \frac{\partial_{jj} V(q_1, \dots, q_N)}{\partial_j V(q_1, \dots, q_N)} \right| \leq \frac{2}{q_j}.$$

Proof. Recall that $V(q_1, \dots, q_N) = e_1' \cdot [(\Sigma^0)^{-1} + C'QC]^{-1} \cdot e_1$. Thus

$$\partial_j V = -e_1' \cdot [(\Sigma^0)^{-1} + C'QC]^{-1} \cdot c_j \cdot c_j' \cdot [(\Sigma^0)^{-1} + C'QC]^{-1} \cdot e_1,$$

and

$$\partial_{jj} V = 2e_1' \cdot [(\Sigma^0)^{-1} + C'QC]^{-1} \cdot c_j \cdot c_j' \cdot [(\Sigma^0)^{-1} + C'QC]^{-1} \cdot c_j \cdot c_j' \cdot [(\Sigma^0)^{-1} + C'QC]^{-1} \cdot e_1.$$

Let $\gamma_j = e_1' \cdot [(\Sigma^0)^{-1} + C'QC]^{-1} \cdot c_j$, which is a number. Then the above becomes

$$\partial_j f = -\gamma_j^2; \quad \partial_{jj} f = 2\gamma_j^2 \cdot c_j' \cdot [(\Sigma^0)^{-1} + C'QC]^{-1} \cdot c_j.$$

⁴⁵To see this, let $s_1, \dots, s_N$ denote the limit signal counts, where $s_i = \infty$ if and only if $i \in \mathcal{A}$. We need to find some signal j such that $V(s_j + 1, s_{-j}) < V(s_j, s_{-j})$. If such a signal does not exist, then all partial derivatives of V at s are zero. Since V is always differentiable (unlike V^*), this would imply that all directional derivatives of V are also zero. By the convexity of V , V must be minimized at s . However, the minimum value of V is zero because there exists a complementary set. This contradicts $V(s) > 0$.

Note that $(\Sigma^0)^{-1} + C'QC \succeq q_j \cdot c_j c_j'$ in matrix norm. Thus the number $c_j' \cdot [(\Sigma^0)^{-1} + C'QC]^{-1} \cdot c_j$ is bounded above by $\frac{1}{q_j}$.⁴⁶ This proves the lemma. $\square$

Since the second derivative is small compared to the first derivative, we deduce that the variance reduction of any *discrete* signal can be approximated by the partial derivative of f . This property is summarized in the following lemma:

Lemma 11. *For any $q_1, \dots, q_N$, we have⁴⁷*

$$V(q) - V(q_j + 1, q_{-j}) \geq \frac{q_j}{q_j + 1} |\partial_j V(q)|.$$

Proof. We will show the more general result:

$$V(q) - V(q_j + x, q_{-j}) \geq \frac{q_j x}{q_j + x} \cdot |\partial_j V(q)|, \forall x \geq 0.$$

This clearly holds at $x = 0$. Differentiating with respect to x , we only need to show

$$-\partial_j V(q_j + x, q_{-j}) \geq \frac{q_j^2}{(q_j + x)^2} |\partial_j V(q)|, \forall x \geq 0.$$

Equivalently, we need to show

$$-(q_j + x)^2 \cdot \partial_j V(q_j + x, q_{-j}) \geq -q_j^2 \cdot \partial_j V(q), \forall x \geq 0.$$

Again, this inequality holds at $x = 0$. Differentiating with respect to x , it becomes

$$-2(q_j + x) \cdot \partial_j V(q_j + x, q_{-j}) - (q_j + x)^2 \cdot \partial_{jj} V(q_j + x, q_{-j}) \geq 0.$$

This is exactly the result of Lemma 10. $\square$

A.6.3 Lower Bound on Variance Reduction

Our next result lower bounds the directional derivative of V along the “optimal” direction λ^* :

Lemma 12. *For any $q_1, \dots, q_N$, we have $|\partial_{\lambda^*} V(q)| \geq \frac{V(q)^2}{\phi(\mathcal{S}^*)^2}$.*

⁴⁶Formally, we need to show that for any $\varepsilon > 0$, the number $c_j' [c_j c_j' + \varepsilon I_K]^{-1} c_j$ is at most 1. Using the trace identify $\text{tr}(AB) = \text{tr}(BA)$, we can rewrite this number as

$$\text{tr}([c_j c_j' + \varepsilon I_K]^{-1} c_j c_j') = \text{tr}(I_K - [c_j c_j' + \varepsilon I_K]^{-1} \varepsilon I_K) = K - \varepsilon \cdot \text{tr}([c_j c_j' + \varepsilon I_K]^{-1}).$$

The matrix $c_j c_j'$ has rank 1, so $K - 1$ of its eigenvalues are zero. Thus the matrix $[c_j c_j' + \varepsilon I_K]^{-1}$ has eigenvalue $1/\varepsilon$ with multiplicity $K - 1$, and the remaining eigenvalue is positive. This implies $\varepsilon \cdot \text{tr}([c_j c_j' + \varepsilon I_K]^{-1}) > K - 1$, and then the above display yields $c_j' \cdot [(\Sigma^0)^{-1} + C'QC]^{-1} \cdot c_j < 1$ as desired.

⁴⁷Note that the convexity of V gives $V(q) - V(q_j + 1, q_{-j}) \leq |\partial_j V(q)|$. This lemma provides a converse that we need for the subsequent analysis.

Proof. To compute this directional derivative, we think of agents acquiring signals in fractional amounts, where a fraction of a signal is just the same signal with precision multiplied by that fraction. Consider an agent who draws λ_i^* realizations of each signal i . Then he essentially obtains the following signals:

$$Y_i = \langle c_i, \theta \rangle + \mathcal{N}\left(0, \frac{1}{\lambda_i^*}\right), \forall i.$$

This is equivalent to

$$\lambda_i^* Y_i = \langle \lambda_i^* c_i, \theta \rangle + \mathcal{N}(0, \lambda_i^*), \forall i.$$

Such an agent receives at least as much information as the sum of these signals:

$$\sum_i \lambda_i^* Y_i = \sum_i \langle \lambda_i^* c_i, \theta \rangle + \sum_i \mathcal{N}(0, \lambda_i^*) = \frac{\omega}{\phi(\mathcal{S}^*)} + \mathcal{N}(0, 1).$$

Hence the agent's posterior precision about ω (which is the inverse of his posterior variance V) must increase by at least $\frac{1}{\phi(\mathcal{S}^*)^2}$ along the direction λ^* . The chain rule of differentiation yields the lemma. $\square$

We can now bound the variance reduction at late periods:

Lemma 13. *Fix any $q_1, \dots, q_N$. Suppose L is a positive number such that $(\Sigma^0)^{-1} + C'QC \succeq Lc_j c_j'$ holds for each signal $j \in \mathcal{S}^*$. Then we have*

$$\min_{j \in \mathcal{S}^*} V(q_j + 1, q_{-j}) \leq V(q) - \frac{L}{L+1} \cdot \frac{V(q)^2}{\phi(\mathcal{S}^*)^2}.$$

Proof. Fix any signal $j \in \mathcal{S}^*$. Using the condition $(\Sigma^0)^{-1} + C'QC \succeq Lc_j c_j'$, we can deduce the following variant of Lemma 11:⁴⁸

$$V(q) - V(q_j + 1, q_{-j}) \geq \frac{L}{L+1} |\partial_j V(q)|.$$

Since V is always differentiable, $\partial_{\lambda^*} V(q)$ is a convex combination of the partial derivatives of V .⁴⁹ Thus $\max_{j \in \mathcal{S}^*} |\partial_j V(q)| \geq |\partial_{\lambda^*} V(q)|$. These inequalities, and Lemma 12, complete the proof. $\square$

A.6.4 Proof of the Restated Theorem 2 “If” Part

We will show $t \cdot V(m(t)) \rightarrow \phi(\mathcal{S}^*)^2$, so that society eventually approximates the optimal speed of learning. Since λ^* is the unique minimizer of V^* , this will imply the desired conclusion $\frac{m(t)}{t} \rightarrow \lambda^*$ via the second half of Proposition 2 part (a).

⁴⁸Even though we are not guaranteed $q_j \geq L$, we can modify the prior and signal counts such that the precision matrix $(\Sigma^0)^{-1} + C'QC$ is unchanged, and signal j has been observed at least L times. This is possible thanks to the condition $(\Sigma^0)^{-1} + C'QC \succeq Lc_j c_j'$. Then, applying Lemma 11 to this modified problem yields the result here.

⁴⁹While this may be a surprising contrast with V^* , the difference arises because the formula for V always involves a full-rank prior covariance matrix, whereas its asymptotic variant V^* corresponds to a flat prior.

To estimate $V(m(t))$, we note that for any fixed L , society's acquisitions $m(t)$ eventually satisfy the condition $(\Sigma^0)^{-1} + C'QC \succeq Lc_jc_j'$. This is due to our assumption that the signals observed infinitely often span $\mathbb{R}^K$, which implies that $C'QC$ becomes arbitrarily large in matrix norm. Hence, we can apply Lemma 13 to find that

$$V(m(t+1)) \leq V(m(t)) - \frac{L}{L+1} \cdot \frac{V(m(t))^2}{\phi(\mathcal{S}^*)^2}$$

for all $t \geq t_0$, where t_0 depends only on L .

We introduce the auxiliary function $g(t) = \frac{V(m(t))}{\phi(\mathcal{S}^*)^2}$. Then the above simplifies to

$$g(t+1) \leq g(t) - \frac{L}{L+1} g(t)^2.$$

Inverting both sides, we have

$$\frac{1}{g(t+1)} \geq \frac{1}{g(t)(1 - \frac{L}{L+1}g(t))} = \frac{1}{g(t)} + \frac{\frac{L}{L+1}}{1 - \frac{L}{L+1}g(t)} \geq \frac{1}{g(t)} + \frac{L}{L+1}. \quad (11)$$

This holds for all $t \geq t_0$. Thus by induction, $\frac{1}{g(t)} \geq \frac{L}{L+1}(t - t_0)$ and so $g(t) \leq \frac{L+1}{L(t-t_0)}$. Going back to the posterior variance function V , this implies

$$V(m(t)) \leq \frac{L+1}{L} \cdot \frac{\phi(\mathcal{S}^*)^2}{t - t_0}. \quad (12)$$

Hence, by choosing L sufficiently large and then considering large t , we find that society's speed of learning is arbitrarily close to the optimal speed $\phi(\mathcal{S}^*)^2$. This completes the proof.

We comment that the above argument leaves open the possibility that some signals outside of $\mathcal{S}^*$ are observed *infinitely often*, yet with *zero long-run frequency*. In Online Appendix B, we show this does not happen.

A.7 Proofs for Interventions (Section 9)

A.7.1 Proof of Proposition 4

Given any history of observations, an agent can always allocate his B observations as follows: He draws $\lfloor B \cdot \lambda_i^* \rfloor$ realizations of each signal i , and samples arbitrarily if there is any capacity remaining. Here $\lfloor \cdot \rfloor$ denotes the floor function.

Fix any $\varepsilon > 0$. If B is sufficiently large, then the above strategy acquires at least $(1 - \varepsilon) \cdot B \cdot \lambda_i^*$ observations of each signal i . Adapting the proof of Lemma 12, we see that the agent's posterior precision about ω must increase by $\frac{(1-\varepsilon)B}{\phi(\mathcal{S}^*)^2}$ under this strategy. Thus the same must hold for his optimal strategy, so that society's posterior precision at time t is at least $\frac{(1-\varepsilon)Bt}{\phi(\mathcal{S}^*)^2}$. This implies that average precision per signal is at least $\frac{1-\varepsilon}{\phi(\mathcal{S}^*)^2}$, which can be arbitrarily close to the optimal precision $\text{val}([N]) = \frac{1}{\phi(\mathcal{S}^*)^2}$ with appropriate choice of ε .

Since λ^* is the unique minimizer of V^* , society's long-run frequencies must be close to λ^* . In particular, with ε sufficiently small, we can ensure that each signal in $\mathcal{S}^*$ are observed with positive frequencies. The restated Theorem 2 “if” part extends to the current setting and implies that society's long-run frequency vector must be λ^* . This yields the proposition.⁵⁰

A.7.2 Proof of Proposition 5

Suppose without loss that the best complementary set $\mathcal{S}^*$ is $\{1, \dots, k\}$. By taking a linear transformation, we further assume each of the first k signals only involves ω and the first $k-1$ confounding terms $b_1, \dots, b_{k-1}$. We will show that whenever $k-1$ sufficiently precise signals are provided about each of these confounding terms, the long-run frequency vector converges to λ^* regardless of the prior.

Fix any positive real number L . Since the $k-1$ free signals are very precise, it is as if the prior precision matrix (after taking into account these free signals) satisfies

$$(\Sigma^0)^{-1} \succeq L^2 \sum_{i=2}^k \Delta_{ii}$$

where Δ_{ii} is the $K \times K$ matrix that has one at the (i, i) entry and zero otherwise. Recall also that society eventually learns ω . Thus at some late period t_0 , society's acquisitions must satisfy

$$C'QC \succeq L^2 \Delta_{11}.$$

Adding up the above two displays, we have

$$(\Sigma^0)^{-1} + C'QC \succeq L^2 \sum_{i=1}^k \Delta_{ii} \succeq Lc_j c_j', \forall 1 \leq j \leq k.$$

The last inequality uses the fact that each c_j only involves the first k coordinates.

Now this is exactly the condition we need in order to apply Lemma 13: Crucially, whether or not the condition is met for signals j outside of $\mathcal{S}^*$ does not affect the argument there. Thus we can follow the proof of the restated Theorem 2 “if” part to deduce (12). That is, for fixed L and corresponding free information, society's long-run precision per signal is at least $\frac{L}{(L+1)\phi(\mathcal{S}^*)^2}$. This can be made arbitrarily close to the optimal average precision. Identical to the previous proof, we deduce that for large L , society's long-run frequency vector must be close to λ^* . The restated Theorem 2 “if” part allows us to conclude that the frequency is exactly λ^* .

⁵⁰This proof also suggests that how small ε (and how large B) need to be depends on the distance between the optimal speed of learning and the “second-best” speed of learning from any other complementary set. Intuitively, in order to achieve long-run efficient learning, agents need to allocate B observations in the best set to approximate the optimal frequencies. If another set of signals offers a speed of learning that is only slightly worse, we will need B sufficiently large for the approximately optimal frequencies in the best set to beat this other set.

B Proofs for the Autocorrelated Model (Section 8.2)

B.1 Proof of Theorem 3

We work with the transformed model such that the signals in $\mathcal{S}$ become the first k transformed states $\tilde{\theta}_1, \dots, \tilde{\theta}_k$. The payoff-relevant state becomes a certain linear combination $w_1\tilde{\theta}_1 + \dots + w_k\tilde{\theta}_k$ with positive weights $w_1, \dots, w_k$. Choose M so that the innovations corresponding to the transformed states are independent from each other. In other words, $\tilde{M}$ (the transformed version of M) is given by $\text{diag}(\frac{x}{w_1}, \dots, \frac{x}{w_k}, y_{k+1}, \dots, y_K)$. Here x is a *small* positive number, while $y_{k+1}, \dots, y_K$ are *large* positive numbers. We further choose $\Sigma^0 = M$, which is the stable belief without learning.

With these choices, it is clear that if all agents only sample from $\mathcal{S}$, society's beliefs about the transformed states remain independent at every period. Let v_i^{t-1} denote the prior variance of $\tilde{\theta}_i^t$ at the beginning of period t (before the signal acquisition in that period). Then as long as agent t would continue to sample a signal $\tilde{\theta}_j + \mathcal{N}(0, 1)$ in $\mathcal{S}$, these prior variances would evolve as follows: $v_i^0 = \frac{x}{w_i}$ for $1 \leq i \leq k$ and $v_i^0 = y_i$ for $i > k$. And for $t \geq 1$,

$$v_i^t = \begin{cases} \alpha \cdot v_i^{t-1} + (1 - \alpha)\tilde{M}_{ii}, & \text{if } i \neq j; \\ \alpha \cdot \frac{v_i^{t-1}}{1 + v_i^{t-1}} + (1 - \alpha)\tilde{M}_{ii} & \text{if } i = j. \end{cases}$$

The particular signal j maximizes the reduction in the posterior variance of $\omega^t = \sum_{i=1}^k w_i \tilde{\theta}_i^t$. That is, $j \in \arg\max_{1 \leq i \leq k} \frac{(w_i \cdot v_i^{t-1})^2}{1 + v_i^{t-1}}$.

By induction, it is clear that $v_i^t \leq \tilde{M}_{ii}$ holds for all pairs i, t , with equality for $i > k$. Thus at the beginning of each period t , assuming that all previous agents have sampled from $\mathcal{S}$, agent t 's prior uncertainties about $\tilde{\theta}_1, \dots, \tilde{\theta}_k$ are small while his uncertainties about $\tilde{\theta}_{k+1}, \dots, \tilde{\theta}_K$ are large. As such, our previous proof for the existence of learning traps with persistent states carries over, and we deduce that agent t continues to observe from $\mathcal{S}$.

Note that for α close to 1, agents will sample each signal i with frequency close to $\frac{w_i}{w_1 + \dots + w_k}$. It follows that the prior variances v_i^t approximately satisfy the following fixed-point equation:

$$v_i = \alpha \cdot \left(v_i - \frac{w_i}{w_1 + \dots + w_k} \cdot \frac{v_i^2}{1 + v_i} \right) + (1 - \alpha) \cdot \frac{x}{w_i}.$$

This yields the first-order approximation $v_i^t \sim \frac{\sqrt{(1-\alpha)x \cdot (w_1 + \dots + w_k)}}{w_i}$ as $\alpha \rightarrow 1$ and $t \rightarrow \infty$. The posterior variance of ω^t is therefore approximated as

$$\sum_{i=1}^k w_i^2 \cdot v_i^t \sim \sqrt{(1 - \alpha)x \cdot \left(\sum_{i=1}^k w_i \right)^3}.$$

This is exactly $\sqrt{(1 - \alpha) \left(\frac{M_{11}}{\text{val}(\mathcal{S})} \right)}$ since $M_{11} = \sum_{i=1}^k w_i^2 \cdot \tilde{M}_{ii} = x \cdot \sum_{i=1}^k w_i$ and $\text{val}(\mathcal{S}) = \frac{1}{\phi(\mathcal{S})^2} = \frac{1}{(\sum_{i=1}^k w_i)^2}$. We thus deduce the payoff estimate in part (1) of the theorem.

It remains to prove part (2). For that we just need to show society can achieve posterior variances smaller than $(1 + \varepsilon) \cdot \sqrt{(1 - \alpha) \left(\frac{M_{11}}{\text{val}(\mathcal{S}^*)} \right)}$ at *every* late period. In fact, we show below that myopically choosing from the best set $\mathcal{S}^*$ achieves this.

Let V^t denote society's prior covariance matrix at the beginning of period $t + 1$, under this alternative sampling strategy. Note that for α close to 1, each signal in $\mathcal{S}^*$ is observed with positive frequency. Thus, for any $L > 0$ we have

$$[V^t]^{-1} \succeq L \cdot c_j c_j', \quad \forall j \in \mathcal{S}^*$$

for α close to 1 and t large.

Now, take $v(t) := [V^t]_{11}$ to be the prior variance of ω^{t+1} . Then myopic sampling together with Lemma 13 implies the *posterior* variance of ω^{t+1} is bounded above by

$$v(t) - \frac{L}{L + 1} \cdot \frac{v(t)^2}{\phi(\mathcal{S}^*)^2}.$$

Together with the innovation terms given by M , the next prior variance $v(t+1)$ admits the following upper bound:

$$v(t+1) \leq \alpha \cdot \left(v(t) - \frac{L}{L + 1} \cdot \frac{v(t)^2}{\phi(\mathcal{S}^*)^2} \right) + (1 - \alpha) \cdot M_{11}.$$

Clearly, this implies

$$v(t+1) < v(t) + (1 - \alpha) \cdot M_{11} - \frac{L}{L + 1} \cdot \frac{v(t)^2}{\phi(\mathcal{S}^*)^2}.$$

As a consequence, $\limsup_{t \rightarrow \infty} v(t) \leq \sqrt{1 + 1/L} \cdot \sqrt{(1 - \alpha) M_{11} \cdot \phi(\mathcal{S}^*)^2}$. Since $\text{val}(\mathcal{S}^*) = \frac{1}{\phi(\mathcal{S}^*)^2}$, this yields the desired estimate if we choose $L > \frac{1}{\varepsilon}$ in the first place. We have thus completed the proof of Theorem 3.

B.2 Proof of Proposition 3

The environment in Example 6 is equivalent to one with three signals $\frac{\omega+b}{2}, \frac{\omega-b}{2}$ and $\frac{1}{L}\omega$, each with standard Gaussian noise (just let $b = \omega + 2b_1$). We assume L is large, so that the best complementary set consists of the latter two signals.

For the autocorrelated model, we choose $M = \Sigma^0 = \text{diag}(x, x)$ with $x \geq L^2$ (this is the covariance matrix for the innovations associated with ω and b). Then assuming that all previous agents have chosen the third (unbiased) signal, agent t 's prior variance of b^t remains $x \geq L^2$. As such, he (and in fact each agent) continues to observe the third signal. In this case the prior variance v^t about ω^{t+1} evolves according to

$$v^t = \alpha \cdot \frac{L^2 \cdot v^{t-1}}{L^2 + v^{t-1}} + (1 - \alpha)x.$$

It is not difficult to show that v^t must converge to the (positive) fixed point of the above equation. Let us in particular take $\alpha = 1 - \frac{1}{L^3}$ and $x = L^2$, then the long-run prior variance v

solves $v = \frac{(L^2 - \frac{1}{L})v}{L^2 + v} + \frac{1}{L}$. This yields exactly that $v = \sqrt{L}$. Hence long-run posterior variance is $\frac{L^2 \cdot v}{L^2 + v} > \sqrt{L}/2$, which implies $\limsup_{\delta \rightarrow 1} U_\delta^M \leq -\sqrt{L}/2$.

Let us turn to the optimal sampling strategy. Write $\tilde{\theta}_1 = \frac{\omega + b}{2}$ and $\tilde{\theta}_2 = \frac{\omega - b}{2}$. In this transformed model, $\tilde{M} = \tilde{\Sigma}^0 = \text{diag}(\frac{x}{2}, \frac{x}{2})$, and the payoff-relevant state is the sum of $\tilde{\theta}_1$ and $\tilde{\theta}_2$. Consider now a strategy that observes the first two signals alternatively. Then the beliefs about $\tilde{\theta}_1$ and $\tilde{\theta}_2$ remain independent (as in $\tilde{M}$ and $\tilde{\Sigma}^0$), and their variances evolve as follows: $v_1^0 = v_2^0 = \frac{x}{2}$; in odd periods t

$$v_1^t = \alpha \cdot \frac{v_1^{t-1}}{1 + v_1^{t-1}} + (1 - \alpha) \frac{x}{2} \text{ and } v_2^t = \alpha \cdot v_2^{t-1} + (1 - \alpha) \frac{x}{2},$$

and symmetrically for even t .

These imply that for odd t , v_1^t converges to v_1 and v_2^t converges to v_2 below (while for even t $v_1^t \rightarrow v_2$ and $v_2^t \rightarrow v_1$):

$$v_1 = \alpha \cdot \frac{\alpha v_1 + (1 - \alpha)x/2}{1 + \alpha v_2 + (1 - \alpha)x/2} + (1 - \alpha)x/2;$$

$$v_2 = \alpha^2 \cdot \frac{v_2}{1 + v_2} + (1 - \alpha^2) \cdot \frac{x}{2}.$$

From the second equation, we obtain $(1 - \alpha^2)(\frac{x}{2} - v_2) = \alpha^2 \cdot \frac{(v_2)^2}{1 + v_2}$. With $\alpha = 1 - \frac{1}{L^3}$ and $x = L^2$, it follows that

$$v_2 = (1 + o(1)) \frac{1}{\sqrt{L}}.$$

where $o(1)$ is a term that vanishes as $L \rightarrow \infty$. Thus we also have

$$v_1 = \alpha \frac{v_2}{1 + v_2} + (1 - \alpha) \frac{x}{2} = (1 + o(1)) \frac{1}{\sqrt{L}}.$$

Hence under this alternating sampling strategy, long-run posterior variances about $\tilde{\theta}_1$ and $\tilde{\theta}_2$ are both bounded above by $\frac{2}{\sqrt{L}}$. Since $\omega = \tilde{\theta}_1 + \tilde{\theta}_2$, we conclude that $\liminf_{\delta \rightarrow 1} U_\delta^{SP} \geq -\frac{4}{\sqrt{L}}$. Choosing L large proves the proposition.

References

- Ali, Nageeb.** 2018. “Herding with Costly Information.” *Journal of Economic Theory*, 175: 713–720.
- Athey, Susan, and Armin Schmutzler.** 1995. “Product and Process Flexibility in an Innovative Environment.” *RAND Journal of Economics*, 26(4): 557–574.
- Banerjee, Abhijit.** 1992. “A Simple Model of Herd Behavior.” *Quarterly Journal of Economics*, 107(3): 797–817.
- Bikhchandani, Sushil, David Hirshleifer, and Ivo Welch.** 1992. “A Theory of Fads, Fashion, Custom, and Cultural Change as Information Cascades.” *Journal of Political Economy*, 100(5): 992–1026.

- Blackwell, David.** 1951. “Proceedings of the Second Berkeley Symposium on Mathematical Statistics and Probability.” Chapter Comparison of Experiments, 93–102. Berkeley and Los Angeles: University of California Press.
- Börgers, Tilman, Angel Hernando-Veciana, and Daniel Krahmer.** 2013. “When Are Signals Complements Or Substitutes.” *Journal of Economic Theory*, 148(1): 165–195.
- Burguet, Roberto, and Xavier Vives.** 2000. “Social Learning and Costly Information.” *Economic Theory*, 15(1): 185–205.
- Chade, Hector, and Edward E. Schlee.** 2002. “Another Look at the Radner-Stiglitz Nonconcavity in the Value of Information.” *Journal of Economic Theory*, 107: 421–452.
- Chade, Hector, and Jan Eeckhout.** 2018. “Matching Information.” *Theoretical Economics*, 13: 377–414.
- Chaloner, Kathryn.** 1984. “Optimal Bayesian Experimental Design for Linear Models.” *The Annals of Statistics*, 12(1): 283–300.
- Chen, Yiling, and Bo Waggoner.** 2016. “Informational Substitutes.” Working Paper.
- Che, Yeon-Koo, and Konrad Mierendorff.** 2019. “Optimal Dynamic Allocation of Attention.” *American Economic Review*. Forthcoming.
- Cunha, Flavio, and James Heckman.** 2007. “The Technology of Skill Formation.” *American Economic Review*, 97(2): 31–47.
- Dasaratha, Krishna, Ben Golub, and Nir Hak.** 2018. “Social Learning in a Dynamic Environment.” Working Paper.
- DeMarzo, Peter, Dimitri Vayanos, and Jeffrey Zwiebel.** 2003. “Persuasion Bias, Social Influence, and Unidimensional Opinions.” *The Quarterly Journal of Economics*.
- Easley, David, and Nicolas M. Kiefer.** 1988. “Controlling a Stochastic Process with Unknown Parameters.” *Econometrica*, 56: 1045–1064.
- Frongillo, Rafael, Grant Schoenebeck, and Omer Tamuz.** 2011. “Social Learning in a Changing World.” *WINE’11 Proceedings of the 7th international conference on Internet and Network Economics*.
- Fudenberg, Drew, Philip Strack, and Tomasz Strzalecki.** 2018. “Speed, Accuracy, and the Optimal Timing of Choices.” *American Economic Review*, 108(12): 3651–3684.
- Gittins, J. C.** 1979. “Bandit processes and dynamic allocation indices.” *Journal of the Royal Statistical Society, Series B*, 148–177.
- Goldstein, Itay, and Liyan Yang.** 2015. “Information Diversity and Complementarities in Trading and Information Acquisition.” *Journal of Finance*, 70(4): 1723–1765.
- Golub, Benjamin, and Matthew Jackson.** 2012. “How Homophily Affects the Speed of Learning and Best-Response Dynamics.” *The Quarterly Journal of Economics*, 127(3): 1287–1338.
- Hann-Caruthers, Wade, Vadim Martynov, and Omer Tamuz.** 2017. “The Speed of Sequential Asymptotic Learning.” Working Paper.
- Hansen, Ole Havard, and Eric N. Torgersen.** 1974. “Comparison of Linear Normal Experiments.” *The Annals of Statistics*, 2: 367–373.
- Harel, Matan, Elchanan Mossel, Philipp Strack, and Omer Tamuz.** 2018. “Groupthink

- and the Failure of Information Aggregation in Large Groups.” Working Paper.
- Jovanovic, Boyan, and Yaw Nyarko.** 1996. “Learning by Doing and the Choice of Technology.” *Econometrica*, 64(6): 1299–1310.
- Liang, Annie, Xiaosheng Mu, and Vasilis Syrgkanis.** 2017. “Optimal and Myopic Information Acquisition.” Working Paper.
- Liang, Annie, Xiaosheng Mu, and Vasilis Syrgkanis.** 2019. “Dynamically Aggregating Diverse Information.” Working Paper.
- Lizzeri, Alessandro, and Marciano Siniscalchi.** 2008. “Parental Guidance and Supervised Learning.” *The Quarterly Journal of Economics*.
- Mayskaya, Tatiana.** 2019. “Dynamic Choice of Information Sources.” Working Paper.
- Milgrom, Paul, and Rober J. Weber.** 1982*a*. “A Theory of Auctions and Competitive Bidding.” *Econometrica*, 50(5): 1089–1122.
- Milgrom, Paul, and Robert J. Weber.** 1982*b*. “The Value of Information in a Sealed-Bid Auction.” *Journal of Mathematical Economics*, 10(1): 105–114.
- Monderer, Dov, and Lloyd Shapley.** 1996. “Potential Games.” *Games and Economic Behavior*, 14(1): 124–143.
- Moscarini, Giuseppe, Marco Ottaviani, and Lones Smith.** 1998. “Social Learning in a Changing World.” *Economic Theory*, 11(3): 657–665.
- Mueller-Frank, Manuel, and Mallesh Pai.** 2016. “Social Learning with Costly Search.” *American Economic Journal: Microeconomics*, 8(1): 83–109.
- Sandholm, William.** 2010. *Population Games and Evolutionary Dynamics*. MIT Press.
- Sethi, Rajiv, and Muhamet Yildiz.** 2016. “Communication with Unknown Perspectives.” *Econometrica*, 84(6): 2029–2069.
- Sethi, Rajiv, and Muhamet Yildiz.** 2019. “Culture and Communication.” Working Paper.
- Smith, Lones, and Peter Sørensen.** 2000. “Pathological Outcomes of Observational Learning.” *Econometrica*, 68(2): 371–398.
- Vives, Xavier.** 1992. “How Fast do Rational Agents Learn?” *Review of Economic Studies*, 60(2): 329–347.
- Vivi Alatas, Abhijit Banerjee, Arun Chandrasekhar Rema Hanna, and Ben Olken.** 2016. “Network Structure and the Aggregation of Information: Theory and Evidence from Indonesia.” *American Economic Review*, 106(7): 1663–1704.
- Wolitzky, Alex.** 2018. “Learning from Others’ Outcomes.” *American Economic Review*, 108: 2763–2801.

Online Appendix

“Complementary Information and Learning Traps”

Annie Liang and Xiaosheng Mu

Appendices

A	Proof of Proposition 2 Part (b)	2
A.1	Long-run Characterization for All δ	2
A.2	Efficiency as $\delta \rightarrow 1$	4
B	Strengthening of Theorem 2 “If” Part	5
B.1	Log Residual Term	5
B.2	Getting Rid of the Log	6
C	Incomplete Learning	8
C.1	Dimension Reduction	8
C.2	“Learnable” Component	8
C.3	Generalization of Results	9
C.4	Example	10
D	Other Extensions	12
D.1	General Payoff Functions	12
D.2	Low Altruism	13
D.3	Multiple Payoff-Relevant States	13
E	Regions of Prior Beliefs	16
F	Example Mentioned in Footnote 34 (Section 9.1)	18
G	Example of a Learning Trap beyond Normality	18
G.1	Proofs of the Claims	20
H	Comparison with Börgers, Hernando-Veciana and Krahmer (2013)	21

A Proof of Proposition 2 Part (b)

We will first generalize the “if” part of Theorem 2 to show that for *any* $\delta \in (0, 1)$ and any prior belief, the social planner’s sampling strategy that maximizes δ -discounted payoff yields frequency vectors that converge over time. Moreover, the limit is the optimal frequency vector associated with some strongly complementary set. Later we will argue that that for δ close to 1, this long-run outcome must be the best complementary set $\mathcal{S}^*$ starting from all priors.

A.1 Long-run Characterization for All δ

Here we prove the following result:

Proposition 7. *Suppose $\delta \in (0, 1)$. Given any prior, let $d_\delta(t)$ denote the vector of signals counts associated with any signal acquisition strategy that maximizes the δ -discounted average payoff. Then $\lim_{t \rightarrow \infty} \frac{d_\delta(t)}{t}$ exists and is equal to $\lambda^{\mathcal{S}}$ for some strongly complementary set $\mathcal{S}$.*

Proof. We follow the proof of Theorem 2 “if” part in Appendix A.6. The same argument there shows that for any $\delta < 1$, any strategy that maximizes δ -discounted payoff must infinitely observe a set of signals that span ω . Therefore it remains to prove the analogue of the restated version of Theorem 2 “if” part.

To do that, let

$$W(t) = (1 - \delta) \sum_{t' \geq t} \delta^{t'-t} \cdot V(d(t'))$$

denote the expected discounted loss from period t onwards; henceforth we fix δ and use $d(t)$ as shorthand for $d_\delta(t)$. Suppose signal acquisitions in the first t periods satisfy $C'QC \succeq Lc_jc'_j$ for each signal $j \in \mathcal{S}^*$, where L is some positive constant. Then we are going to show that

$$\frac{1}{W(t+1)} \geq \frac{1}{W(t)} + \frac{L}{(L+1)\phi(\mathcal{S}^*)^2}. \quad (13)$$

Once this is proved, we can choose L large to show $W(t) \leq \frac{(1+\varepsilon)\phi(\mathcal{S}^*)^2}{t}$ for any $\varepsilon > 0$ and all sufficiently large t . Pick m so that $\delta^m \leq \varepsilon$. Then for $t' \in (t, t+m)$ we have $V(d(t')) \geq \frac{(1-\varepsilon/2)\phi(\mathcal{S}^*)^2}{t'} \geq \frac{(1-\varepsilon)\phi(\mathcal{S}^*)^2}{t}$, so that

$$(1 - \delta) \sum_{t'=t+1}^{t+m-1} \delta^{t'-t} \cdot V(d(t')) \geq (\delta - \delta^m) \cdot \frac{(1 - \varepsilon)\phi(\mathcal{S}^*)^2}{t} \geq \frac{(\delta - \varepsilon)(1 - \varepsilon)\phi(\mathcal{S}^*)^2}{t}.$$

Subtracting this from $W(t)$, we obtain

$$(1 - \delta) \cdot V(d(t)) \leq \frac{(1 + \varepsilon - (\delta - \varepsilon)(1 - \varepsilon))\phi(\mathcal{S}^*)^2}{t}$$

again for t sufficiently large depending on ε . Since ε is arbitrary, we would be able to conclude $t \cdot V(d(t)) \rightarrow \phi(\mathcal{S}^*)^2$, and $\frac{d(t)}{t} \rightarrow \lambda^*$ would follow.

To prove (13), we consider a deviation strategy that chooses signals myopically in every period $t' \geq t+1$. Let the resulting signal count vectors be $\tilde{d}(t')$, and define $\tilde{d}(t) = d(t)$. This deviation provides an upper bound on $W(t+1)$, given by

$$W(t+1) \leq (1-\delta) \sum_{t' \geq t+1} \delta^{t'-t-1} \cdot V(\tilde{d}(t')).$$

Since $W(t) = (1-\delta) \cdot V(d(t)) + \delta \cdot W(t+1)$, we have

$$\frac{1}{W(t+1)} - \frac{1}{W(t)} = \frac{(1-\delta) \cdot (V(d(t)) - W(t+1))}{W(t+1) \cdot ((1-\delta) \cdot V(d(t)) + \delta \cdot W(t+1))},$$

which is decreasing in $W(t+1)$ (holding $V(d(t))$ equal). Thus from the previous upper bound on $W(t+1)$, we obtain that

$$\frac{1}{W(t+1)} - \frac{1}{W(t)} \geq \frac{1}{\sum_{j=0}^{\infty} (1-\delta)\delta^j \cdot V(\tilde{d}(t+1+j))} - \frac{1}{\sum_{j=0}^{\infty} (1-\delta)\delta^j \cdot V(\tilde{d}(t+1+j))} \quad (14)$$

By the assumption that $C'QC \succeq Lc_jc'_j$ after t periods, we can apply (11) to deduce that for each $j \geq 0$,

$$\frac{1}{V(\tilde{d}(t+1+j))} - \frac{1}{V(\tilde{d}(t+j))} \geq \frac{L}{(L+1)\phi(\mathcal{S}^*)^2}.$$

Given this and (14), the desired result (13) follows from the technical lemma below (with $a = \frac{L}{(L+1)\phi(\mathcal{S}^*)^2}$, $x_j = V(\tilde{d}(t+1+j))$, $y_j = V(\tilde{d}(t+j))$ and $\beta_j = (1-\delta)\delta^j$):

Lemma 14. *Suppose a is a positive number. $\{x_j\}_{j=0}^{\infty}, \{y_j\}_{j=0}^{\infty}$ are two sequences of positive numbers such that $\frac{1}{x_j} \geq \frac{1}{y_j} + a$ for each j . Then for any sequence of positive numbers $\{\beta_j\}_{j=0}^{\infty}$ that sum to 1, it holds that*

$$\frac{1}{\sum_{j=0}^{\infty} \beta_j x_j} \geq \frac{1}{\sum_{j=0}^{\infty} \beta_j y_j} + a.$$

To see why this lemma holds, note that it is without loss to assume $\frac{1}{x_j} = \frac{1}{y_j} + a$ holds with equality. Then

$$1 - a \sum_j \beta_j x_j = \sum_j \beta_j (1 - ax_j) = \beta_j \frac{x_j}{y_j}$$

By the Cauchy-Schwarz inequality,

$$\sum_j \beta_j \frac{x_j}{y_j} \geq \frac{1}{\sum_j \beta_j \frac{y_j}{x_j}} = \frac{1}{\sum_j \beta_j (1 + ay_j)} = \frac{1}{1 + a \sum_j \beta_j y_j}.$$

So $1 - a \sum_j \beta_j x_j \geq \frac{1}{1 + a \sum_j \beta_j y_j}$, which is easily seen to be equivalent to $\frac{1}{\sum_j \beta_j x_j} \geq \frac{1}{\sum_j \beta_j y_j} + a$.

Hence Lemma 14 is proved, and so is Proposition 7. $\square$

A.2 Efficiency as $\delta \rightarrow 1$

We now prove that for δ close to 1, the sampling strategy that maximizes δ -discounted payoff must eventually focus on the best complementary set $\mathcal{S}^*$. Recall that V^* is uniquely maximized at λ^* . Thus there exists positive η such that $V^*(\lambda) > (1 + \eta)V^*(\lambda^*)$ whenever λ puts zero frequency on at least one signal in $\mathcal{S}^*$.

Suppose for contradiction that sampling eventually focuses on a strongly complementary set $\mathcal{S}$ different from $\mathcal{S}^*$. Then at large periods t we must have $V(d(t)) > \frac{(1+\eta)\phi(\mathcal{S}^*)^2}{t}$, using the fact that V^* is the asymptotic version of V . As a result, there exists sufficiently large L_0 such that some signal in $\mathcal{S}^*$ is observed less than L_0 times under the optimal strategy for maximizing δ -discounted payoff.⁵¹ Crucially, this L_0 can be chosen independently of δ . As a consequence, under the hypothesis of inefficient long-run outcome, $V(d(t)) > \frac{(1+\eta)\phi(\mathcal{S}^*)^2}{t}$ in fact holds for all $t > \underline{t}$ where $\underline{t}$ is also independent of δ .

Now we fix a positive integer $L > \frac{2}{\eta}$, and consider the following deviation strategy starting in period $\underline{t} + 1$:

1. In periods $\underline{t} + 1$ through $\underline{t} + Lk$, observe each signal in the best set $\mathcal{S}^*$ (of size k) exactly L times, in any order.
2. Starting in period $\underline{t} + Lk + 1$, sample myopically.

Let us study the posterior variance after period $\underline{t} + j$ under such a deviation. For $j \geq Lk + 1$, note that each signal $j \in \mathcal{S}^*$ has been observed at least L times before the period $\underline{t} + Lk + 1$. So $C'QC \succeq Lc_jc_j'$ holds, and we can deduce (similar to (12)) that the posterior variance is at most $(1 + \frac{1}{L}) \cdot \frac{\phi(\mathcal{S}^*)^2}{j - Lk}$. Since $\frac{1}{L} < \frac{\eta}{2}$, there exists $\underline{j}$ (depending on $\eta, \underline{t}, L, k$) such that the posterior variance after period $\underline{t} + j$ is at most $(1 + \eta/2) \frac{\phi(\mathcal{S}^*)^2}{\underline{t} + j}$ for $j > \underline{j}$. Thus the flow payoff *gain* in each such period is at least

$$\frac{\eta}{2} \cdot \frac{\phi(\mathcal{S}^*)^2}{\underline{t} + j}, \quad \forall j > \underline{j}$$

under this deviation strategy.

On the other hand, for $j \leq \underline{j}$ we can trivially bound the posterior variance from above by the prior variance V_0 . This V_0 also serves as an upper bound on the flow payoff *loss* in these periods.

Combining both estimates, we find that the deviation strategy achieves payoff gain of at least

$$\delta^{\underline{t}} \cdot \left(\sum_{j > \underline{j}} \delta^{j-1} \cdot \frac{\eta}{2} \cdot \frac{\phi(\mathcal{S}^*)^2}{\underline{t} + j} - \sum_{j=1}^{\underline{j}} \delta^{j-1} \cdot V_0 \right).$$

Importantly, all other parameters in the above are constants independent of δ . As δ approaches 1, the sum $\sum_{j > \underline{j}} \frac{\delta^{j-1}}{\underline{t} + j}$ approaches a harmonic sum which diverges. Thus for all δ close to 1 the above display is strictly positive, suggesting that the constructed deviation is profitable. This contradiction completes the proof of Proposition 2 part (b).

⁵¹Otherwise, $C'QC \succeq L_0c_jc_j'$ holds at large t , implying a contradicting upper bound on $V(d(t))$ (see the argument in the previous subsection).

B Strengthening of Theorem 2 “If” Part

Here we show the following result, which strengthens the restated Theorem 2 “if” part (see Appendix A.6). It says that any signal observed with zero long-run frequency must in fact be observed only finitely often.

Stronger Version of Theorem 2 “if” part. *Suppose that the signals observed infinitely often span $\mathbb{R}^K$. Then $m_i(t) = \lambda_i^* \cdot t + O(1), \forall i$.*

The proof is divided into two subsections below.

B.1 Log Residual Term

Recall that we have previously shown $m_i(t) \sim \lambda_i^* \cdot t$. We can first improve the estimate of the residual term to $m_i(t) = \lambda_i^* \cdot t + O(\ln t)$. Indeed, Lemma 13 yields that for some constant L and every $t \geq L$,

$$V(m(t+1)) \leq V(m(t)) - \left(1 - \frac{L}{t}\right) \cdot \frac{V(m(t))^2}{\phi(\mathcal{S}^*)^2}. \quad (15)$$

This is because we may apply Lemma 13 with $M = \min_{j \in \mathcal{S}^*} m_j(t)$, which is at least $\frac{t}{L}$.

Let $g(t) = \frac{V(m(t))}{\phi(\mathcal{S}^*)^2}$. Then the above simplifies to

$$g(t+1) \leq g(t) - \left(1 - \frac{L}{t}\right) g(t)^2.$$

Inverting both sides, we have

$$\frac{1}{g(t+1)} \geq \frac{1}{g(t)} + \frac{1 - \frac{L}{t}}{1 - (1 - \frac{L}{t})g(t)} \geq \frac{1}{g(t)} + 1 - \frac{L}{t}. \quad (16)$$

This enables us to deduce

$$\frac{1}{g(t)} \geq \frac{1}{g(L)} + \sum_{x=L}^{t-1} \left(1 - \frac{L}{x}\right) \geq t - O(\ln t).$$

Thus $g(t) \leq \frac{1}{t - O(\ln t)} \leq \frac{1}{t} + O(\frac{\ln t}{t^2})$. That is,

$$V(m(t)) \leq \frac{\phi(\mathcal{S}^*)^2}{t} + O\left(\frac{\ln t}{t^2}\right).$$

Since $t \cdot V(\lambda t)$ approaches $V^*(\lambda)$ at the rate of $\frac{1}{t}$, we have

$$V^*\left(\frac{m(t)}{t}\right) \leq t \cdot V(m(t)) + O\left(\frac{1}{t}\right) \leq \phi(\mathcal{S}^*)^2 + O\left(\frac{\ln t}{t}\right). \quad (17)$$

Suppose $\mathcal{S}^* = \{1, \dots, k\}$. Then the above estimate together with (10) implies $\sum_{j > k} \frac{m_j(t)}{t} = O(\frac{\ln t}{t})$. Hence $m_j(t) = O(\ln t)$ for each signal j outside of the best set.

Now we turn attention to those signals in the best set. We work with the transformed model, as in Appendix A.2.2. Let $\Sigma = \Sigma(t)$ be the posterior covariance matrix after observing $m_j(t)$ observations of each signal $j > k$. Then the posterior covariance matrix after additionally observing $m_i(t)$ observations of each signal $i \leq k$ can be written as $[\Sigma^{-1} + \text{diag}(m_1(t), \dots, m_k(t))]^{-1}$. It follows that

$$V(m(t)) = \beta' \cdot [\Sigma^{-1} + \text{diag}(m_1(t), \dots, m_k(t))]^{-1} \cdot \beta,$$

where $\beta \in \mathbb{R}^k$ is a shorthand for $\beta^{\mathcal{S}^*}$.

Using the formula for matrix derivatives, we have that for $1 \leq i \leq k$,

$$\partial_i V(m(t)) = -\beta' \cdot [\Sigma^{-1} + \text{diag}(m_1(t), \dots, m_k(t))]^{-1} \cdot \Delta_{ii} \cdot [\Sigma^{-1} + \text{diag}(m_1(t), \dots, m_k(t))]^{-1} \cdot \beta. \quad (18)$$

Recall that each $m_i(t)$ is on the order of t , whereas the precision matrix Σ^{-1} is given by

$$\Sigma^{-1} = (\Sigma^0)^{-1} + \sum_{j>k} m_j(t) \cdot c_j c_j' = O(\ln t).$$

Thus the matrix inverse $[\Sigma^{-1} + \text{diag}(m_1(t), \dots, m_k(t))]^{-1}$ can be approximated by $\text{diag}(m_1(t), \dots, m_k(t))^{-1}$ up to a factor of $O(\frac{\ln t}{t})$. Plugging into (18), we deduce

$$\begin{aligned} \partial_i V(m(t)) &= -\beta' \cdot \text{diag}(1/m_1(t), \dots, 1/m_k(t)) \cdot \Delta_{ii} \cdot \text{diag}(1/m_1(t), \dots, 1/m_k(t)) \cdot \beta \cdot (1 + O(\ln t/t)) \\ &= -\left(\frac{\beta_i}{m_i(t)}\right)^2 \cdot \left(1 + O\left(\frac{\ln t}{t}\right)\right). \end{aligned} \quad (19)$$

We now use this to show $m_i(t) \leq \lambda_i^* \cdot t + O(\ln t)$ for each $1 \leq i \leq k$. Suppose for the sake of contradiction that $m_1(t)$ exceeds $\lambda_1^* \cdot t$ by a (big) multiple of $\ln t$. Consider $\tau + 1 \leq t$ to be the last period in which signal 1 was observed. Then $m_1(\tau)$ is larger than $\lambda_1^* \cdot \tau$ by several $\ln \tau$. Since $\sum_{1 \leq i \leq k} m_i(\tau) \leq \tau$, there exists some other signal in the best set, say signal 2, with $m_2(\tau) < \lambda_2^* \cdot \tau$. This implies $\frac{\beta_2}{m_2(\tau)}$ exceeds $\frac{\beta_1}{m_1(\tau)}$ by a factor larger than $O(\frac{\ln \tau}{\tau})$. By (19), we then deduce that $\partial_2 V(m(\tau))$ is more negative than $\partial_1 V(m(\tau))$. But this suggests that the agent in period $\tau + 1$ should not have chosen signal 1, leading to a contradiction.

Hence $m_i(t) \leq \lambda_i^* \cdot t + O(\ln t)$ holds for all signals i , whether $i \leq k$ or $i > k$. Using $\sum_{i=1}^N m_i(t) = t$, we deduce that each $m_i(t)$ is also lower-bounded by $\lambda_i^* \cdot t - O(\ln t)$. This proves $m_i(t) = \lambda_i^* \cdot t + O(\ln t)$ as desired.

B.2 Getting Rid of the Log

In order to remove the $\ln t$ residual term, we need a refined analysis. The reason we ended up with $\ln t$ is because we used (15) and (16) at *each* period t ; the " $\frac{L}{t}$ " term in those equations adds up to $\ln t$. In what follows, instead of quantifying the variance reduction in each period (as we did), we will lower-bound the variance reduction over multiple periods. This will lead to better estimates and enable us to prove $m_i(t) = \lambda_i^* \cdot t + O(1)$.

To give more detail, let $t_1 < t_2 < \dots$ denote the periods in which some signal $j > k$ is chosen. Since $m_j(t) = O(\ln t)$ for each such signal j , $t_l \geq 2^{\varepsilon \cdot l}$ holds for some positive constant ε and each positive integer l . Continuing to let $g(t) = \frac{V(m(t))}{\phi(\mathcal{S}^*)^2}$, our goal is to estimate the difference between $\frac{1}{g(t_{l+1})}$ and $\frac{1}{g(t_l)}$.

Ignoring period t_{l+1} for the moment, we are interested in $\frac{\phi(\mathcal{S}^*)^2}{V(m(t_{l+1}-1))} - \frac{\phi(\mathcal{S}^*)^2}{V(m(t_l))}$, which is just the difference in the *precision* about ω when the division vector changes from $m(t_l)$ to $m(t_{l+1}-1)$. From the proof of Lemma 12, we can estimate this difference if the change were along the direction λ^* :

$$\frac{\phi(\mathcal{S}^*)^2}{V(m(t_l) + \lambda^*(t_{l+1} - 1 - t_l))} - \frac{\phi(\mathcal{S}^*)^2}{V(m(t_l))} \geq t_{l+1} - 1 - t_l. \quad (20)$$

Now, the vector $m(t_{l+1}-1)$ is not exactly equal to $m(t_l) + \lambda^*(t_{l+1}-1-t_l)$, so the above estimate is not directly applicable. However, by our definition of t_l and t_{l+1} , any difference between these vectors must be in the first k signals. In addition, the difference is bounded by $O(\ln t_{l+1})$ by what we have shown. This implies⁵²

$$V(m(t_{l+1}-1)) - V(m(t_l) + \lambda^*(t_{l+1}-1-t_l)) = O\left(\frac{\ln^2 t_{l+1}}{t_{l+1}^3}\right).$$

Since $V(m(t_{l+1}-1))$ is on the order of $\frac{1}{t_{l+1}}$, we thus have (if the constant L is large)

$$\frac{\phi(\mathcal{S}^*)^2}{V(m(t_{l+1}-1))} - \frac{\phi(\mathcal{S}^*)^2}{V(m(t_l) + \lambda^*(t_{l+1}-1-t_l))} \geq -\frac{L \ln^2 t_{l+1}}{t_{l+1}}. \quad (21)$$

(20) and (21) together imply

$$\frac{1}{g(t_{l+1}-1)} \geq \frac{1}{g(t_l)} + (t_{l+1} - 1 - t_l) - \frac{L \ln^2 t_{l+1}}{t_{l+1}}.$$

Finally, we can apply (16) to $t = t_{l+1} - 1$. Altogether we deduce

$$\frac{1}{g(t_{l+1})} \geq \frac{1}{g(t_l)} + (t_{l+1} - t_l) - \frac{2L \ln^2 t_{l+1}}{t_{l+1}}.$$

Now observe that $\sum_l \frac{2L \ln^2 t_{l+1}}{t_{l+1}}$ converges (this is the sense in which our estimates here improve upon (16), where $\frac{L}{t}$ leads to a divergent sum). Thus we are able to conclude

$$\frac{1}{g(t_l)} \geq t_l - O(1), \quad \forall l.$$

In fact, this holds also at periods $t \neq t_l$. Therefore $V(m(t)) \leq \frac{\phi(\mathcal{S}^*)^2}{t} + O(\frac{1}{t^2})$, and

$$V^*\left(\frac{m(t)}{t}\right) \leq t \cdot V(m(t)) + O\left(\frac{1}{t}\right) \leq \phi(\mathcal{S}^*)^2 + O\left(\frac{1}{t}\right). \quad (22)$$

⁵²By the mean-value theorem, the difference can be written as $O(\ln t_{l+1})$ multiplied by a certain directional derivative. Since the coordinates of $m(t_{l+1}-1)$ and of $m(t_l) + \lambda^*(t_{l+1}-1-t_l)$ both sum to $t_{l+1}-1$, this directional derivative has a direction vector whose coordinates sum to zero. Combined with $\partial_i V(m(t)) = -(\frac{\phi(\mathcal{S}^*)^2}{t}) \cdot (1 + O(\frac{\ln t}{t}))$ (which we showed before), this directional derivative has size $O(\frac{\ln t}{t^3})$.

This equation (22) improves upon the previously-derived (17). Hence by (10) again, $m_j(t) = O(1)$ for each signal $j > k$.

Once these signal counts are fixed, we can repeat the argument in the previous subsection to argue that $m_i(t) = \lambda_i^* \cdot t + O(1)$ for $i \leq k$. Indeed, the matrix Σ^{-1} is now bounded by $O(1)$, rather than $O(\ln t)$. Thus from (18) we can obtain

$$\partial_i V(m(t)) = -\frac{\beta_i}{m_i(t)} \cdot \left(1 + O\left(\frac{1}{t}\right)\right) \quad (23)$$

Observe that the current estimate (23) improves upon the previous estimate (19). Hence by the same argument we used after (19), we can now deduce that

$$m_i(t) = \lambda_i^* \cdot t + O(1), \quad \forall i.$$

This completes the proof.

C Incomplete Learning

C.1 Dimension Reduction

Our results extend to situations where ω *cannot* be identified from the available sources. To see this, suppose that the signal coefficient vectors $c_1, \dots, c_N$ span a $k - 1$ dimensional subspace. Moreover, without loss assume that $c_1, \dots, c_{k-1}$ are linearly independent (and thus span this subspace). Re-define the confounding variables to be the following linear combinations of the original states:

$$\tilde{b}_i = \langle c_i, \theta \rangle, \quad \forall 1 \leq i \leq k - 1.$$

By assumption, each of the N signals can then be written as a linear combination of $\tilde{b}_1, \dots, \tilde{b}_{k-1}$ plus noise. On the other hand, ω cannot be written as such a linear combination since it is not identified. We can thus work in this linearly transformed model, with new state vector $\tilde{\theta} = (\omega, \tilde{b}_1, \dots, \tilde{b}_{k-1})'$ having dimension $k \leq K$. Note that the original prior covariance matrix $\Sigma^0 \in \mathbb{R}^{K \times K}$ induces a transformed prior covariance matrix $\tilde{\Sigma}^0 \in \mathbb{R}^{k \times k}$.

C.2 “Learnable” Component

In order to reduce this problem into one where the payoff-relevant state is identified, we will decompose ω into its *learnable* and *un-learnable* components. That is, we will write

$$\omega = \tilde{\omega} + \omega^\perp$$

where $\tilde{\omega}$ is a linear combination of $\tilde{b}_1, \dots, \tilde{b}_{k-1}$ (thus “learnable”), and $\omega^\perp$ is orthogonal to $\tilde{b} := (\tilde{b}_1, \dots, \tilde{b}_{k-1})'$ according to the prior $\tilde{\Sigma}^0$ (thus “un-learnable”). Such a decomposition exists and is

unique. Formally, we seek $\tilde{\omega} = \gamma' \cdot \tilde{b}$ for some $\gamma \in \mathbb{R}^{k-1}$, such that $\text{Cov}(\tilde{b}, \omega - \tilde{\omega}) = 0$. Such a vector γ is uniquely determined by

$$\text{Cov}(\tilde{b}, \tilde{b}) \cdot \gamma = \text{Cov}(\tilde{b}, \omega). \quad (24)$$

In this equation, $\text{Cov}(\tilde{b}, \tilde{b})$ represents the $(k-1) \times (k-1)$ prior covariance matrix of $\tilde{b}$ (i.e., the bottom-right submatrix of $\tilde{\Sigma}^0$). Similarly, $\text{Cov}(\tilde{b}, \omega)$ is the $(k-1) \times 1$ vector of covariances between $\tilde{b}_i$ and ω (i.e., the bottom-left submatrix of $\tilde{\Sigma}^0$).

Since the random variable $\omega^\perp = \omega - \tilde{\omega}$ is independent from $\tilde{b}_1, \dots, \tilde{b}_{k-1}$, it is also independent from any linear combination of these variables. Hence, uncertainty about $\omega^\perp$ cannot be reduced upon observation of any of the available signals. It follows that minimizing the variance of ω (the sum of the variances of $\tilde{\omega}$ and of $\omega^\perp$) is equivalent to minimizing the variance of $\tilde{\omega}$. Thus agents only seek to learn about $\tilde{\omega}$, and the problem is as if $\tilde{\omega}$ were the payoff-relevant state. This returns us to the case where the payoff-relevant state is identified. We can then define (strongly) complementary sets with respect to learning about $\tilde{\omega}$, as well as the efficient set $\mathcal{S}^*$.

C.3 Generalization of Results

We emphasize that the payoff-relevant state $\tilde{\omega}$ obtained in this way depends on the prior covariance matrix $\tilde{\Sigma}^0$. As a result, which sets are complementary in this more general setting are dependent on the prior belief over ω and $\tilde{b}$, since they are defined w.r.t. $\tilde{\omega}$.

Nonetheless, our results directly generalize for any fixed $\tilde{\omega}$, in the following way.

Definition 5. Fix nonzero $\gamma \in \mathbb{R}^{k-1}$ and $\tilde{\omega} = \gamma' \cdot \tilde{b}$.⁵³ Define the set of $\tilde{\omega}$ -learnable priors to include all prior covariance matrices over $(\omega, \tilde{b})$ that satisfy (24).

Then Theorem 2 asserts: Fix $\tilde{\omega}$ such that different complementary sets (w.r.t. $\tilde{\omega}$) have distinct informational values (Assumption 2). Then as the prior belief varies within the set of $\tilde{\omega}$ -learnable priors, strongly complementary sets (w.r.t. $\tilde{\omega}$) are the only possible long-run observation sets.

Theorem 1 also generalizes, so long as we replace “ K ” by “ $k-1$ ” in the theorem statement and in the Strong Linear Independence assumption (Assumption 3). This modification reflects the dimension reduction we performed earlier.

We note that Assumption 2 is more than necessary for these results to hold. As the proof of Theorem 2 shows, it is sufficient to assume that *within each subspace, a unique complementary set maximizes val* (i.e., minimizes ϕ). Using the derivations in the proof of Lemma 8, this latter assumption is satisfied (for all γ and $\tilde{\omega} = \gamma' \cdot \tilde{b}$) whenever the signal coefficient vectors $c_1, \dots, c_N$ have the following *generic* property:

Assumption 4. For every set of linearly independent signal coefficient vectors $c_{i_1}, \dots, c_{i_m}$ and every other c_j that can be (uniquely) written as $c_j = \sum_{l=1}^m \alpha_l \cdot c_{i_l}$, it holds that

$$\pm \alpha_1 \pm \alpha_2 \pm \dots \pm \alpha_m \neq 1$$

for all (2^m) choices of pluses and minuses.

⁵³If $\gamma = 0$, $\omega = \omega^\perp$ cannot be learned at all. We will henceforth rule out such priors.

Below we maintain this assumption.

While the above generalization of Theorem 2 considers prior beliefs that are $\tilde{\omega}$ -learnable for a fixed $\tilde{\omega}$, we can take a step further and ask which sets can be observed in the long run *across all prior beliefs over* $(\omega, \tilde{b})$ (so that $\tilde{\omega}$ also varies). The answer builds on the following definition inspired by Lemma 8:

Definition 6. *The set $\mathcal{S}$ is locally best if its signal coefficient vectors $c_{i_1}, \dots, c_{i_m}$ are linearly independent and have the following property: There exists a set of choices of pluses and minuses $\delta_1, \dots, \delta_m \in \{-1, 1\}$, such that for all signal coefficient vectors $c_j \notin \mathcal{S}$ that can be written as $c_j = \sum_{l=1}^m \alpha_l \cdot c_{i_l}$ for some (unique) $(\alpha_1, \dots, \alpha_m) \in \mathbb{R}^m$, it holds that*

$$|\delta_1 \alpha_{j1} + \dots + \delta_m \alpha_{jm}| < 1. \quad (25)$$

In words: any signal outside of $\mathcal{S}$ that is spanned by signals in $\mathcal{S}$ (i.e., can be written as a linear combination of signals in $\mathcal{S}$) must be expressible via “small” weights $(\alpha_1, \dots, \alpha_m)$, in the sense described by (25). The rough intuition is that the “smaller” are these weights (as if the noise ϵ_j is scaled up), the less informative is the signal $X_j \notin \mathcal{S}$ compared to the signals in $\mathcal{S}$. As formalized in the result below, this is precisely what guarantees $\mathcal{S}$ to have the highest informational value in its subspace (thus strongly complementary by Lemma 1).

Proposition 8. *Under Assumption 4, a set of sources $\mathcal{S}$ are eventually exclusively observed starting from some set of prior beliefs over $(\omega, \tilde{b})$ if and only if it is “locally best” as defined above.*

Proof. It is equivalent to prove that $\mathcal{S}$ is strongly complementary w.r.t. *some* $\tilde{\omega}$ if and only if it is locally best. Suppose $\mathcal{S}$ is strongly complementary w.r.t. $\tilde{\omega}$, then it is at least complementary. So some linear combination $\sum_{l=1}^m \beta_l X_{i_l}$ of the signals in $\mathcal{S}$ produces an unbiased estimate of $\tilde{\omega}$. Let δ_l be the sign of β_l for $1 \leq l \leq m$, then the fact that $\mathcal{S}$ is strongly complementary together with Lemma 8 (when applied to the subspace spanned by $\mathcal{S}$) yields the above inequality (25). So $\mathcal{S}$ is locally best.

Conversely, suppose $\mathcal{S}$ is locally best. Let $\delta_1, \dots, \delta_m \in \{-1, 1\}$ be given as in the definition. We now define $\tilde{\omega}$ to be $\sum_{l=1}^m \delta_l X_{i_l}$ with noise terms removed. Reversing the proof of Lemma 8, we deduce that $\phi(\mathcal{S}) < \phi(\mathcal{S}')$ whenever $|\mathcal{S} - \mathcal{S}'| = |\mathcal{S}' - \mathcal{S}| = 1$ and $\mathcal{S}'$ belongs to the subspace $\overline{\mathcal{S}}$. Note that $\phi(\mathcal{S}) < \phi(\mathcal{S}')$ also holds if such a set $\mathcal{S}'$ is not in this subspace, since in that case $\phi(\mathcal{S}') = \infty$.⁵⁴ Hence $\phi(\mathcal{S}) < \phi(\mathcal{S}')$ always holds, and it follows from Proposition 6 that $\text{val}(\mathcal{S}) > \text{val}(\mathcal{S}')$. $\mathcal{S}$ is then strongly complementary w.r.t. $\tilde{\omega}$ by definition. $\square$

C.4 Example

We provide an example to illustrate the above analysis.

⁵⁴Note that $c_{i_1}, c_{i_2}, \dots, c_{i_l}$ uniquely span $\tilde{\omega}$. So $c_{i_2}, \dots, c_{i_l}$ and another vector c_j (replacing c_{i_1}) span $\tilde{\omega}$ only if c_j is spanned by $c_{i_2}, \dots, c_{i_l}$ and $\tilde{\omega}$. Thus c_j must be a linear combination of $c_{i_1}, c_{i_2}, \dots, c_{i_l}$.

Example 7. The available sources are $X_1 = \omega + b_1 + b_2 + \varepsilon_1$, $X_2 = b_1 + \varepsilon_2$, and $X_3 = \omega + b_1 + b_3 + \varepsilon_3$. Note that ω cannot be completely learned, even given infinite observations of all three signals. Define $\tilde{b}_1 = \omega + b_1 + b_2$, $\tilde{b}_2 = b_1$ and $\tilde{b}_3 = \omega + b_1 + b_3$.

- (a) First consider the simplest prior belief over (ω, b_1, b_2, b_3) : $\Sigma^0 = I$. Decomposing ω into its learnable and un-learnable components, we get

$$\tilde{\omega} = \frac{2\omega + b_2 + b_3}{3} = \frac{\tilde{b}_1 - \tilde{b}_2 + \tilde{b}_3}{3},$$

which can be completely learned, and

$$\omega^\perp = \omega - \tilde{\omega} = \frac{\omega - b_2 - b_3}{3},$$

which is orthogonal to $\tilde{b}$ according to the prior. Because the prior Σ^0 is very simple, we can check orthogonality directly without having to compute $\tilde{\Sigma}^0$ and γ (which, as shown above, is equal to $(1/3, -1/3, 1/3)$).

Since agents cannot reduce uncertainty about $\omega^\perp$ by observing any of the available signals, the problem is equivalent to one in which the payoff-relevant state were given instead by $\tilde{\omega}$. In this case there is a unique complementary set w.r.t. learning about $\tilde{\omega}$, and it is the whole set $\{X_1, X_2, X_3\}$. Thus, starting from this prior belief (or any $\tilde{\omega}$ -learnable prior), agents eventually observe all three signals with equal frequencies.

- (b) As the prior belief Σ^0 varies, however, different long-run outcomes can emerge. Indeed, as the three available signals are linearly independent, condition (25) is trivially satisfied because $c_j \notin \mathcal{S}$ cannot be written as a linear combination of $c_{i_1}, \dots, c_{i_m}$. Thus every non-empty subset is “locally best,” and can be exclusively observed in the long run under some prior.

To illustrate, suppose we want to look for priors such that $\{X_1, X_2\}$ will be eventually observed with equal frequencies, whereas X_3 is not eventually observed. This suggests⁵⁵

$$\tilde{\omega} = \tilde{b}_1 - \tilde{b}_2 = \omega + b_1 + b_2 - b_1 = \omega + b_2.$$

If this $\tilde{\omega}$ were the learnable component, then $\{X_1, X_2\}$ would be the only complementary set, and agents would indeed achieve the desired long-run frequencies.

For completeness, we provide an example of such a $\tilde{\omega}$ -learnable prior. For this construction, note that $\omega^\perp$ should be $\omega - \tilde{\omega} = -b_2$. So the set of $\tilde{\omega}$ -learnable priors are those covariance matrices Σ^0 such that b_2 is orthogonal to $\omega + b_1 + b_2$, b_1 , as well as $\omega + b_1 + b_3$. Specifically, one such prior is

$$\begin{pmatrix} \omega \\ b_1 \\ b_2 \\ b_3 \end{pmatrix} \sim \mathcal{N} \left(\begin{pmatrix} \mu_1 \\ \mu_2 \\ \mu_3 \\ \mu_4 \end{pmatrix}, \begin{pmatrix} 3 & 0 & -1 & 0 \\ 0 & 1 & 0 & 0 \\ -1 & 0 & 1 & 1 \\ 0 & 0 & 1 & 3 \end{pmatrix} \right).$$

for arbitrary $\mu_1, \mu_2, \mu_3, \mu_4$.

⁵⁵Another possibility is to have $\tilde{\omega} = \tilde{b}_1 + \tilde{b}_2$, which involves a slightly more complex computation.

D Other Extensions

D.1 General Payoff Functions

Our main results extend when each agent t chooses an action to maximize an arbitrary individual payoff function $u_t(a_t, \omega)$ (we have so far restricted to $u_t(a_t, \omega) = -(a_t - \omega)^2$). Since in the Gaussian environment, the signal that minimizes the posterior variance of ω Blackwell-dominates every other signal (Hansen and Torgersen, 1974), each agent’s signal acquisition remains unchanged.⁵⁶

However, the interpretation of the optimal benchmark (that we defined in Section 5) is more limited. Specifically, while the optimal frequency vector λ^* can still be interpreted as maximizing information revelation (part (a) of Proposition 2), the relationship to the social planner problem (part (b) of Proposition 2) may fail. We comment on this possibility below.

We first note that the quadratic loss payoff assumed in our main model belongs to a class of “prediction problems”, where every agent’s payoff function $u(a, \omega)$ is the same and depends only on $|a - \omega|$. Part (b) of Proposition 2 does generalize to some other prediction problems; for example, our proof applies to any payoff function of the form $u(a, \omega) = |a - \omega|^\gamma$, with exponent $\gamma \in (0, 2]$.

Nonetheless, even restricting to prediction problems, part (b) of Proposition 2 does *not* hold in general. For a counterexample, consider $u(a, \omega) = -\mathbf{1}_{\{|a - \omega| > 1\}}$, which punishes the agent for any prediction that differs from the true state by more than 1.⁵⁷ Intuitively, the payoff gain from further information decreases sharply (indeed, exponentially) with the amount of information that has already been acquired. Thus, even with a forward-looking objective function, the range of future payoffs is limited and each agent cares mostly to maximize his own payoff. This results in an optimal sampling strategy that resembles myopic behavior, and differs from the rule that would maximize speed of learning.

The above counterexample illustrates the difficulty in estimating the value of information with an arbitrary payoff function. In order to make intertemporal payoff comparisons, we need to know how much payoff is gained/lost when the posterior variance is decreased/increased by a certain amount. This can be challenging in general, see Chade and Schlee (2002) for a related discussion.⁵⁸

Finally, while it may not be necessary to assume that agents have the same payoff function, part (b) of Proposition 2 can only hold under some restrictions on how the payoff functions differ. Otherwise, suppose for example that payoffs take the form $-\alpha_t(a_t - \omega)^2$, where α_t decreases ex-

⁵⁶To be fully rigorous, we also need a mild regularity condition on individual payoff functions to ensure that the signal minimizing the posterior variance yields *strictly* higher payoffs than other signals. The following condition is sufficient: For every t , any variance $\sigma^2 > 0$ and any action a^* , there exists a positive Lebesgue measure of μ for which a^* does *not* maximize $\mathbb{E}[u_t(a, \omega) \mid \omega \sim \mathcal{N}(\mu, \sigma^2)]$. That is, we require that for every belief variance, the expected value of ω should affect the optimal action to take. This rules out cases where some actions are dominant regardless of how much is learned.

⁵⁷We thank Alex Wolitzky for this example.

⁵⁸This difficulty becomes more salient if we try to go beyond prediction problems: The value of information in that case will depend on signal realizations.

ponentially fast. Then even with the δ -discounted objective, the social planner puts most of the weight on earlier agents, making it optimal to acquire signals myopically.

D.2 Low Altruism

We have assumed that each agent is short-lived and cares only to maximize the accuracy of his own prediction of the payoff-relevant state. Now suppose that agents are slightly altruistic; that is, each agent t chooses a signal as well as an action a_t to maximize discounted payoffs $\mathbb{E} \left[\sum_{t' \geq t} \rho^{t'-t} \cdot (a_{t'} - \omega)^2 \right]$, assuming that future agents will behave similarly. Note that $\rho = 0$ returns our main model. Below we show that for ρ sufficiently small, the existence of learning traps extends to this setting.⁵⁹

Suppose signals $1, \dots, k$ are strongly complementary. We want to show that for sufficiently small $\rho > 0$, there exist priors given which agents with discount factor ρ always observe these signals in equilibrium. We follow the construction in Appendix A.5. The added difficulty here is to show that if any agent ever chooses a signal $j > k$, the payoff loss in that period (relative to myopically choosing among the first k signals) is at least a constant fraction of possible payoff gains in future periods. Once this is proved, then for sufficiently small ρ such a deviation is not profitable.

Suppose that agents sample only from the first k signals in the first $t-1$ periods, with frequencies close to λ^* . Then, the posterior variances $V_{11}, \dots, V_{kk}$ (which are also the prior for period t) are on the order of $\frac{1}{t}$. Thus, following the computation in Appendix A.5, we can show that for some positive constant ξ (independent of t), the variance reduction of ω by any signal $j > k$ is at least $\frac{\xi}{t^2}$ smaller than the variance reduction by signal 1. This is the amount of payoff loss in period t under a deviation to signal j .

Such a deviation could improve the posterior variance in future periods. But even for the best continuation strategy, the posterior variance in period $t+m$ could at most be reduced by $O(\frac{m}{t^2})$.⁶⁰ Thus if we choose ξ to be small enough, the payoff gain in each period $t+m$ is bounded above by $\frac{m}{\xi t^2}$. Note that for ρ sufficiently small,

$$-\frac{\xi}{t^2} + \sum_{m \geq 1} \rho^m \cdot \frac{m}{\xi t^2} < 0.$$

Hence the deviation is not profitable and the proof is complete.

D.3 Multiple Payoff-Relevant States

In our main model, only one of the K persistent states is payoff-relevant. Consider a generalization in which each agent predicts (the same) $r \leq K$ unknown states and his payoff is determined via a

⁵⁹The other half of Theorem 2 also extends: Proposition 7 shows that strongly complementary sets are the only possible long-run outcomes (for any ρ).

⁶⁰This is because over m periods, the increase in the precision matrix is at most linear in m .

weighted sum of quadratic losses. We show here that our main results extend to this setting. As before, let $V(q_1, \dots, q_N)$ denote this weighted posterior variance as a function of the signal counts. V^* is the normalized, asymptotic version of V .

We assume that V^* is uniquely minimized at some frequency vector λ^* . Part (a) of Proposition 2 extends and implies that λ^* maximizes speed of learning. Unlike the case of $r = 1$, this optimal frequency vector generally involves more than K signals if $r > 1$.⁶¹ We are not aware of any simple method to characterize λ^* .

Nonetheless, we can generalize the notion of “complementary sets” as follows: A set of signals $\mathcal{S}$ is complementary if both of the following properties hold:

1. each of the r payoff-relevant states is spanned by $\mathcal{S}$;
2. the optimal frequency vector supported on $\mathcal{S}$ puts positive weight on *each* signal in $\mathcal{S}$.

Similarly, we say that a complementary set $\mathcal{S}$ is “strongly complementary” if it is best in its subspace: the optimal frequency vector supported on $\bar{\mathcal{S}}$ only puts positive weights on signals in $\mathcal{S}$. When $r = 1$, these definitions agree with our main model.

By this definition, the existence of learning traps readily extends: For suitable prior beliefs, the marginal value of each signal in $\mathcal{S}$ persistently exceeds the marginal value of each signal in $\bar{\mathcal{S}} - \mathcal{S}$. Since the marginal values of the remaining signals (those outside of the subspace) can be made very low by imposing large prior uncertainty on the relevant confounding terms, we deduce that society exclusively observes from the strongly complementary set $\mathcal{S}$.

We mention that the “if” part of Theorem 2 also generalizes. For that we need a different proof, since there is no obvious analogue of Lemma 12 (and thus of Lemma 13) when $r > 1$. Instead, we prove the restated Theorem 2 “if” part in Appendix A.6 as follows: When society infinitely samples a set that spans $\mathbb{R}^K$, the marginal value of each signal j can be approximated by its asymptotic version:

$$\partial_i V(q_1, \dots, q_N) \sim \frac{1}{t^2} \cdot \partial_i V^*\left(\frac{q_1}{t}, \dots, \frac{q_N}{t}\right).$$

Together with Lemma 11, this shows that the myopic signal choice j in any sufficient late period must *almost* minimize the partial derivative of V^* , in the following sense:

Lemma 15. *For any $\varepsilon > 0$, there exists sufficiently large $t(\varepsilon)$ such that if signal j is observed in any period $t + 1$ later than $t(\varepsilon)$, then*

$$\partial_j V^*\left(\frac{m(t)}{t}\right) \leq (1 - \varepsilon) \min_{1 \leq i \leq N} \partial_i V^*\left(\frac{m(t)}{t}\right).$$

Consider society’s frequency vectors $\lambda(t) = \frac{m(t)}{t} \in \Delta^{N-1}$. Then they evolve according to

$$\lambda(t + 1) = \frac{t}{t + 1} \lambda(t) + \frac{1}{t + 1} e_j.$$

⁶¹A theorem of Chaloner (1984) shows that λ^* is supported on at most $\frac{r(2K+1-r)}{2}$ signals.

whenever j is the signal choice in period $t + 1$. So the frequencies $\lambda(t)$ move in the direction of e_j , which is the direction where V^* decreases almost the fastest. This suggests that the evolution of $\lambda(t)$ over time resembles the gradient descent dynamics. As such, we can expect that the value of $V^*(\lambda(t))$ roughly decreases over time, and that eventually $\lambda(t)$ approaches $\lambda^* = \operatorname{argmin} V^*$.

To formalize this argument, we have (for fixed $\varepsilon > 0$ and sufficiently large t)

$$\begin{aligned}
V^*(\lambda(t+1)) &= V^*\left(\frac{t}{t+1}\lambda(t) + \frac{1}{t+1}e_j\right) \\
&= V^*\left(\frac{t}{t+1}\lambda(t)\right) + \frac{1}{t+1} \cdot \partial_j V^*\left(\frac{t}{t+1}\lambda(t)\right) + O\left(\frac{1}{(t+1)^2} \cdot \partial_{jj} V^*\left(\frac{t}{t+1}\lambda(t)\right)\right) \\
&\leq V^*\left(\frac{t}{t+1}\lambda(t)\right) + \frac{1-\varepsilon}{t+1} \cdot \partial_j V^*\left(\frac{t}{t+1}\lambda(t)\right) \\
&= \frac{t+1}{t} \cdot V^*(\lambda(t)) + \frac{(1-\varepsilon)(t+1)}{t^2} \cdot \partial_j V^*(\lambda(t)) \\
&\leq V^*(\lambda(t)) + \frac{1}{t} \cdot V^*(\lambda(t)) + \frac{1-2\varepsilon}{t} \cdot \min_{1 \leq i \leq N} \partial_i V^*(\lambda(t)).
\end{aligned} \tag{26}$$

The first inequality uses Lemma 10, the next equality uses the homogeneity of V^* , and the last inequality uses Lemma 15.

Write $\lambda = \lambda(t)$ for short. Note that V^* is differentiable at λ , since $\lambda_i(t) > 0$ for a set of signals that spans the entire space. Thus the convexity of V^* yields

$$V^*(\lambda^*) \geq V^*(\lambda) + \sum_{i=1}^N (\lambda_i^* - \lambda_i) \cdot \partial_i V^*(\lambda).$$

The homogeneity of V^* implies $\sum_{i=1}^N \lambda_i \cdot \partial_i V^*(\lambda) = -V^*(\lambda)$. This enables us to rewrite the preceding inequality as

$$\sum_{i=1}^N \lambda_i^* \cdot \partial_i V^*(\lambda) \leq V^*(\lambda^*) - 2V^*(\lambda).$$

Thus, in particular,

$$\min_{1 \leq i \leq N} \partial_i V^*(\lambda(t)) \leq V^*(\lambda^*) - 2V^*(\lambda). \tag{27}$$

Combining (26) and (27), we have for all large t :

$$V^*(\lambda(t+1)) \leq V^*(\lambda(t)) + \frac{1}{t} \cdot [(1-2\varepsilon) \cdot V^*(\lambda^*) - (1-4\varepsilon) \cdot V^*(\lambda(t))]. \tag{28}$$

Now, suppose (for contradiction) that $V^*(\lambda(t)) > (1+4\varepsilon) \cdot V^*(\lambda^*)$ holds for all large t . Then (28) would imply $V^*(\lambda(t+1)) \leq V^*(\lambda(t)) - \frac{\varepsilon \cdot V^*(\lambda^*)}{t}$. But since the harmonic series diverges, $V^*(\lambda(t))$ would eventually decrease to be negative, which is impossible. Thus

$$V^*(\lambda(t)) \leq (1+4\varepsilon) \cdot V^*(\lambda^*)$$

must hold for *some* large t . By (28), the same is true at all future periods. But since ε is arbitrary, the above inequality proves that $V^*(\lambda(t)) \rightarrow V^*(\lambda^*)$. Hence $\lambda(t) \rightarrow \lambda^*$, completing the proof of Theorem 2 for multiple payoff-relevant states.

E Regions of Prior Beliefs

Here we analyze a pair of examples to illustrate which prior beliefs would lead to which strongly complementary set as the long-run outcome.

E.0.1 Revisiting Section 2.1

First consider the example in Section 2.1, with signals $X_1 = 3\omega + b_1 + \varepsilon_1$, $X_2 = b_1 + \varepsilon_2$, $X_3 = \omega + \varepsilon_3$. For simplicity, we restrict attention to priors beliefs that are independent across ω and b_1 . We have shown that if the prior variance of b_1 is larger than 8, every agent observes the last signal and society gets stuck in a learning trap.

As a converse, we now show that whenever the prior variance of b_1 is *smaller* than 8, every agent chooses from the first two signals and learns efficiently in the long run. It is clear that the first agent observes X_1 . The rest of the argument is based on the following claim: After any number $k \geq 1$ of observations of X_1 and any number $l \geq 0$ of observations of X_2 , the next agent finds *either* signal X_1 *or* signal X_2 to be more valuable than X_3 .

We first prove this claim assuming $l = 0$. Let τ_ω and τ_b be the prior precisions about ω and b , with $\tau_b > \frac{1}{8}$ by assumption. After $k + 1$ observations of X_1 , society's posterior precision matrix is

$$P = \begin{pmatrix} \tau_\omega + 9(k+1) & 3(k+1) \\ 3(k+1) & \tau_b + k + 1 \end{pmatrix}$$

So the posterior variance of ω is $[P^{-1}]_{11}$, which is

$$f(k) = \frac{\tau_b + k + 1}{\tau_\omega \tau_b + 9(k+1)\tau_b + (k+1)\tau_\omega}.$$

If society had observed X_1 k times and X_2 once, posterior variance is

$$g(k) = \frac{\tau_b + k + 1}{\tau_\omega \tau_b + 9k\tau_b + (k+1)\tau_\omega + 9k}.$$

If society had observed X_1 k times and X_3 once, posterior variance would be

$$h(k) = \frac{\tau_b + k}{\tau_\omega \tau_b + (9k+1)\tau_b + k\tau_\omega + k}.$$

Now we need to check $\min\{f(k), g(k)\} < h(k)$. The comparison reduces to

$$\frac{\tau_b + k + 1}{\tau_b + k} < \frac{\max\{\tau_\omega \tau_b + 9(k+1)\tau_b + (k+1)\tau_\omega, \tau_\omega \tau_b + 9k\tau_b + (k+1)\tau_\omega + 9k\}}{\tau_\omega \tau_b + (9k+1)\tau_b + k\tau_\omega + k},$$

which further simplifies to

$$\frac{1}{\tau_b + k} < \frac{\max\{\tau_\omega + 8\tau_b - k, \tau_\omega - \tau_b + 8k\}}{\tau_\omega \tau_b + (9k+1)\tau_b + k\tau_\omega + k}.$$

Note that $\max\{\tau_\omega + 8\tau_b - k, \tau_\omega - \tau_b + 8k\} > \frac{2}{9}(\tau_\omega + 8\tau_b - k) + \frac{7}{9}(\tau_\omega - \tau_b + 8k) = \tau_\omega + \tau_b + 6k$. The above inequality (in the last display) thus follows from $(\tau_b + k)(\tau_\omega + \tau_b + 6k) > \tau_\omega \tau_b + (9k + 1)\tau_b + k\tau_\omega + k$, which is equivalent to the simple inequality $\tau_b^2 + 6k^2 - k > (2k + 1)\tau_b$.

Next we consider the case with $l \geq 1$ and show that whenever society has observed signals X_1 and X_2 both at least once, no future agent will choose X_3 . In this case we can apply Lemma 13 with $L = 1$ and $\mathcal{S}^* = \{X_1, X_2\}$. Note that $\phi(\mathcal{S}^*) = \frac{2}{3}$, so myopic sampling from this set leads to posterior variance at most

$$V(m(t)) - \frac{L}{L+1} \cdot \frac{V(m(t))^2}{\phi(\mathcal{S}^*)^2} = V(m(t)) - \frac{9}{8} \cdot V(m(t))^2 < V(m(t)) - V(m(t))^2.$$

On the other hand, observing signal X_3 would lead to posterior variance $\frac{V(m(t))}{1+V(m(t))}$, which is larger than $V(m(t)) \cdot (1 - V(m(t)))$. This proves the claim.

E.0.2 More Complex Dynamics

Here we study a more involved example, with signals

$$X_1 = 3\omega + b_1 + \varepsilon_1$$

$$X_2 = b_1 + \varepsilon_2$$

$$X_3 = 2\omega + b_2 + \varepsilon_3$$

$$X_4 = b_2 + \varepsilon_4$$

Note that the first two signals are the same as before, but X_3 and X_4 now involve another confounding term b_2 . The (strongly) complementary sets are $\{X_1, X_2\}$ and $\{X_3, X_4\}$, which partition the whole set as in Sethi and Yildiz (2019). We restrict attention to prior beliefs such that b_1 is independent from ω and b_2 .

We will show that learning trap is possible if and only if the prior variance of b_1 is larger than 8. On the one hand, if prior uncertainty about b_1 is large, we can specify prior beliefs over ω and b_2 such that $2\omega + b_2$ and b_2 are *independent and have the same small variance*. Given such a prior, the reduction in posterior variance of ω by either signal X_3 or X_4 is close to the reduction by the unbiased signal $\omega + \varepsilon$ (see the proof of Theorem 2). This allows us to show that agents alternate between observing X_3 and X_4 , leading to a learning trap.

On the other hand, if the prior variance of b_1 is smaller than 8, we claim that society eventually focuses on the efficient set $\{X_1, X_2\}$. To see this, we assume for contradiction that the long-run outcome is $\{X_3, X_4\}$. Then at late periods, posterior beliefs have the property that $2\omega + b_2$ and b_2 are approximately independent and have approximately the same small variance. This means both signals X_3 and X_4 have approximately the same marginal value as the unbiased signal $\omega + \varepsilon$. But then we return to the situation in Section 2.1, where future agents find it optimal to choose from $\{X_1, X_2\}$ instead (as analyzed in the previous section). This contradiction proves the claim.

Note however that unlike in Section 2.1, it may take a long time for society to eventually switch to the best set. For example, with prior variances 1, 7, $\frac{1}{100}$ of ω, b_1, b_2 respectively, the first 89 agents

observe X_3 and signal X_1 or X_2 is only observed afterwards. Moreover, switching between different complementary sets can also happen more than once: With a standard Gaussian prior of (ω, b_1, b_2) , society's signal path begins with $X_1 X_3 X_1 X_2 \dots$. These more complex dynamics make it difficult to pinpoint the long-run outcome as a function of the prior.

F Example Mentioned in Footnote 34 (Section 9.1)

Suppose the available signals are

$$\begin{aligned} X_1 &= 10x + \varepsilon_1 \\ X_2 &= 10y + \varepsilon_2 \\ X_3 &= 4x + 5y + 10b \\ X_4 &= 8x + 6y - 20b \end{aligned}$$

where $\omega = x + y$ and b is a payoff-irrelevant unknown. Set the prior to be

$$(x, y, b)' \sim \mathcal{N} \left(\begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 0.1 & 0 & 0 \\ 0 & 0.1 & 0 \\ 0 & 0 & 0.039 \end{pmatrix} \right).$$

It can be computed that agents observe only the signals X_1 and X_2 , although the set $\{X_3, X_4\}$ is optimal with $\phi(\{X_1, X_2\}) = 1/5 > 3/16 = \phi(\{X_3, X_4\})$. Thus, the set $\{X_1, X_2\}$ constitutes a learning trap for this problem. But if each signal choice were to produce 10 independent realizations, agents starting from the above prior would observe only the signals X_3 and X_4 . This breaks the learning trap.

G Example of a Learning Trap beyond Normality

The existence of learning traps is not special to the assumption of normally-distributed states and signals. We provide here a simple example demonstrating a learning trap in an environment with binary states and signals, and leave to future work the interesting question of what general conditions are sufficient.

Suppose there is a payoff-relevant state $\omega \in \{\omega_1, \omega_2\}$, which agents can learn about from any of three available signals. The first signal X_1 is described by the following information structure:

$$\begin{array}{ccc} X_1 = a & X_1 = b & \\ \omega = \omega_1 & p & 1 - p \\ \omega = \omega_2 & 1 - p & p \end{array}$$

where p itself is unknown and equals either $p_1 = 0.9$ or $p_2 = 0.1$. A second signal X_2 provides information about p , and is described by

$$\begin{array}{ccc} & X_2 = a & X_2 = b \\ p = p_1 & 0.9 & 0.1 \\ p = p_2 & 0.1 & 0.9 \end{array}$$

Finally, agents have access to a signal with information structure

$$\begin{array}{ccc} & X_3 = a & X_3 = b \\ \omega = \omega_1 & 2/3 & 1/3 \\ \omega = \omega_2 & 1/3 & 2/3 \end{array}$$

As in our earlier example from Section 2.1, repeated acquisition of X_3 is sufficient for agents to learn ω , but this path of acquisitions is *inefficient* in the sense that it leads to a sub-optimal long-run speed of learning. Specifically, agents can learn faster by alternating between acquisition of X_1 and X_2 , as the following result shows (see the next subsection for its proof):

Claim 2. *The pair of signals (X_1, X_2) strictly Blackwell-dominates the pair (X_3, X_3) , so that observing each of X_1 and X_2 once is more informative than two independent realizations of X_3 .*

Recall from Blackwell (1951) that if an experiment P is more informative another experiment Q , then for every $n \geq 1$, drawing n conditionally independent samples from P is also more informative than n samples from Q . Thus Claim 2 implies that with $2n$ total observations to allocate across the sources X_1, X_2 and X_3 , an even division between X_1 and X_2 is better than repeated acquisition of X_3 . And if the total number of observations is an odd number, the optimal allocation should at most acquire one observation of X_3 . These imply that repeated and exclusive observation of X_3 is not an efficient long-run outcome.

Nonetheless, we show next that there is an open set of prior beliefs given which all agents acquire signal X_3 , so that the set $\{X_3\}$ constitutes a “learning trap.” The logic is similar to the example in Section 2.1. Suppose the prior belief is such that ω and p are independent, and that p is equally likely to be $p_1 = 0.9$ or $p_2 = 0.1$. Due to independence, signal X_2 alone is uninformative about ω . Moreover, the uniform prior on p implies that observation of X_1 alone does not change the agent’s belief about ω . Thus the first agent myopically acquires X_3 . But since this acquisition maintains independence and the uniform prior on p , *every* agent acquires X_3 starting from this prior (regardless of signal realizations).

The result below generalizes this argument to an open set of priors. To state the result, we note that the belief over (ω, p) can be summarized by four numbers:

$$u = \mathbb{P}\{\omega_1, p_1\}, \quad v = \mathbb{P}\{\omega_1, p_2\}, \quad y = \mathbb{P}\{\omega_2, p_1\}, \quad z = \mathbb{P}\{\omega_2, p_2\}.$$

Claim 3. *Suppose the prior belief satisfies*

$$\frac{1}{2} < \frac{u}{v} < 2 \quad \text{and} \quad \frac{1}{2} < \frac{y}{z} < 2. \tag{29}$$

Then every agent chooses to acquire X_3 .

We conjecture that from all prior beliefs, agents will almost surely eventually concentrate their acquisitions on the set $\{X_3\}$ or $\{X_1, X_2\}$. These long-run outcomes constitute “complementary sets” in the sense that repeated acquisition of signals from these sets allows for complete learning of ω (almost surely), and moreover all signals in these sets are required to learn ω .⁶² Indeed, the sets are “strongly complementary” in the sense that swapping out a single signal cannot improve the long-run rate of learning. These observations are suggestive that some of our key notions and results may be generalized beyond the normal informational environment that we consider in the main text.

G.1 Proofs of the Claims

Proof of Claim 2. Note that both pairs of signals share the signal space $\{aa, ab, ba, bb\}$. Their signal structures can be summarized as the following row-stochastic matrices, where the first one corresponds to (X_1, X_2) and the second one corresponds to (X_3, X_3) :

	<i>aa</i>	<i>ab</i>	<i>ba</i>	<i>bb</i>		<i>aa</i>	<i>ab</i>	<i>ba</i>	<i>bb</i>
ω_1, p_1	0.81	0.09	0.09	0.01	ω_1, p_1	4/9	2/9	2/9	1/9
ω_1, p_2	0.01	0.09	0.09	0.81	ω_1, p_2	4/9	2/9	2/9	1/9
ω_2, p_1	0.09	0.01	0.81	0.09	ω_2, p_1	1/9	2/9	2/9	4/9
ω_2, p_2	0.09	0.81	0.01	0.09	ω_2, p_2	1/9	2/9	2/9	4/9

Let A denote the first matrix and B denote the second matrix. By direct computation,

$$A^{-1} \cdot B = \begin{array}{cc} & \begin{array}{cccc} aa & ab & ba & bb \end{array} \\ \begin{array}{c} aa \\ ab \\ ba \\ bb \end{array} & \begin{array}{cccc} \frac{155}{288} & 2/9 & 2/9 & \frac{5}{288} \\ \frac{5}{288} & 2/9 & 2/9 & \frac{155}{288} \\ \frac{5}{288} & 2/9 & 2/9 & \frac{155}{288} \\ \frac{155}{288} & 2/9 & 2/9 & \frac{5}{288} \end{array} \end{array}$$

Crucially, all entries are positive. Thus this row-stochastic matrix specifies a garbling that takes the pair (X_1, X_2) to (X_3, X_3) . $\square$

Proof of Claim 3. Consider observing the realization $X_3 = a$. By Bayes’ rule, the posterior belief is summarized by the four numbers

$$(u_{3a}, v_{3a}, y_{3a}, z_{3a}) = \frac{1}{2u + 2v + y + z} \cdot (2u, 2v, y, z).$$

Similarly the realization $X_3 = b$ would lead to the posterior belief

$$(u_{3b}, v_{3b}, y_{3b}, z_{3b}) = \frac{1}{u + v + 2y + 2z} \cdot (u, v, 2y, 2z).$$

⁶²We believe that the sets are also complementary as per our Definition 2, using a suitably generalized notion of informational value. This is left for future work.

Note that the ratios $\frac{u}{v}$ and $\frac{y}{z}$ are unchanged after either signal realization; that is, the condition (29) is preserved after acquisitions of X_3 . Thus, it suffices to show that under this condition, the *first* agent prefers signal X_3 to both X_1 and X_2 .

To do this, we will compare the expected payoffs from acquiring each of the three signals. Note that each signal has two possible realizations and leads to two possible posterior beliefs about ω . We show below that the two posterior beliefs induced by X_3 are *more extreme* than the posterior beliefs induced by any realization of X_1 or X_2 . This implies that X_3 is better than X_1 and X_2 for *every* decision problem based on the binary state ω (since the indirect utility function is convex in the posterior belief).

More formally, observe from the preceding calculations that the signal X_3 induces two possible likelihood ratios between $\omega = \omega_1$ and $\omega = \omega_2$, which are $\frac{2u+2v}{y+z}$ and $\frac{u+v}{2y+2z}$. On the other hand, the different realizations of signal X_1 or X_2 induce likelihood ratios $\frac{9u+v}{y+9z}$, $\frac{u+9v}{9y+z}$, $\frac{9u+v}{9y+z}$ and $\frac{u+9v}{y+9z}$. We will show that all four of these likelihood ratios lie strictly between $\frac{u+v}{2y+2z}$ and $\frac{2u+2v}{y+z}$.

We focus on the first of the four, $\frac{9u+v}{y+9z}$, although essentially the same argument applies to the others. The lower bound

$$\frac{9u+v}{y+9z} > \frac{u+v}{2y+2z}$$

is equivalent to $17uy + vy + 9uz > 7vz$. This follows easily from the assumption that $u > v/2$ and $y > z/2$. The upper bound

$$\frac{9u+v}{y+9z} < \frac{2u+2v}{y+z}$$

is equivalent to $7uy < 17vz + vy + 9uz$, which similarly follows from $v > u/2$ and $z > y/2$. This completes the proof. $\square$

H Comparison with Börger, Hernando-Veciana and Krahmer (2013)

Our definition of complementary sets mirrors the constructions in Börger, Hernando-Veciana and Krahmer (2013) for complementary pairs of signals, but differ in a few key aspects:

First, Definition 2 is for sets of signals, while Börger, Hernando-Veciana and Krahmer (2013) focus on pairs. Our generalization to sets can be understood in either of two ways. The proposed Definition 2 requires each *pair of subsets* that partition the whole set to be complementary. In this way, we generalize from “two complementary signals” to “two complementary sets.” For example, the sources $X_1 = \omega + b_1 + \varepsilon_1$, $X_2 = b_1 + b_2 + \varepsilon_2$, and $X_3 = b_2 + \varepsilon_3$ are complementary, since access to the set $\{X_1, X_2\}$ is complementary to access to $\{X_3\}$ (likewise for the other combinations).

Alternatively, one might consider a set to be complementary if *all* of the pieces combine to enhance the whole. For example, the above sources X_1, X_2, X_3 are complementary, since the presence of each is critical to enhancing the value of the others (ω can only be learned by observing all three sources). Proposition 1 shows these two perspectives are formally equivalent.

Another difference from [Börger, Hernando-Veciana and Krahmer \(2013\)](#) is that we consider complementary *sources* as opposed to complementary *signal observations*. Our definition does not ask whether *a single observation* of some signal improves the (marginal) value of another signal observation. Rather, we ask whether *access* to some sources improves the value of access to others, where the social planner can optimally allocate many observations across the available sources.

Finally, [Börger, Hernando-Veciana and Krahmer \(2013\)](#) value information based on its contribution to decision problems, while the *val* function that we use is more statistical in nature (based on asymptotic improvements to belief precision).⁶³ Like the definition in [Börger, Hernando-Veciana and Krahmer \(2013\)](#), however, our notion of complementarity does turn out to be uniform across prior beliefs.

⁶³That being said, Proposition 1 shows that our definition is robust to monotone transformations of the *val* function.