

Learning with Heterogeneous Misspecified Models: Characterization and Robustness*

J. Aislinn Bohren[†]

Daniel N. Hauser[‡]

July 29, 2021

First version: May 15, 2017

This paper develops a general framework to study how misinterpreting information impacts learning. Our main result is a simple criterion to characterize long-run beliefs based on the underlying form of misspecification. We present this characterization in the context of social learning, then highlight how it applies to other learning environments, including individual learning. A key contribution is that our characterization applies to settings with model heterogeneity and provides conditions for entrenched disagreement. Our characterization can be used to determine whether a representative agent approach is valid in the face of heterogeneity, study how differing levels of bias or unawareness of others' biases impact learning, and explore whether the impact of a bias is sensitive to parametric specification or the source of information. This unified framework synthesizes insights gleaned from previously studied forms of misspecification and provides novel insights in specific applications, as we demonstrate in settings with partisan bias, overreaction, naive learning, and level-k reasoning.

KEYWORDS: Model misspecification, Social learning

JEL: C73, D83

*This paper was previously circulated under the titles “Bounded Rationality and Learning: A Framework and a Robustness Result,” “Social Learning with Model Misspecification: A Framework and a Robustness Result,” and “Social Learning with Model Misspecification: A Framework and a Characterization.” We thank Nageeb Ali, Mira Frick, Drew Fudenberg, Alex Imas, Brian Kovak, Ryota Iijima, Shuya Li, George Mailath, Margaret Meyer, Pauli Murto, Wojciech Olszewski, Pietro Ortoleva, Ali Polat, Andrew Postlewaite, Andrea Prat, Yuval Salant, Larry Samuelson, Joel Sobel, Ran Spiegler, Philipp Strack, Juuso Toikka, Juuso Välimäki, Rakesh Vohra, Yuichi Yamamoto and conference and seminar participants for helpful comments and suggestions. Cuimin Ba provided excellent research assistance. Bohren gratefully acknowledges financial support from NSF grant SES-1851629 and the briq Institute.

[†]University of Pennsylvania; Email: abohren@sas.upenn.edu

[‡]Aalto University School of Business and Helsinki Graduate School of Economics; Email: daniel.hauser@aalto.fi

1 Introduction

How do individuals learn when they misinterpret information? The literature on misspecified learning typically takes the following approach: fix an incorrect, or *misspecified*, model—such as overreaction to signals or a failure to account for correlated information—and explore how it impacts the long-run beliefs about the state. We know from this literature that a misspecified model may lead to *incorrect learning*, where beliefs converge to the wrong state, *cyclical learning*, where beliefs do not converge, *entrenched disagreement*, where agents with different models become certain of different states, and *path-dependent learning*, where multiple limit beliefs can arise—for example, correct and incorrect learning.¹

In deriving these insights, studies usually consider a parameterized misspecified model that captures the cognitive bias or heuristic of interest and assume that all individuals have an identical level of this same bias.² But multiple parameterizations can often capture a given cognitive error, and these different parameterizations may yield different predictions about asymptotic learning. Further, there may be model heterogeneity, either because agents exhibit varying levels of the same bias, have fundamentally distinct biases, or use different heuristics. This raises the question of whether it is valid to use a representative agent approach or consider a single bias in isolation. Finally, a given form of misspecification may have a different impact on learning depending on whether the source of information is private, public, or social.

This paper develops a general framework to study how misinterpreting information impacts learning. A central contribution of this framework is the ability to allow for model heterogeneity. Our main result is a simple criterion to characterize long-run beliefs and behavior based on the underlying form of misspecification. We present this characterization in the context of a social learning environment in which individuals observe a private signal and the action choices of predecessors and critically, have misspecified models of how to interpret these sources. We then highlight how our characterization applies to other learning environments, including active individual learning settings and settings with different sources of information (e.g. public signals, social outcomes).

This characterization provides a deeper understanding of how misspecification influences learning and can be used to address the issues raised above. Specifically, it can determine whether a class of misspecified models are *robust*—in that the learning predictions are not sensitive to parametric specification and similar levels of a bias lead to similar learning outcomes—without needing to analyze variations on a case-by-case basis. Such robustness ensures that knowledge of the exact parametric form and level of a bias are not necessary to accurately predict its impact on learning. When agents exhibit varying levels of the same bias, the characterization can be used to evaluate whether a representative agent model is

¹Berk (1966) showed that model misspecification can lead to incorrect learning. When information depends on beliefs—as is the case when agents learn from their peers, their models vary with the history, or their actions influence future signals—misspecification can also give rise to cyclical learning (e.g. Nyarko (1991)) and path-dependent learning (e.g. (Rabin and Schrag 1999)), while when agents have heterogeneous models, entrenched disagreement can emerge (e.g. Gagnon-Bartsch (2016)). See the Related Literature section for additional references.

²Notable exceptions include Ortleva and Snowberg (2015), which allows for heterogeneous levels of correlation neglect, and Gagnon-Bartsch (2016); Frick, Iijima, and Ishii (2019), in which agents exhibit the false consensus effect about other agents’ preferences or population characteristics.

a good approximation, and in the case where it is not, to determine how such heterogeneity impacts learning. When multiple biases coexist, it can be used to study how agents with different models influence each others’ learning. Using a common framework to encompass multiple learning environments provides insight into how the impact of misspecification varies with the source of information.

Our framework captures a rich array of ways in which individuals are biased when processing information and interpreting others’ choices. Depending on the context, the empirical literature in psychology and economics has documented that individuals exhibit behavioral biases such as systematically overreacting or underreacting to new information, slanting information towards a preferred state (e.g. partisan bias), selectively weighting information (e.g. confirmation bias), incorrectly aggregating correlated information, or misperceiving others’ preferences and beliefs (e.g. false consensus effect, pluralistic ignorance), and use simplifying decision rules such as the counting and social-circle heuristics.³ Our framework represents these and other cognitive biases and heuristics as misspecified models where individuals have incorrect models of the signal distribution, others’ preferences, and how others interpret information. By studying the impact of different biases within a unified framework, our characterization can synthesize the insights gleaned from papers that focus on a single bias and can link seemingly distinct biases that have a similar effect on behavior.⁴

The details of our framework are as follows. In the social learning environment, a sequence of individuals learn about a binary state from the actions of their predecessors and a private signal, then select an action; each agent’s payoff depends on this action and the state. In the individual learning environment, a single myopic agent observes a sequence of signals and acts repeatedly; future information can depend on the history, which can capture, for example, active learning or a history-dependent model of inference. In both settings, an agent’s type specifies preferences and a model of inference. This includes a subjective belief about the signal distribution, which determines how she interprets signals, and a subjective belief about the type distribution, which determines how she interprets actions via her beliefs about others’ preferences and models of inference. Agents are Bayesian learners with respect to their subjective distributions. Model misspecification refers to the case where these subjective distributions *differ* from the true distributions.

In order to derive meaningful predictions, our framework requires additional structure on how agents interpret signals and actions. First, we focus on *aligned* type spaces in which all agents share a common ordinal ranking of signal realizations and action choices in terms of which are stronger evidence for a given state. Second, we focus on environments that

³See [Section 2.3](#) for relevant citations and the details of how our framework models these biases. Biases may be due to systematic errors or may emerge endogenously due to cognitive limitations (e.g. bounded memory ([Wilson 2014](#))) or belief-based utility (e.g. desire to appear competent, anticipatory utility ([Brunnermeier and Parker 2005](#); [Kőszegi 2006](#); [Gottlieb 2015](#))). The context of the learning setting determines which biases are of first order relevance.

⁴Our framework nests several prior behavioral models of learning. [Section 4.2](#) shows how our framework nests naive learning in [Bohren \(2016\)](#). [Online Appendix E](#) demonstrates how it nests the parameterization of confirmation bias in [Rabin and Schrag \(1999\)](#) and the non-Bayesian learning rule in [Epstein, Noor, and Sandroni \(2010\)](#). More generally, our framework can be used to study heuristics and biases that reduce to Markovian updating rules. It cannot nest heuristics that reduce to non-Markovian updating rules or are calibrated based on equilibrium objects (e.g. the analogy-based expectation equilibria in [Guarino and Jehiel \(2013\)](#)).

have *uniformly informative actions*, in that for each state, there is an action that occurs with higher probability in this state at all possible beliefs. When all agents have a correctly specified model, these assumptions ensure that learning is almost surely *correct*, in that beliefs converge to the realized state—they rule out confounded learning and informational herds (Banerjee 1992; Bikhchandani, Hirshleifer, and Welch 1992; Smith and Sørensen 2000).

This framework resolves several challenges that arise when incorporating model heterogeneity. First, model heterogeneity can lead to complicated higher-order beliefs. For example, when an agent believes that others have a misspecified signal distribution, it is also necessary to specify what the agent believes these misspecified agents believe about others, and so on. In our framework, types serve as a modeling tool to represent agents’ hierarchies of beliefs; higher order beliefs are fully determined by the subjective type distributions. Second, heterogeneity leads to a more complex learning process, as agents with different models have different beliefs following the same history. These beliefs determine their action choices. Therefore, when agents learn from others’ actions, the informational content of the history depends on a vector of beliefs. Our framework provides substantial added structure and tractability for analyzing this multidimensional belief process.

Using this framework, we explore when the asymptotic learning outcomes described above—correct, incorrect, cyclical, entrenched disagreement and path-dependent—arise. Our main result (Theorem 4) characterizes the set of learning outcomes that arise with positive probability based on two expressions that are straightforward to derive from the primitives of the misspecification: (i) the difference between the Kullback-Leibler divergence from a type’s perceived action distribution in each state at a candidate learning outcome to the true action distribution in the realized state at this learning outcome; and (ii) an ordering over the type space—*maximal accessibility*—based on each type’s perceived action distributions at certain beliefs (i.e. all types have a degenerate belief on one of the states).

To establish Theorem 4, we first determine whether a belief is locally stable in that the belief process converges to it with positive probability from nearby beliefs. Building on techniques used in Smith and Sørensen (2000) and Bohren (2016), we use the first expression described above to derive a necessary and sufficient condition for local stability (Theorem 1). Intuitively, at a given vector of beliefs, a type’s belief process moves towards the state that is more likely to generate the observed history at that vector. The difference between the Kullback-Leibler divergences at a degenerate belief determines whether this is the case in a neighborhood of the belief, and therefore, whether the belief is locally stable. We also show that non-degenerate beliefs cannot be locally stable—this stems from our focus on aligned environments with uniformly informative actions. Therefore, each type is certain about the state at any locally stable belief. The set of locally stable beliefs correspond to the set of strict Berk-Nash equilibria.⁵

We next determine whether a locally stable learning outcome is globally stable, in that beliefs converge to it with positive probability from any initial belief. This step follows immediately from local stability for the learning outcomes in which all types have correct

⁵Berk-Nash equilibrium is a solution concept for agents with misspecified models developed in Esponda and Pouzo (2016). In our framework, this means each agent plays optimally given her model and belief about the state, where her belief is the best fitting stationary belief in terms of the Kullback-Leibler divergence under the equilibrium strategy profile. By strict we mean equilibria in which there is a unique best-fitting stationary belief at a given equilibrium strategy profile.

or incorrect learning ([Theorem 2](#)). But in learning outcomes in which types have different limit beliefs, i.e. disagreement outcomes, even if the outcome is locally stable, it may not be possible to separate the beliefs of different types and push them to a neighborhood of the disagreement outcome. Maximal accessibility—the second expression described above—is sufficient to do this ([Theorem 3](#)).

Taken together, these local and global stability results establish [Theorem 4](#). An important feature of our characterization is that the two expressions we outline only need to be verified at a *finite* set of beliefs—that is, the set of certain beliefs. When the informational content of the history depends on the belief for each type, in principle, the asymptotic properties of beliefs could depend on the dynamics of beliefs across the entire belief space. Therefore, this feature significantly simplifies the calculations required to use the characterization in specific settings. Given a particular form of misspecification, the expressions are straightforward to verify.

From this characterization, we see that model heterogeneity has several important implications for learning that are distinct from settings with a single type. First, entrenched disagreement can arise within a population that observes a common history (see [Section 4.3](#) for an application). This arises despite our focus on aligned type spaces, which ensures that agents have a common interpretation of the relative informational content of signal realizations and action choices. Therefore, model heterogeneity provides an explanation for how connected populations observing shared sources can perpetually disagree. Second, cognitive biases impact the learning of agents who are not inherently biased but are misspecified due to their unawareness of others’ biases. Such unawareness can have an equally severe impact on learning as the bias itself (see [Example 2](#)).

We use our characterization to explore whether learning predictions are robust. We show that, except for knife-edge cases, nearby misspecified environments will have the same set of learning outcomes ([Theorem 5](#)). Further, learning is almost surely correct when agents have approximately correct models ([Theorem 6](#)). These results strengthen the applicability of correctly specified environments to real-world settings with mild biases. They also establish that small errors on the part of a researcher in modeling or measuring biases will not significantly alter the predicted learning outcomes. In contrast, [Frick, Iijima, and Ishii \(2020a,b\)](#) show that correctly specified environments are not robust in settings with either private actions and an infinite state space or that violate our uniformly informative actions assumption.⁶

Given our characterization, we return to the issues raised in the second paragraph in the context of specific biases. In [Section 4.1](#), we demonstrate that overreaction has a qualitatively different impact based on whether agents learn from a private or social source: when agents learn from their peers, it can lead to cyclical learning, while when individuals learn directly from signals, learning is almost surely correct. In contrast, [Epstein, Noor, and Sandroni \(2008\)](#) find that a different parameterization of overreaction leads to incorrect learning when agents observe signals directly, suggesting that overreaction is sensitive to modeling choice. In [Section 4.2](#), we show that a representative agent model is a good approximation when agents fail to account for redundant information at a similar level, but when there is sufficient

⁶[Madarász and Prat \(2016\)](#) find a failure of robustness in an agency setting, which stems from the interaction between misspecification and incentives.

heterogeneity in their bias, the representative agent model will underestimate the set of parameters that lead to correct learning. In [Section 4.3](#), we show how agents using different levels of reasoning impact each others’ learning outcomes. The presence of higher level agents can lead to different learning outcomes for level-2 agents (i.e. naive learners) relative to settings that consider the impact of naive learning in isolation (e.g. [Eyster and Rabin \(2010\)](#); [Bohren \(2016\)](#)).

Related Literature. As discussed above, stable long-run beliefs in our characterization are strict Berk-Nash equilibria ([Esponda and Pouzo 2016](#)). [Arrow and Green \(1973\)](#) provided the first equilibrium framework that explicitly distinguished between the true model and agents’ subjective models in a setting with a misspecified oligopolist. Other solution concepts for specific forms of misspecification include cursed equilibrium ([Eyster and Rabin 2005](#)), analogy-based expectation equilibrium ([Jehiel 2005](#)), and personal equilibrium ([Spiegler 2016](#)).

A rich literature explores which learning outcomes arise for specific forms of misspecification. In an individual learning setting, selective attention ([Schwartzstein 2014](#)) and misattribution of reference dependence ([Bushong and Gagnon-Bartsch 2019](#)) lead to incorrect learning almost surely, confirmation bias ([Rabin and Schrag 1999](#)) and overreaction to signals ([Epstein et al. 2010](#)) lead to correct and incorrect learning, overconfidence in one’s ability leads to inefficiently low effort ([Heidhues, Köszegi, and Strack 2018](#)), and misspecified prior beliefs ([Nyarko 1991](#); [Fudenberg, Romanyuk, and Strack 2017](#)) lead to cyclical learning. In contrast, underreaction to signals leads to correct learning almost surely ([Epstein et al. 2010](#)).

Turning to social learning, in the canonical binary state sequential environment, underestimating redundant information leads to correct and incorrect learning ([Eyster and Rabin 2010](#); [Bohren 2016](#)), while overestimating redundant information leads to cyclical learning ([Bohren 2016](#)). In a variation of this canonical environment, underestimating redundant information leads to incorrect learning almost surely or cyclical learning ([Gagnon-Bartsch and Rabin 2016](#)). Misinterpreting others’ preferences ([Frick et al. 2020a](#)) and the gambler’s fallacy ([He 2020](#)) lead to incorrect learning almost surely; misinterpreting others’ preferences can also lead to cyclical learning ([Gagnon-Bartsch 2016](#); [Bohren and Hauser 2019a](#)) or entrenched disagreement ([Gagnon-Bartsch 2016](#)).⁷ In contrast, coarse reasoning ([Guarino and Jehiel 2013](#))—which also results in underestimating redundant information—or a linear updating heuristic that puts sufficient weight on agents’ own signals ([Jadbabaie, Molavi, Sandroni, and Tahbaz-salehi 2012](#)) lead to correct learning almost surely. By capturing multiple forms of model misspecification and learning environments within the same framework, our analysis provides a tool to unify some of these insights.

A recent set of papers explore convergence in more general misspecified learning environments.⁸ For the most part, this work focuses on active individual learning settings. [Fudenberg et al. \(2017\)](#) characterize long-run beliefs for an agent who learns about a bi-

⁷In other settings, correlation neglect leads to inefficient risk-taking ([Levy and Razin 2015](#)) and ideological extremeness ([Ortoleva and Snowberg 2015](#))

⁸An older statistics literature on model misspecification characterizes limiting beliefs in terms of the Kullback-Leibler divergence ([Berk 1966](#); [Shalizi 2009](#)). These papers do not apply to active and social learning settings, as the signal process is exogenous, or to settings where an agent’s model varies with the history, as the model is fixed across time.

nary state from a diffusion process with drift that depends on the state and current action. They use this characterization to illustrate how learning outcomes can differ for myopic and patient misspecified agents. [Esponda and Pouzo \(2019\)](#) characterize steady state behavior for a class of Markov decision problems, while [Heidhues, Köszegi, and Strack \(2019\)](#) derive convergence results in a setting with Gaussian signals and state. [Fudenberg, Lanzani, and Strack \(2020\)](#); [Esponda, Pouzo, and Yamamoto \(2019\)](#) characterize properties of the limiting action distribution when the agent is non-myopic and the state space is infinite. The former show that if actions converge, then they must converge to a refinement of Berk-Nash equilibrium. The latter characterize the long-run action distribution in terms of the solutions to a generalization of a differential equation, providing insight into which action distributions arise when actions fail to converge. In a social learning setting, [Molavi, Tahbaz-Salehi, and Jadbabaie \(2018\)](#) study information aggregation when agents share beliefs on a network and treat neighbors’ current beliefs as sufficient statistics for the history. They nest common rules to aggregate beliefs on a network, including the canonical DeGroot model. Our paper complements this work by focusing on the asymptotic properties of social learning environments in which agents use heuristics or have biases that can be captured by misspecified Bayesian updating and these misspecified models may differ across agents. [Frick et al. \(2020b\)](#) build on our results to explore convergence in settings that have a finite number of states and can violate our uniformly informative property. The technical challenges that arise when there are more than two states are similar to those that arise from model heterogeneity in our setting. Relaxing uniform informativeness necessitates new methods to characterize local stability.

A complementary literature proposes explanations for how misspecification can persist despite infinite data that contradicts the model. [Gagnon-Bartsch, Rabin, and Schwartzstein \(2018\)](#) show that limited attention causes agents to ignore information that would overturn their models. [Kominers, Mu, and Peysakhovich \(2018\)](#) consider a setting where updating via Bayes rule is costly, and therefore, agents do not always update after observing new information. In [Ba \(2021\)](#), an agent switches to an alternative model only if the alternative model fits the data significantly better than the status quo model.⁹

The paper proceeds as follows. [Section 2](#) sets up the model, [Section 3](#) presents the analysis, [Section 4](#) develops several applications, and [Section 5](#) concludes. Proofs for [Section 3](#) are in [Appendix A](#) and proofs for [Section 4](#) are in [Online Appendix C](#).

2 A General Framework

We first introduce a general framework for social learning, and then discuss how to adapt it to individual learning. We conclude the section with several examples of settings that our framework captures. A reader who prefers to skip the microfoundation for the learning environment can jump to the reduced form stochastic process we analyze in [Section 3](#).

⁹A related set of papers provide a foundation for non-Bayesian updating and model misspecification. [Ortoleva \(2012\)](#) axiomatizes non-Bayesian updating rules in which agents deviate from Bayes rule when reacting to “unexpected” news and [Cripps \(2018\)](#) axiomatizes rules that are independent of how information is partitioned. [Frick et al. \(2019\)](#) show that the false consensus effect can arise when agents’ beliefs are derived from local interactions in an assortative society.

2.1 The Model: Social Learning

States and Actions. Nature selects one of two payoff-relevant states of the world $\omega \in \{L, R\}$ at the beginning of the game according to prior $p_0 \equiv \Pr(\omega = R) \in (0, 1)$. A countably infinite set of agents $t = 1, 2, \dots$ act sequentially and choose an action $\tilde{a}_t$ from a finite set $\mathcal{A}$ with $M \equiv |\mathcal{A}| \geq 2$ actions.¹⁰ Let $h_t \equiv (\tilde{a}_1, \dots, \tilde{a}_{t-1})$ denote the publicly observable action history.

Signals. Agents learn about the state from private information and the actions of other agents. Given state ω , agent t observes signal $\tilde{s}_t$ in $[0, 1]$ governed by conditional c.d.f. F^ω , independently of the signals of other agents. No signal realization perfectly reveals the state: F^L and F^R are mutually absolutely continuous with common support $\mathcal{S}$. Therefore, there exists a positive finite Radon-Nikodym derivative dF^R/dF^L . At least some signal realizations are *informative*, which rules out $dF^R/dF^L = 1$ almost surely. As is conventional, normalize the signal realization to be the posterior probability that the state is R following a neutral prior, i.e. $s = 1/(1 + dF^L/dF^R(s))$ for all $s \in \mathcal{S}$.

Types. Each agent has a privately observed type $\tilde{\theta}_t \in \Theta$ drawn independently from distribution $\pi \in \Delta(\Theta)$, where $\Theta \equiv (\theta_1, \dots, \theta_n)$ is a non-empty finite set. A type specifies preferences and a model of inference.¹¹

Preferences. Type θ_i earns payoff $u_i(a, \omega)$ from choosing action a in state ω , where $u_i : \mathcal{A} \times \{L, R\} \rightarrow \mathbb{R}$. Given probability p that the state is R , the type chooses the action that maximizes its expected payoff $(1 - p)u_i(a, L) + pu_i(a, R)$. Assume that at least two actions are not weakly dominated, no two actions yield the same payoff in both states, and no action is optimal at a single belief. Without loss of generality, assume no action is dominated for all types.

Model of Inference. A type interprets signals and actions using its subjective model of the world. Type θ_i has a subjective signal distribution in each state, represented as conditional c.d.f. $\hat{F}_i^\omega$ in state ω , and subjective type distribution $\hat{\pi}_i \in \Delta(\Theta)$.¹² It believes that no signal realization perfectly reveals the state: $\hat{F}_i^L$ and $\hat{F}_i^R$ are mutually absolutely continuous, with full support on $\mathcal{S}$ to ensure realized signals are consistent with θ_i 's model. Given signal realization s , let $\hat{s}_i(s) \equiv 1/(1 + d\hat{F}_i^L/d\hat{F}_i^R(s))$ denote θ_i 's subjective belief that the state is R following a neutral prior.¹³ All types share common prior p_0 that the state is R . It is straightforward to allow for type-specific prior beliefs or a model of inference that varies with the belief about the state (see [Online Appendix E](#) for the latter extension). Agents do not update their models of inference—we take these models as fixed and explore how they impact learning about features that are directly payoff-relevant, i.e. the state.

¹⁰We maintain the convention that a_i or a corresponds to an arbitrary element of $\mathcal{A}$ and $\tilde{a}_t$ corresponds to a random variable with support $\mathcal{A}$, with analogous notation for subsequent random variables.

¹¹While we assume types are private, suitably defining the action space and preferences so that each type chooses distinct actions can render types observable.

¹²We implicitly restrict attention to forms of signal misspecification in which signal realizations that map to the same true posterior also map to the same subjective posterior. In this case, it is without loss of generality to define subjective signal distributions with respect to a signal space normalized to correspond to private beliefs ([Bohren and Hauser 2021b](#)).

¹³We can also take $\hat{s}$ as a primitive: for any strictly increasing function $\hat{s} : \mathcal{S} \rightarrow [0, 1]$ with $\hat{s}(\inf \mathcal{S}) < 1/2$ and $\hat{s}(\sup \mathcal{S}) > 1/2$, there exists a pair of mutually absolutely continuous probability measures with full support on $\mathcal{S}$ that are represented by $\hat{s}$ (see ([Bohren and Hauser 2021b](#))).

A *correctly specified* type has a correct model of inference, $(\hat{F}_i^L, \hat{F}_i^R) = (F^L, F^R)$ and $\hat{\pi}_i = \pi$, while a *misspecified* type has an incorrect model, $(\hat{F}_i^L, \hat{F}_i^R) \neq (F^L, F^R)$ and/or $\hat{\pi}_i \neq \pi$. We group types into three categories. A *noise* type does not learn from its signal or the action history: it believes that signals are uninformative, $d\hat{F}_i^L/d\hat{F}_i^R = 1$ almost surely, and all agents are noise types, $\hat{\pi}_i(\Theta_N) = 1$, where Θ_N denotes the set of noise types. An *autarkic* type learns from its signal but not the action history: it believes that signals are informative and all agents are noise types. To avoid the case in which an autarkic type is observationally equivalent to a noise type, assume that an autarkic type has preferences such that there are at least two strictly optimal actions on the set of posterior beliefs that arise from its subjective signal distribution. A *social* type learns from its signal and the action history: it believes that both are informative. Let Θ be ordered such that the first k types are social, denoted by $\Theta_S \equiv (\theta_1, \dots, \theta_k)$, and the remaining $n - k$ types are noise or autarkic, denoted by Θ_A and Θ_N , respectively. A learning environment (Θ, π) is *correctly specified* if all social types are correctly specified and otherwise is *misspecified*.

When agents have different models of inference, this can lead to complicated higher-order beliefs. For example, when an agent believes that other agents have a misspecified signal distribution, we also need to model what the agent believes these misspecified agents believe about others. In our framework, higher-order beliefs are fully determined by the subjective type distributions. If type θ_i believes that all agents are type θ_j , then θ_j 's subjective type distribution $\hat{\pi}_j$ captures θ_i 's second order beliefs, the subjective type distributions of the types in the support of $\hat{\pi}_j$ capture third order beliefs, and so on. Therefore, in addition to describing the types that actually exist, Θ may contain types that serve as a tool to represent hierarchies of beliefs—in other words, types that occur with positive probability under a type's subjective distribution but with probability zero under the true distribution.¹⁴

Aligned Environments. In order to derive meaningful predictions, our framework requires some structure on how agents interpret signals and actions. We focus on environments that are *aligned*, in that it is common knowledge that agents have the same ordinal ranking of signal realizations and action choices in terms of which are stronger evidence for state R . For the signal, this corresponds to subjective signal distributions that satisfy the following assumption.

Assumption 1 (Aligned Subjective Signals). *For each $\theta_i \in \Theta$, the subjective signal distribution is either aligned with the true signal distribution, i.e. for any $s, s' \in \mathcal{S}$ such that $s > s'$, then $\hat{s}_i(s) > \hat{s}_i(s')$ or uninformative, i.e. $\hat{s}_i(s) = 1/2$ for all $s \in \mathcal{S}$.*

In other words, for any two signal realizations s and s' , if s is stronger evidence for state R than s' under the true measure, then s is also stronger evidence for state R than s' under the subjective measure. We make one exception to allow for types who believe the signal is uninformative. Types can differ in the perceived *strength* of signal realizations—both relative to other types and to the true distribution. For example, all agents can believe that lung cancer is stronger evidence that smoking is harmful than shortness of breath, but differ in their perceived strength of this evidence. [Assumption 1](#) implies common knowledge that

¹⁴[Mertens and Zamir \(1985\)](#) construct the universal type space, which is the set of hierarchies of beliefs that satisfy certain consistency requirements. Finiteness combined with subsequent restrictions we impose on Θ restrict the set of belief hierarchies we analyze to a subset of the universal type space.

signals are aligned, since all agents believe that other agents have a type in Θ , and so on.

For actions, we assume that, when the state is known, each type has the same ordinal ranking over its undominated actions. Types may have different sets of undominated actions, and therefore, choose different actions when the state is known.

Assumption 2 (Aligned Preferences). *The set of types Θ have aligned preferences, in that there exists a complete order $\succ$ on $\mathcal{A}$ such that if $a \succ a'$, then for each $\theta_i \in \Theta$, either $u_i(a, R) > u_i(a', R)$ or a is dominated for θ_i .¹⁵*

For example, all agents prefer a risky asset in one state and a safe asset in the other, but differ in the belief at which they are willing to start investing in the risky asset or some agents prefer less risky assets across all beliefs about the state. Given [Assumption 2](#), we maintain a complete order over the action space by relative preference in state R .¹⁶ Fixing such an order, index actions $\mathcal{A} \equiv (a_1, \dots, a_M)$ so that $a_m \succ a_l$ iff $m > l$.

While signal and preference alignment are not necessary, they are a simple yet general set of restrictions that allow us to derive sharp predictions in a broad class of learning environments. These restrictions do rule out some natural economic settings—for example, some versions of horizontally differentiated environments (e.g. the horizontally differentiated preferences in [Gagnon-Bartsch \(2016\)](#)). In [Section 3.6](#), we discuss how our techniques can be applied to environments that are not aligned.

Informative Actions and Consistent Histories. We focus on environments that are *uniformly informative*, in that for each state there is an action that occurs and is perceived to occur with higher probability in this state regardless of the history. Since an autarkic type believes its signal is informative and does not observe the history, its action choice depends on its signal realization regardless of the history. Therefore, as we show in [Section 3.1](#), the following simple condition—combined with aligned signals and preferences—is sufficient to establish that a_1 is *uniformly informative* of state L —that is, it occurs with higher probability in state L at all possible beliefs—and similarly, a_M is uniformly informative of state R .

Assumption 3 (Informative Actions). *For actions $a \in \{a_1, a_M\}$, there exists an autarkic type $\theta_j \in \Theta_A$ with $\pi(\theta_j) > 0$ that plays a with positive probability, and each social type $\theta_i \in \Theta_S$ believes that such an autarkic type exists.*

Alternative assumptions can also establish uniform informativeness. Our analysis carries through unchanged provided at least one action is uniformly informative of each state. Further, uniform informativeness does not need to hold with respect to actions—for example, our analysis also applies if there is an alternative source that is uniformly informative. We discuss this further in [Section 3.1](#).

We also focus on settings in which the realized history is consistent with each type’s model of inference. To rule out the possibility that a type observes what it believes to be a zero probability history, we assume that social types believe that there is an autarkic or noise type that plays each action with positive probability.

¹⁵For any undominated actions a and a' , if $u_i(a, R) > u_i(a', R)$, then $u_i(a, L) < u_i(a', L)$. Therefore, if $a \succ a'$, then for each $\theta_i \in \Theta$, either $u_i(a', L) > u_i(a, L)$ or a' is dominated for θ_i .

¹⁶This order may not be unique since [Assumption 2](#) places no restriction on how a type ranks its dominated actions or actions that are optimal for a single type.

Assumption 4 (Consistent Histories). For each $a \in \mathcal{A}$ and for each social type $\theta_i \in \Theta_S$, there exists an autarkic or noise type $\theta_j \in \Theta_A \cup \Theta_N$ with $\hat{\pi}_i(\theta_j) > 0$ that plays a with positive probability.

This ensures that each social type believes that all histories are on the equilibrium path. Any learning environment can be slightly perturbed so that it satisfies [Assumptions 3](#) and [4](#) by adding such an autarkic or noise type with arbitrarily small probability.

Timing. At time t , agent t draws its type $\tilde{\theta}_t$ and observes the history h_t and signal $\tilde{s}_t$, then chooses action $\tilde{a}_t$. Then the history is updated to h_{t+1} .

In [Section 3.6](#), we discuss possible extensions to our framework, including misaligned type spaces, heterogeneous signal distributions, and state-dependent type distributions.

2.2 The Model: Individual Learning

Our framework can capture misspecified learning with a single long-run agent by modifying the learning environment so that the signal process is public. Suppose $\mathcal{S}$ is finite and otherwise maintain the same assumptions on the true signal process as in [Section 2.1](#). A single type has subjective signal distributions $\hat{F}^L$ and $\hat{F}^R$ that are mutually absolutely continuous with full support on $\mathcal{S}$. Replace [Assumption 3](#) with the assumption that signals are perceived as informative, $d\hat{F}^R/d\hat{F}^L \neq 1$. When there is a single type, [Assumptions 1](#), [2](#) and [4](#) are unnecessary and the type distribution is trivial. Allowing the perceived signal distributions $\hat{F}^L$ and $\hat{F}^R$ to vary with the belief about the state captures cognitive biases such as *confirmation bias*, nesting [Rabin and Schrag \(1999\)](#), and certain forms of *under-/overreaction*, nesting [Epstein et al. \(2010\)](#) (see [Online Appendix E](#)). Allowing the true distributions F^L and F^R to vary with the belief about the state captures an active learning model in which action choices influence information. In [Section 3.1](#), we show how this individual learning set-up reduces to a belief process that satisfies the same properties as the social learning set-up, and therefore, we can also apply our learning characterization to this setting.

2.3 Examples

The following examples demonstrate how our framework can capture many information-processing biases, heuristics, and other misspecified models of inference.

Partisan Bias. A type systematically slants evidence towards one state ([Bartels 2002](#); [Jerit and Barabas 2012](#)). For example, an R-partisan type interprets all signal realizations as being stronger evidence for state R than is actually the case, $\hat{s}(s) = s^\nu$.

Under-/Overreaction. A type under- or overreacts to signals ([Moore and Healy 2008](#); [Moore, Tenney, and Haran 2015](#); [Angrisani, Guarino, Jehiel, and Kitagawa 2020](#)). For example, $\frac{\hat{s}(s)}{1-\hat{s}(s)} = (\frac{s}{1-s})^\nu$, where $\nu \in [0, 1)$ corresponds to underreaction and $\nu \in (1, \infty)$ corresponds to overreaction.

Correlation Neglect/Naive Learning. A type underestimates the correlation in the actions of prior agents: the true share of autarkic types is $\pi(\Theta_A)$, but the type believes that the share of autarkic types is $\hat{\pi}(\Theta_A) > \pi(\Theta_A)$ ([Eyster and Rabin 2010](#); [Bohren 2016](#); [Enke and Zimmermann 2019](#); [Eyster, Rabin, and Weizsäcker 2020](#)). The *counting heuristic*, where agents simply count actions to form beliefs, provides a foundation for this bias ([Ungeheuer and Weber 2020](#)).

Level-k/Cognitive Hierarchy. Level-1 is an autarkic type that believes all agents are noise

types; level-2 believes all agents are level-1 and interprets each prior action as reflecting an independent private signal; level-3 believes all agents are level-2, and so on (Costa-Gomes and Crawford 2006; Penczynski 2017). The cognitive hierarchy model is similar, but allows for richer beliefs over types (Camerer, Ho, and Chong 2004).

False Consensus Effect. A type overweights the likelihood that others have similar preferences or models of inference (Ross, Greene, and House 1977; Marks and Miller 1987; Gagnon-Bartsch 2016). For example, there are two types with preferences $u_1 \neq u_2$. Both types believe all agents share their preferences, $\hat{\pi}_1(\theta_1) = 1$ and $\hat{\pi}_2(\theta_2) = 1$.

Pluralistic Ignorance. A type underweights the likelihood that others have similar preferences or models of inference (Miller and McFarland 1987, 1991). For example, all types correctly interpret signals but believe others overreact.

A model of inference that depends on the belief about the state can capture confirmation biases such as *selective exposure*, *selective perception* and *selective recall* (Nickerson 1998); see Online Appendix E for this extension. Type-specific signal distributions can capture biases that involve interpersonal comparisons of the quality of information, such as *overconfidence* in the accuracy of one’s information relative to others (Moore and Healy 2008) and overestimating the precision of signals from agents who have similar preferences or models i.e. the *social circle heuristic* (Pachur, Rieskamp, and Hertwig 2004); see Section 3.6 for details.

3 Asymptotic Learning

We study asymptotic learning outcomes—the long-run beliefs about the state—for social types. Our main result characterizes how these long-run beliefs depend on *two expressions* that are straightforward to derive from the primitives of the model.

3.1 Belief Updating

We first characterize how each type updates its belief about the state.

Individual Decision Problem. Consider an agent of type θ_i who observes history h and private signal realization s . The agent uses her model of inference to compute the probability of h in each state, $\hat{P}_i(h|\omega)$, and applies Bayes rule to form the likelihood ratio

$$\lambda_i(h) \equiv \frac{\hat{P}_i(R|h)}{\hat{P}_i(L|h)} = \left(\frac{p_0}{1-p_0} \right) \frac{\hat{P}_i(h|R)}{\hat{P}_i(h|L)}$$

that the state is R versus L . By Assumption 4, θ_i believes that h occurs with positive probability and therefore Bayes rule can be used to update beliefs. If the agent is an autarkic or noise type, she believes that the history is uninformative, implying $\lambda_i(h) = p_0/(1-p_0)$ for all h . The agent then uses her subjective signal distribution $\hat{s}_i(s)$ to update to private posterior belief $\lambda_i(h)\hat{s}_i(s)/(1-\hat{s}_i(s))$ and chooses the action that maximizes her expected payoff with respect to this belief. The following lemma represents each type’s decision rule as a set of signal cutoffs.

Lemma 1 (Decision Rule). Assume Assumptions 1 and 2. For each $\theta_i \in \Theta$ and $\lambda \in [0, \infty]$, there exist signal cutoffs $0 = \bar{s}_{i,0}(\lambda) \leq \bar{s}_{i,1}(\lambda) \leq \dots \leq \bar{s}_{i,M}(\lambda) = 1$ such that an agent of type θ_i chooses action a_m at likelihood ratio λ iff $\bar{s}_{i,m-1}(\lambda) \neq \bar{s}_{i,m}(\lambda)$ and she observes private

signal realization $s \in (\bar{s}_{i,m-1}(\lambda), \bar{s}_{i,m}(\lambda)]$, with a closed interval if $\bar{s}_{i,m-1}(\lambda) = 0$.¹⁷

Interpreting Action Histories. These signal cutoffs determine how agents interpret past action choices. A social type $\theta_i \in \Theta_S$ believes type θ_j with likelihood ratio λ_j chooses a_m with probability $\hat{F}_i^\omega(\bar{s}_{j,m}(\lambda_j)) - \hat{F}_i^\omega(\bar{s}_{j,m-1}(\lambda_j))$ in state ω , i.e. the subjective probability θ_i assigns to θ_j observing a signal realization in the interval $(\bar{s}_{j,m-1}(\lambda_j), \bar{s}_{j,m}(\lambda_j)]$ that lead to choice a_m . Therefore, given likelihood ratios $\boldsymbol{\lambda} \equiv (\lambda_1, \dots, \lambda_k) \in [0, \infty]^k$ for social types and $\lambda_j = p_0/(1 - p_0)$ for autarkic or noise types, θ_i believes that a_m is chosen with probability

$$\hat{\psi}_i(a_m|\omega, \boldsymbol{\lambda}) \equiv \sum_{j=1}^n \hat{\pi}_i(\theta_j) (\hat{F}_i^\omega(\bar{s}_{j,m}(\lambda_j)) - \hat{F}_i^\omega(\bar{s}_{j,m-1}(\lambda_j))) \quad (1)$$

in state ω , i.e. the subjective probability θ_i assigns to each type choosing a_m weighted by θ_i 's subjective type distribution. Note that the probability of a_m varies with both the realized state and the belief about the state. The true probability of a_m is analogous, substituting the true signal and type distributions,

$$\psi(a_m|\omega, \boldsymbol{\lambda}) \equiv \sum_{j=1}^n \pi(\theta_j) (F^\omega(\bar{s}_{j,m}(\lambda_j)) - F^\omega(\bar{s}_{j,m-1}(\lambda_j))). \quad (2)$$

When $\hat{\psi}_i(a_m|\omega, \boldsymbol{\lambda}) \neq \psi(a_m|\omega, \boldsymbol{\lambda})$, misspecification introduces a wedge between the subjective and true probability of observing each action.

The following lemma establishes several key properties of $\psi(a|\omega, \boldsymbol{\lambda})$ and $\hat{\psi}_i(a|\omega, \boldsymbol{\lambda})$. First, all social types perceive action a_1 as *uniformly informative* of state L —that is, more likely in state L at all values of the likelihood ratio—and perceive a_M as uniformly informative of state R . This follows from [Assumptions 1 and 2](#), which rule out confounded learning, and [Assumption 3](#), which rules out informational herds.¹⁸ Additionally, $\hat{\psi}_i(a|\omega, \boldsymbol{\lambda})$ and $\psi(a|\omega, \boldsymbol{\lambda})$ are continuous with respect to $\boldsymbol{\lambda}$ at certainty and no action is perceived to perfectly reveal the state.

Lemma 2. *Assume [Assumptions 1 to 4](#). For all $\theta_i \in \Theta_S$, action a_1 (a_M) is perceived as uniformly informative of state L (state R),*

$$\sup_{\boldsymbol{\lambda} \in [0, \infty]^k} \frac{\hat{\psi}_i(a_1|R, \boldsymbol{\lambda})}{\hat{\psi}_i(a_1|L, \boldsymbol{\lambda})} < 1 \text{ and } \inf_{\boldsymbol{\lambda} \in [0, \infty]^k} \frac{\hat{\psi}_i(a_M|R, \boldsymbol{\lambda})}{\hat{\psi}_i(a_M|L, \boldsymbol{\lambda})} > 1,$$

¹⁷Throughout the paper, we work with the extended real number line $\mathbb{R} \cup \{-\infty, \infty\}$ to allow $\lambda = \infty$ to denote the belief at which an agent places probability one on state R .

¹⁸Confounded learning corresponds to an interior belief at which actions are uninformative even though each type acts based on private information. When preferences are not aligned, such as $u_1 = \mathbb{1}_{a=\omega}$ and $u_2 = \mathbb{1}_{a \neq \omega}$, the aggregate probability that an action is chosen can be the same in each state even though different types choose actions with different probabilities ([Smith and Sørensen 2000](#)). Model heterogeneity can lead to the same issue when it is not common knowledge that signals and preferences are aligned. An informational herd corresponds to an interior belief at which all types choose the same action regardless of their private information ([Banerjee 1992](#); [Bikhchandani et al. 1992](#)). Unbounded private signals also rule out informational herds ([Smith and Sørensen 2000](#)). While signals may be unbounded in our setting, our learning characterization requires the stronger notion of uniformly informative actions.

these actions occur with positive probability in each state uniformly across $\lambda \in [0, \infty]^k$, and $\lambda \mapsto \hat{\psi}_i(a|\omega, \lambda)$ and $\lambda \mapsto \psi(a|\omega, \lambda)$ are continuous at $\lambda \in \{0, \infty\}^k$ for $(a, \omega) \in \mathcal{A} \times \{L, R\}$. Further, each action $a \in \mathcal{A}$ is boundedly informative i.e. there exists an $\varepsilon > 0$ such that $\hat{\psi}_i(a|R, \lambda)/\hat{\psi}_i(a|L, \lambda) \in [\varepsilon, 1/\varepsilon]$ for $\lambda \in [0, \infty]^k$.

From this point forward, our analysis is derived in terms of $\psi(a|\omega, \lambda)$ and $\hat{\psi}_i(a|\omega, \lambda)$. Therefore, one could abstract from the microfoundations for individual and social learning presented in Section 2 and work directly with a reduced form learning model represented by a state and belief-dependent stochastic process on a finite set $\mathcal{A}$. Our subsequent analysis holds for any such process that satisfies the properties derived in Lemma 2 (uniform informativeness can hold at any two elements in $\mathcal{A}$ —it does not necessarily need to hold at a_1 and a_M) and a consistency requirement that all observed information is in the support of agents’ subjective distributions.¹⁹ This allows our learning characterization to be immediately applied to large class of dynamic decision problems with model misspecification—including active individual learning models with myopic agents, learning from alternative sources of information (i.e. stochastic outcomes, public signals), learning from multiple sources of information, and models of inference that vary with agents’ beliefs about the state.²⁰ While some common economic settings violate uniform informativeness—namely, social learning with unbounded private signals and individual learning with costly information acquisition—the property is restored if there is an additional source of information or if all realizations of the information process are a stochastic function of choices and the state (e.g. $\mathcal{A}$ corresponds to stochastic outcomes as in Bohren and Hauser (2019a)).

The Likelihood Ratio Process. From Eq. (1), we derive how the likelihood ratios for social types depend on the history. Each $\theta_i \in \Theta_S$ initially has likelihood ratio $\lambda_i(h_1) = p_0/(1 - p_0) \in (0, \infty)$. At any history h_t with $t > 1$,

$$\lambda_i(h_t) = \left(\frac{p_0}{1 - p_0} \right) \prod_{\tau=1}^{t-1} \frac{\hat{\psi}_i(\tilde{a}_\tau|R, \lambda(h_\tau))}{\hat{\psi}_i(\tilde{a}_\tau|L, \lambda(h_\tau))}.$$

The process is recursive: given $\lambda(h_t)$ and $\tilde{a}_t$, $\lambda_i(h_{t+1}) = \lambda_i(h_t) \frac{\hat{\psi}_i(\tilde{a}_t|R, \lambda(h_t))}{\hat{\psi}_i(\tilde{a}_t|L, \lambda(h_t))}$. Therefore, $\lambda_t \equiv \lambda(h_t)$ is sufficient for the history and we suppress the dependence on h . The behavior of $\langle \lambda_t \rangle_{t=1}^\infty$ determines the learning dynamics for each social type—where $\hat{\psi}_i(a|L, \lambda)$ and $\hat{\psi}_i(a|R, \lambda)$ determine the *update* to the likelihood ratio, and $\psi(a|\omega, \lambda)$ in the realized state determines the *probability* of this update. Characterizing the behavior of $\langle \lambda_t \rangle_{t=1}^\infty$ is challenging: the process is an equilibrium object with nonlinear transitions that depend on its current value, and in contrast to correctly specified environments, it is generally not a martingale.

We close with an example, which we use throughout the paper to illustrate our results.

Example 1 (Partisan Bias). Suppose agents systematically slant signal realizations towards

¹⁹Formally, for all $\theta_i \in \Theta_S$ and $(a, \omega, \omega') \in \mathcal{A} \times \{L, R\}^2$, $\hat{\psi}_i(a|\omega, \lambda) = 0$ iff $\psi(a|\omega', \lambda) = 0$.

²⁰Note that when ψ and $\hat{\psi}_i$ are independent of λ , our set-up is a special case of Berk (1966). While in principle, our analysis could be applied to non-myopic active learning problems, in practice, it would be a significant challenge to verify the properties derived in Lemma 2 as $\hat{\psi}(\cdot|\omega, \lambda)$ would depend on the solution to a dynamic optimization problem. Therefore, our framework does not easily apply to settings such as that studied in Fudenberg et al. (2017).

state R . There are two types of agents: θ_1 is social and θ_2 is autarkic, with share $\pi(\theta_1) \in (0, 1)$ of social types. Both have an identical level of partisan bias parameterized by $\hat{F}_i^\omega(s) = F^\omega(s^\nu)$ for $\nu \in (0, 1)$ and $\mathcal{S} = [0, 1]$, which results in private belief $\hat{s}_i(s) = s^\nu$. The social type has a correctly specified type distribution, $\hat{\pi}_1 = \pi$. Both types seek to choose the action that matches the state, $\mathcal{A} = \{L, R\}$ and $u_i(a, \omega) = \mathbb{1}_{a=\omega}$, and start with prior $p_0 = 0.5$. When the social type has likelihood ratio λ and observes signal realization s , it updates to private belief $\lambda \left(\frac{s^\nu}{1-s^\nu} \right)$. It chooses action L if this private belief is less than one, which occurs for signal realizations less than $\bar{s}_{1,1} = 1/(1 + \lambda)^{1/\nu}$. Similarly, the autarkic type chooses L for signal realizations less than $\bar{s}_{2,1} = 0.5^{1/\nu}$. The social type believes social types choose L with probability $\hat{F}_1^\omega(1/(1 + \lambda)^{1/\nu}) = F^\omega(1/(1 + \lambda))$ and autarkic types choose L with probability $\hat{F}_1^\omega(0.5^{1/\nu}) = F^\omega(0.5)$. At any $\lambda \in (0, \infty)$, this is greater than the true probabilities $F^\omega(1/(1 + \lambda)^{1/\nu})$ and $F^\omega(0.5^{1/\nu})$, respectively, that each type chooses L . Therefore, the social type overestimates the frequency of L actions.

3.2 Stationary Beliefs and Learning Outcomes

At a stationary belief, the likelihood ratio remains constant for any action that occurs with positive probability.

Definition 1 (Stationary). Belief $\lambda^* \in [0, \infty]^k$ is stationary if for all $a \in \mathcal{A}$, either (i) $\psi(a|\omega, \lambda^*) = 0$ or (ii) $\lambda^* = \lambda^* \left(\frac{\hat{\psi}_i(a|R, \lambda^*)}{\hat{\psi}_i(a|L, \lambda^*)} \right)$ for all $\theta_i \in \Theta_S$.

From Lemma 2, actions a_1 and a_M are uniformly informative across the belief space. Therefore, the set of stationary beliefs corresponds to the set of certain beliefs in which each type places probability one on a single state.

Lemma 3 (Stationary Beliefs). Assume Assumptions 1 to 4. The set of stationary beliefs is $\{0, \infty\}^k$. Given likelihood ratio process $\langle \lambda_t \rangle_{t=1}^\infty$, for any belief $\lambda^* \in [0, \infty]^k$, if $\lambda_i^* \in (0, \infty)$ for some $\theta_i \in \Theta_S$, then $\Pr(\lambda_t \rightarrow \lambda^*) = 0$.

These stationary beliefs are the candidate limit points of the likelihood ratio: if the likelihood ratio converges for all types, then it must converge to $\lambda^* \in \{0, \infty\}^k$. This rules out incomplete learning, i.e. $\lambda_{i,t}$ converges to an interior belief for some type.

We define asymptotic learning outcomes relative to the set of stationary beliefs.

Definition 2 (Learning Outcomes). In state L , correct learning (for type θ_i) denotes the event where $\lambda_t \rightarrow 0^k$ ($\lambda_{i,t} \rightarrow 0$), incorrect learning (for type θ_i) denotes the event where $\lambda_t \rightarrow \infty^k$ ($\lambda_{i,t} \rightarrow \infty$), entrenched disagreement denotes the event where $\lambda_t \rightarrow \{0, \infty\}^k \setminus \{0^k, \infty^k\}$, cyclical learning (for type θ_i) denotes the event where λ_t ($\lambda_{i,t}$) does not converge, and mixed learning denotes the event where $\lambda_{i,t}$ converges for some social types but not others. The definitions are analogous in state R .

When all social types have the same limit belief, we refer to this as an *agreement* outcome. *Entrenched disagreement* occurs when different types converge to different limit beliefs; throughout the paper, when we say ‘disagreement’ we are referring to ‘entrenched disagreement’.²¹ Learning is *complete* if correct learning occurs almost surely and is *path-dependent* if multiple learning outcomes arise with positive probability—for example, correct

²¹Types’ beliefs can also differ when beliefs do not converge i.e. cyclical or mixed learning.

and incorrect learning. In correctly specified environments, the Martingale Convergence Theorem rules out incorrect, cyclical, and mixed learning, and entrenched disagreement. This is not the case in misspecified environments.

3.3 Stability of Learning Outcomes

In this section, we derive conditions for the likelihood ratio to converge to each stationary belief with positive probability. To do so, we first characterize the behavior of the likelihood ratio when it is in a neighborhood of a stationary belief. Building on results on the local stability of nonlinear stochastic difference equations developed in [Smith and Sørensen \(2000\)](#), we establish necessary and sufficient conditions for the likelihood ratio to converge to this stationary belief with positive probability from nearby beliefs, which we refer to as *local stability* ([Theorem 1](#)). We then determine when the likelihood ratio converges to a locally stable belief with positive probability from any initial belief, which we refer to as *global stability* ([Theorems 2 and 3](#)). This ensures that our characterization holds independently of the initial belief.

Our approach builds on techniques used in [Bohren \(2016\)](#) to characterize asymptotic learning outcomes when there is a single type with a misspecified model of the share of a-tarkic types. Our key technical innovations are to allow for multiple types and to characterize conditions for entrenched disagreement. Relative to [Bohren \(2016\)](#), establishing the global stability of disagreement outcomes and belief convergence with multiple types requires novel and different techniques.

Local Stability. A learning outcome is *locally stable* if the likelihood ratio converges to it with positive probability from nearby beliefs and is *unstable* if the likelihood ratio almost surely does not converge to it.

Definition 3 (Local Stability). λ^* is locally stable if there exists an $\varepsilon > 0$ and neighborhood $B_\varepsilon(\lambda^*)$ such that $Pr(\lambda_t \rightarrow \lambda^*) > 0$ for $\lambda_1 \in B_\varepsilon(\lambda^*)$ and is unstable if $Pr(\lambda_t \rightarrow \lambda^*) = 0$ for all $\lambda_1 \in (0, \infty)^k$.

Our characterization of local stability depends on the sign of the average update from an action, weighted by the true probability of each action. In state ω and at belief λ , this is equal to

$$\gamma_i(\omega, \lambda) \equiv \sum_{a \in \mathcal{A}} \psi(a|\omega, \lambda) \log \left(\frac{\hat{\psi}_i(a|R, \lambda)}{\hat{\psi}_i(a|L, \lambda)} \right)$$

for social type θ_i . This expression has two natural interpretations. First, at interior beliefs, it corresponds to the expected change in the log likelihood ratio. Second, it is the difference between (i) the Kullback-Leibler divergence from type θ_i 's subjective action distribution in state L , $\hat{\psi}_i(\cdot|L, \lambda)$, to the true action distribution in state ω , $\psi(\cdot|\omega, \lambda)$, and (ii) the Kullback-Leibler divergence from θ_i 's subjective action distribution in state R , $\hat{\psi}_i(\cdot|R, \lambda)$, to the true action distribution in state ω , $\psi(\cdot|\omega, \lambda)$. If θ_i 's subjective action distribution in state L is closer to the true action distribution than θ_i 's subjective action distribution in state R , then this difference is negative and θ_i 's log likelihood ratio moves towards state L in expectation; otherwise, it moves towards state R .

We show that the sign of each component of $\gamma(\omega, \lambda) \equiv (\gamma_i(\omega, \lambda))_{i=1}^k$ at a stationary belief

determines whether the belief is locally stable. Let

$$\Lambda_i(\omega) \equiv \{\boldsymbol{\lambda} \in \{0, \infty\}^k \mid \gamma_i(\omega, \boldsymbol{\lambda}) < 0 \text{ if } \lambda_i = 0 \text{ and } \gamma_i(\omega, \boldsymbol{\lambda}) > 0 \text{ if } \lambda_i = \infty\}, \quad (3)$$

denote the set of certain beliefs at which $\gamma_i(\omega, \boldsymbol{\lambda})$ is negative if social type θ_i believes the state is L and positive if θ_i believes the state is R . This is the first expression for our learning characterization. [Theorem 1](#) establishes that all beliefs in $\Lambda(\omega) \equiv \cap_{i=1}^k \Lambda_i(\omega)$ are locally stable and, subject to a minor technical condition, beliefs not in $\Lambda(\omega)$ are unstable.

Theorem 1 (Locally Stable Beliefs). *Assume [Assumptions 1 to 4](#). Then $\boldsymbol{\lambda}^* \in \{0, \infty\}^k$ is locally stable in state ω if $\boldsymbol{\lambda}^* \in \Lambda(\omega)$ and is unstable if $\boldsymbol{\lambda}^* \notin \Lambda(\omega)$ and $\gamma_i(\omega, \boldsymbol{\lambda}^*) \neq 0$ for some θ_i with $\boldsymbol{\lambda}^* \notin \Lambda_i(\omega)$. All $\boldsymbol{\lambda}^* \in [0, \infty]^k \setminus \{0, \infty\}^k$ are unstable.*

The intuition is as follows. Consider certain belief $\boldsymbol{\lambda}^* \in \{0, \infty\}^k$ and suppose $\boldsymbol{\lambda}^* \in \Lambda(\omega)$. By the continuity of $\gamma_i(\omega, \boldsymbol{\lambda})$ at $\boldsymbol{\lambda}^*$, if $\lambda_i^* = 0$, then $\gamma_i(\omega, \boldsymbol{\lambda}) < 0$ in a neighborhood of $\boldsymbol{\lambda}^*$ and if $\lambda_i^* = \infty$, then $\gamma_i(\omega, \boldsymbol{\lambda}) > 0$. Therefore, in expectation, the log likelihood ratio moves towards $\boldsymbol{\lambda}^*$ from nearby beliefs. By similar reasoning, if $\boldsymbol{\lambda}^* \notin \Lambda(\omega)$, the log likelihood ratio moves away from $\boldsymbol{\lambda}^*$ at nearby beliefs (provided $\gamma_i(\omega, \boldsymbol{\lambda}^*) \neq 0$ for some θ_i with $\boldsymbol{\lambda}^* \notin \Lambda_i(\omega)$).

The set $\Lambda(\omega)$ has a natural interpretation: it corresponds to the set of strict Berk-Nash equilibria ([Esponda and Pouzo 2016](#)).²² At each $\boldsymbol{\lambda}^* \in \Lambda(\omega)$, each social type places probability one on the state that generates a perceived action distribution closest to the observed action distribution—that is, the state that minimizes the Kullback-Leibler divergence from the type’s perceived action distribution to the true action distribution when all types act optimally with respect to $\boldsymbol{\lambda}^*$. Given the strict inequalities in [Eq. \(3\)](#), this state is unique for each type and the equilibrium is strict. Therefore, [Theorem 1](#) establishes that all strict Berk-Nash equilibria are locally stable. [Esponda et al. \(2019\)](#); [Fudenberg et al. \(2020\)](#) establish a similar result in an individual learning setting that allows for non-myopic agents and a more general state space.

When $\gamma_i(\omega, \boldsymbol{\lambda}^*) = 0$ for a type, the perceived action distributions in each state are equally close to the true action distribution. At a stationary belief $\boldsymbol{\lambda}^* \in \{0, \infty\}^k$, if $\gamma_i(\omega, \boldsymbol{\lambda}^*) = 0$ for all social types with $\boldsymbol{\lambda}^* \notin \Lambda_i(\omega)$, we cannot determine whether $\boldsymbol{\lambda}^*$ is stable or unstable from the sign of $\gamma_i(\omega, \boldsymbol{\lambda}^*)$ (these correspond to weak Berk-Nash equilibria). Conditions for stability in this case significantly complicate the analysis without adding much economic insight. Going forward, we focus on learning environments in which this does not occur.

Definition 4 (Identified at Certainty). *A learning environment is identified at certainty if $\gamma_i(\omega, \boldsymbol{\lambda}) \neq 0$ for all $\theta_i \in \Theta_S$, $\boldsymbol{\lambda} \in \{0, \infty\}^k$ and $\omega \in \{L, R\}$.*

The set of learning environments that are identified at certainty is *generic* in the set of environments that satisfy [Assumptions 1 to 4](#). Note that all correctly specified environments that satisfy [Assumptions 1 to 4](#) are identified at certainty. When the learning environment is identified at certainty, [Theorem 1](#) simplifies to the following corollary.

Corollary 1. *Assume [Assumptions 1 to 4](#). If the learning environment is identified at certainty, then $\boldsymbol{\lambda}^* \in \{0, \infty\}^k$ is locally stable in state ω if and only if $\boldsymbol{\lambda}^* \in \Lambda(\omega)$. All*

²²By strict we mean equilibria in which there is a unique vector of stationary beliefs that minimizes the Kullback-Leibler divergence from the subjective to the true action distribution at the given equilibrium strategy profile.

$\lambda^* \in [0, \infty]^k \setminus \{0, \infty\}^k$ are unstable.

In other words, if the likelihood ratio converges for all social types, then it must converge to a limit random variable whose support lies in $\Lambda(\omega)$. This reduces characterizing local stability to determining the sign of $\gamma_i(\omega, \lambda)$ for each social type at certain beliefs. It is straightforward to do so in applications, as demonstrated in the following example.

Example 1 (Partisan Bias, cont.). From the perceived and true probabilities of each action derived in [Section 3.1](#), $\gamma_1(L, 0)$ is equal to

$$\underbrace{(\pi(\theta_1) + \pi(\theta_2)F^L(.5^{\frac{1}{\nu}})) \log \frac{\pi(\theta_1) + \pi(\theta_2)F^R(.5)}{\pi(\theta_1) + \pi(\theta_2)F^L(.5)}}_{L\text{-action}} + \underbrace{\pi(\theta_2)(1 - F^L(.5^{\frac{1}{\nu}})) \log \frac{1 - F^R(.5)}{1 - F^L(.5)}}_{R\text{-action}}.$$

The construction of $\gamma_1(L, \infty)$ is analogous. When the bias is small (ν is close to one), both $\gamma_1(L, 0)$ and $\gamma_1(L, \infty)$ are negative, and $\Lambda(L) = \{0\}$. As the bias grows, R actions occur more frequently and both expressions increase. For intermediate levels of bias, $\gamma_1(L, 0)$ is positive, $\gamma_1(L, \infty)$ is negative, and $\Lambda(L) = \emptyset$. When the bias is sufficiently large, both expressions are positive and $\Lambda(L) = \{\infty\}$. See [Online Appendix B.1](#) for this derivation.

Global Stability. We are interested in a characterization of asymptotic learning that is independent of the initial belief. Therefore, we need a stronger notion of stability than local stability. We say a learning outcome is *globally stable* if the likelihood ratio converges to this outcome with positive probability, from *any* initial belief.

Definition 5 (Global Stability). λ^* is globally stable if for any initial belief $\lambda_1 \in (0, \infty)^k$, $Pr(\lambda_t \rightarrow \lambda^*) > 0$.

Clearly, the set of globally stable learning outcomes is a subset of the set of locally stable learning outcomes. It remains to establish when local stability implies global stability.

Aligned signals and preferences ([Assumptions 1](#) and [2](#)) guarantee that there exist actions that move the beliefs of all types in the same direction. Therefore, from any initial belief, it is possible to construct a finite sequence of actions that occur with positive probability and move the likelihood ratio arbitrarily close to an agreement outcome. Once the likelihood ratio is in a neighborhood of the agreement outcome, local stability establishes convergence. Therefore, global stability follows immediately from local stability for an *agreement* outcome.

Theorem 2 (Global Stability of Agreement). Consider a learning environment that is identified at certainty and satisfies [Assumptions 1](#) to [4](#). Agreement outcome $\lambda^* \in \{0^k, \infty^k\}$ is globally stable in state ω if and only if $\lambda^* \in \Lambda(\omega)$.²³

Given [Theorem 2](#), deriving $\Lambda(\omega)$ is the only calculation necessary to determine whether correct or incorrect learning occur with positive probability: these learning outcomes occur with positive probability if and only if the corresponding limit beliefs are in $\Lambda(\omega)$. Further, this result fully characterizes global stability when there is a single social type.

Global stability does not immediately follow from local stability for *disagreement* outcomes, as it may not be possible to separate the beliefs of different types and reach a

²³When a learning environment is not identified at certainty, our proof establishes that any agreement outcome $\lambda^* \in \Lambda(\omega)$ is globally stable in state ω .

neighborhood of the disagreement outcome. For example, the beliefs of two similar types can remain close together when starting from a common prior, even if disagreement is possible when they start with very different priors. The second expression for our learning characterization—*maximal accessibility*—provides a sufficient condition to separate beliefs and push the likelihood ratio arbitrarily close to a given disagreement outcome. We first define a partial order on how types update following a_1 , which decreases the likelihood ratio, and a_M , which increases it.

Definition 6 (Maximal R-order). *Define the maximal R-order $\succeq_{\lambda}$ over Θ at belief λ as $\theta_i \succeq_{\lambda} \theta_j$ iff θ_i interprets a_1 and a_M as stronger evidence of state R than θ_j ,*

$$\frac{\hat{\psi}_j(a|R, \lambda)}{\hat{\psi}_j(a|L, \lambda)} \leq \frac{\hat{\psi}_i(a|R, \lambda)}{\hat{\psi}_i(a|L, \lambda)} \quad (4)$$

for $a \in \{a_1, a_M\}$, with strict order $\succ_{\lambda}$ if Eq. (4) holds with strict inequality for at least one action $a \in \{a_1, a_M\}$.

We use this order to define maximal accessibility. As the number of possible disagreement outcomes increases with the number of social types, so does the complexity of such a property; we present the case of two social types here and relegate the case of more than two social types to [Online Appendix D](#).

Definition 7 (Maximal Accessibility ($k = 2$)). *Disagreement outcome $(0, \infty)$ is maximally accessible if $\theta_2 \succ_{(0,0)} \theta_1$ or $\theta_2 \succ_{(\infty, \infty)} \theta_1$, and disagreement outcome $(\infty, 0)$ is maximally accessible if $\theta_1 \succ_{(0,0)} \theta_2$ or $\theta_1 \succ_{(\infty, \infty)} \theta_2$.*

It is straightforward to verify maximal accessibility in applications by evaluating Eq. (4) at beliefs $(0, 0)$ or (∞, ∞) .

When a disagreement outcome is maximally accessible, for any initial belief, there exists a finite sequence of actions that moves beliefs to a neighborhood of the disagreement outcome. To see this, suppose $\theta_2 \succ_{(0,0)} \theta_1$. As discussed above, the likelihood ratio enters a neighborhood of $(0, 0)$ with positive probability from any initial belief. Given $\theta_2 \succ_{(0,0)} \theta_1$, we can construct a sequence of a_1 and a_M actions that decrease θ_1 's beliefs and increase θ_2 's beliefs in a neighborhood of $(0, 0)$. We show that this guarantees that there exists a finite sequence of actions that occurs with positive probability and moves beliefs from a neighborhood of $(0, 0)$ to a neighborhood of $(0, \infty)$. Once the likelihood ratio is sufficiently close to the disagreement outcome, local stability establishes convergence. Therefore, the global stability of a disagreement outcome follows from maximal accessibility and local stability.²⁴

Theorem 3 (Global Stability of Disagreement ($k = 2$)). *Consider a learning environment that satisfies [Assumptions 1 to 4](#). If disagreement outcome $\lambda^* \in \{(0, \infty), (\infty, 0)\}$ is in $\Lambda(\omega)$*

²⁴Maximal accessibility is simple and tractable, but it can be restrictive, especially in models with large action spaces. In [Theorem 7](#) in [Appendix A.3](#), we establish a more general condition that uses all actions to separate beliefs, which we call *separability* ([Definition 8](#)). It is more cumbersome to verify but holds for a larger set of learning environments. Another simple sufficient condition to separate beliefs is: $(0, 0) \in \Lambda_1(\omega) \setminus \Lambda_2(\omega)$ or $(\infty, \infty) \in \Lambda_2(\omega) \setminus \Lambda_1(\omega)$ for $(0, \infty)$, with an analogous condition for $(\infty, 0)$. It can be directly verified from the construction of $\Lambda(\omega)$ but will not hold when both agreement outcomes are locally stable.

and maximally accessible, then λ^* is globally stable in state ω .

See [Section 4.3](#) for an application that uses maximal accessibility.

Mixed Learning. Next, we establish conditions to rule out mixed learning. Suppose $\omega = L$ and consider the mixed outcome in which θ_1 has correct learning and θ_2 has cyclical learning. If either $(0, 0) \in \Lambda_2(L)$ or $(0, \infty) \in \Lambda_2(L)$, then $\langle \lambda_{2,t} \rangle$ converges with positive probability when $\lambda_1^* = 0$, and hence, almost surely cannot oscillate infinitely often. Therefore, in order for this mixed outcome to arise with positive probability, it must be that $(0, 0) \notin \Lambda_2(L)$ and $(0, \infty) \notin \Lambda_2(L)$, i.e. in a neighborhood of $(0, 0)$ or $(0, \infty)$, θ_2 's beliefs drift away from the outcome. Generalizing this intuition, let $\Lambda_M(\omega)$ denote the set of mixed outcomes in which there are no locally stable beliefs for the non-convergent types in state ω . When $k = 2$, this corresponds to

$$\Lambda_M(\omega) \equiv \{(\lambda_i^*, \theta_i) \in \{0, \infty\} \times \{\theta_1, \theta_2\} \mid (\lambda_i^*, 0) \notin \Lambda_{-i}(\omega) \text{ and } (\lambda_i^*, \infty) \notin \Lambda_{-i}(\omega)\}, \quad (5)$$

where (λ_i^*, θ_i) denotes the mixed outcome in which θ_i 's beliefs converge to λ_i^* and θ_{-i} 's beliefs do not converge. As in the case of disagreement, mixed learning is more involved when there are more than two social types; we relegate this case to [Online Appendix D](#) and define $\Lambda_M(\omega) = \emptyset$ when $k = 1$. The following result establishes that if a mixed outcome is not in $\Lambda_M(\omega)$, then it almost surely does not occur.

Lemma 4 (Unstable Mixed Outcomes ($k = 2$)). *Consider a learning environment that is identified at certainty and satisfies [Assumptions 1 to 4](#). If $(\lambda_i^*, \theta_i) \notin \Lambda_M(\omega)$, then $\Pr(\lambda_{i,t} \rightarrow \lambda_i^* \text{ and } \lambda_{-i,t} \text{ does not converge}) = 0$ in state ω .*

Therefore, $\Lambda_M(\omega) = \emptyset$ rules out mixed learning in state ω .²⁵ It is straightforward to derive $\Lambda_M(\omega)$ from $\Lambda_i(\omega)$. See [Example 2](#) and [Sections 4.2](#) and [4.3](#) for applications that show $\Lambda_M(\omega)$ is empty.

[Lemma 4](#) does not determine whether a mixed outcome in $\Lambda_M(\omega)$ arises with positive probability. Doing so presents a novel challenge relative to stationary learning outcomes, as it requires characterizing the movement of the convergent type's belief across all possible beliefs for the non-convergent type. We leave this question for future research.

3.4 Learning Characterization

We use the stability results in the previous section to characterize the set of asymptotic learning outcomes in each state. The final piece of the characterization involves showing when the likelihood ratio converges almost surely for all social types ([Lemma 7](#) in [Appendix A.5](#)). This establishes our main result.

Theorem 4 (Learning Characterization ($k \leq 2$)). *Consider a learning environment that is identified at certainty and satisfies [Assumptions 1 to 4](#). When the state is L :*

- (i) **Correct** learning occurs with positive probability if and only if $0^k \in \Lambda(L)$.
- (ii) **Incorrect** learning occurs with positive probability if and only if $\infty^k \in \Lambda(L)$.

²⁵A simple sufficient condition for $\Lambda_M(\omega) = \emptyset$ is that both agreement outcomes or both disagreement outcomes are in $\Lambda(\omega)$. When $k > 2$, an analogous condition rules out mixed outcomes in which one type has cyclical learning. Ruling out mixed outcomes in which more than one type has cyclical learning requires joint conditions on $\Lambda_i(\omega)$ for the non-convergent types.

- (iii) **Entrenched Disagreement** occurs with positive probability if $\Lambda(L)$ contains a maximally accessible disagreement outcome and almost surely does not occur if $\Lambda(L)$ contains no disagreement outcome. Each maximally accessible disagreement outcome in $\Lambda(L)$ occurs with positive probability.
- (iv) **Cyclical Learning** occurs almost surely if and only if $\Lambda(L) = \emptyset$ when $k = 1$. When $k = 2$, cyclical learning occurs almost surely for both social types if $\Lambda(L) = \emptyset$ and $\Lambda_M(L) = \emptyset$, occurs almost surely for at least one social type if $\Lambda(L) = \emptyset$, and almost surely does not occur if $\Lambda(L)$ contains an agreement outcome or maximally accessible disagreement outcome and $\Lambda_M(L) = \emptyset$.

An analogous result holds in state R .

See [Online Appendix D](#) for the analogous result for $k > 2$ social types, using the generalized definitions of maximal accessibility and $\Lambda_M(\omega)$.

The conditions for correct and incorrect learning are tight: these learning outcomes arise if and only if the respective limit beliefs are in $\Lambda(\omega)$. For a disagreement outcome, there is a wedge between the sufficient conditions for it to arise—maximal accessibility—and not arise—instability.²⁶ This wedge disappears if all locally stable disagreement outcomes are maximally accessible (see [Section 4.3](#) for an example.) If multiple learning outcomes occur with positive probability, then learning is path-dependent—agents become certain of different states following different histories (again see [Section 4.3](#)). [Theorem 4](#) also determines when learning is complete, as stated in the following corollary (using the expanded definition of $\Lambda_M(\omega)$ in [Online Appendix D](#) when $k > 2$.)

Corollary 2 (Complete Learning). *If $\Lambda(L) = \{0^k\}$ and $\Lambda_M(L) = \emptyset$, correct learning occurs almost surely in state L . An analogous result holds in state R .*

It follows immediately from the martingale property of the likelihood ratio that these conditions are satisfied in a correctly specified environment.²⁷ As we show in [Example 1](#) at the end of this subsection, they can also hold in misspecified environments.

An important feature of [Theorem 4](#) is that the characterization requires calculations at a *finite* set of beliefs—in particular, determining $\Lambda(\omega)$, $\Lambda_M(\omega)$ and maximal accessibility only requires computing $\psi(a|\omega, \lambda)$ and $\hat{\psi}(a|\omega, \lambda)$ at stationary beliefs $\lambda \in \{0, \infty\}^k$. Since action choices depend on beliefs, in principle, determining the asymptotic properties of the likelihood ratio could require characterizing its behavior across the entire belief space. The fact that this is not necessary makes [Theorem 4](#) straightforward to use in applications.

Several economic insights emerge from [Theorem 4](#). First, belief convergence forces action convergence: each type eventually settles on an action if and only if its beliefs converge. It follows that the limit action choice is efficient if and only if learning is correct. If learning is incorrect for a type, the type will choose inefficient actions infinitely often, and if learning is cyclical, the type will choose both efficient and inefficient actions infinitely often.²⁸

²⁶[Theorem 7](#) in [Appendix A.3](#) presents a weaker condition—separability—to establish the global stability of disagreement outcomes.

²⁷The likelihood ratio is a martingale in state L and the log function is concave. Together with [Assumptions 2](#) and [3](#), this implies $\gamma_i(L, \lambda) < 0$ for all $\lambda \in [0, \infty]^k$, so $\Lambda(L) = \{0^k\}$ and $\Lambda_M(L) = \emptyset$.

²⁸In the proof of [Theorem 4](#), we show that if the likelihood ratio for a type does not converge, then it enters a neighborhood of each certain belief infinitely often.

Second, model misspecification gives rise to two potential sources of entrenched disagreement in society. Model heterogeneity can lead to entrenched disagreement *within* a population due to differing interpretations of a common history. A signal realization that, for instance, a vaccine is safe or a politician is corrupt can cause agents to update in opposite directions based on their belief about the credibility of the source. This arises even though preferences and signals are aligned, so that agents have a common interpretation of whether one signal realization or action choice is relatively more likely in a given state than another. Therefore, model heterogeneity can explain how connected populations observing shared sources can perpetually disagree.²⁹ Path-dependent learning can also lead to entrenched disagreement, but *across* populations that observe different histories rather than within populations that have different models.³⁰ This can explain how separate populations with similar models can come to have polarized ingrained views. When both sources are present, then within-population disagreement can vary across populations depending on whether the observed history has a common or polarizing interpretation—in other words, the order in which information arrives will impact the level of disagreement within a population. For example, in a cognitive hierarchy learning model, agreement emerges following some histories while disagreement emerges following others (see the working paper version of this article for this analysis (Bohren and Hauser 2021a)).

While path-dependent learning—and hence, across-population disagreement—can also occur in correctly specified social learning environments (Banerjee 1992; Bikhchandani et al. 1992), learning is incomplete at all but at most one possible limit beliefs (the degenerate belief on the correct state). Since agents remain uncertain about the state, disagreement between populations can be easily resolved by introducing a common source. In contrast, misspecified learning gives rise to path-dependent learning with multiple degenerate limit beliefs. As different populations come to place high probability on different states, it becomes increasingly difficult to reconcile prolonged disagreement with common information.

The following examples demonstrate how to apply Theorem 4.

Example 1 (Partisan Bias, cont.). Applying Theorem 4 to the characterization of $\Lambda(L)$ above establishes that correct learning occurs almost surely for mild partisan bias, cyclical learning arises almost surely for intermediate levels, and incorrect learning occurs almost surely for severe partisan bias (see Proposition 5 in Online Appendix B.1.)

Example 2 (Partisan Bias and Unawareness). In the presence of model heterogeneity, agents have a more complex inference problem—in order to be correctly specified, an agent must know the form and frequency of different biases in the population. In this example, we show that when a type accurately interprets signals but does not account for others’ partisan bias, it can be just as wrong as the partisan types. Augment Example 1 to include non-partisan agents who have correctly specified signal distributions. Analogous to the partisan types, there is a social and an autarkic non-partisan type. Neither social type is aware that some agents have a different subjective signal distribution. An analogous derivation of $\Lambda(L)$ establishes that the non-partisan social type has the same long-run learning outcome as the partisan

²⁹For example, Gagnon-Bartsch (2016) show that taste projection can lead to entrenched disagreement.

³⁰Earlier work establishing that path-dependent learning with multiple degenerate limit beliefs arises for specific misspecified models includes Rabin and Schrag (1999) (confirmation bias), Epstein et al. (2010) (overreaction) and Eyster and Rabin (2010); Bohren (2016) (naive social learning).

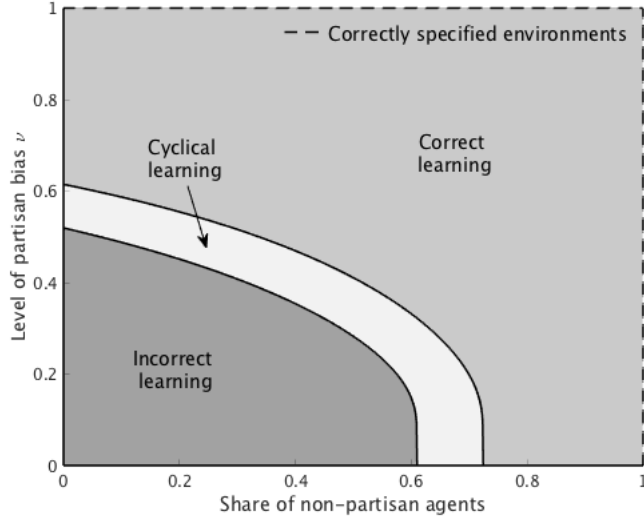


FIGURE 1. Partisan Bias and Unawareness
 $(\omega = L, F^L(s) = 2s - s^2, F^R(s) = s^2, \text{share autarkic types}=0.1)$

social type—correct learning occurs almost surely for both types when there is a small share of partisan agents or mild bias, cyclical learning arises almost surely for intermediate levels, and incorrect learning occurs almost surely when there is a large share of very biased partisan agents. In other words, the presence of a large share of unaccounted for partisan types can prevent types who correctly interpret signals from making efficient choices. On the other hand, the presence of many agents who correctly interpret signals can help even severely biased agents adopt the efficient action. *Fig. 1* illustrates these three learning regions. See [Online Appendix B.2](#) for the analysis.

3.5 Robustness of Learning.

We next establish that learning is generally robust to the details of the environment. Specifically, correctly specified environments are robust to some misspecification in that learning is complete when social types have approximately correct models. Misspecified environments are also robust, in that except for knife-edge cases that separate different learning regions, nearby misspecified environments will have the same set of learning outcomes. To explore robustness, we first fix a learning environment and show that the same learning outcomes arise in environments with sufficiently similar equilibrium distributions of actions. We then use this result to establish that learning is complete when social types' models of inference are sufficiently close to correctly specified. Taken together, these results strengthen the applicability of correctly specified environments to real-world settings with mild biases and establish that small errors in measuring or forecasting more severe biases will not significantly alter the predicted learning outcomes.

For the first result, we focus on learning environments in which [Theorem 4](#) fully characterizes learning outcomes—that is, environments that satisfy [Assumptions 1](#) to [4](#), are identified at certainty, and in which all locally stable disagreement outcomes are maximally accessible and mixed learning almost surely does not arise, $\Lambda_M(L) = \Lambda_M(R) = \emptyset$ (using the general-

ized definitions of maximal accessibility, $\Lambda_M(\omega)$, and the property that either $\Lambda(\omega) = \emptyset$ or $\mathcal{G}(\omega)$ has no cycles from [Online Appendix D](#) when $k > 2$). We refer to such environments as *regular*.³¹ Fixing a regular learning environment, [Theorem 5](#) establishes that any learning environment with sufficiently similar action distributions has the same set of long-run learning outcomes, where we use the total variation distance to measure the closeness of distributions.³²

Theorem 5 (Robustness). *Let (Θ^*, π^*) be a regular learning environment with set of stable learning outcomes $\Lambda^*(\omega)$ in state ω . There exists a $\delta > 0$ such that any learning environment (Θ, π) that satisfies [Assumptions 1 to 4](#), has the same number of social types, and is sufficiently close to (Θ^*, π^*) in terms of the true and perceived distributions over actions in that $\|\psi^*(\cdot|\omega, \boldsymbol{\lambda}) - \psi(\cdot|\omega, \boldsymbol{\lambda})\| < \delta$, $\|\hat{\psi}_i^*(\cdot|L, \boldsymbol{\lambda}) - \hat{\psi}_i(\cdot|L, \boldsymbol{\lambda})\| < \delta$ and $\|\hat{\psi}_i^*(\cdot|R, \boldsymbol{\lambda}) - \hat{\psi}_i(\cdot|R, \boldsymbol{\lambda})\| < \delta$ for all $\boldsymbol{\lambda} \in \{0, \infty\}^k$ and $i = 1, \dots, k$, has the same set of long-run learning outcomes as (Θ^*, π^*) in state ω i.e. given sets $\Lambda(\omega)$ and $\Lambda_M(\omega)$ for (Θ, π) , (i) $\Lambda(\omega) = \Lambda^*(\omega)$; (ii) $\Lambda_M(\omega) = \emptyset$; and (iii) all disagreement outcomes in $\Lambda(\omega)$ are maximally accessible.*

This result follows from the continuity of $\gamma(\omega, \boldsymbol{\lambda})$ in ψ and $\hat{\psi}_i$, which implies that, provided $\gamma_i(\omega, \boldsymbol{\lambda}) \neq 0$, it does not change sign at belief $\boldsymbol{\lambda}$ when ψ and $\hat{\psi}_i$ are perturbed. Therefore, the locally stable set $\Lambda(\omega)$ remains the same, the mixed outcome set $\Lambda_M(\omega)$ remains empty, and locally stable disagreement outcomes remain globally stable. For example, in [Fig. 1](#) we see that for any environment that is identified at certainty (i.e. does not lie on the two blue lines dividing the learning regions), nearby environments have the same learning outcome.

Correctly specified environments are regular and have complete learning. Therefore, fixing a correctly specified environment, [Theorem 5](#) establishes that learning is complete in misspecified environments with similar action distributions. We use this to establish robustness when social types' models of inference are approximately correct. In order to compare models of inference without needing to define a complicated metric over types, we vary social types' models of inference while holding fixed other aspects of the learning environment—specifically, the preferences of social types and the preferences and models of autarkic and noise types. Formally, (Θ, π) is *structurally equivalent* to (Θ^*, π^*) if $|\Theta_S| = |\Theta_S^*|$, $u_i = u_i^*$ for $i = 1, \dots, k$, $\Theta_A = \Theta_A^*$, $\Theta_N = \Theta_N^*$, and $\pi(\theta_i) = \pi^*(\theta_i^*)$ for $i = 1, \dots, n$.³³ [Theorem 6](#) establishes that learning is complete when social types have subjective type and signal distributions close enough to the true distributions, where again we use the total

³¹Within the set of environments that satisfy [Assumptions 1 to 4](#) and are identified at certainty, this includes all correctly specified environments and all environments with a single social type, as well as the environments with multiple social types in [Example 2](#) and [Sections 4.2](#) and [4.3](#). Robustness also holds locally for environments with $\Lambda_M(\omega) \neq \emptyset$ or disagreement outcomes that are not globally stable.

³²The total variation distance between distributions ψ and ψ' is $\|\psi(\cdot|\omega, \boldsymbol{\lambda}) - \psi'(\cdot|\omega, \boldsymbol{\lambda})\| \equiv \max_{A \subset \mathcal{A}} |\sum_{a \in A} (\psi(a|\omega, \boldsymbol{\lambda}) - \psi'(a|\omega, \boldsymbol{\lambda}))|$. Pinsker's inequality implies that our results also hold for the Kullback-Leibler divergence, which has been used to study robustness in subsequent work ([Frick et al. 2020a,b](#)).

³³Recall that a correctly specified environment requires all social types to be correctly specified, but autarkic types may be misspecified. Therefore, misspecified environments have structurally equivalent correctly specified environments.

variation distance to measure the closeness of distributions.^{34,35}

Theorem 6 (Robustness). *Let (Θ^*, π^*) be a correctly specified environment that satisfies Assumptions 2 and 3. There exists a $\delta > 0$ such that in any structurally equivalent misspecified environment (Θ, π) that satisfies Assumptions 1 to 4 and in which social types have sufficiently correct models of inference in that $\|\hat{\pi}_i - \pi\| < \delta$, $\|\hat{F}_i^L - F^L\| < \delta$ and $\|\hat{F}_i^R - F^R\| < \delta$ for all $\theta_i \in \Theta_S$, learning is complete.*

A similar result holds for an individual type that has an approximately correct model of inference, regardless of the degree to which other types are misspecified. Therefore, agents do not need to know exactly how their misspecified peers behave in order to accurately learn from their choices. Beyond Theorems 5 and 6, we can use Theorem 4 for a precise characterization of how large perturbations to the environment can be before they alter the set of learning outcomes. For example, in Fig. 1, we see that if the share of non-partisan types is 0.8, then learning is complete for any level of bias $\nu > 0$, whereas if the share of non-partisan types is 0.5, then learning is complete when $\nu > 0.4$.

These robustness results may not seem surprising, since Bayes rule is continuous. But in an infinite horizon setting, nearby models with small per-period differences in belief updating have the potential to aggregate to very different limit beliefs. For example, if agents’ models are equidistant from the truth in either state at a certain belief—and therefore, the environment is not identified at certainty—then arbitrarily small perturbations alter the set of learning outcomes that arise. In Fig. 1, this is the case for the environments described by the parameters (q, ν) that trace out the boundaries dividing the different learning regions. At these parameters, there is an abrupt shift from complete learning to cyclical learning or cyclical learning to incorrect learning. But this corresponds to a measure zero set.

More generally, uniformly informative actions ensure that identification at certainty is a generic property of the learning environments we consider—and therefore, so is robustness. Further, it ensures that all correctly specified environments are identified at certainty, and therefore, all correctly specified environments are robust. The same holds true in Bohren (2016), whose robustness result is a special case of Theorem 6. In contrast, when actions are not uniformly informative, Frick et al. (2020b) show that a failure of robustness occurs in a correctly specified environment that is not identified at certainty. Identification at certainty fails because actions (or signals) are perceived to be uninformative at certainty—in contrast to our framework, in which individual actions are perceived to be informative but identification at certainty fails because the Kullback-Leibler divergence from a type’s perceived action distribution to the true action distribution is the same in each state (a

³⁴In a slight abuse of notation, we simultaneously let $\|\cdot\|$ denote the total variation distance between two probability measures over the type space, i.e. $\|\pi - \pi'\| \equiv \sup_{X \subset \{1, \dots, n\}} |\sum_{i \in X} (\pi(\theta_i) - \pi'(\theta'_i))|$, and the distance between two signal c.d.f.s with common support, i.e. $\|F - F'\| \equiv \sup_{X \in \mathcal{B}_S} |P_F(X) - P_{F'}(X)|$, where $\mathcal{B}_S$ denotes the Borel σ -algebra on $\mathcal{S}$ and $P_F(X)$ denotes the probability of $X \in \mathcal{B}_S$ under the measure described by c.d.f. F .

³⁵An analogous result holds in settings where agents are very wrong about the type distribution, as long as the types that they believe exist are “close” to the actual types. For example, neither type in Example 2 allows for the possibility of the other—their subjective type distributions have non-overlapping supports. But learning is complete when the partisan type is not too biased, as the partisan type is sufficiently close to the non-partisan type. By defining a suitable measure of distance between types, one could use Theorem 5 to derive an analogous result to Theorem 6 for such settings.

property that cannot arise in a correctly specified model).

Frick et al. (2020a) also demonstrate a failure of robustness in a social learning setting with an infinite state space and privately observed actions. In this environment, action choices are sensitive to small amounts of misspecification and agents with almost correct models can come to place probability one on an incorrect state. Together with our results, this demonstrates that the details of the learning environment are important to consider when exploring robustness. If one wishes to design a robust learning environment, then adjusting these details may be an important lever.

Robustness relates to whether a learning outcome is a weak or strict Berk-Nash equilibrium. In our setting, correct learning is a strict and unique Berk-Nash equilibrium in correctly specified environments. This ensures that correct learning is also the unique Berk-Nash equilibrium in nearby misspecified environments. In contrast, when a correctly specified environment is not identified at certainty—as in Smith and Sørensen (2000); Frick et al. (2020b)—correct learning is a weak Berk-Nash equilibrium and therefore is not robust.

3.6 Discussion

Focus on Asymptotic Learning. We focus on how misspecification affects long-run learning. When using this approach, an important question is whether the long-run is economically relevant. For incorrect learning, cyclical learning, or disagreement, showing these outcomes arise asymptotically establishes that agents are bounded away from efficiency or agreement, *irrespective* of the amount of information that agents observe or the rate of learning. Therefore, the source of these inefficiencies is not a lack of sufficient information to learn the state or the slow pace at which this information arrives. Long-run results also highlight an important distinction between inefficient action choices due to incomplete learning in correctly specified environments versus incorrect learning in misspecified environments. Incomplete learning is fragile and a herd can be overturned by a relatively uninformative piece of information at any point in time. In contrast, when incorrect learning arises, more informative interventions are required to overturn longer incorrect herds.

When correct learning arises asymptotically leaves open important questions such as how quickly actions converge to efficiency. The expression $\gamma(\lambda, \omega)$ also determines the asymptotic rate of learning; we leave further study of this to future work.

Extensions. We outline several possible extensions to the learning framework.

Misaligned Type Spaces. Set $\Lambda(\omega)$ also characterizes locally stable beliefs in misaligned environments. Beyond ruling out confounded learning, the aligned assumptions are used to establish the global stability of agreement—they guarantee that there are actions that uniformly move social types’ beliefs towards both agreement outcomes. If a misaligned environment has an action that uniformly moves all social types’ beliefs towards a stationary belief, then our method can also be used to establish the global stability of this belief.

Signal and Type Distributions. It is a straightforward extension to allow types to observe signals from different distributions, to believe that other types observe signals from different distributions, or to allow the true and/or subjective type distributions to depend on the state. Simply augment the definition of a type to include the additional primitives. The extension to heterogeneous signal distributions allows us to model biases that involve interpersonal comparisons related to the quality of information. For instance, a natural way to model overconfidence is with a type that correctly interprets its own signal but believes other

agents observe signals from a less informative distribution. It is also straightforward to allow multiple pieces of information to arrive at the same time—for example, a finite number of agents act simultaneously or all agents also observe a public signal process (for the latter, see an earlier working paper version of this article (Bohren and Hauser 2019b)). As long as the model reduces to a belief process that satisfies the conditions outlined in Lemma 2, our characterization applies.

Action and State Space. For technical convenience, we assume that the action space is finite. Allowing for a continuous action space would not qualitatively change the analysis. Similar techniques to those we use can be used to analyze a finite state space with more than two states, with the caveats that the definition of an aligned environment is more complicated and the notation is more cumbersome. We use results pertaining to stochastic difference equations in our analysis, which means that generalizing to an infinite state space requires different techniques.

4 Applications

We present three applications to demonstrate how our general framework can be used to address the issues raised in the introduction. First, we show that overreaction—a form of signal misspecification—has a fundamentally different impact when individuals learn from social versus private sources. Second, we examine whether a representative agent model is a good approximation in a setting with heterogeneous levels of naive learning. Finally, we show that entrenched disagreement arises and agreement almost surely does not in a level- k social learning model where different types have fundamentally distinct models of inference. All proofs for the results in this section are in Online Appendix C.

4.1 Overreaction: Individual versus Social Learning

This application demonstrates how our characterization can be used to determine whether the impact of a bias differs when agents learn from private versus social sources. We explore this question in a setting in which individuals overreact to signals.³⁶ We show that overreaction interacts with social learning to create long-run inefficiencies that are not present when agents learn directly from signals. In particular, cyclical learning arises for sufficiently severe levels of overreaction when agents learn from a social source, whereas learning is complete *regardless* of the severity of the bias when agents learn directly from signals. We conclude with a discussion of other models of overreaction and highlight how these different parameterizations differentially impact learning.

We model overreaction as a type who forms beliefs as if it has observed the same signal realization multiple times. The type believes that the private signal is distributed according to $\hat{F}^\omega(s) = F^\omega(\frac{s^\nu}{(1-s)^\nu + s^\nu})$ in state ω , where F^ω is continuous with support $\mathcal{S} = [0, 1]$ and $\nu \in (1, \infty)$ captures the degree of overreaction. This leads to private belief $\frac{\hat{s}(s)}{1-\hat{s}(s)} = (\frac{s}{1-s})^\nu$ following signal realization s . For example, if $\nu = 2$, the signal is double counted—the private belief following realization s corresponds to the correct belief following two realizations s —independent of the direction and strength of the signal realization.

Benchmark: Individual Learning. Given that the level of overreaction is independent of the realized signal, when type θ_1 learns directly from signals, the bias does not alter the

³⁶Overreaction has been widely documented empirically; see Section 2.3 for citations.

sign of $\gamma_1(\omega, \boldsymbol{\lambda})$: as in the correctly specified model, $\gamma_1(\omega, \boldsymbol{\lambda}) < 0$ for all $\boldsymbol{\lambda}$ and learning is complete for any level of overreaction. This follows almost directly from [Berk \(1966\)](#).

Observation 1. *If type θ_1 observes signals directly, learning is complete for any $\nu \in [1, \infty)$.*

Social Learning. When signals are filtered through other agents' actions, the induced level of overreaction to an action depends on the observed action. Therefore, the bias can alter the sign of $\gamma_i(\omega, \boldsymbol{\lambda})$, and hence, the set of learning outcomes. We illustrate this in a learning environment with symmetric preferences and signal distributions, so that there are no inherent asymmetries in the underlying overreaction to signals or the mapping from beliefs to actions. Suppose there are two types of agents: θ_1 is social and θ_2 is autarkic, with $\pi(\theta_1) \in (0, 1)$. Both types have the same level of overreaction, as outlined above, face a decision problem with four actions, $a \in \{a_1, a_2, a_3, a_4\}$ (ordered according to increasing preference in state R) and a symmetric signal, $F^L(s) = 1 - F^R(1 - s)$, and have the same symmetric preferences, i.e. if a_1 is optimal at belief p that the state is R , then a_4 is optimal at belief $1 - p$, and similarly for a_2 and a_3 . Given this symmetry, the agent's decision-rule can be represented as a cutoff rule $p^* \in (0, 1/2)$ such that the agent chooses a_1 if $p \leq p^*$, a_2 if $p \in (p^*, 0.5]$, a_3 if $p \in (0.5, 1 - p^*]$ and a_4 if $p \in (1 - p^*, 1]$. To ensure that moderate actions a_2 and a_3 are chosen for a sufficiently large window of beliefs, assume p^* is sufficiently small so that $\frac{F^L(p^*) - F^R(p^*)}{\log F^L(p^*) - \log F^R(p^*)} < F^R(.5)$. To close the model, assume that θ_1 has a correct subjective type distribution and both types have prior $p_0 = 1/2$.

The following result establishes that social learning causes overreaction to interfere with asymptotic learning. In particular, sufficiently severe overreaction leads to cyclical learning.

Proposition 1. *There exists a cutoff $\bar{\pi} \in (0, 1)$ on the share of social types such that: (i) if $\pi(\theta_1) > \bar{\pi}$, then there exists a cutoff $\bar{\nu}(\pi(\theta_1)) \in (1, \infty)$, which is decreasing in $\pi(\theta_1)$, such that cyclical learning arises almost surely if $\nu > \bar{\nu}(\pi(\theta_1))$ and learning is complete if $\nu < \bar{\nu}(\pi(\theta_1))$; (ii) if $\pi(\theta_1) < \bar{\pi}$, then learning is complete for all $\nu \in [1, \infty)$.*

When beliefs are close to zero, action a_1 is chosen for a larger set of signal realizations than a_4 . Therefore, the overreaction is asymmetric: it is stronger with respect to contradictory a_4 actions than confirmatory a_1 actions. This pulls beliefs away from zero. Similarly, when beliefs are close to infinity, overreaction is stronger with respect to contradictory a_1 actions. For sufficiently severe overreaction, this gives rise to cyclical learning, as illustrated in [Fig. 2](#).

Other Parameterizations. Other approaches to modeling overreaction include [Epstein et al. \(2010\)](#), who model overreaction as a linear updating rule that places negative weight on the prior belief and a weight above one on the correctly specified posterior belief, and [Bushong and Gagnon-Bartsch \(2019\)](#), where an agent underestimates the extent of her reference dependence, which leads her to overreact when recalling past outcomes.³⁷ These different parameterizations are of consequence: when an agent learns directly from signals, [Epstein et al. \(2010\)](#)'s parameterization of overreaction alters updating in an asymmetric way and sufficiently severe overreaction leads to the possibility of incorrect learning. Similarly, incorrect learning can arise in [Bushong and Gagnon-Bartsch \(2019\)](#) when the agent is loss averse, so that she overreacts asymmetrically to losses and gains. In contrast, if the agent is not loss averse, the overreaction is symmetric and, as in our individual learning setting, this

³⁷In [Online Appendix E](#), we show how our framework nests [Epstein et al. \(2010\)](#).

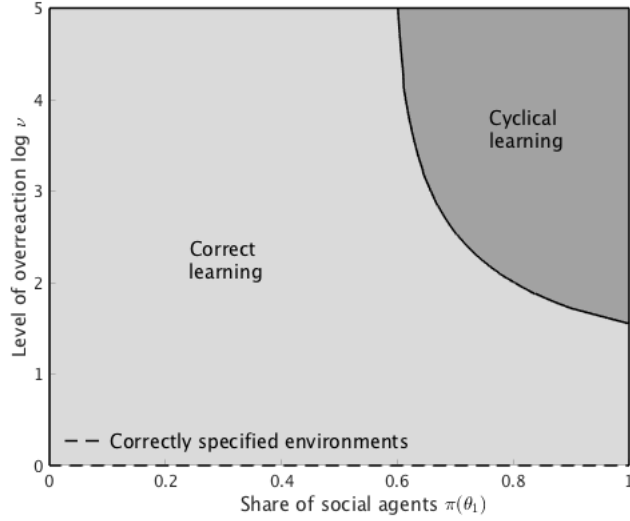


FIGURE 2. Overreaction to social information
 $(F^L = 2s - s^2, F^R = s^2, p^* = .25.)$

does not interfere with learning.

4.2 Naive Learning with Model Heterogeneity

As discussed in the introduction, papers that study model misspecification generally assume that all agents have the same form and level of misspecification. This can be viewed as a representative agent approach, which significantly simplifies the analysis. However, the empirical literature has shown that even when agents have similar biases, they will exhibit it in differing degrees. As this application demonstrates, our characterization can determine whether a representative agent approach is valid in the face of heterogeneity in the sense that the long-run behavior of a representative agent approximates the long-run behavior of heterogeneous agents. This provides a template for evaluating the representative agent approach that is straightforward to apply to other forms of misspecification.

We explore this question in a setting in which agents are naive learners who overestimate the private information reflected in actions. We compare learning in a setting in which agents have heterogeneous levels of naivete to a representative agent setting in which a single type has a level of naivete equal to the average naivete of the population (the latter is a special case of [Bohren \(2016\)](#)). We show that when heterogeneity is small, this representative agent model is a good approximation of the underlying heterogeneous environment in that both environments have the same learning outcomes, whereas when heterogeneity is large, incorrect learning arises with positive probability in the corresponding representative agent model even though both types almost surely learn the correct state.

As in [Bohren \(2016\)](#), we model naive learning as a misspecified belief about the share of autarkic types. Let θ_A denote the autarkic type and assume $\pi(\theta_A) \in (0, 1)$. To capture model heterogeneity, suppose there are two social types, θ_1 and θ_2 , that occur with equal probability, $\pi(\theta_1) = \pi(\theta_2)$. Both social types overestimate the share of autarkic types, with type θ_2 having a more severe bias, $\pi(\theta_A) < \hat{\pi}_1(\theta_A) \leq \hat{\pi}_2(\theta_A) \leq 1$. This leads agents to underestimate the correlation between prior actions. In the representative agent setting,

a single social type believes that the autarkic type occurs with probability $\hat{\pi} \in (\pi(\theta_A), 1]$. We compare heterogeneous environments with biases $(\hat{\pi}_1(\theta_A), \hat{\pi}_2(\theta_A))$ to the corresponding representative agent environment with a bias equal to the average bias in the heterogeneous setting, i.e. $\hat{\pi} = (\hat{\pi}_1(\theta_A) + \hat{\pi}_2(\theta_A))/2$. To close the model, assume that each agent faces a binary decision problem in which she earns a payoff of one from choosing the action that matches the state, $\mathcal{A} = \{L, R\}$ and $u(a, \omega) = \mathbb{1}_{a=\omega}$, all types correctly interpret private signals, social types have correct beliefs about the relative frequency of each social type, and all types have common prior $p_0 = 1/2$.

We first show that the representative agent model is a good approximation when heterogeneity is sufficiently small.

Proposition 2. *Generically, for any average bias $\hat{\pi} \in (\pi(\theta_A), 1]$, there exists an $\varepsilon > 0$ such that if heterogeneity is sufficiently small, $|\hat{\pi}_1(\theta_A) - \hat{\pi}_2(\theta_A)| < \varepsilon$, then the heterogeneous and representative agent settings have the same set of long-run learning outcomes.*

This result illustrates the robustness of misspecified environments discussed in [Theorem 5](#).

Next, we explore how heterogeneity affects learning. It is a priori unclear whether heterogeneity will facilitate or hinder learning, compared to the representative agent model. The type with milder misspecification may facilitate learning by counteracting the type with more severe misspecification, or the type with the more severe misspecification may distort information in a way that hinders learning for both types. The following result establishes that the first effect dominates and heterogeneity facilitates learning.

Proposition 3. *Suppose the signal distribution is symmetric, $F^L(s) = 1 - F^R(1 - s)$. If learning is almost surely correct in the representative agent model with $\hat{\pi} \in (\pi(\theta_A), 1]$, then learning is almost surely correct in the heterogeneous model for all $\hat{\pi}_1(\theta_A) \in (\pi(\theta_A), 1]$ and $\hat{\pi}_2(\theta_A) \in (\hat{\pi}_1(\theta_A), 1]$ such that $(\hat{\pi}_1(\theta_A) + \hat{\pi}_2(\theta_A))/2 = \hat{\pi}$, and if incorrect learning occurs with positive probability in the heterogeneous model with $\hat{\pi}_1(\theta_A) \in (\pi(\theta_A), 1]$ and $\hat{\pi}_2(\theta_A) \in (\hat{\pi}_1(\theta_A), 1]$, then incorrect learning occurs with positive probability in the representative agent model with $\hat{\pi} \equiv (\hat{\pi}_1(\theta_A) + \hat{\pi}_2(\theta_A))/2$. For all $\hat{\pi}_1(\theta_A) \in (\pi(\theta_A), 1]$ and $\hat{\pi}_2(\theta_A) \in [\hat{\pi}_1(\theta_A), 1]$, almost surely learning is either correct or incorrect.*

Type θ_1 is more adept at correcting for correlated information, and as a result, asymptotically adopts the inefficient action with lower probability than θ_2 . In turn, this helps θ_2 learn the true state. Actions from θ_1 confirm the state and θ_2 overestimates the private information reflected in these actions. This reduces the probability that θ_2 herds on an inefficient action.³⁸ [Fig. 3](#) illustrates these learning regions.

This characterization allows us to precisely determine how much heterogeneity can be present before the representative agent model is no longer a good approximation. For the parameters considered in [Fig. 3](#), when the average bias is low or high—for example, 0.4 or 0.8—any feasible level of heterogeneity results in the same set of learning outcomes as the corresponding representative agent model. At intermediate levels of the average bias—for example, 0.6—sufficient heterogeneity yields different learning outcomes than the represen-

³⁸Heterogeneity does not always improve learning. If heterogeneity leads to fundamentally different biases—for example, if one type overestimates the correlation in prior actions and the other type underestimates it—then sufficient heterogeneity will interfere with long-run learning, even when the average bias is close to the truth (e.g. $\hat{\pi} \approx \pi(\theta_A)$) and learning is complete in the representative agent model.

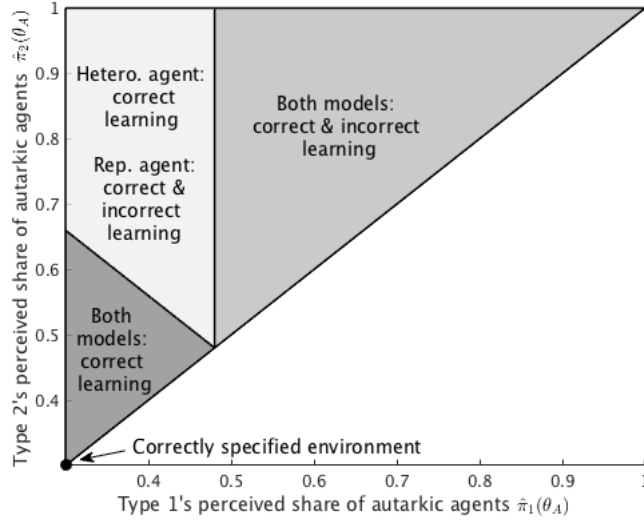


FIGURE 3. Naive Learning with a Representative Agent
 $(\pi(\theta_A) = .3, F^L(s) = 2s - s^2, F^R(s) = s^2)$

tative agent model. In the knife-edge case of an average bias of 0.48, the representative agent model is not identified at certainty—representative agent models with $\hat{\pi} < 0.48$ are robust to any level of heterogeneity while representative agent models with $\hat{\pi}$ slightly above 0.48 are only robust to a very small level of heterogeneity.

Proposition 3 has important implications for policy interventions aimed at mitigating inefficient choices. Suppose a social planner wishes to intervene if and only if agents face the possibility of incorrect learning. The planner measures the average level of bias in the population and uses a representative agent approach to determine whether to intervene. Given Proposition 3, this method will result in *overintervention*, in that there are levels of bias at which the planner will intervene even though learning is almost surely correct. However, underintervention will not be an issue, as the planner will never fail to intervene when incorrect learning arises.

4.3 Entrenched Disagreement in a Level-k Learning Model

As discussed in the introduction, a central contribution of our framework is the ability to allow for model heterogeneity—either due to varying levels of the same bias, as illustrated in Section 4.2, or due to fundamentally distinct biases. This section demonstrates how our characterization can be used to explore whether agents with distinct models influence each others’ learning and to determine when entrenched disagreement emerges. We illustrate this in the context of a social learning setting where agents use level-k reasoning. We show that entrenched disagreement emerges as a robust feature of this setting. Further, the presence of agents who use level-3 reasoning alters the learning outcomes of agents who use level-2 reasoning. This contrasts with settings that consider level-2 reasoning in isolation (Eyster and Rabin 2010; Bohren 2016), highlighting the potential error in predicting learning outcomes without properly accounting for the interaction between different models of inference.

Level-k models describe how boundedly rational agents draw inference in strategic set-

tings (Costa-Gomes and Crawford 2006). Agents are characterized by their “depth” of reasoning, where higher levels use progressively more sophisticated reasoning. Applying this model of inference to a social learning setting, each level corresponds to a type with a misspecified model of the strategic link between prior actions and private signals. Level-0 is a noise type that chooses an action without learning from its signal or the actions of others, i.e. $\hat{s}_0(s) = 1/2$ and $\hat{\pi}_0(\theta_0) = 1$. Higher levels accurately learn from their signals but misinterpret actions, which is captured by a misspecified type distribution. Level-1 is an autarkic type who acts solely based on its private signal, i.e. it believes all agents are noise types, $\hat{\pi}_1(\theta_0) = 1$. Level-2 fails to account for redundant information in prior actions, i.e. it believes all agents are autarkic types, $\hat{\pi}_2(\theta_1) = 1$.³⁹ Level-3 understands that prior actions contain redundant information, but does not allow for the possibility that other agents also account for this, i.e. it believes almost all other agents are level-2, $\hat{\pi}_3(\theta_2) = 1 - \varepsilon$ for some small $\varepsilon > 0$ (for technical reasons, we assume this type places arbitrarily small probability on the level-1 type, $\hat{\pi}_3(\theta_1) = \varepsilon$).⁴⁰ As is customary, the level-0 type anchors the model of level-1 but does not actually exist in the population, $\pi(\theta_0) = 0$. To close the model, assume that each agent faces a binary decision problem in which she earns a payoff of one from choosing the action that matches the state, $\mathcal{A} = \{L, R\}$ and $u(a, \omega) = \mathbb{1}_{a=\omega}$, the level-1 type occurs with positive probability, $\pi(\theta_1) \in (0, 1)$, there are no level-4 or higher types, and all types have common prior $p_0 = 1/2$.⁴¹ Note that a correctly specified environment is not a special case of this set-up, as no type allows for the existence of its own type.

Although depth of reasoning models feature prominently in the empirical literature on social learning (Kübler and Weizsäcker 2004; Penczynski 2017), it has been relatively unexplored in the corresponding theoretical literature—despite the interest in naive learning—as characterizing learning outcomes is significantly more complex when agents learn in different ways.⁴² The addition of level-3 adds two complications relative to a naive learning model with only level-2. First, it is necessary to characterize learning outcomes for multiple types simultaneously. Second, the presence of level-3 affects the learning of level-2, even though level-2 is not aware of level-3.

Proposition 4 establishes that either entrenched disagreement or cyclical learning almost surely arise, depending on the true distribution over types. Strikingly, agreement almost surely does not arise for any distribution over types.

Proposition 4. *There exists an $\bar{\varepsilon} > 0$ such that if $\varepsilon \in (0, \bar{\varepsilon})$, then either learning is cyclical almost surely or entrenched disagreement occurs almost surely. For $\varepsilon \in (0, \bar{\varepsilon})$, there exists*

³⁹A level-2 type is analogous to the “BRTNI” agents in Eyster and Rabin (2010) and the “naive Bayesians” in Hung and Plott (2001). The naive learners in Bohren (2016) are a modified level-2 type that allows for the possibility of other level-2 agents. In Eyster and Rabin (2010), all agents have the same model—they are all level-2—while Bohren (2016), level-1 and level-2 types both occur with positive probability.

⁴⁰The exact parameterization of the level- k model, i.e. $\varepsilon = 0$, violates Assumptions 3 and 4. An alternative similar model that can be captured with this set-up is the cognitive hierarchy model (Camerer et al. 2004), where level-3 places non-trivial probability on level-1 and level-2 types. We explore this alternative parameterization in the working paper version of this article (Bohren and Hauser 2021a).

⁴¹Our framework can allow for higher levels, but empirical studies rarely find evidence of such reasoning. For example, in a social learning experiment, Penczynski (2017) finds that most agents’ behavior is consistent with level-1, 2 or 3 types, with a modal type of level-2.

⁴²As far as we know, no theoretical papers characterize asymptotic learning in a level- k framework. While naive learners are analogous to the level-2 type, existing models focus on the case where all agents are level-2.

a cutoff $\bar{\pi}_3 \in (0, 1)$ such that if $\pi(\theta_3) > \bar{\pi}_3$, then almost surely learning is cyclical, there exists a cutoff $\bar{\pi}_2 \in (0, 1)$ such that if $\pi(\theta_2) > \bar{\pi}_2$, then both disagreement outcomes arise with positive probability, and there exists a cutoff $\bar{\pi}_1 \in (0, 1)$ such that if $\pi(\theta_1) > \bar{\pi}_1$, then the disagreement outcome in which level-2 learns the correct state and level-3 learns the incorrect state arises almost surely.

In contrast, cyclical learning does not arise in settings that examine level-2 reasoning in isolation (Eyster and Rabin 2010; Bohren 2016). Given the empirical evidence for both level-2 and level-3 reasoning in social learning settings (Penczynski 2017), the presence of both models of inference is important to take into account.

Disagreement is driven by level-2’s imitation of the more frequent action and level-3’s anti-imitation in order to correct for level-2’s overreaction. If a large share of agents are level-1, then level-2’s model is close to correct and almost surely level-2 agents learn the correct state. Therefore, the disagreement outcome in which level-2 agents learn the correct state and level-3 agents learn the incorrect state arises almost surely. Note that in this case, a higher level of reasoning performs strictly *worse* than a lower level of reasoning. Otherwise, if a large share of agents are level-2, then both disagreement outcomes emerge and learning is path-dependent. Therefore, two similar populations who learn from different action histories may converge to different forms of disagreement. If a large share of agents are level-3, then neither disagreement outcome is stable and learning is almost surely cyclical. Intuitively, near beliefs $(0, \infty)$ where level-2 agents choose L and level-3 agents choose R , a level-2 agent overreacts to the more frequently chosen R action, pulling her belief away from state L , and similarly for the other disagreement outcome. Correct and incorrect learning almost surely do not arise because the level-2 and level-3 agents have models that interpret actions in opposite ways, which prevents agreement.

Fig. 4 illustrates the learning regions for the level-k model, including the thresholds described in Proposition 4.⁴³ Penczynski (2017)’s estimate of the type distribution lies in the gray region in which both disagreement outcomes arise with positive probability.

5 Conclusion

We develop a general framework to study learning with model misspecification, which captures many information-processing biases and heuristics of interest in economic decision-making. A key contribution of our framework is the ability to allow for model heterogeneity, in which agents exhibit different levels of a bias or have distinct biases. Our main result characterizes the set of asymptotic learning outcomes based on two expressions that are straightforward to derive from the underlying form of misspecification. This characterization provides a unified way to compare different forms of misspecification that have been previously studied and yields new insights about forms of misspecification that have not been previously explored. The characterization also provides a rationale for entrenched disagreement, in which agents with different models converge to different certain beliefs despite observing a common history. Our results yield insights into how the source of information (i.e. social versus private) impacts learning, whether learning predictions are sensitive to

⁴³In Online Appendix C.3, we show analytically that the qualitative features Fig. 4 hold for the full characterization of learning outcomes across $(\pi(\theta_1), \pi(\theta_2), \pi(\theta_3)) \in \Delta^2$; for expositional clarity, we state a partial characterization in Proposition 4.

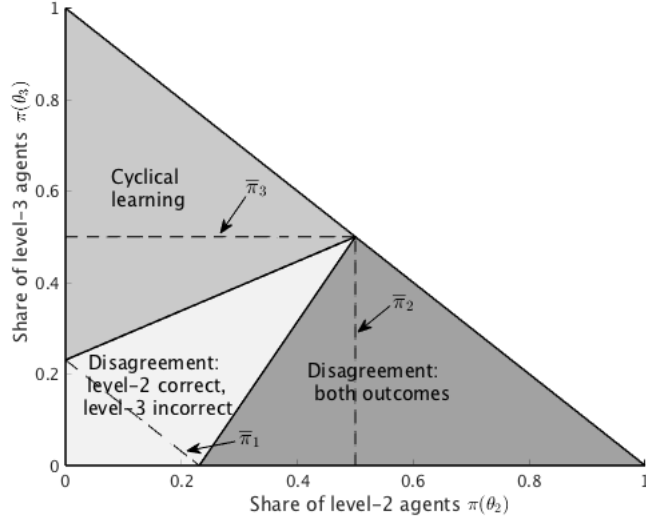


FIGURE 4. Entrenched Disagreement in Level-k Learning
 $(\omega = L, F^L = \frac{10}{3}(s - .5s^2) - .6, F^R(s) = \frac{5}{3}(s^2 - .04))$

different parameterizations of a bias, and when a set of learning outcomes is robust to varying levels of misspecification. In the presence of model heterogeneity, our results can also be used to explore how different biases interact and to determine whether a representative agent approach generates accurate learning predictions.

References

- ANGRISANI, M., A. GUARINO, P. JEHL, AND T. KITAGAWA (2020): “Information Redundancy Neglect versus Overconfidence: A Social Learning Experiment,” *American Economic Journal: Microeconomics*.
- ARROW, K. J. AND J. R. GREEN (1973): “Notes on Expectations Equilibria in Bayesian Settings,” *Institute for Mathematical Studies in the Social Sciences Working Papers*.
- BA, C. (2021): “Robust Model Misspecification and Paradigm Shift,” Mimeo.
- BANERJEE, A. (1992): “A Simple Model of Herd Behavior,” *The Quarterly Journal of Economics*, 107, 797–817.
- BARTELS, L. M. (2002): “Beyond the running tally: Partisan bias in political perceptions,” *Political Behavior*, 24, 117–150.
- BERK, R. H. (1966): “Limiting Behavior of Posterior Distributions When the Model is Incorrect,” *The Annals of Mathematical Statistics*, 37, 51–58.
- BIKHCHANDANI, S., D. HIRSHLEIFER, AND I. WELCH (1992): “A Theory of Fads, Fashion, Custom, and Cultural Change as Informational Cascades,” *The Journal of Political Economy*, 100, 992–1026.
- BOHREN, A. (2016): “Informational Herding with Model Misspecification,” *Journal of Economic Theory*, 222–247.
- BOHREN, J. A. AND D. N. HAUSER (2019a): “Misinterpreting Social Outcomes and Information Campaigns,” Mimeo.
- (2019b): “Social Learning with Model Misspecification: A Framework and a Characterization,” *PIER Working Paper 18-017*, available at SSRN: <https://ssrn.com/abstract=3236842> or <http://dx.doi.org/10.2139/ssrn.3236842>.

- (2021a): “Learning with Heterogeneous Misspecified Models: Characterization and Robustness,” *PIER Working Paper 21-005*, available at SSRN: <https://ssrn.com/abstract=3783840> or <http://dx.doi.org/10.2139/ssrn.3783840>.
- (2021b): “Posteriors as Signals in Misspecified Models,” *Mimeo*.
- BRUNNERMEIER, M. K. AND J. A. PARKER (2005): “Optimal expectations,” *The American Economic Review*, 95, 1092–1118.
- BUSHONG, B. AND T. GAGNON-BARTSCH (2019): “Learning with Misattribution of Reference Dependence,” .
- CAMERER, C. F., T.-H. HO, AND J.-K. CHONG (2004): “A Cognitive Hierarchy Model of Games,” *The Quarterly Journal of Economics*, 119, 861–898.
- COSTA-GOMES, M. AND V. CRAWFORD (2006): “Cognition and Behavior in Two-Person Guessing Games: An Experimental Study,” *American Economic Review*, 96, 1737 – 1768.
- CRIPPS, M. W. (2018): “Divisible Updating,” .
- ENKE, B. AND F. ZIMMERMANN (2019): “Correlation Neglect in Belief Formation,” *The Review of Economic Studies*, 86, 313–332.
- EPSTEIN, L., J. NOOR, AND A. SANDRONI (2008): “Non-Bayesian updating: a theoretical framework,” *Theoretical Economics*, 3, 193–229.
- EPSTEIN, L. G., J. NOOR, AND A. SANDRONI (2010): “Non-Bayesian Learning,” *The B.E. Journal of Theoretical Economics*, 10.
- ESPONDA, I. AND D. POUZO (2016): “Berk-Nash Equilibrium: A Framework for Modeling Agents with Misspecified Models,” *Econometrica*, 84, 1093–1130.
- (2019): “Equilibrium in Misspecified Markov Decision Processes,” *Unpublished Manuscript*, available at <https://arxiv.org/abs/1502.06901>.
- ESPONDA, I., D. POUZO, AND Y. YAMAMOTO (2019): “Asymptotic Behavior of Bayesian Learners with Misspecified Models,” *Unpublished Manuscript*, available at <http://arxiv.org/abs/1904.08551>.
- EYSTER, E. AND M. RABIN (2005): “Cursed Equilibrium,” *Econometrica*, 73, 1623–1672.
- (2010): “Naive Herding in Rich-Information Settings,” *American Economic Journal: Microeconomics*, 2, 221–243.
- EYSTER, E., M. RABIN, AND G. WEIZSÄCKER (2020): “An Experiment on Social Mislearning,” *Unpublished Manuscript*, available at <https://ideas.repec.org/p/rco/dpaper/73.html>.
- FRICK, M., R. IJIMA, AND Y. ISHII (2019): “Dispersed Behavior and Perceptions in Assortative Societies,” *Cowles Foundation Discussion Paper No. 2128R*, available at <https://ssrn.com/abstract=3362249>.
- (2020a): “Misinterpreting others and the fragility of social learning,” *Econometrica*, 88, 2281–2328.
- (2020b): “Stability and Robustness in Misspecified Learning Models,” *Cowles Foundation Discussion Paper No. 2235*, available at <https://ssrn.com/abstract=3600633>.
- FUDENBERG, D., G. LANZANI, AND P. STRACK (2020): “Limits Points of Endogenous Misspecified Learning,” *Unpublished Manuscript*, available at https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3553363.
- FUDENBERG, D., G. ROMANYUK, AND P. STRACK (2017): “Active learning with a misspecified prior,” *Theoretical Economics*, 12, 1155–1189.
- GAGNON-BARTSCH, T. (2016): “Taste Projection in Models of Social Learning,” *Mimeo*.
- GAGNON-BARTSCH, T. AND M. RABIN (2016): “Naive Social Learning, Mislearning and Unlearning,” .
- GAGNON-BARTSCH, T., M. RABIN, AND J. SCHWARTZSTEIN (2018): “Channeled attention and

- stable errors,” *HBS Working paper 18-108*.
- GOTTLIEB, D. (2015): “Will you never learn? Self deception and biases in information processing,” *Unpublished Manuscript*, 1–46.
- GUARINO, A. AND P. JEHIEL (2013): “Social Learning with Coarse Inference,” *American Economic Journal: Microeconomics*, 5, 147–174.
- HE, K. (2020): “Mislearning from Censored Data: The Gambler’s Fallacy in Optimal-Stopping Problems,” *Unpublished Manuscript*, available at <http://arxiv.org/abs/1803.08170>.
- HEIDHUES, P., B. KŐSZEGI, AND P. STRACK (2018): “Unrealistic Expectations and Misguided Learning,” *Econometrica*, 86, 1159–1214.
- (2019): “Convergence in Misspecified Learning Models with Endogenous Actions,” *Unpublished Manuscript*, available at https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3312968.
- HUNG, A. A. AND C. R. PLOTT (2001): “Information Cascades: Replication and an Extension to Majority Rule and Conformity-Rewarding Institutions,” *American Economic Review*, 91, 1508–1520.
- JADBABAIE, A., P. MOLAVI, A. SANDRONI, AND A. TAHBAZ-SALEHI (2012): “Non-Bayesian social learning,” *Games and Economic Behavior*, 76, 210–225.
- JEHIEL, P. (2005): “Analogy-based expectation equilibrium,” *Journal of Economic Theory*, 123, 81–104.
- JERIT, J. AND J. BARABAS (2012): “Partisan perceptual bias and the information environment,” *The Journal of Politics*, 74, 672–684.
- KOMINERS, S. D., X. MU, AND A. PEYSAKHOVICH (2018): “Paying (for) attention: The impact of information processing costs on bayesian inference,” *Unpublished Manuscript*, available at https://papers.ssrn.com/sol3/papers.cfm?abstract_id=2857978.
- KŐSZEGI, B. (2006): “Ego utility, overconfidence, and task choice,” *Journal of the European Economic Association*, 4, 673–707.
- KÜBLER, D. AND G. WEIZSÄCKER (2004): “Limited depth of reasoning and failure of cascade formation in the laboratory,” *The Review of Economic Studies*, 71, 425–441.
- LEVY, G. AND R. RAZIN (2015): “Correlation neglect, voting behavior, and information aggregation,” *American Economic Review*, 105, 1634–45.
- MADARÁSZ, K. AND A. PRAT (2016): “Sellers with misspecified models,” *The Review of Economic Studies*.
- MARKS, G. AND N. MILLER (1987): “Ten years of research on the false-consensus effect: An empirical and theoretical review,” *Psychological Bulletin*, 102, 72.
- MERTENS, J.-F. AND S. ZAMIR (1985): “Formulation of Bayesian analysis for games with incomplete information,” *International Journal of Game Theory*, 14, 1–29.
- MILLER, D. T. AND C. MCFARLAND (1987): “Pluralistic ignorance: When similarity is interpreted as dissimilarity,” *Journal of Personality and Social Psychology*, 53, 298.
- (1991): “When social comparison goes awry: The case of pluralistic ignorance,” in *Social Comparison: Contemporary Theory and Research*, ed. by J. Suls and T. A. Wills, Lawrence Erlbaum Associates, Inc, 287–313.
- MOLAVI, P., A. TAHBAZ-SALEHI, AND A. JADBABAIE (2018): “A Theory of Non-Bayesian Social Learning,” *Econometrica*, 86, 445–490.
- MOORE, D. A. AND P. J. HEALY (2008): “The trouble with overconfidence,” *Psychological Review*, 115, 502–517.
- MOORE, D. A., E. R. TENNEY, AND U. HARAN (2015): “Overprecision in Judgment,” in *The Wiley Blackwell Handbook of Judgment and Decision Making*, John Wiley & Sons, Ltd, chap. 6, 182–209.

- NICKERSON, R. S. (1998): “Confirmation Bias: A Ubiquitous Phenomenon in Many Guises,” *Review of General Psychology*, 2, 175–220.
- NYARKO, Y. (1991): “Learning in Misspecified Models and the Possibility of Cycles,” *Journal of Economic Theory*, 55, 416–427.
- ORTOLEVA, P. (2012): “Modeling the Change of Paradigm: Non-Bayesian Reactions to Unexpected News,” *American Economic Review*, 102, 2410–2436.
- ORTOLEVA, P. AND E. SNOWBERG (2015): “Overconfidence in political behavior,” *American Economic Review*, 105, 504–535.
- PACHUR, T., J. RIESKAMP, AND R. HERTWIG (2004): “The Social Circle Heuristic: Fast and Frugal Decisions Based on Small Samples,” *Proceedings of the Annual Meeting of the Cognitive Science Society*.
- PENCZYNSKI, S. P. (2017): “The nature of social learning: Experimental evidence,” *European Economic Review*, 94, 148–165.
- RABIN, M. AND J. L. SCHRAG (1999): “First Impressions Matter: A Model of Confirmatory Bias,” *The Quarterly Journal of Economics*, 114, 37–82.
- ROSS, L., D. GREENE, AND P. HOUSE (1977): “The “false consensus effect”: An egocentric bias in social perception and attribution processes,” *Journal of experimental social psychology*, 13, 279–301.
- SCHWARTZSTEIN, J. (2014): “Selective Attention and Learning,” *Journal of the European Economic Association*, 12, 1423–1452.
- SHALIZI, C. R. (2009): “Dynamics of Bayesian updating with dependent data and misspecified models,” *Electronic Journal of Statistics*, 3, 1039–1074.
- SMITH, L. AND P. N. SØRENSEN (2000): “Pathological Outcomes Observational Learning,” *Econometrica*, 68, 371–398.
- SPIEGLER, R. (2016): “Bayesian networks and boundedly rational expectations,” *Quarterly Journal of Economics*, 131, 1243–1290.
- UNGEHEUER, M. AND M. WEBER (2020): “The Perception of Dependence, Investment Decisions, and Stock Prices,” *Unpublished Manuscript*, available at https://papers.ssrn.com/sol3/papers.cfm?abstract_id=2739130.
- WILSON, A. (2014): “Bounded Memory and Biases in Information Processing,” *Econometrica*, 82, 2257–2294.

A Proofs from Section 3

Throughout this section, we use the following notation. Given $\varepsilon > 0$ and $\kappa \in \{1, \dots, k\}$, define a neighborhood of $\boldsymbol{\lambda} \in \{0, \infty\}^\kappa$ as $B_\varepsilon(\boldsymbol{\lambda}^*) \equiv \{\boldsymbol{\lambda} \in [0, \infty]^\kappa \mid \lambda_i \in [0, \varepsilon] \text{ if } \lambda_i^* = 0 \text{ and } \lambda_i \in (1/\varepsilon, \infty] \text{ if } \lambda_i^* = \infty\}$. Let $\text{int}(B_\varepsilon(\boldsymbol{\lambda}^*)) \equiv B_\varepsilon(\boldsymbol{\lambda}^*) \cap (0, \infty)^\kappa$.

A.1 Proofs of Lemmas 1 to 3

Proof of Lemma 1. Let $p_i(\lambda_i(h), s) \equiv \hat{P}_i(R|h, s)$ denote the private belief of type θ_i following history h and signal realization s . Given $\lambda \in (0, \infty)$ and $s \in \mathcal{S}$, consider type θ_i ’s choice at private belief $p_i(\lambda, s)$. Recall that actions $(a_1, \dots, a_M)$ are ordered by relative preference in state R . Since no two actions yield the same payoff in both states, no action is optimal at a single belief, and preferences are aligned ([Assumption 2](#)), there exist belief thresholds $0 = \bar{p}_{i,0} \leq \bar{p}_{i,1} \leq \dots \leq \bar{p}_{i,M} = 1$ such that we can partition the belief space $[0, 1]$ into a finite set of closed intervals, with action a_m optimal at $p_i(\lambda, s)$ if $p_i(\lambda, s) \in [\bar{p}_{i,m-1}, \bar{p}_{i,m}]$ and $\bar{p}_{i,m-1} \neq \bar{p}_{i,m}$, and a_m never optimal iff $\bar{p}_{i,m-1} = \bar{p}_{i,m}$. Without loss of generality, assume the

tie-breaking rule is to choose the optimal action with the lower index at each interior cutoff $\bar{p}_{i,m} \in (0, 1)$, i.e. if $\bar{p}_{i,m-1} \neq \bar{p}_{i,m}$, choose a_m over any other optimal action $a_{m+\kappa}$ at belief $\bar{p}_{i,m}$. Since there are at least two undominated actions, there are at least two intervals with a non-empty interior. Since signals are aligned ([Assumption 1](#)), when $\theta_i \in \Theta_S \cup \Theta_A$, $p_i(\lambda, s)$ is strictly increasing in s for all $\lambda \in (0, \infty)$. Therefore, for each $\lambda \in (0, \infty)$, we can define the decision rule with respect to signal cutoffs $0 \equiv \bar{s}_{i,0}(\lambda) \leq \bar{s}_{i,1}(\lambda) \leq \dots \leq \bar{s}_{i,M}(\lambda) \equiv 1$, where $\bar{s}_{i,m}(\lambda) \equiv \inf\{\bar{s} \in [0, 1] : p_i(\lambda, s) > p_{i,m} \text{ for all } s \in \mathcal{S} \text{ s.t. } s > \bar{s}\}$ for $m = 1, \dots, M-1$, so that the agent chooses action a_m iff $\bar{s}_{i,m-1}(\lambda) \neq \bar{s}_{i,m}(\lambda)$ and she observes signal realization $s \in (\bar{s}_{i,m-1}(\lambda), \bar{s}_{i,m}(\lambda)]$, with a closed interval if $\bar{s}_{i,m-1}(\lambda) = 0$. When $\theta_i \in \Theta_N$, $p_i(\lambda, s)$ is constant with respect to s for all $\lambda \in [0, \infty]$. Therefore, the same set of actions are optimal at all signal realizations, and generically, this set is a singleton. We define an analogous set of signal cutoffs as follows: if a_m is the optimal action with the lowest index at $\lambda \in (0, \infty)$, then $\bar{s}_{i,0}(\lambda) = \dots = \bar{s}_{i,m-1}(\lambda) = 0$, $\bar{s}_{i,m}(\lambda) = \dots = \bar{s}_{i,M-1}(\lambda) = \sup \mathcal{S}$ and $\bar{s}_{i,M}(\lambda) = 1$. The case of $\lambda \in \{0, \infty\}$ for $\theta_i \in \Theta_S \cup \Theta_A$ is similar to the noise type: a single action is optimal at all signal realizations and we can define the cutoffs analogously. $\square$

Proof of Lemma 2. Let $\bar{s}_{i,m}(\lambda)$ denote the cutoff rules defined in [Lemma 1](#) for $m = 0, \dots, M$ and $\theta_i \in \Theta$. To establish the uniform bound, we first show that $\hat{\psi}_i(a_1|R, \lambda) \leq \hat{\psi}_i(a_1|L, \lambda)$ for all $\lambda \in [0, \infty]^k$ and social types $\theta_i \in \Theta_S$, and then show that this inequality is strict. Fix $\lambda \in [0, \infty]^k$ and consider how social type θ_i updates its beliefs following a_1 . By [Lemma 1](#), $\theta_i \in \Theta_S$ believes that type $\theta_j \in \Theta$ plays a_1 with probability $\hat{F}_i^\omega(\bar{s}_{j,1}(\lambda_j))$. By [Lemma A.1](#) in [Smith and Sørensen \(2000\)](#), $F^R(s) \leq F^L(s)$, with strict inequality except when both terms are 0 or 1. Since a social type believes the signal is informative, by [Assumption 1](#), this is also true for $\hat{F}_i^R(s)$ and $\hat{F}_i^L(s)$. Therefore, $\hat{F}_i^R(\bar{s}_{j,1}(\lambda_j)) \leq \hat{F}_i^L(\bar{s}_{j,1}(\lambda_j))$. This implies $\hat{\psi}_i(a_1|R, \lambda) \leq \hat{\psi}_i(a_1|L, \lambda)$, since $\hat{\psi}_i(a|\omega, \lambda)$ is a convex combination of $\hat{F}_i^\omega(\bar{s}_{j,1}(\lambda_j))$ for each $\theta_j \in \Theta$. To see that the inequality is strict, recall that autarkic types have a likelihood ratio that is constant and equal to $p_0/(1-p_0)$. Therefore, for social type θ_i ,

$$\begin{aligned} \frac{\hat{\psi}_i(a_1|R, \lambda)}{\hat{\psi}_i(a_1|L, \lambda)} &= \frac{\sum_{\theta_j \in \Theta_A} \hat{\pi}_i(\theta_j) \hat{F}_i^R(\bar{s}_{j,1}(\frac{p_0}{1-p_0})) + \sum_{\theta_j \in \Theta_S} \hat{\pi}_i(\theta_j) \hat{F}_i^R(\bar{s}_{j,1}(\lambda_j))}{\sum_{\theta_j \in \Theta_A} \hat{\pi}_i(\theta_j) \hat{F}_i^L(\bar{s}_{j,1}(\frac{p_0}{1-p_0})) + \sum_{\theta_j \in \Theta_A} \hat{\pi}_i(\theta_j) \hat{F}_i^L(\bar{s}_{j,1}(\lambda_j))} \\ &\leq \frac{\sum_{\theta_j \in \Theta_A} \hat{\pi}_i(\theta_j) \hat{F}_i^R(\bar{s}_{j,1}(\frac{p_0}{1-p_0})) + \sum_{\theta_j \in \Theta_S} \hat{\pi}_i(\theta_j) \hat{F}_i^L(\bar{s}_{j,1}(\lambda_j))}{\sum_{\theta_j \in \Theta_A} \hat{\pi}_i(\theta_j) \hat{F}_i^L(\bar{s}_{j,1}(\frac{p_0}{1-p_0})) + \sum_{\theta_j \in \Theta_S} \hat{\pi}_i(\theta_j) \hat{F}_i^L(\bar{s}_{j,1}(\lambda_j))} \\ &\leq \frac{\sum_{\theta_j \in \Theta_A} \hat{\pi}_i(\theta_j) \hat{F}_i^R(\bar{s}_{j,1}(\frac{p_0}{1-p_0})) + \hat{\pi}_i(\Theta_S)}{\sum_{\theta_j \in \Theta_A} \hat{\pi}_i(\theta_j) \hat{F}_i^L(\bar{s}_{j,1}(\frac{p_0}{1-p_0})) + \hat{\pi}_i(\Theta_S)} < 1 \end{aligned} \quad (6)$$

where the first line follows by definition, the second line follows from $\hat{F}_i^R(s) \leq \hat{F}_i^L(s)$, the third line follows from $\sum_{\theta_j \in \Theta_S} \hat{\pi}_i(\theta_j) \hat{F}_i^L(s) \leq \hat{\pi}_i(\Theta_S)$, and the bound of one follows from [Assumption 3](#), which ensures there exists at least one autarkic type $\theta_j \in \Theta_A$ with $\hat{\pi}_i(\theta_j) > 0$ and $\bar{s}_{j,1}(\frac{p_0}{1-p_0})$ such that $\hat{F}_i^R(\bar{s}_{j,1}(\frac{p_0}{1-p_0})) < \hat{F}_i^L(\bar{s}_{j,1}(\frac{p_0}{1-p_0}))$. Therefore, $\hat{\psi}_i(a_1|R, \lambda)/\hat{\psi}_i(a_1|L, \lambda)$ is uniformly bounded away from one. Similar logic holds for action a_M .

Fix $\theta_i \in \Theta_S$ and $a \in \mathcal{A}$. From [Assumption 4](#), $\hat{\psi}_i(a|\omega, \lambda) > 0$ for each $\omega \in \{L, R\}$ and $\lambda \in [0, \infty]^k$. Further, $\hat{\psi}_i(a|\omega, \lambda)$ is bounded below by the perceived probability that autarkic or noise types play a in state ω , which is independent of λ . Therefore, $\hat{\psi}_i(a|\omega, \lambda)$

is uniformly bounded away from zero for $\lambda \in [0, \infty]^k$. It follows that a is perceived to be boundedly informative, i.e. there exists an $\varepsilon > 0$ such that $\hat{\psi}_i(a|R, \lambda)/\hat{\psi}_i(a|L, \lambda) \in [\varepsilon, 1/\varepsilon]$ for all $\lambda \in [0, \infty]^k$. Similar reasoning uses [Assumption 3](#) to establish that $\psi(a_1|\omega, \lambda)$ and $\psi(a_M|\omega, \lambda)$ are uniformly bounded away from zero for $\lambda \in [0, \infty]^k$ in each state.

To establish the continuity of $\psi(a_m|\omega, \lambda)$ at $\lambda \in \{0, \infty\}^k$, from [Eq. \(2\)](#), it is sufficient to show that $\lambda \mapsto F^\omega(\bar{s}_{i,m}(\lambda))$ is continuous at $\lambda \in \{0, \infty\}$ for all $m = 0, \dots, M$, $\theta_i \in \Theta_S$ and $\omega \in \{L, R\}$. Consider continuity at $\lambda = 0$. Fix an $\varepsilon > 0$. Since a perfectly informative signal realization does not occur with positive probability and is not perceived to occur with positive probability, there exists an $\tilde{s} \in \mathcal{S}$ such that $\tilde{s} < 1$, $1 - F^\omega(\tilde{s}) < \varepsilon$ for each $\omega \in \{L, R\}$, and $\hat{s}_i(\tilde{s}) \notin \{0, 1\}$ for each $\theta_i \in \Theta_S$.⁴⁴ From [Lemma 1](#), it follows that for each $m = 1, \dots, M-1$ and $\theta_i \in \Theta_S$, $\bar{s}_{i,m}(0) \in \{0, \sup \mathcal{S}\}$, with $\bar{s}_{i,m}(0) = \sup \mathcal{S}$ if and only if $\bar{s}_{i,m}(\lambda) > 0$ for some $\lambda \in (0, \infty)$. From the properties of Bayes Rule, there exists a $\delta > 0$ such that if $\lambda < \delta$, then for each $m = 1, \dots, M-1$ and $\theta_i \in \Theta_S$ such that $\bar{s}_{i,m}(0) = \sup \mathcal{S}$, $\bar{s}_{i,m}(\lambda) \geq \tilde{s}$. Therefore, $|F^\omega(\bar{s}_{i,m}(\lambda)) - F^\omega(\bar{s}_{i,m}(0))| \leq |F^\omega(\tilde{s}) - 1| < \varepsilon$ for $\omega \in \{L, R\}$. For $m = 1, \dots, M-1$ and $\theta_i \in \Theta_S$ such that $\bar{s}_{i,m}(0) = 0$, it follows from [Lemma 1](#) that $\bar{s}_{i,m}(\lambda) = 0$ for all $\lambda \in [0, \infty]$ and therefore, $|F^\omega(\bar{s}_{i,m}(\lambda)) - F^\omega(\bar{s}_{i,m}(0))| = 0$ for $\omega \in \{L, R\}$. By definition, $\bar{s}_{i,0}(\lambda) = 0$ and $\bar{s}_{i,M}(\lambda) = 1$ for all $\lambda \in [0, \infty]$, and so $|F^\omega(\bar{s}_{i,m}(\lambda)) - F^\omega(\bar{s}_{i,m}(0))| = 0$ for $m \in \{0, M\}$ and $\omega \in \{L, R\}$. Taken together, this establishes the continuity of $\lambda \mapsto F^\omega(\bar{s}_{i,m}(\lambda))$ at $\lambda = 0$ for all $m = 0, \dots, M$, $\theta_i \in \Theta_S$ and $\omega \in \{L, R\}$. Analogous reasoning establishes continuity at $\lambda = \infty$. To establish continuity for $\hat{\psi}_i(a_m|\omega, \lambda)$, it is sufficient to show that $\lambda \mapsto \hat{F}_i^\omega(\bar{s}_{j,m}(\lambda))$ is continuous at $\lambda \in \{0, \infty\}$ for all $m = 0, \dots, M$, $\theta_j \in \Theta_S$ and $\omega \in \{L, R\}$. This follows from identical logic to the case of $\psi(a_m|\omega, \lambda)$. $\square$

Proof of Lemma 3. At a stationary belief $\lambda^* \in [0, \infty]^k$, $\lambda^* = \lambda^* \frac{\hat{\psi}_i(a|R, \lambda^*)}{\hat{\psi}_i(a|L, \lambda^*)}$ for all a such that $\psi(a|\omega, \lambda^*) > 0$. Trivially, this is satisfied for all $\lambda^* \in \{0, \infty\}^k$, independent of $\psi(a|\omega, \lambda^*)$, and these beliefs are stationary. It remains to be determined whether it is satisfied for any interior beliefs $\lambda^* \in (0, \infty)^k$. Suppose $\lambda^* \in (0, \infty)^k$. By [Assumption 3](#), $\psi(a_1|\omega, \lambda^*) > 0$ for each $\omega \in \{L, R\}$. By [Lemma 2](#), $\frac{\hat{\psi}_i(a_1|R, \lambda^*)}{\hat{\psi}_i(a_1|L, \lambda^*)} < 1$. Therefore, this does not hold for a_1 and λ^* cannot be stationary.

Suppose beliefs converge to a non-stationary belief $\lambda^* \in [0, \infty]^k \setminus \{0, \infty\}^k$ with positive probability. By [Lemma 2](#), following action $\tilde{a}_t = a_M$, $\lambda_{i,t+1} - \lambda_{i,t}$ is bounded uniformly away from zero for all social types $\theta_i \in \Theta_S$. For sufficiently small $\varepsilon > 0$, when $\lambda_t \in B_\varepsilon(\lambda^*)$, $\lambda_{i,t+1} \notin B_\varepsilon(\lambda^*)$ for any type with an interior belief $\lambda_{i,t} \in (0, \infty)$. The probability $Pr(\exists t < T \text{ s.t. } \tilde{a}_t = a_M)$ converges to one as $T \rightarrow \infty$. Therefore, the likelihood ratio almost surely leaves $B_\varepsilon(\lambda^*)$. $\square$

A.2 Local Stability

Proof of Theorem 1. Without loss of generality, consider the case where $\omega = L$.

Part 1. We first show that if $\lambda^* \in \Lambda(L)$, then λ^* is locally stable. Consider $\lambda^* = 0^k$ and suppose $\gamma_i(L, 0^k) < 0$ for all social types $\theta_i \in \Theta_S$. By the continuity property established in

⁴⁴If there did not exist such an $\tilde{s}$, then $1 - F^\omega(s) \geq \varepsilon$ for all $s \in [0, 1)$. Given $F^\omega(1) = 1$, this implies that there is a mass point at $s = 1$, i.e. a perfectly informative signal realization that occurs with positive probability. This is a contradiction.

Lemma 2, there exists an $\varepsilon > 0$ such that for neighborhood $B_\varepsilon(0^k) \equiv [0, \varepsilon]^k$,

$$\sum_{a \in \mathcal{A}} \psi(a|L, 0^k) g_i(a) < 0, \quad (7)$$

for all $\theta_i \in \Theta_S$, where $g_i(a) \equiv \sup_{\lambda \in B_\varepsilon(0^k)} \log \frac{\hat{\psi}_i(a|R, \lambda)}{\hat{\psi}_i(a|L, \lambda)}$ denotes the maximal update for type θ_i following action a across all beliefs in $B_\varepsilon(0^k)$. Let $\mathbf{g}(a) \equiv (g_1(a), \dots, g_k(a))$ denote the vector of maximal updates for action a , $\bar{g}_i \equiv \max_{a \in \mathcal{A}} g_i(a)$ denote the maximal update across all actions, and $\bar{\mathbf{g}} \equiv (\bar{g}_1, \dots, \bar{g}_k)$ denote the corresponding vector of maximal updates.

For any $\delta > 0$, choose $\varepsilon_\delta \in (0, \varepsilon]$ such that $\sup_{\lambda \in [0, \varepsilon_\delta]^k} |\psi(a|L, \lambda) - \psi(a|L, 0^k)| < \delta$. By **Lemma 2**, $\psi(a|L, \lambda)$ is continuous at $\lambda = 0^k$, so such an ε_δ exists. Let $\psi_\delta(a) \equiv \inf_{\lambda \in [0, \varepsilon_\delta]^k} \psi(a|L, \lambda)$. Define a linear system $\langle \lambda_{\delta, t} \rangle_{t=1}^\infty$ as follows: $\lambda_{\delta, 1} = \lambda_1$, and for each $a \in \mathcal{A}$, whenever $\tilde{a}_t = a$ and $\lambda_t \in [0, \varepsilon_\delta]^k$, let $\log \lambda_{\delta, t+1} = \log \lambda_{\delta, t} + \mathbf{g}(a)$ with probability $\psi_\delta(a)/\psi(a|L, \lambda_t)$ and $\log \lambda_{\delta, t+1} = \log \lambda_{\delta, t} + \bar{\mathbf{g}}$ otherwise. When $\omega = L$, $\psi_\delta(a)$ is the joint probability of $\tilde{a}_t = a$ and the former event (i.e. $\psi(a|L, \lambda_t) * \psi_\delta(a)/\psi(a|L, \lambda_t)$), while $\bar{\psi}_\delta \equiv 1 - \sum_{a \in \mathcal{A}} \psi_\delta(a)$ is the cumulative probability of the latter event across all actions. To complete the construction, if $\lambda_t \notin [0, \varepsilon_\delta]^k$ then independent of the realization of $\tilde{a}_t$ let $\log \lambda_{\delta, t+1} = \log \lambda_{\delta, t} + \mathbf{g}(a)$ with probability $\psi_\delta(a)$ for each $a \in \mathcal{A}$ and $\log \lambda_{\delta, t+1} = \log \lambda_{\delta, t} + \bar{\mathbf{g}}$ otherwise. The transition probabilities and updates to $\langle \lambda_{\delta, t} \rangle$ are i.i.d. By construction, when $\lambda_t \in [0, \varepsilon_\delta]^k$, the update to the process $\langle \lambda_{\delta, t} \rangle$ is larger than the update to the process $\langle \lambda_t \rangle$ (this follows from $\log \frac{\hat{\psi}_i(a|R, \lambda_t)}{\hat{\psi}_i(a|L, \lambda_t)} \leq g_i(a) \leq \bar{g}_i$ for all $\theta_i \in \Theta_S$). It follows that if $\lambda_t \leq \lambda_{\delta, t}$ and $\lambda_t \in [0, \varepsilon_\delta]^k$, then $\lambda_{t+1} \leq \lambda_{\delta, t+1}$.

From a straightforward adaptation of Lemma C.1 of **Smith and Sørensen (2000)** to a multidimensional Markov process, if

$$\bar{\psi}_\delta \bar{g}_i + \sum_{a \in \mathcal{A}} \psi_\delta(a) g_i(a) < 0 \quad (8)$$

for all $\theta_i \in \Theta_S$, then almost surely $\lim_{t \rightarrow \infty} \lambda_{\delta, t} = 0^k$. By **Eq. (7)**, **Eq. (8)** holds for sufficiently small δ (this follows from $\lambda \mapsto \psi(a|\omega, \lambda)$ continuous at 0^k (**Lemma 2**) and $\lim_{\delta \rightarrow 0} \psi_\delta(a) = \psi(a|L, 0^k)$, which implies $\lim_{\delta \rightarrow 0} \bar{\psi}_\delta = 0$). Let $\delta_1 > 0$ denote an upper bound such that **Eq. (8)** holds for all $\delta < \delta_1$.

Fix $\delta \in (0, \delta_1)$ and choose a corresponding $\varepsilon_\delta > 0$ as defined above. Conditional on the event that $\lambda_t \in [0, \varepsilon_\delta]^k$ for all $t \geq 1$, $\langle \lambda_t \rangle$ is bounded above by a stochastic process that converges to zero almost surely. Therefore, conditional on this event, $\lim_{t \rightarrow \infty} \lambda_t = 0^k$ almost surely. We next show that there exists an $\varepsilon^* \in (0, \varepsilon_\delta]$ such that the probability of this event is uniformly bounded away from zero across $[0, \varepsilon^*]^k$. For any $\lambda \in (0, \infty)^k$, when $\lambda_1 = \lambda$, $\lim_{t \rightarrow \infty} \lambda_{\delta, t} = 0^k$ almost surely, and therefore, $Pr(\cup_t \cap_{s \geq t} \{\lambda_{\delta, s} \in [0, \varepsilon_\delta]^k\} | \lambda_1 = \lambda) = 1$. It follows that there exists a finite $\tau(\lambda) \geq 1$ such that $Pr(\lambda_{\delta, t} \in [0, \varepsilon_\delta]^k \forall t \geq \tau(\lambda) | \lambda_1 = \lambda) > 0$. Since $\tau(\lambda) < \infty$ and $g_i(a)$ and $\bar{g}_i$ are bounded for all $a \in \mathcal{A}$ and $\theta_i \in \Theta_S$, $\text{supp } \lambda_{\delta, \tau(\lambda)} \subset (0, \infty)^k$. Since the system is linear, for any $\lambda' \in \text{supp } \lambda_{\delta, \tau(\lambda)}$ such that $Pr(\lambda_{\delta, t} \in [0, \varepsilon_\delta]^k \forall t \geq \tau(\lambda) | \lambda_{\delta, \tau(\lambda)} = \lambda') > 0$, $Pr(\lambda_{\delta, t} \in [0, \varepsilon_\delta]^k \forall t \geq 1 | \lambda_1 = \lambda') > 0$. Further, since $\langle \lambda_{\delta, t} \rangle$ has i.i.d. transitions, $Pr(\lambda_{\delta, t} \in [0, \varepsilon_\delta]^k \forall t \geq 1 | \lambda_1 = \lambda'') \geq Pr(\lambda_{\delta, t} \in [0, \varepsilon_\delta]^k \forall t \geq 1 | \lambda_1 = \lambda') > 0$ for $\lambda'' \leq \lambda'$. Therefore, there exists an $\varepsilon^* \in (0, \varepsilon_\delta]$ such that for $\lambda_1 \in (0, \varepsilon^*]^k$, with probability

uniformly bounded away from zero, $\lambda_t \in [0, \varepsilon_\delta]^k$ for all $t \geq 1$. Therefore, for $\lambda_1 \in (0, \varepsilon^*]^k$, $\lim_{t \rightarrow \infty} \lambda_t = 0^k$ with probability uniformly bounded away from zero i.e. 0^k is locally stable.⁴⁵

The proofs for the other stationary beliefs $\lambda^* \in \{0, \infty\}^k$ are analogous. If $\lambda_i^* = \infty$, substitute λ_i^{-1} for type θ_i and modify the transition rules accordingly.

Part 2. We next show that $\lambda^* \in \{0, \infty\}^k$ is unstable if $\lambda^* \notin \Lambda(L)$ and $\gamma_i(\omega, \lambda^*) \neq 0$ for some θ_i with $\lambda^* \notin \Lambda_i(L)$. Without loss of generality, order the types so that the first κ types correspond to $\lambda_i^* = 0$ and the latter $k - \kappa$ types correspond to $\lambda_i^* = \infty$. Suppose $\lambda_1^* = 0$ but $\gamma_1(L, \lambda^*) > 0$. By the continuity property established in Lemma 2, there exists an $\varepsilon > 0$ such that for neighborhood $B_\varepsilon(\lambda^*) \equiv [0, \varepsilon]^\kappa \times [1/\varepsilon, \infty]^{k-\kappa}$,

$$\sum_{a \in \mathcal{A}} \psi(a|L, \lambda^*) g_1(a) > 0, \quad (9)$$

where $g_1(a) \equiv \inf_{\lambda \in B_\varepsilon(\lambda^*)} \log \frac{\hat{\psi}_1(a|R, \lambda)}{\hat{\psi}_1(a|L, \lambda)}$ denotes the minimal update for type θ_1 following action a across $B_\varepsilon(\lambda^*)$. Let $\underline{g}_1 \equiv \min_{a \in \mathcal{A}} g_1(a)$ denote the minimal update across all actions.

For any $\delta \in (0, \varepsilon]$, let $\psi_\delta(a) \equiv \inf_{\lambda \in B_\delta(\lambda^*)} \psi(a|L, \lambda)$ and $\tau_1(\delta) \equiv \min\{\tau | \lambda_t \in B_\delta(\lambda^*)\}$ denote the first time $\langle \lambda_t \rangle$ enters $B_\delta(\lambda^*)$. Define a linear system $\langle \lambda_{\delta,t} \rangle_{t=1}^\infty$ as follows: $\lambda_{\delta,t} = \lambda_{1,t}$ for $t \leq \tau_1(\delta)$, where $\lambda_{1,t}$ denotes the likelihood ratio for θ_1 at time t . For $t > \tau_1(\delta)$, for each $a \in \mathcal{A}$, whenever $\tilde{a}_t = a$ and $\lambda_t \in B_\delta(\lambda^*)$, $\log \lambda_{\delta,t+1} = \log \lambda_{\delta,t} + g_1(a)$ with probability $\psi_\delta(a)/\psi(a|L, \lambda_t)$ and $\log \lambda_{\delta,t+1} = \log \lambda_{\delta,t} + \underline{g}_1$ otherwise. When $\omega = L$, $\psi_\delta(a)$ is the joint probability of $\tilde{a}_t$ and the former event, while $\psi_\delta \equiv 1 - \sum_{a \in \mathcal{A}} \psi_\delta(a)$ is the cumulative probability of the latter event across all actions. To complete the construction, for $t > \tau_1(\delta)$, if $\lambda_t \notin B_\delta(\lambda^*)$ then independent of $\tilde{a}_t$, $\log \lambda_{\delta,t+1} = \log \lambda_{\delta,t} + g_1(a)$ with probability $\psi_\delta(a)$ for each $a \in \mathcal{A}$ and $\log \lambda_{\delta,t+1} = \log \lambda_{\delta,t} + \underline{g}_1$ otherwise. By construction, if $\lambda_{1,t} \geq \lambda_{\delta,t}$ and $\lambda_t \in B_\delta(\lambda^*)$, then $\lambda_{1,t+1} \geq \lambda_{\delta,t+1}$. The increments $(\log \lambda_{\delta,t+1} - \log \lambda_{\delta,t})_{t=\tau_1(\delta)}^\infty$ form an i.i.d. process with expectation equal to

$$\underline{\psi}_\delta \underline{g}_1 + \sum_{a \in \mathcal{A}} \psi_\delta(a) g_1(a). \quad (10)$$

By Eq. (9), Eq. (10) is strictly positive for sufficiently small δ (this follows from $\lambda \mapsto \psi(a|\omega, \lambda)$ continuous at λ^* (Lemma 2) and $\lim_{\delta \rightarrow 0} \underline{\psi}_\delta = 0$). Let $\delta_1 \in (0, \varepsilon]$ denote an upper bound such that Eq. (10) is strictly positive for all $\delta < \delta_1$.

Fix $\delta \in (0, \delta_1)$. Conditional on the event $\tau_1(\delta) = T$ for some $T < \infty$, by the Law of Large Numbers, $\frac{1}{t-T} \sum_{s=T}^{t-1} (\log \lambda_{\delta,s+1} - \log \lambda_{\delta,s})$ converges to Eq. (10) almost surely, which is

⁴⁵One can apply an identical argument to any component i of the likelihood ratio process to show that when $0^k \in \Lambda_i(\omega)$, for small enough $\varepsilon > 0$ we can construct an analogous one-dimensional process that starts at $\lambda_{i,1}$ and, for some $\varepsilon^* \in (0, \varepsilon]$, remains in $B_\varepsilon(0)$ with probability uniformly bounded away from zero across $\lambda_{i,1} \in [0, \varepsilon^*]$. Further, this process bounds $\langle \lambda_{i,t} \rangle$ from above provided $\lambda_s \in B_\varepsilon(0^k)$ for all $s < t$. This establishes that if $\langle \lambda_t \rangle$ almost surely exits $B_\varepsilon(0^k)$ and $0^k \in \Lambda_i(\omega)$, then with probability uniformly bounded away from zero, $\langle \lambda_{j,t} \rangle$ exits $B_\varepsilon(0)$ for some $j \neq i$. A similar observation holds for the other stationary beliefs. We use this observation in the proof of Lemma 4.

positive. Given $\lambda_{1,1} \in (0, \infty)$ and [Lemma 2](#), $\lambda_{1,T} \in (0, \infty)$. Therefore,

$$\lim_{t \rightarrow \infty} \log \lambda_{\delta,t} = \lim_{t \rightarrow \infty} (\log \lambda_{1,T} + \sum_{s=T}^{t-1} (\log \lambda_{\delta,s+1} - \log \lambda_{\delta,s})) = \infty \text{ a.s.}$$

This implies that if $\langle \lambda_t \rangle$ remains in $B_\delta(\lambda^*)$ for all $t > T$, then almost surely $\lim_{t \rightarrow \infty} \log \lambda_{1,t} \geq \lim_{t \rightarrow \infty} \log \lambda_{\delta,t} = \infty$, which is a contradiction. Therefore, conditional on $\tau_1(\delta) = T$, almost surely $\tau_2(\delta) \equiv \min\{\tau > \tau_1(\delta) | \lambda_t \notin B_\delta(\lambda^*)\}$ is finite. Since the log likelihood ratio process is linear, the same reasoning establishes that whenever $\langle \lambda_t \rangle$ enters $B_\delta(\lambda^*)$, it almost surely exits. Therefore, letting $\tau(\delta) \equiv \min\{\tau | \lambda_t \in B_\delta(\lambda^*) \forall t \geq \tau\}$ be the first time $\langle \lambda_t \rangle$ enters $B_\delta(\lambda^*)$ and never exits, it follows that conditional on $\tau_1(\delta) = T$, $\tau(\delta) = \infty$ almost surely. Given $\tau(\delta) \geq \tau_1(\delta)$, conditional on $\tau_1(\delta) = \infty$, $\tau(\delta) = \infty$. Together, this establishes that $\tau(\delta) = \infty$ almost surely, and therefore, $Pr(\lambda_t \rightarrow \lambda^*) = 0$. The logic is analogous when $\lambda_1^* = \infty$ and $\gamma_1(L, \lambda^*) < 0$.

Part 3. Finally, [Lemma 3](#) established that for $\lambda^* \in [0, \infty]^k \setminus \{0, \infty\}^k$, $Pr(\lambda_t \rightarrow \lambda^*) = 0$ for all $\lambda_1 \in (0, \infty)^k$. Therefore, such λ^* are unstable. $\square$

A.3 Global Stability

Preliminary Notation. From [Theorem 1](#), for each $\lambda^* \in \Lambda(\omega)$, there exists an $\varepsilon > 0$ and a *stable* neighborhood $B_\varepsilon(\lambda^*)$ such that $Pr(\lambda_t \rightarrow \lambda^* | \lambda_1)$ is uniformly bounded away from zero across $\lambda_1 \in B_\varepsilon(\lambda^*)$. Also, for each $\lambda^* \in \{0, \infty\}^k \setminus \Lambda(\omega)$, there exists an $\varepsilon > 0$ and an *unstable* neighborhood $B_\varepsilon(\lambda^*)$ such that when $\lambda_1 \in \text{int}(B_\varepsilon(\lambda^*))$, $\langle \lambda_t \rangle$ almost surely leaves $B_\varepsilon(\lambda^*)$. Additionally, for each $\varepsilon' > 0$ such that $B_{\varepsilon'}(\lambda^*)$ is an *unstable* neighborhood, there exists an $\varepsilon \in (0, \varepsilon')$ such that for each θ_i with $\lambda^* \in \Lambda_i(\lambda^*)$, if $\langle \lambda_{i,t} \rangle$ is in $B_\varepsilon(\lambda_i^*)$, then as long as $\langle \lambda_t \rangle$ is in $B_{\varepsilon'}(\lambda^*)$, $\langle \lambda_{i,t} \rangle$ is bounded above by a process that with probability uniformly bounded away from zero (i) converges to λ_i^* and (ii) almost surely does not leave $B_{\varepsilon'}(\lambda_i^*)$. Moreover, given the continuity at certainty property established in [Lemma 2](#), for each $\lambda^* \in \{0, \infty\}^k$ and each $\varepsilon' > 0$ such that $B_{\varepsilon'}(\lambda^*)$ is a stable or unstable neighborhood of $\lambda^* \in \{0, \infty\}^k \setminus \Lambda^*$, we can select a sufficiently small $\varepsilon > 0$ such that $B_\varepsilon(\lambda^*)$ is a stable or unstable neighborhood and if $\lambda_t \in B_\varepsilon(\lambda^*)$, then for each θ_i such that $(\lambda')_i \neq (\lambda^*)_i$, $Pr(\lambda_{i,t+1} \in B_{\varepsilon'}((\lambda')_i)) = 0$. This also implies that $Pr(\lambda_{t+1} \in B_{\varepsilon'}(\lambda')) = 0$. Fixing state ω , choose an $\varepsilon(\lambda^*) \in (0, 1)$ for each $\lambda^* \in \{0, \infty\}^k$ such that the set of neighborhoods $\{B_{\varepsilon(\lambda^*)}(\lambda^*)\}_{\lambda^* \in \{0, \infty\}^k}$ satisfies these properties. Define

$$E \equiv \min_{\lambda^* \in \{0, \infty\}^k} -\log \varepsilon(\lambda^*). \quad (11)$$

Note $E \in (0, \infty)$. If $\log \lambda_i \in \mathbb{R} \cup \{-\infty, \infty\} \setminus [-E, E]$ for each $\theta_i \in \Theta_S$, then λ is contained in one of these stable or unstable neighborhoods. Let

$$B_E(\lambda^*) \equiv \{\lambda \in [0, \infty]^k | \log \lambda_i < -E \text{ if } \lambda_i^* = 0 \text{ and } \log \lambda_i > E \text{ if } \lambda_i^* = \infty\} \quad (12)$$

denote the corresponding neighborhood for each stationary λ^* , where in a slight abuse of notation, when using E we switch from the neighborhood subscript denoting the bound for the likelihood ratio to denoting the bound for the log likelihood ratio to simplify future notation. Let $\mathcal{B} \equiv \cup_{\lambda^* \in \Lambda(\omega)} B_E(\lambda^*)$ denote the union of the stable neighborhoods and let $\mathcal{B}_U \equiv \cup_{\lambda^* \in \{0, \infty\}^k \setminus \Lambda(\omega)} B_E(\lambda^*)$ denote the union of the unstable neighborhoods. We will use

these neighborhoods in subsequent proofs.

Proof of Theorem 2. Suppose the agreement outcome $0^k \in \Lambda(\omega)$ is locally stable. By Lemma 2, a_1 occurs with positive probability and observing a_1 decreases the likelihood ratio. Given initial likelihood ratio $\lambda_1 \in (0, \infty)^k$, let N be the minimum number of consecutive a_1 actions required for the likelihood ratio to reach the stable neighborhood of 0^k defined in Eq. (12), i.e. $\lambda_{N+1} \in B_E(0^k)$. By Lemma 2, the change in the likelihood ratio following a_1 is bounded away from zero. Therefore, $N < \infty$. Further, given a_1 occurs with positive probability each period, the probability of a_1 occurring N times is strictly positive. Let $\tau_1 \equiv \min\{t | \lambda_t \in B_E(0^k)\}$ be the first time that $\langle \lambda_t \rangle$ enters $B_E(0^k)$, $\tau_2 \equiv \min\{t > \tau_1 | \lambda_t \notin B_E(0^k)\}$ be the first that $\langle \lambda_t \rangle$ leaves $B_E(0^k)$ after entering, and $\tau_3 \equiv \min\{\tau | \lambda_t \in B_E(0^k) \forall t \geq \tau\}$ be the first time the likelihood ratio enters $B_E(0^k)$ and never leaves. We know that $Pr(\tau_1 < \infty) > 0$, since the probability of transitioning from λ_1 to $B_E(0^k)$ is bounded below by the probability of observing N consecutive a_1 actions starting in period one. Also, $Pr(\tau_2 = \infty) > 0$, since by local stability, when the likelihood ratio is in $B_E(0^k)$, with positive probability, it never leaves. Therefore, $Pr(\tau_3 < \infty) > Pr(\tau_1 < \infty \wedge \tau_2 = \infty) > 0$. Therefore, with positive probability, the likelihood ratio eventually enters and remains in $B_E(0^k)$. By Theorem 1, if the likelihood ratio remains in $B_E(0^k)$ for all t , beliefs almost surely converge to 0^k . Therefore, if $0^k \in \Lambda(\omega)$, then from any initial belief $\lambda_1 \in (0, \infty)^k$, $Pr(\lambda_t \rightarrow 0^k) > 0$. The proof for agreement outcome ∞^k is analogous. $\square$

Global Stability of Disagreement. We first state a result that uses a much weaker but more complicated to verify condition called *separability* to establish the global stability of a disagreement outcome (see Theorem 7 below). Starting from any interior initial belief, separability uses all of the actions to separate the beliefs of the different types and reach a neighborhood of the disagreement outcome. Therefore, together with local stability, the separability condition implies global stability. We then prove Theorem 3 by showing that maximal accessibility—which is easier to verify but only uses actions a_1 and a_M to separate beliefs—implies separability.

Define $\Psi(\lambda)$ as the matrix consisting of the log of the ratios of the perceived action probabilities in each state at belief λ , where each row corresponds to the ratios for social type θ_i and each column corresponds to the ratios for action a_m ,

$$\Psi(\lambda)_{im} \equiv \log \frac{\hat{\psi}_i(a_m | R, \lambda)}{\hat{\psi}_i(a_m | L, \lambda)}, \quad (13)$$

and define the submatrix

$$\Psi[\theta_i, \theta_j; a_1, a_M](\lambda) \equiv \begin{pmatrix} \log \frac{\hat{\psi}_i(a_1 | R, \lambda)}{\hat{\psi}_i(a_1 | L, \lambda)} & \log \frac{\hat{\psi}_i(a_M | R, \lambda)}{\hat{\psi}_i(a_M | L, \lambda)} \\ \log \frac{\hat{\psi}_j(a_1 | R, \lambda)}{\hat{\psi}_j(a_1 | L, \lambda)} & \log \frac{\hat{\psi}_j(a_M | R, \lambda)}{\hat{\psi}_j(a_M | L, \lambda)} \end{pmatrix} \quad (14)$$

as these ratios for social types θ_i and θ_j from actions a_1 and a_M . We use $\Psi(\lambda)$ to define separability (Definition 8) and use $\Psi[\theta_i, \theta_j; a_1, a_M](\lambda)$ to show that maximal accessibility implies separability.

Given a belief $\lambda^* = (0^\kappa, \infty^{k-\kappa})$, we say λ^* is separable at zero for type θ_κ if there exists a finite sequence of actions that are on average more likely in state R than state L for θ_κ

and the types with $\lambda_i^* = \infty$, and are on average more likely in state L than state R for the remaining $\kappa - 1$ types with $\lambda_i^* = 0$. In other words, in a neighborhood of λ^* , there exists a finite sequence of actions that will decrease the beliefs of types $(\theta_1, \dots, \theta_{\kappa-1})$ and increase the beliefs of types $(\theta_\kappa, \dots, \theta_k)$. Separability at infinity is similar—in a neighborhood of λ^* , there exists a finite sequence of actions that will on average decrease the beliefs of types $(\theta_1, \dots, \theta_{\kappa+1})$ and increase the beliefs of types $(\theta_{\kappa+2}, \dots, \theta_k)$. The following definition formalizes this notion for an arbitrary stationary belief.

Definition 8 (Separability ($k \geq 2$)). (i) Belief $\lambda^* \in \{0, \infty\}^k \setminus \infty^k$ is separable at zero for type θ_i with $\lambda_i^* = 0$ if there exist vectors $c \in [0, \infty)^{|A|}$ and $G \in \mathbb{R}^k$ with $G_i > 0$, $G_j > 0$ for all j with $\lambda_j^* = \infty$ and $G_j < 0$ for all $j \neq i$ with $\lambda_j^* = 0$, such that $\Psi(\lambda^*) \cdot c = G$; (ii) Belief $\lambda^* \in \{0, \infty\}^k \setminus 0^k$ is separable at infinity for type θ_i with $\lambda_i^* = \infty$ if there exist vectors $c \in (0, \infty)^{|A|}$ and $G \in \mathbb{R}^k$ with $G_i < 0$, $G_j > 0$ for all $j \neq i$ with $\lambda_j^* = \infty$ and $G_j < 0$ for all j with $\lambda_j^* = 0$, such that $\Psi(\lambda^*) \cdot c = G$.

The following result shows that separability can be used to establish the global stability of a disagreement outcome for the case of two social types. [Theorem 7'](#) in [Online Appendix D](#) presents an analogous result for the case of more than two social types.

Theorem 7 (Global Stability of Disagreement ($k = 2$)). Consider a learning environment that is identified at certainty and satisfies [Assumptions 1](#) to [4](#). If $(0, \infty) \in \Lambda(\omega)$ and $(0, 0)$ is separable at zero for θ_2 or (∞, ∞) is separable at infinity for θ_1 , then $(0, \infty)$ is globally stable in state ω . Similarly, if $(\infty, 0) \in \Lambda(\omega)$ and $(0, 0)$ is separable at zero for θ_1 or (∞, ∞) is separable at infinity for θ_2 , then $(\infty, 0)$ is globally stable in state ω .

We use [Lemmas 5](#) and [6](#) to establish [Theorem 7](#). We say a belief $\lambda_2^* \in \{0, \infty\}^k$ is adjacent to a belief $\lambda_1^* \in \{0, \infty\}^k$ if all but one component of the vectors are equal, i.e. there exists exactly one $i = 1, \dots, k$ such that $(\lambda_2^*)_i \neq (\lambda_1^*)_i$ (where the subscript of λ^* is used to index different stationary beliefs). In other words, all but one social type have the same beliefs in λ_1^* and λ_2^* . We say a belief $\lambda_2^* \in \{0, \infty\}^k$ is adjacently accessible from adjacent belief $\lambda_1^* \in \{0, \infty\}^k$ if, given the likelihood ratio process is at an interior value in a neighborhood of λ_1^* , then with positive probability it enters a neighborhood of λ_2^* in finite time.

Definition 9 (Adjacently Accessible ($k \geq 2$)). Belief $\lambda_2^* \in \{0, \infty\}^k$ is adjacently accessible from adjacent belief $\lambda_1^* \in \{0, \infty\}^k$ if for any $\varepsilon_2 > 0$, there exists an $\varepsilon_1 > 0$ such that for any $\lambda \in \text{int}(B_{\varepsilon_1}(\lambda_1^*))$, there exists a $\tau(\lambda) < \infty$ such that if $\lambda_t = \lambda$, then $\Pr(\lambda_{t+\tau(\lambda)} \in \text{int}(B_{\varepsilon_2}(\lambda_2^*))) > 0$.

The following lemma establishes that separability can be used to establish adjacent accessibility for any number of social types.

Lemma 5 (Adjacently Accessible ($k \geq 2$)). Consider a learning environment that satisfies [Assumptions 1](#) to [4](#). If $\lambda_1^* \in \{0, \infty\}^k \setminus \infty^k$ with $(\lambda_1^*)_i = 0$ is separable at zero for θ_i , then adjacent belief λ_2^* with $(\lambda_2^*)_i = \infty$ is adjacently accessible from λ_1^* . If $\lambda_1^* \in \{0, \infty\}^k \setminus 0^k$ with $(\lambda_1^*)_i = \infty$ is separable at infinity for θ_i , then adjacent belief λ_2^* with $(\lambda_2^*)_i = 0$ is adjacently accessible from λ_1^* .

Proof. Given $\kappa \in \{1, \dots, k\}$, let $\lambda_1^* = (0^\kappa, \infty^{k-\kappa})$, $\lambda_2^* = (0^{\kappa-1}, \infty^{k-\kappa+1})$ and suppose λ_1^* is separable at zero for θ_κ . We will show that for any $\varepsilon_2 > 0$, there exists an $\varepsilon_1 > 0$

such that for any $\lambda \in \text{int}(B_{\varepsilon_1}(\lambda_1^*))$, there exists a $\tau(\lambda) < \infty$ such that if $\lambda_t = \lambda$, then $Pr(\lambda_{t+\tau(\lambda)} \in \text{int}(B_{\varepsilon_2}(\lambda_2^*))) > 0$.⁴⁶ Since the log likelihood ratio process is linear, we can show this for $\lambda_1 = \lambda$ and it immediately follows that it holds when $\lambda_t = \lambda$.

We first define several pieces of notation. Given $\varepsilon > 0$, recall that $B_\varepsilon(\lambda_1^*) \equiv [0, \varepsilon)^\kappa \times (1/\varepsilon, \infty)^{k-\kappa}$ denotes an ε -neighborhood of λ_1^* . Let $\ell(\varepsilon) \equiv -\log \varepsilon$. Then $[-\infty, -\ell(\varepsilon))^\kappa \times (\ell(\varepsilon), \infty]^{k-\kappa}$ denotes the corresponding neighborhood of $\log \lambda_1^*$. Let

$$g_{\varepsilon,i}(a) \equiv \begin{cases} \inf_{\lambda \in B_\varepsilon(\lambda_1^*)} \log \frac{\hat{\psi}_i(a|R, \lambda)}{\hat{\psi}_i(a|L, \lambda)} & i \geq \kappa \\ \sup_{\lambda \in B_\varepsilon(\lambda_1^*)} \log \frac{\hat{\psi}_i(a|R, \lambda)}{\hat{\psi}_i(a|L, \lambda)} & i < \kappa \end{cases}$$

denote the smallest (largest) update to the log likelihood ratio for type θ_i with $i \geq \kappa$ ($i < \kappa$) following action a when the likelihood ratio is in the neighborhood $B_\varepsilon(\lambda_1^*)$, and let

$$\bar{g}_{\varepsilon,\kappa}(a) \equiv \sup_{\lambda \in B_\varepsilon(\lambda_1^*)} \log \frac{\hat{\psi}_\kappa(a|R, \lambda)}{\hat{\psi}_\kappa(a|L, \lambda)}$$

denote the largest update to the log likelihood ratio for type θ_κ following action a when the likelihood ratio is in the neighborhood $B_\varepsilon(\lambda_1^*)$. By Lemma 2, for any $\varepsilon > 0$, $g_{\varepsilon,i}(a)$ and $g_{\varepsilon,\kappa}(a)$ are bounded for all $i = 1, \dots, k$ and $a \in \mathcal{A}$.

We next define a set of processes that we use to separate the beliefs of types $(\theta_1, \dots, \theta_{\kappa-1})$ and type θ_κ . By λ_1^* separable at zero for θ_κ , there exist vectors $c \in [0, \infty)^k$ and $G \in \mathbb{R}^k$ that satisfy the separability condition. Moreover, since the rationals are dense in the reals, there exists vector $c \in [0, \infty)^k$ of rational numbers and vector $G \in \mathbb{R}^k$ that satisfies the separability condition. Therefore, there exists an $\varepsilon_3 > 0$ and finite integers $c_a \geq 0$ for each $a \in \mathcal{A}$ such that

$$G_i \equiv \sum_{a \in \mathcal{A}} c_a g_{\varepsilon_3,i}(a), \quad (15)$$

with $G_i > 0$ for all $i \geq \kappa$ and $G_i < 0$ for all $i < \kappa$. Let

$$\bar{G}_\kappa \equiv \sum_{a \in \mathcal{A}} c_a \bar{g}_{\varepsilon_3,\kappa}(a) \quad (16)$$

and note $\bar{G}_\kappa \geq G_\kappa > 0$. Next we define processes $\xi_{i,t} \equiv \sum_{s=1}^{t-1} g_{\varepsilon_3,i}(a_s)$ and $\bar{\xi}_{\kappa,t} \equiv \sum_{s=1}^{t-1} \bar{g}_{\varepsilon_3,\kappa}(a_s)$ for $t > 1$ and $\xi_{i,1} = \bar{\xi}_{\kappa,1} = 0$. Consider an action sequence with c_a realizations of each a , starting with the action that minimizes $\bar{g}_{\varepsilon_3,\kappa}(a)$, followed by the action that leads to the second lowest $\bar{g}_{\varepsilon_3,\kappa}(a)$, and so on. Following this action sequence, at time $\tau_1 \equiv \sum_{a \in \mathcal{A}} c_a + 1 < \infty$, the process $\xi_{i,\tau_1} = G_i$ by Eq. (15) and $\bar{\xi}_{\kappa,\tau_1} = \bar{G}_\kappa$ by Eq. (16). For $i \geq \kappa$, $G_i > 0$, and therefore, $\xi_{i,\tau_1} > 0$, while for $i < \kappa$, $G_i < 0$, and therefore, $\xi_{i,\tau_1} < 0$. Since $g_{\varepsilon,i}(a)$ and $g_{\varepsilon,\kappa}(a)$ are bounded, there exists a $\underline{K} > 0$ such that $\xi_{i,t} \geq -\underline{K}$ for all $t < \tau_1$ and $i > \kappa$, and there exists a $\bar{K} > 0$ such that $\xi_{i,t} < \bar{K}$ for all $t < \tau_1$ and $i < \kappa$. Therefore, for any $K > 0$, there exists an $N_K < \infty$ such that following N_K repetitions of this action sequence, at time $\tau_K \equiv N_K \sum_{a \in \mathcal{A}} c_a + 1 < \infty$, the following properties hold:

⁴⁶Note that the subscript of λ^* is used to index different stationary beliefs, while the subscript of λ is used to denote the likelihood ratio process at a given time i.e. λ_t is the value of the process at time t .

- (i) $\xi_{i,\tau_K} < -K$ for all $i < \kappa$,
- (ii) $\xi_{i,\tau_K} > 0$ for all $i \geq \kappa$,
- (iii) For all $t < \tau_K$, $\xi_{i,t} \leq \bar{K}$ for all $i < \kappa$ and $\xi_{i,t} \geq -\underline{K}$ for all $i > \kappa$,
- (iv) $\bar{\xi}_{\kappa,t} \leq N_K \bar{G}_\kappa$ for all $t \leq \tau_K$, with equality at $t = \tau_K$ (the inequality holds due to the order of the action sequence described above).

In summary, following N_K repetitions of the action sequence, the processes $\langle \xi_{i,t} \rangle$ for $i < \kappa$ and the process $\langle \xi_{\kappa,t} \rangle$ are separated by at least K , and at all t during the sequence, the processes $\langle \xi_{i,t} \rangle$ for $i < \kappa$ are bounded above by $\bar{K}$ and the processes $\langle \xi_{i,t} \rangle$ for $i > \kappa$ are bounded below by $-\underline{K}$.

These processes bound the updates to the likelihood ratio while it remains in an ε_3 -neighborhood of λ_1^* : when $\lambda_s \in B_{\varepsilon_3}(\lambda_1^*)$ for all $s \leq t$, then $\log \lambda_{i,t} - \log \lambda_{i,1} \leq \xi_{i,t} \leq \bar{K}$ for $i < \kappa$, $\xi_{\kappa,t} \leq \log \lambda_{\kappa,t} - \log \lambda_{\kappa,1} \leq \bar{\xi}_{\kappa,t}$, and $-\underline{K} \leq \xi_{i,t} \leq \log \lambda_{i,t} - \log \lambda_{i,1}$ for $i > \kappa$. Further, if the likelihood ratio remains in an ε_3 -neighborhood of λ_1^* for all $t \leq \tau_K$, then at time τ_K the update to the log likelihood ratios of types $(\theta_1, \dots, \theta_{\kappa-1})$ and type θ_κ are separated by K , i.e. $\log \lambda_{\kappa,\tau_K} - \log \lambda_{\kappa,1} - (\log \lambda_{i,\tau_K} - \log \lambda_{i,1}) > K$ for $i < \kappa$.

Using these processes, we establish local accessibility in three steps. Fix an $\varepsilon_3 > 0$ that satisfies Eq. (15), with corresponding processes $\xi_{i,t}$ and $\bar{\xi}_{\kappa,t}$, constants c_a , G_i and $\bar{G}_\kappa$, and bounds $\underline{K}$ and $\bar{K}$ defined above. Fix $\varepsilon_2 \in (0, \varepsilon_3)$. Given ε_2 and ε_3 , there exists an $N_2 < \infty$ such that for any $\log \lambda_{\kappa,t} > -\ell(\varepsilon_3)$, following N_2 realizations of a_M , $\log \lambda_{\kappa,t+N_2} > \ell(\varepsilon_2)$ (recall $\ell(\varepsilon) \equiv -\log \varepsilon$). Let $K_2 > 0$ be the most that the log likelihood ratio increases for any type θ_i with $i < \kappa$ across all beliefs in $B_{\varepsilon_2}(\lambda_1^*)$ following these N_2 actions. Note $K_2 < \infty$ since $N_2 < \infty$ and no type believes a_M is perfectly informative. Fix $K \geq \bar{K} + K_2$ and, given the corresponding N_K , choose $\varepsilon_1 > 0$ such that for $\lambda \in B_{\varepsilon_1}(\lambda_1^*)$, $\log \lambda_i < -\ell(\varepsilon_2) - \max(\bar{K}, N_K \bar{G}_\kappa)$ for $i \leq \kappa$ and $\log \lambda_i > \ell(\varepsilon_2) + \underline{K}$ for $i > \kappa$. Note $\varepsilon_1 < \varepsilon_2$. Choose an arbitrary $\lambda \in \text{int}(B_{\varepsilon_1}(\lambda_1^*))$ and suppose $\lambda_1 = \lambda$.

Step 1. Repeat the sequence of c_a realizations of each action a (ordered as described above) N_K times, where N_K and c_a are finite as established above. It follows from the chosen ε_1 and properties (iii) and (iv) above that λ_t remains in $\text{int}(B_{\varepsilon_2}(\lambda_1^*))$ for all $t \leq \tau_K$, where $\tau_K \equiv N_K \sum_A c_a + 1$ as defined above, and from properties (i) and (ii) that $\log \lambda_{i,\tau_K} < \log \lambda_{i,1} - K$ for $i < \kappa$, $\log \lambda_{\kappa,\tau_K} \in (\log \lambda_{\kappa,1}, -\ell(\varepsilon_2))$, and $\log \lambda_{i,\tau_K} > \ell(\varepsilon_1)$ for $i > \kappa$.

Step 2. Continue repeating this sequence of c_a realizations of each action a until the first action such that $\log \lambda_{\kappa,t} > -\ell(\varepsilon_3)$. It is possible to do this with a finite number of actions $n(\lambda_{\kappa,1})$ since $\lambda_{\kappa,1} \in (0, \infty)$ and $\log \lambda_{\kappa,\tau_K} > \log \lambda_{\kappa,1}$. As shown in Step 1, the likelihood ratios for $i \neq \kappa$ remain in $\text{int}(B_{\varepsilon_2}(\lambda_1^*))$ after every action in this sequence. For $i < \kappa$, given $\log \lambda_{i,\tau_K} < \log \lambda_{i,1} - K < -\ell(\varepsilon_1) - K$ following Step 1, it follows that $\log \lambda_{i,t} < -\ell(\varepsilon_1) - K + \bar{K}$ for all $t \in \{\tau_K + 1, \dots, \tau_K + n(\lambda_{\kappa,1})\}$. We use this observation in the following step.

Step 3. Repeat N_2 realizations of a_M so that $\log \lambda_{\kappa,\tau_K+n(\lambda_{\kappa,1})+N_2} > \ell(\varepsilon_2)$. Since a_M increases the likelihood ratio, beliefs remain in $\text{int}(B_{\varepsilon_2}(\lambda_1^*))$ for $i > \kappa$. For $i < \kappa$, given $\log \lambda_{i,\tau_K+n(\lambda_{\kappa,1})} < -\ell(\varepsilon_1) - K + \bar{K}$, after N_2 realizations of a_M , $\log \lambda_{i,t} < -\ell(\varepsilon_1) - K + \bar{K} + K_2$. Given $K \geq \bar{K} + K_2$, beliefs remain in $\text{int}(B_{\varepsilon_2}(\lambda_1^*))$ for $i < \kappa$.

Following these three steps, the likelihood ratio is in $\text{int}(B_{\varepsilon_2}(\lambda_2^*))$. Each step requires a finite number of actions that occur with positive probability. Therefore, given ε_1 and ε_2 defined above, for any $\lambda \in \text{int}(B_{\varepsilon_1}(\lambda_1^*))$, there exists a $\tau(\lambda) < \infty$ such that if $\lambda_1 = \lambda$,

then $Pr(\lambda_{1+\tau} \in \text{int}(B_{\varepsilon_2}(\lambda_2^*))) > 0$. The cases of the other disagreement outcomes $\lambda_1^* \in \{0, \infty\}^k \setminus \infty^k$ that are separable at zero are analogous, as are the cases of $\lambda_1^* \in \{0, \infty\}^k \setminus 0^k$ that are separable at infinity. $\square$

We say an belief $\lambda^* \in \{0, \infty\}^k$ is accessible if, from any interior initial belief, with positive probability the likelihood ratio process enters a neighborhood of λ^* in finite time.

Definition 10 (Accessible ($k \geq 2$)). A belief $\lambda^* \in \{0, \infty\}^k$ is accessible if for any initial belief $\lambda_1 \in (0, \infty)^k$ and any $\varepsilon > 0$, there exists a $\tau < \infty$ such that $Pr(\lambda_\tau \in B_\varepsilon(\lambda^*)) > 0$.

We next show that adjacent accessibility can be used to establish accessibility.

Lemma 6 (Accessible Disagreement ($k \geq 2$)). Consider a learning environment that satisfies Assumptions 1 to 4. Disagreement outcome $\lambda^* \in \{0, \infty\}^k \setminus \{0^k, \infty^k\}$ is accessible if there exists a finite sequence of stationary beliefs $\lambda_1^*, \lambda_2^*, \dots, \lambda_L^* = \lambda^*$, with $\lambda_1^* \in \{0^k, \infty^k\}$ and λ_{l+1}^* adjacently accessible from adjacent belief λ_l^* for $l = 1, \dots, L-1$.

Proof. Given disagreement outcome λ^* , suppose there exists a finite sequence of stationary beliefs $\lambda_1^*, \lambda_2^* \dots \lambda_L^*$, with $\lambda_1^* \in \{0^k, \infty^k\}$, λ_{l+1}^* adjacently accessible from adjacent belief λ_l^* for $l = 1, \dots, L-1$ and $\lambda_L^* = \lambda^*$. By definition of adjacently accessible, for any $\varepsilon_L > 0$, there exists an $\varepsilon_{L-1} > 0$ and $\tau_L < \infty$ such that if $\lambda_t \in \text{int}(B_{\varepsilon_{L-1}}(\lambda_{L-1}^*))$, then $Pr(\lambda_{t+\tau_L} \in B_{\varepsilon_L}(\lambda_L^*)) > 0$. Iterating back to λ_1^* , for any $\varepsilon_L > 0$, there exists an $\varepsilon_1 > 0$ and $\tau_2 < \infty$ such that if $\lambda_t \in \text{int}(B_{\varepsilon_1}(\lambda_1^*))$, then $Pr(\lambda_{t+\sum_{l=2}^L \tau_l} \in B_{\varepsilon_L}(\lambda_L^*)) > 0$. Consider agreement outcome $\lambda_1^* \in \{0^k, \infty^k\}$. By Theorem 2, for any initial belief $\lambda_1 \in (0, \infty)^k$ and any $\varepsilon_1 > 0$, there exists a finite sequence of τ_1 actions that occur with positive probability such that following this sequence, $\lambda_{\tau_1+1} \in \text{int}(B_{\varepsilon_1}(\lambda_1^*))$. Therefore, starting from any initial belief, $Pr(\lambda_{\tau_1+1} \in \text{int}(B_{\varepsilon_1}(\lambda_1^*))) > 0$. Therefore, for any $\varepsilon_L > 0$ and initial belief $\lambda_1 \in (0, \infty)^k$, $Pr(\lambda_\tau \in B_{\varepsilon_L}(\lambda_L^*)) > 0$, where $\tau \equiv \sum_{l=1}^L \tau_l + 1 < \infty$ since each $\tau_l < \infty$. By definition, this means that $\lambda^* = \lambda_L^*$ is accessible. $\square$

Accessibility and local stability together imply global stability. Therefore, by Lemmas 5 and 6, the separability of an agreement outcome combined with the local stability of an adjacently accessible disagreement outcome establishes that the disagreement outcome is globally stable.

Proof of Theorem 7. Suppose $(0, \infty) \in \Lambda(\omega)$ and either $(0, 0)$ is separable at zero for θ_2 or (∞, ∞) is separable at infinity for θ_1 . By Lemma 5, $(0, \infty)$ is adjacently accessible from $(0, 0)$. It follows from Lemma 6 that $(0, \infty)$ is accessible. Fix initial belief $\lambda_1 \in (0, \infty)^2$ and choose $\varepsilon < e^{-E}$, where E is defined in Eq. (11). By accessibility, there exists a finite sequence ξ of N actions that occurs with positive probability, such that following ξ , $\lambda_{N+1} \in B_\varepsilon(0, \infty)$. From $(0, \infty) \in \Lambda(\omega)$, $Pr(\lambda_t \rightarrow (0, \infty) | h = \xi) > 0$. Therefore, from any initial belief, $Pr(\lambda_t \rightarrow (0, \infty)) > 0$, which implies that $(0, \infty)$ is globally stable. The proof of $(\infty, 0)$ is analogous. $\square$

Finally, to prove Theorem 3, we show that maximal accessibility implies the conditions for separability outlined in Theorem 7.

Proof of Theorem 3. Suppose $\theta_2 \succ_{(0,0)} \theta_1$. We show that this implies $(0,0)$ is separable at zero for θ_2 . Since $\theta_2 \succ_{(0,0)} \theta_1$, the submatrix $\Psi[\theta_2, \theta_1; a_1, a_M](0,0)$ defined in Eq. (14) has a positive determinant. Therefore, there exists a $c \in \mathbb{R}_+^2$ that solves

$$\Psi[\theta_2, \theta_1; a_1, a_M](0,0) \cdot c = \begin{pmatrix} 1 \\ 0 \end{pmatrix}.$$

By continuity, there exists a perturbation of c to $\tilde{c} \in \mathbb{R}_+^2$ such that

$$\Psi[\theta_2, \theta_1; a_1, a_M](0,0) \cdot \tilde{c} = \begin{pmatrix} G_2 \\ G_1 \end{pmatrix},$$

where $G_1 < 0$ and $G_2 > 0$. Therefore, by Definition 8, $(0,0)$ is separable at zero for θ_2 , since the definition holds for vector $c' \in (0, \infty)^{|A|}$ with $c'_1 = c_1$, $c'_M = c_2$ and $c'_i = 0$ otherwise. The case where $\theta_2 \succ_{(\infty, \infty)} \theta_1$ is analogous, as is the proof for $(\infty, 0)$. $\square$

A.4 Mixed Learning

Proof of Lemma 4. Fix state ω and consider the mixed learning outcome $(0, \theta_1)$ in which θ_1 's belief converges to zero and θ_2 's belief doesn't converge. Suppose $(0, \theta_1) \notin \Lambda_M(\omega)$, i.e. $(0,0) \in \Lambda_2(\omega)$ or $(0, \infty) \in \Lambda_2(\omega)$. Without loss of generality, consider the case where $(0,0) \in \Lambda_2(\omega)$. Let $\varepsilon \in (0, e^{-E})$, where E is defined in Eq. (11). Suppose $\lambda_{1,1} \in B_\varepsilon(0)$, and let $\tau \equiv \min\{t | \lambda_{1,t} \notin B_\varepsilon(0)\}$ be the first time that θ_1 's belief leaves a neighborhood of zero. We will show that almost surely, either (i) $\tau < \infty$ or (ii) $\langle \lambda_t \rangle$ converges to $(0,0)$ or $(0, \infty)$. By the linearity of the likelihood ratio process, this implies the same holds whenever $\lambda_{1,t} \in B_\varepsilon(0)$, and therefore, almost surely $(0, \theta_1)$ does not occur.

Whenever $\langle \lambda_t \rangle$ is in $B_\varepsilon(0) \times [\varepsilon, 1/\varepsilon]$, $\langle \lambda_t \rangle$ almost surely exits $B_\varepsilon(0) \times [\varepsilon, 1/\varepsilon]$. Therefore, when $\lambda_1 \in B_\varepsilon(0) \times (0, \infty)$, almost surely either $\lambda_{1,t}$ exits $B_\varepsilon(0)$ or $\langle \lambda_t \rangle$ is in $B_\varepsilon(0,0)$ or $B_\varepsilon(0, \infty)$ infinitely often,

$$Pr(\tau < \infty \text{ or } \lambda_t \in B_\varepsilon(0,0) \cup B_\varepsilon(0, \infty) \text{ i.o.}) = 1. \quad (17)$$

We next determine how the behavior of $\langle \lambda_t \rangle$ in $B_\varepsilon(0,0)$ and $B_\varepsilon(0, \infty)$ depends on $\Lambda_1(\omega)$ and $\Lambda_2(\omega)$. The following properties follow from the proof of Theorem 1. (i) Suppose $(0,0) \in \Lambda_1(\omega)$. If $\langle \lambda_t \rangle$ enters $B_\varepsilon(0,0)$, then with probability bounded away from zero across $B_\varepsilon(0,0)$, $\langle \lambda_t \rangle$ converges to $(0,0)$ and almost surely, $\langle \lambda_t \rangle$ either converges to $(0,0)$ or exits $B_\varepsilon(0,0)$. Therefore, if $\langle \lambda_t \rangle$ is in $B_\varepsilon(0,0)$ infinitely often, then $\langle \lambda_t \rangle$ almost surely converges to $(0,0)$. (ii) Suppose $(0,0) \in \Lambda_2(\omega) \setminus \Lambda_1(\omega)$. If $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(0,0))$, then $\langle \lambda_t \rangle$ exits $B_\varepsilon(0,0)$ almost surely and with probability uniformly bounded away from zero, $\langle \lambda_{1,t} \rangle$ exits $B_\varepsilon(0)$ (this follows from Footnote 45). Therefore, if $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(0,0))$ infinitely often, then $\langle \lambda_{1,t} \rangle$ exits $B_\varepsilon(0)$ almost surely, and hence, $\tau < \infty$ almost surely. Combined with Eq. (17), it follows that almost surely either $\tau < \infty$ or $\langle \lambda_t \rangle$ enters $B_\varepsilon(0, \infty)$ infinitely often, $Pr(\tau < \infty \text{ or } \lambda_t \in B_\varepsilon(0, \infty) \text{ i.o.}) = 1$. (iii) When $(0, \infty) \in \Lambda_2(\omega)$, the behavior of $\langle \lambda_t \rangle$ in a neighborhood of $(0, \infty)$ is similar to the case of $(0,0)$. From property (i), if $(0, \infty) \in \Lambda_1(\omega)$, then if $\langle \lambda_t \rangle$ is in $B_\varepsilon(0, \infty)$ infinitely often, $\langle \lambda_t \rangle$ almost surely converges to $(0, \infty)$. From property (ii), if $(0, \infty) \notin \Lambda_1(\omega)$, then $Pr(\tau < \infty \text{ or } \lambda_t \in B_\varepsilon(0,0) \text{ i.o.}) = 1$. (iv) When $(0, \infty) \notin \Lambda_2(\omega)$, then if $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(0, \infty))$, it exits $B_\varepsilon(0, \infty)$ almost surely. We use

these properties to establish the claim outlined in the first paragraph for four possible cases.

Case (i): Suppose $(0, 0) \in \Lambda(\omega)$ and $(0, \infty) \in \Lambda(\omega)$. If $\langle \lambda_t \rangle$ is in $B_\varepsilon(0, 0) \cup B_\varepsilon(0, \infty)$ infinitely often, then since both $(0, 0)$ and $(0, \infty)$ are locally stable, $\langle \lambda_t \rangle$ almost surely converges to $(0, 0)$ or $(0, \infty)$. Combined with Eq. (17), this implies that almost surely either $\tau < \infty$ or $\langle \lambda_t \rangle$ converges to $(0, 0)$ or $(0, \infty)$.

Case (ii): Suppose $(0, 0) \in \Lambda_2(\omega) \setminus \Lambda_1(\omega)$ and $(0, \infty) \in \Lambda(\omega)$. As established in property (ii), $\Pr(\tau < \infty \text{ or } \lambda_t \in B_\varepsilon(0, \infty) \text{ i.o.}) = 1$. If $\langle \lambda_t \rangle$ is in $B_\varepsilon(0, \infty)$ infinitely often, then since it is locally stable, $\langle \lambda_t \rangle$ almost surely converges to $(0, \infty)$. Therefore, almost surely either $\tau < \infty$ or $\langle \lambda_t \rangle$ converges to $(0, \infty)$. If $(0, 0) \in \Lambda(\omega)$ and $(0, \infty) \in \Lambda_2(\omega) \setminus \Lambda_1(\omega)$, then by analogous reasoning, almost surely either $\tau < \infty$ or $\langle \lambda_t \rangle$ converges to $(0, 0)$.

Case (iii): Suppose $(0, 0) \in \Lambda(\omega)$ and $(0, \infty) \notin \Lambda_2(\omega)$. Then by property (iv), if $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(0, \infty))$, it exits $B_\varepsilon(0, \infty)$ almost surely. First suppose $\langle \lambda_t \rangle$ enters $B_\varepsilon(0) \times [\varepsilon, 1/\varepsilon]$ from $B_\varepsilon(0, \infty)$ with positive probability. Then with probability uniformly bounded away from zero across $B_\varepsilon(0) \times [\varepsilon, 1/\varepsilon]$, either $\langle \lambda_t \rangle$ enters $B_\varepsilon(0, 0)$ or $\langle \lambda_{1,t} \rangle$ exits $B_\varepsilon(0)$ (and hence, $\tau < \infty$).⁴⁷ When $\langle \lambda_t \rangle$ enters $B_\varepsilon(0, 0)$, then with probability uniformly bounded away from zero across $B_\varepsilon(0, 0)$, $\langle \lambda_t \rangle$ converges to $(0, 0)$. Together, this establishes that if $\langle \lambda_t \rangle$ is in $B_\varepsilon(0, \infty)$ infinitely often and $\langle \lambda_t \rangle$ enters $B_\varepsilon(0) \times [\varepsilon, 1/\varepsilon]$ from $B_\varepsilon(0, \infty)$ with positive probability, then almost surely either $\tau < \infty$ or $\langle \lambda_t \rangle$ converges to $(0, 0)$. Next suppose $\langle \lambda_t \rangle$ enters $B_\varepsilon(0) \times [\varepsilon, 1/\varepsilon]$ from $B_\varepsilon(0, \infty)$ with probability zero. Then when $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(0, \infty))$, almost surely $\tau < \infty$. Taken together, this establishes that if $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(0, \infty))$ infinitely often, then almost surely either $\tau < \infty$ or $\langle \lambda_t \rangle$ converges to $(0, 0)$. We already established that if $\langle \lambda_t \rangle$ enters $B_\varepsilon(0, 0)$ infinitely often, then $\langle \lambda_t \rangle$ converges to $(0, 0)$ almost surely. It follows from Eq. (17) that almost surely either $\tau < \infty$ or $\langle \lambda_t \rangle$ converges to $(0, 0)$.

Case (iv): Suppose $(0, 0) \in \Lambda_2(\omega) \setminus \Lambda_1(\omega)$ and $(0, \infty) \notin \Lambda(\omega)$. When $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(0, 0))$, with probability uniformly bounded away from zero, $\langle \lambda_{1,t} \rangle$ exits $B_\varepsilon(0)$. Therefore, if $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(0, 0))$ infinitely often, then $\tau < \infty$ almost surely. If $(0, \infty) \in \Lambda_2(\omega) \setminus \Lambda_1(\omega)$, then the same holds for $B_\varepsilon(0, \infty)$ and it follows from Eq. (17) that $\tau < \infty$ almost surely. If $(0, \infty) \notin \Lambda_2(\omega)$, then by property (iv), if $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(0, \infty))$, it exits $B_\varepsilon(0, \infty)$ almost surely. First suppose $\langle \lambda_t \rangle$ enters $B_\varepsilon(0) \times [\varepsilon, 1/\varepsilon]$ from $B_\varepsilon(0, \infty)$ with positive probability. By similar reasoning to Case (iii), when $\langle \lambda_t \rangle$ enters $B_\varepsilon(0) \times [\varepsilon, 1/\varepsilon]$, with probability uniformly bounded away from zero across $B_\varepsilon(0) \times [\varepsilon, 1/\varepsilon]$, $\langle \lambda_t \rangle$ enters $B_\varepsilon(0, 0)$ or $\langle \lambda_{1,t} \rangle$ exits $B_\varepsilon(0)$. When $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(0, 0))$, then with probability uniformly bounded away from zero, $\langle \lambda_{1,t} \rangle$ exits $B_\varepsilon(0)$. Together, this establishes that if $\langle \lambda_t \rangle$ enters $B_\varepsilon(0, \infty)$ infinitely often and $\langle \lambda_t \rangle$ enters $B_\varepsilon(0) \times [\varepsilon, 1/\varepsilon]$ from $B_\varepsilon(0, \infty)$ with positive probability, then almost surely $\tau < \infty$. If $\langle \lambda_t \rangle$ enters $B_\varepsilon(0) \times [\varepsilon, 1/\varepsilon]$ from $B_\varepsilon(0, \infty)$ with probability zero, then when $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(0, \infty))$, almost surely $\tau < \infty$. Taken together, this establishes that if $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(0, \infty))$ infinitely often, then $\tau < \infty$ almost surely. We already established that if $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(0, 0))$ infinitely often, then $\tau < \infty$ almost surely. It follows from Eq. (17) that $\tau < \infty$ almost surely.

Together, these cases establish the claim outlined above when $(0, 0) \in \Lambda_2(\omega)$. The case

⁴⁷This follows from Lemma 2, which guarantees that there exists a finite sequence of a_1 actions that occurs with positive probability such that following this sequence, $\langle \lambda_t \rangle$ enters $B_\varepsilon(0, 0) \cup [\varepsilon, \infty) \times (0, 1/\varepsilon]$ from any starting belief in $B_\varepsilon(0) \times [\varepsilon, 1/\varepsilon]$. An analogous property follows from a finite sequence of a_M actions when the role of $(0, 0)$ and $(0, \infty)$ are switched.

of $(0, \infty) \in \Lambda_2(\omega)$ and the other mixed outcomes are analogous. $\square$

A.5 Learning Characterization

Lemma 7 (Belief Convergence ($k \leq 2$)). *Consider a learning environment that is identified at certainty and satisfies Assumptions 1 to 4. If (i) $k = 1$ and $\Lambda(\omega) \neq \emptyset$ or (ii) $k = 2$, $\Lambda(\omega)$ contains an agreement outcome or maximally accessible disagreement outcome and $\Lambda_M(\omega) = \emptyset$, then for any initial belief $\lambda_1 \in (0, \infty)^k$, there exists a random variable λ_∞ with $\text{supp}(\lambda_\infty) = \Lambda(\omega)$ such that $\lambda_t \rightarrow \lambda_\infty$ almost surely in state ω .*

Proof. Fix state ω and consider the case of $k = 2$. Recall that $\mathcal{B}$ is the set of stable neighborhoods and $\mathcal{B}_U$ is the set of unstable neighborhoods defined in Eq. (12). We define a similar set of neighborhoods for mixed outcomes. Consider $(0, \theta_1)$. Let $\varepsilon \equiv e^{-E}$ (where E is defined in Eq. (11)) and choose $\varepsilon_1 \in (0, \varepsilon)$ such that there exist a $\tau < \infty$ and $\tau' < \infty$ such that for any $\lambda_1 \in B_{\varepsilon_1}(0) \times [\varepsilon, 1/\varepsilon]$, following τ realizations of action a_M , $\lambda_{\tau+1} \in B_\varepsilon(0, \infty)$ and following τ' realizations of action a_1 , $\lambda_{\tau'+1} \in B_\varepsilon(0, 0)$. Define $B_{\varepsilon, \varepsilon_1}^M(0, \theta_1) \equiv B_\varepsilon(0, 0) \cup B_\varepsilon(0, \infty) \cup (B_{\varepsilon_1}(0) \times [\varepsilon, 1/\varepsilon])$ as the mixed neighborhood of $(0, \theta_1)$, with $\text{int}(B_{\varepsilon, \varepsilon_1}^M(0, \theta_1)) \equiv B_{\varepsilon, \varepsilon_1}^M(0, \theta_1) \cap (0, \infty)^2$. By construction, with probability uniformly bounded away from zero across $B_{\varepsilon_1}(0) \times [\varepsilon, 1/\varepsilon]$, $\langle \lambda_t \rangle$ enters $B_\varepsilon(0, 0)$ and $B_\varepsilon(0, \infty)$ in finite time from $B_{\varepsilon_1}(0) \times [\varepsilon, 1/\varepsilon]$. Given analogous definitions for the mixed neighborhoods of (∞, θ_1) , $(0, \theta_2)$, and (∞, θ_2) with corresponding $\varepsilon_2, \varepsilon_3$, and ε_4 , let $\mathcal{B}_M$ denote the union of these mixed neighborhoods, $\text{int}(\mathcal{B}_M) \equiv \mathcal{B}_M \cap (0, \infty)^2$, and $\mathcal{I} \equiv [0, \infty]^2 \setminus \mathcal{B}_M$.

Suppose $\Lambda(\omega)$ contains an agreement outcome or maximally accessible disagreement outcome and $\Lambda_M(\omega)$ is empty. Taken together, previous results have already shown that almost surely, either $\langle \lambda_t \rangle$ does not converge for both types or $\langle \lambda_t \rangle$ converges to a learning outcome in $\Lambda(\omega)$. This follows from Lemma 3 which shows that $\langle \lambda_t \rangle$ does not converge to a non-stationary belief, Theorem 1 which establishes that almost surely $\langle \lambda_t \rangle$ does not converge to an unstable stationary belief, and Lemma 4 which establishes that almost surely $\langle \lambda_t \rangle$ does not converge for one type. Therefore, it remains to show that $\langle \lambda_t \rangle$ almost surely converges.

We first establish the following claim: if $\lambda_1 \in \text{int}(\mathcal{B}_M)$, then with probability uniformly bounded away from zero across $\text{int}(\mathcal{B}_M)$, $\langle \lambda_t \rangle$ enters $\mathcal{I} \cup \mathcal{B}$ in finite time. Again by the linearity of the log likelihood ratio, this implies that the same holds when $\lambda_t \in \text{int}(\mathcal{B}_M)$. Without loss of generality, suppose $(\infty, \infty) \in \Lambda(\omega)$. If $\lambda_1 \in (B_{\varepsilon_2}(\infty) \times [\varepsilon, 1/\varepsilon]) \cup ([\varepsilon, 1/\varepsilon] \times B_{\varepsilon_4}(\infty)) \cup B_\varepsilon(\infty, \infty)$, the claim follows from the definitions of ε_2 and ε_4 and $B_\varepsilon(\infty, \infty) \subset \mathcal{B}$. Therefore, suppose $\lambda_1 \in \text{int}(B_{\varepsilon, \varepsilon_1}^M(0, \theta_1) \cup B_{\varepsilon, \varepsilon_3}^M(0, \theta_2))$.

Case (i): Suppose $\Lambda(\omega) = \{(\infty, \infty)\}$, so $\mathcal{B} = \{B_\varepsilon(\infty, \infty)\}$. Then by $\Lambda_M(\omega) = \emptyset$, either $(0, \infty) \in \Lambda_2(\omega) \setminus \Lambda_1(\omega)$ or $(\infty, 0) \in \Lambda_1(\omega) \setminus \Lambda_2(\omega)$. Without loss of generality suppose $(0, \infty) \in \Lambda_2(\omega) \setminus \Lambda_1(\omega)$. By the logic in Lemma 4, with probability uniformly bounded away from zero across $\text{int}(B_{\varepsilon, \varepsilon_1}^M(0, \theta_1) \setminus B_\varepsilon(0, 0))$, $\langle \lambda_t \rangle$ enters $\mathcal{I} \cup B_\varepsilon(\infty, \infty) \cup (B_{\varepsilon_2}(\infty) \times [\varepsilon, 1/\varepsilon]) \cup ([\varepsilon, 1/\varepsilon] \times B_{\varepsilon_4}(\infty))$ from $\text{int}(B_{\varepsilon, \varepsilon_1}^M(0, \theta_1) \setminus B_\varepsilon(0, 0))$. This implies $\langle \lambda_t \rangle$ enters $\mathcal{I} \cup \mathcal{B}$ from $\text{int}(B_{\varepsilon, \varepsilon_1}^M(0, \theta_1) \setminus B_\varepsilon(0, 0))$ with probability uniformly bounded away from zero across $\text{int}(B_{\varepsilon, \varepsilon_1}^M(0, \theta_1) \setminus B_\varepsilon(0, 0))$. Additionally, if $\lambda_1 \in \text{int}(B_{\varepsilon, \varepsilon_3}^M(0, \theta_2))$, $\langle \lambda_{2,t} \rangle$ almost surely exits $B_\varepsilon(0)$, which implies $\langle \lambda_t \rangle$ almost surely enters $(0, \infty)^2 \setminus B_{\varepsilon, \varepsilon_3}^M(0, \theta_2)$. Since we have already established that $\langle \lambda_t \rangle$ enters $\mathcal{I} \cup \mathcal{B}$ from $\text{int}(\mathcal{B}_M \setminus B_{\varepsilon, \varepsilon_3}^M(0, \theta_2))$ with probability bounded away from zero across $\text{int}(\mathcal{B}_M \setminus B_{\varepsilon, \varepsilon_3}^M(0, \theta_2))$, it follows that this also holds for $\text{int}(\mathcal{B}_M)$.

Case (ii): Suppose $(\infty, \infty) \in \Lambda(\omega)$ and $(0, \infty) \in \Lambda(\omega)$. If $\lambda_1 \in \text{int}(B_{\varepsilon, \varepsilon_1}^M(0, \theta_1) \setminus$

$B_{\varepsilon, \varepsilon_3}^M(0, \theta_2)$), the claim follows from the definition of ε_1 and $B_\varepsilon(0, \infty) \subset \mathcal{B}$. If $\lambda_1 \in \text{int}(B_{\varepsilon, \varepsilon_3}^M(0, \theta_2))$, then by similar reasoning to the previous case, $\langle \lambda_t \rangle$ almost surely enters $(0, \infty)^2 \setminus (B_{\varepsilon, \varepsilon_3}^M(0, \theta_2) \setminus \mathcal{B})$. Since we have already established that $\langle \lambda_t \rangle$ enters $\mathcal{I} \cup \mathcal{B}$ from $\text{int}(\mathcal{B}_M \setminus B_{\varepsilon, \varepsilon_3}^M(0, \theta_2))$ with probability bounded away from zero across $\text{int}(\mathcal{B}_M \setminus B_{\varepsilon, \varepsilon_3}^M(0, \theta_2))$, it follows that this also holds across $\text{int}(\mathcal{B}_M)$. The case of $\{(\infty, \infty), (\infty, 0)\} \in \Lambda(\omega)$ is analogous.

Case (iii): Suppose $\Lambda(\omega) = \{(0, 0), (\infty, \infty)\}$. If $\lambda_1 \in (B_{\varepsilon, \varepsilon_1}^M(0, \theta_1) \cup B_{\varepsilon, \varepsilon_3}^M(0, \theta_2)) \setminus (B_\varepsilon(\infty, 0) \cup B_\varepsilon(0, \infty))$, the claim follows from the definitions of ε_1 and ε_3 combined with $B_\varepsilon(0, 0) \subset \mathcal{B}$. If $\lambda_1 \in \text{int}(B_\varepsilon(0, \infty))$, $\langle \lambda_t \rangle$ almost surely enters $\text{int}(\mathcal{B}_M \setminus (B_\varepsilon(\infty, 0) \cup B_\varepsilon(0, \infty))) \cup \mathcal{I}$. When it enters $\text{int}(\mathcal{B}_M \setminus (B_\varepsilon(\infty, 0) \cup B_\varepsilon(0, \infty)))$, prior reasoning established that it enters $\mathcal{I} \cup \mathcal{B}$ with probability uniformly bounded away from zero. The reasoning is analogous for $\lambda_1 \in \text{int}(B_\varepsilon(\infty, 0))$. It follows that the claim holds across $\text{int}(\mathcal{B}_M)$.

Taken together, these three cases establish the claim when $(\infty, \infty) \in \Lambda(\omega)$. The logic is analogous for the other cases. To complete the proof, let $\tau_1 \equiv \min\{t | \lambda_t \in \mathcal{B}\}$ be the first time that $\langle \lambda_t \rangle$ enters $\mathcal{B}$. By Lemma 2 and Theorem 7, there exists a finite sequence of actions that occurs with positive probability such that for any $\lambda_1 \in \mathcal{I}$, following this sequence $\langle \lambda_t \rangle$ is in $\mathcal{B}$. Therefore, the probability of entering $\mathcal{B}$ is uniformly bounded away from zero across $\mathcal{I}$. Moreover, as shown in the claim above, the probability of entering $\mathcal{I} \cup \mathcal{B}$ is uniformly bounded away from zero across $\text{int}(\mathcal{B}_M)$. Given $\mathcal{I} \cup \text{int}(\mathcal{B}_M) = (0, \infty)^k$, this implies $\Pr(\lambda_t \in \mathcal{B} \text{ i.o.}) = 1$. But if $\langle \lambda_t \rangle$ is in a neighborhood of a locally stable belief infinitely often, then almost surely $\langle \lambda_t \rangle$ converges. The proof for $k = 1$ is analogous, without needing to consider disagreement or mixed outcomes. $\square$

Proof of Theorem 4. Fix state ω . Parts (i) and (ii) follow from the local and global stability of agreement outcomes (Theorems 1 and 2). Part (iii) follows from the local and global stability of disagreement outcomes (Theorems 1 and 3). For part (iv), Lemma 3 rules out convergence to non-stationary beliefs, Theorem 1 rules out convergence to stationary outcomes that are not locally stable, and Lemma 4 rules out convergence to a mixed learning outcome when $\Lambda_M(\omega) = \emptyset$. Therefore, if $\Lambda(\omega) = \emptyset$, there are no locally stable learning outcomes and almost surely the likelihood ratio does not converge for at least one social type. If additionally $\Lambda_M(\omega) = \emptyset$, then almost surely the likelihood ratio does not converge for any social type. The final statement in part (iv) follows from Lemma 7, which establishes when the likelihood ratio almost surely converges. $\square$

A.6 Robustness

Proof of Theorem 5. Fix a regular learning environment (Θ^*, π^*) . Let $\psi^*(\cdot | \omega, \lambda)$ and $\hat{\psi}_i^*(\cdot | \omega, \lambda)$ denote the true and perceived action distributions in this environment, and analogously for $\gamma^*(\omega, \lambda)$, $\Lambda^*(\omega)$, $\Lambda_M^*(\omega)$ and $\mathcal{G}^*(\omega)$ (when $k > 2$, see Eq. (20) in Online Appendix D for the generalized definition of $\Lambda_M(\omega)$ and Definition 15 for the definition of $\mathcal{G}(\omega)$). Throughout the proof, restrict attention to learning environments (Θ, π) that have the same number of social types as (Θ^*, π^*) and satisfy Assumptions 1 to 4. We first consider local stability of nearby learning environments. The mapping $(\psi(a | \omega, \lambda), \hat{\psi}_i(a | \omega, \lambda)) \mapsto \gamma_i(\omega, \lambda)$ is continuous. By definition of identified at certainty, the sign of $\gamma_i^*(\omega, \lambda)$ is strictly positive or negative at $\lambda \in \{0, \infty\}^k$. Therefore, there exists a $\delta_1(\omega) > 0$ such that in any learning environment (Θ, π) that is sufficiently close to (Θ^*, π^*) in that $\|\psi^*(\cdot | \omega, \lambda) - \psi(\cdot | \omega, \lambda)\| < \delta_1(\omega)$, $\|\hat{\psi}_i^*(\cdot | L, \lambda) - \hat{\psi}_i(\cdot | L, \lambda)\| < \delta_1(\omega)$ and $\|\hat{\psi}_i^*(\cdot | R, \lambda) - \hat{\psi}_i(\cdot | R, \lambda)\| < \delta_1(\omega)$ for all $\lambda \in \{0, \infty\}^k$.

and $i = 1, \dots, k$, continuity implies that $\gamma_i(\omega, \boldsymbol{\lambda})$ has the same sign as $\gamma_i^*(\omega, \boldsymbol{\lambda})$ for all $\boldsymbol{\lambda} \in \{0, \infty\}^k$ and $i = 1, \dots, k$. This implies $\Lambda_i(\omega) = \Lambda_i^*(\omega)$ for $i = 1, \dots, k$, and therefore, $\Lambda(\omega) = \Lambda^*(\omega)$, so (Θ, π) has the same set of locally stable outcomes as (Θ^*, π^*) in state ω . Given $\Lambda_M(\omega)$ and $\mathcal{G}(\omega)$ are constructed from $\Lambda_i(\omega)$, in any such learning environment, $\Lambda_M(\omega) = \Lambda_M^*(\omega)$ and $\mathcal{G}(\omega) = \mathcal{G}^*(\omega)$ follows from $\Lambda_i(\omega) = \Lambda_i^*(\omega)$ for $i = 1, \dots, k$. By definition of regular, given $\Lambda_M^*(\omega) = \emptyset$, this implies $\Lambda_M(\omega) = \emptyset$.

Finally, consider the global stability of locally stable disagreement outcomes in a learning environment with $\Lambda(\omega) = \Lambda^*(\omega)$. Suppose $k = 2$ and $(0, \infty) \in \Lambda(\omega)$. Then $(0, \infty) \in \Lambda^*(\omega)$ and by definition of regular, $(0, \infty)$ is maximally accessible in (Θ^*, π^*) , i.e. either $\theta_2^* \succ_{(0,0)} \theta_1^*$ or $\theta_2^* \succ_{(\infty,\infty)} \theta_1^*$. Suppose $\theta_2^* \succ_{(0,0)} \theta_1^*$. By the proof of [Theorem 3](#), this implies $(0, 0)$ is separable at zero for θ_2^* , so there exists a vector $G = (G_1, G_2)'$ with $G_1 < 0$ and $G_2 > 0$ and a vector $c \in \mathbb{R}_+^{|A|}$ such that $\Psi^*(0, 0) \cdot c = G$, where Ψ^* is the matrix defined in [Eq. \(13\)](#) for (Θ^*, π^*) . Since $\Psi(0, 0) \cdot c$ is continuous in $\hat{\psi}_i$, there exists a $\delta_{(0,\infty)} > 0$ such that in any learning environment (Θ, π) that is sufficiently close to (Θ^*, π^*) in that $\|\hat{\psi}_i(\cdot|L, \boldsymbol{\lambda}) - \psi^*(\cdot|L, \boldsymbol{\lambda})\| < \delta_{(0,\infty)}$ and $\|\hat{\psi}_i(\cdot|R, \boldsymbol{\lambda}) - \psi^*(\cdot|R, \boldsymbol{\lambda})\| < \delta_{(0,\infty)}$ for all $\boldsymbol{\lambda} \in \{0, \infty\}^2$ and $i = 1, 2$, the expressions $(\Psi(0, 0) \cdot c)_1 < 0$ and $(\Psi(0, 0) \cdot c)_2 > 0$, where Ψ is the matrix defined in [Eq. \(13\)](#) for (Θ, π) . Therefore, in any such learning environment, $(0, 0)$ is separable at zero for θ_2 .⁴⁸ By [Theorem 7](#), this implies $(0, \infty)$ is globally stable in state ω for (Θ, π) . The case of $\theta_2^* \succ_{(\infty,\infty)} \theta_1^*$ is analogous, as is the proof for disagreement outcome $(\infty, 0)$ using some $\delta_{(\infty,0)} > 0$. When $k > 2$, a similar argument shows that a disagreement outcome $\boldsymbol{\lambda}^*$ that is locally stable and maximally accessible in (Θ^*, π^*) is locally stable and separable in sufficiently close learning environments, where close is defined relative to some $\delta_{\boldsymbol{\lambda}^*} > 0$, and therefore, globally stable (see [Online Appendix D](#) for the expanded definition of maximally accessible when $k > 2$). Let $\delta_2(\omega) \equiv \min_{\boldsymbol{\lambda}^* \in \{\Lambda^*(\omega) \setminus \{0^k, \infty^k\}\}} \delta_{\boldsymbol{\lambda}^*}$ denote the minimum $\delta_{\boldsymbol{\lambda}^*}$ across all locally stable disagreement outcomes in state ω for (Θ^*, π^*) . Then any learning environment (Θ, π) that has $\Lambda(\omega) = \Lambda^*(\omega)$, $\|\hat{\psi}_i^*(\cdot|L, \boldsymbol{\lambda}) - \hat{\psi}_i(\cdot|L, \boldsymbol{\lambda})\| < \delta_2(\omega)$ and $\|\hat{\psi}_i^*(\cdot|R, \boldsymbol{\lambda}) - \hat{\psi}_i(\cdot|R, \boldsymbol{\lambda})\| < \delta_2(\omega)$ for all $\boldsymbol{\lambda} \in \{0, \infty\}^k$ and $i = 1, \dots, k$ has the same set of globally stable outcomes as (Θ^*, π^*) in state ω .

Let $\delta \equiv \min\{\delta_1(\omega), \delta_2(\omega)\}$. Then any learning environment (Θ, π) with $\|\psi^*(\cdot|\omega, \boldsymbol{\lambda}) - \psi(\cdot|\omega, \boldsymbol{\lambda})\| < \delta$, $\|\hat{\psi}_i^*(\cdot|L, \boldsymbol{\lambda}) - \hat{\psi}_i(\cdot|L, \boldsymbol{\lambda})\| < \delta$ and $\|\hat{\psi}_i^*(\cdot|R, \boldsymbol{\lambda}) - \hat{\psi}_i(\cdot|R, \boldsymbol{\lambda})\| < \delta$ for all $\boldsymbol{\lambda} \in \{0, \infty\}^k$ and $i = 1, \dots, k$ has the same set of long-run learning outcomes as (Θ^*, π^*) in state ω . $\square$

Proof of Theorem 6. Fix a correctly specified environment (Θ^*, π^*) that satisfies [Assumptions 2](#) and [3](#). Let $\psi^*(a|\omega, \boldsymbol{\lambda})$ denote the distribution over actions in this environment. In a correctly specified environment, $\hat{\psi}_i^*(a|\omega, \boldsymbol{\lambda}) = \psi^*(a|\omega, \boldsymbol{\lambda})$ for $i = 1, \dots, k$. By [Corollary 2](#), learning is complete in (Θ^*, π^*) . Further, correctly specified environments are regular. Throughout the proof, restrict attention to learning environments (Θ, π) that are structurally equivalent to (Θ^*, π^*) and satisfy [Assumptions 1](#) to [4](#).

We first construct $\hat{\psi}_i(a|\omega, \boldsymbol{\lambda})$ and $\psi(a|\omega, \boldsymbol{\lambda})$ for such a (Θ, π) . Let $\lambda_0 \equiv p_0/(1 - p_0)$. From the decision rules constructed in [Lemma 1](#), an autarkic or noise type $\theta_j \in \Theta_A \cup \Theta_N$ chooses action a_m if $\bar{s}_{j,m-1}(\lambda_0) \neq \bar{s}_{j,m}(\lambda_0)$ and it observes a signal realization $s \in$

⁴⁸Note that $\boldsymbol{\lambda}$ maximally accessible in (Θ^*, π^*) does not imply that it is maximally accessible in (Θ, π) , as the strict maximal R-order can have one weak inequality.

$(\bar{s}_{j,m-1}(\lambda_0), \bar{s}_{j,m}(\lambda_0)]$, with a closed interval if $\bar{s}_{i,m-1}(\lambda_0) = 0$.⁴⁹ Note that $\theta_j^* \in \Theta_A^* \cup \Theta_N^*$ has the same signal cutoffs as θ_j , i.e. $\bar{s}_{j,m}^* = \bar{s}_{j,m}$ for $m = 1, \dots, M$, since $\Theta_A \cup \Theta_N = \Theta_A^* \cup \Theta_N^*$ by definition of structurally equivalent. At any belief $\lambda \in \{0, \infty\}^k$, social type $\theta_j \in \Theta_S$ has a unique optimal action that it plays for all signal realizations, independent of its model of inference $(\hat{F}_j^L, \hat{F}_j^R, \hat{\pi}_j)$. Let $\alpha_j(\lambda)$ denote this optimal action. Note that $\theta_j^* \in \Theta_S^*$ has the same optimal action, $\alpha_j^*(\lambda) = \alpha_j(\lambda)$, since it has the same preferences as θ_j by definition of structurally equivalent. Social type $\theta_i \in \Theta_S$ believes autarkic or noise type $\theta_j \in \Theta_A \cup \Theta_N$ chooses action a_m with probability $\hat{F}_i^\omega(\bar{s}_{j,m}(\lambda_0)) - \hat{F}_i^\omega(\bar{s}_{j,m-1}(\lambda_0))$. It believes $\theta_j \in \Theta_S$ chooses $\alpha_j(\lambda)$ with probability one independent of its model of inference $(\hat{F}_i^L, \hat{F}_i^R, \hat{\pi}_i)$. Therefore, for any $\lambda \in \{0, \infty\}^k$, type θ_i believes a_m is chosen with probability

$$\hat{\psi}_i(a_m | \omega, \lambda) = \sum_{\theta_j \in \Theta_A \cup \Theta_N} \hat{\pi}_i(\theta_j) (\hat{F}_i^\omega(\bar{s}_{j,m}(\lambda_0)) - \hat{F}_i^\omega(\bar{s}_{j,m-1}(\lambda_0))) + \sum_{\theta_j \in \Theta_S} \hat{\pi}_i(\theta_j) \mathbb{1}_{\alpha_j(\lambda) = a_m}. \quad (18)$$

This is continuous in $\hat{\pi}_i$ and $\hat{F}_i^\omega$ under the total variation norm, and it is independent of $(\hat{F}_j^L, \hat{F}_j^R, \hat{\pi}_j)$ for $\theta_j \in \Theta_S \setminus \{\theta_i\}$. Similarly, for any $\lambda \in \{0, \infty\}^k$, the true probability of a_m is

$$\begin{aligned} \psi(a_m | \omega, \lambda) &= \sum_{\theta_j \in \Theta_A \cup \Theta_N} \pi(\theta_j) (F^\omega(\bar{s}_{j,m}(\lambda_0)) - F^\omega(\bar{s}_{j,m-1}(\lambda_0))) + \sum_{\theta_j \in \Theta_S} \pi(\theta_j) \mathbb{1}_{\alpha_j(\lambda) = a_m} \\ &= \sum_{\theta_j^* \in \Theta_A^* \cup \Theta_N^*} \pi^*(\theta_j^*) (F^\omega(\bar{s}_{j,m}^*(\lambda_0)) - F^\omega(\bar{s}_{j,m-1}^*(\lambda_0))) + \sum_{\theta_j^* \in \Theta_S^*} \pi^*(\theta_j^*) \mathbb{1}_{\alpha_j^*(\lambda) = a_m} \\ &= \psi^*(a_m | \omega, \lambda) \end{aligned} \quad (19)$$

where the second equality follows from $\Theta_A \cup \Theta_N = \Theta_A^* \cup \Theta_N^*$ and $\pi(\theta_j) = \pi^*(\theta_j^*)$ by definition of structurally equivalent, $\alpha_j(\lambda) = \alpha_j^*(\lambda)$ for $j = 1, \dots, k$ as shown above, and $\bar{s}_{j,m} = \bar{s}_{j,m}^*$ for $m = 1, \dots, M$ and $j = k+1, \dots, n$ as shown above.

From Eq. (18) continuous in $(\hat{\pi}_i, \hat{F}_i^\omega)$ and Eq. (19), it follows that for any $\varepsilon > 0$, there exists a $\delta(\varepsilon) > 0$ such that any environment (Θ, π) with $\|\hat{\pi}_i - \pi\| < \delta(\varepsilon)$, $\|\hat{F}_i^L - F^L\| < \delta(\varepsilon)$ and $\|\hat{F}_i^R - F^R\| < \delta(\varepsilon)$ for all $\theta_i \in \Theta_S$ satisfies $\|\psi^*(\cdot | L, \lambda) - \hat{\psi}_i(\cdot | L, \lambda)\| < \varepsilon$ and $\|\psi^*(\cdot | R, \lambda) - \hat{\psi}_i(\cdot | R, \lambda)\| < \varepsilon$ for all $\lambda \in \{0, \infty\}^k$ and $i = 1, \dots, k$. Further, from Eq. (19), $\|\psi^*(\cdot | \omega, \lambda) - \psi(\cdot | \omega, \lambda)\| = 0$ for all $\lambda \in \{0, \infty\}^k$ and $\omega \in \{L, R\}$ in any such model. Given this, by Theorem 5, choosing ε small enough establishes that any such learning environment has the same set of long-run learning outcomes as (Θ^*, π^*) in both states. Since learning is complete in both states in (Θ^*, π^*) , this establishes that learning is complete in both states in any learning environment with $\|\hat{\pi}_i - \pi\|$, $\|\hat{F}_i^L - F^L\|$ and $\|\hat{F}_i^R - F^R\|$ sufficiently small for all social types. $\square$

⁴⁹This proof maintains our convention for breaking indifference outlined in Lemma 1, i.e. the agent chooses the optimal action with the lowest index. Our robustness result holds across all equilibria, i.e. it also holds for each of the finite other possible ways to resolve indifference. We omit this analysis, as it requires cumbersome notation without added conceptual insight.

Learning with Heterogeneous Misspecified Models: Characterization and Robustness

Online Appendix

J. Aislinn Bohren

Daniel N. Hauser

July 29, 2021

First version: May 15, 2017

B Derivation of Examples from Section 3

B.1 Example 1: Partisan Bias

Signals and preferences are aligned ([Assumptions 1 and 2](#)) since both types have the same subjective signal distributions and preferences. The autarkic type θ_2 plays both actions with positive probability and the social type θ_1 places positive probability on θ_2 , which establishes that [Assumption 3](#) holds. [Assumption 4](#) is redundant in a binary action decision problem, since [Assumption 3](#) guarantees that the social type believes that the autarkic type plays both actions with positive probability. For technical convenience, we assume that the signal distributions are continuous and symmetric, $F^R(s) = 1 - F^L(1 - s)$.

From the action probabilities derived in [Section 3.1](#), at likelihood ratio λ_1 , type θ_1 believes action L occurs with probability $\hat{\psi}_1(L|\omega, \lambda_1) = \pi(\theta_1)F^\omega(1/(1 + \lambda_1)) + \pi(\theta_2)F^\omega(.5)$, whereas the true probability of action L is $\psi(L|\omega, \lambda_1) = \pi(\theta_1)F^\omega((1/(1 + \lambda_1))^{1/\nu}) + \pi(\theta_2)F^\omega(.5^{1/\nu})$. The construction of $\gamma_1(L, 0)$ in [Section 3.3](#) follows from evaluating these expressions at $\lambda_1 = 0$. Similarly, the construction of $\gamma_1(L, \infty)$ follows from evaluating these expressions at $\lambda_1 = \infty$,

$$\pi(\theta_2)F^L(.5^{1/\nu}) \log \frac{F^R(.5)}{F^L(.5)} + (\pi(\theta_1) + \pi(\theta_2)(1 - F^L(.5^{1/\nu}))) \log \frac{\pi(\theta_1) + \pi(\theta_2)(1 - F^R(.5))}{\pi(\theta_1) + \pi(\theta_2)(1 - F^L(.5))}.$$

We next characterize how $\Lambda(\omega)$ depends on ν . We write $\gamma_1(\omega, \boldsymbol{\lambda}; \nu)$ and $\Lambda(\omega; \nu)$ to make this dependence on ν explicit. To simplify notation, define $\alpha_\nu \equiv F^L(.5^{1/\nu})$ as the probability that type θ_2 chooses an L action in state L and $\pi_A \equiv \pi(\theta_2)$ as the probability of the autarkic type. By symmetry, $F^R(.5) = 1 - F^L(.5) = 1 - \alpha_1$ and by definition of a probability measure, $\pi(\theta_1) = 1 - \pi_A$. Also note that F^L strictly increasing implies that α_ν is strictly increasing in ν , symmetry implies that $\alpha_1 > 1/2$, and F^L continuous implies α_ν is continuous in ν .

First consider $\omega = L$. To determine whether incorrect learning arises, i.e. whether $\infty \in \Lambda(L; \nu)$, we need to determine the sign of

$$\gamma_1(L, \infty; \nu) = \pi_A \alpha_\nu \log \frac{1 - \alpha_1}{\alpha_1} + (1 - \pi_A \alpha_\nu) \log \frac{1 - \pi_A(1 - \alpha_1)}{1 - \pi_A \alpha_1}.$$

Since $\alpha_1 > 1/2$, the update from an L action is negative, $\log \frac{1-\alpha_1}{\alpha_1} < 0$ and the update from an R action is positive, $\log \frac{1-\pi_A(1-\alpha_1)}{1-\pi_A\alpha_1} > 0$. Note both terms are independent of ν . Since α_ν is strictly increasing in ν , the probability of an L action, $\pi_A\alpha_\nu$, is strictly increasing in ν and the probability of an R action, $1 - \pi_A\alpha_\nu$, is strictly decreasing in ν . Therefore, $\gamma_1(L, \infty; \nu)$ is strictly decreasing in ν . At $\nu = 1$, $\gamma_1(L, \infty; 1) < 0$ by the concavity of the log operator. At $\nu = 0$, $\alpha_0 = 0$ and therefore $\gamma_1(L, \infty; 0) = \log \frac{1-\pi_A(1-\alpha_1)}{1-\pi_A\alpha_1} > 0$. Given $\gamma_1(L, \infty; \nu)$ is continuous in ν , there exists a cutoff $\nu_1 \in (0, 1)$ such that for $\nu < \nu_1$, $\gamma_1(L, \infty; \nu) > 0$ and $\infty \in \Lambda(L; \nu)$ and for $\nu > \nu_1$, $\gamma_1(L, \infty; \nu) < 0$ and $\infty \notin \Lambda(L; \nu)$.

To determine whether correct learning arises, i.e. whether $0 \in \Lambda(L; \nu)$, we need to determine the sign of

$$\gamma_1(L, 0; \nu) = (1 - \pi_A(1 - \alpha_\nu)) \log \frac{1 - \pi_A\alpha_1}{1 - \pi_A(1 - \alpha_1)} + \pi_A(1 - \alpha_\nu) \log \frac{\alpha_1}{1 - \alpha_1}.$$

As in the previous case, the update from an L action is negative and the probability of an L action is strictly increasing in ν , while the update from an R action is positive and the probability of an R action is strictly decreasing in ν . Therefore, $\gamma_1(L, 0; \nu)$ is strictly decreasing in ν . At $\nu = 1$, $\gamma_1(L, 0; 1) < 0$ by the concavity of the log operator. At $\nu = 0$, $\alpha_0 = 0$ and therefore,

$$\begin{aligned} \gamma_1(L, 0; 0) &= (1 - \pi_A) \log \frac{1 - \pi_A\alpha_1}{1 - \pi_A(1 - \alpha_1)} + \pi_A \log \frac{\alpha_1}{1 - \alpha_1} \\ &\geq (1 - \pi_A\alpha_1) \log \frac{1 - \pi_A\alpha_1}{1 - \pi_A(1 - \alpha_1)} + \pi_A\alpha_1 \log \frac{\alpha_1}{1 - \alpha_1} \\ &= \gamma_1(R, 0; 1) > 0. \end{aligned}$$

Given $\gamma_1(L, 0; \nu)$ is continuous in ν , there exists a cutoff $\nu_2 \in (0, 1)$ such that for $\nu < \nu_2$, $\gamma_1(L, 0; \nu) > 0$ and $0 \notin \Lambda(L; \nu)$ and for $\nu > \nu_2$, $\gamma_1(L, 0; \nu) < 0$ and $0 \in \Lambda(L; \nu)$.

Finally we show that $\nu_1 < \nu_2$. Note

$$\gamma_1(L, \infty; \nu) - \gamma_1(L, \infty; 1) = \pi_A(\alpha_\nu - \alpha_1) \left(\log \frac{1 - \alpha_1}{\alpha_1} - \log \frac{1 - \pi_A + \pi_A\alpha_1}{1 - \pi_A\alpha_1} \right)$$

and by the symmetry of the signal distributions, $\gamma_1(L, 0; \nu) - \gamma_1(L, 0; 1) = \gamma_1(L, \infty; \nu) - \gamma_1(L, \infty; 1)$. Moreover $\gamma_1(L, 0; 1) - \gamma_1(L, \infty; 1)$ is zero at $\pi_A = 0$ and $\pi_A = 1$, and concave in π_A since the second derivative is

$$\frac{(1 - 2\alpha_1)\pi_A}{(\pi_A(1 - \alpha_1) + (1 - \pi_A))^2(\pi_A\alpha_1 + 1 - \pi_A)^2} \leq 0.$$

Therefore, $0 \notin \Lambda(\omega; \nu)$ before $\infty \in \Lambda(\omega; \nu)$. This implies that $\Lambda(L; \nu) = \{\infty\}$ for $\nu \in (0, \nu_1)$, $\Lambda(L; \nu) = \emptyset$ for $\nu \in (\nu_1, \nu_2)$, and $\Lambda(L; \nu) = \{0\}$ for $\nu \in (\nu_2, 1]$.

Next consider $\omega = R$. Then $\gamma(R, \infty; 1) > 0$ and $\gamma(R, 0; 1) > 0$, since only correct learning can occur at $\nu = 1$. The only change in the above expressions is that now the true probabilities of each action are taken with respect to state R rather than state L . Therefore, the comparative statics are similar to the comparative statics in state L : $\gamma_1(R, 0; \nu)$ and

$\gamma_1(R, \infty; \nu)$ are decreasing in ν . Therefore, $\gamma_1(R, 0; \nu) > 0$ implies $0 \notin \Lambda(R; \nu)$ for all $\nu \in (0, 1]$. Similarly, $\gamma_1(R, \infty; \nu) > 0$ implies $\infty \in \Lambda(R; \nu)$ for all $\nu \in (0, 1]$. Therefore, $\Lambda(R; \nu) = \{\infty\}$ for all $\nu \in (0, 1]$.

When there is a single social type, mixed learning and disagreement are trivially not possible. By [Theorem 4](#), the characterization of the locally stable set fully determines asymptotic learning outcomes. This leads to the following proposition, the proof of which follows from the construction of $\Lambda(\omega; \nu)$ above.

Proposition 5 (Partisan Bias). *When $\omega = L$, there exist unique cutoffs $0 < \nu_1 < \nu_2 < 1$ such that (i) if $\nu \in (\nu_2, 1]$, then almost surely learning is correct; (ii) if $\nu \in (\nu_1, \nu_2)$, then almost surely learning is cyclical; (iii) if $\nu \in (0, \nu_1)$, then almost surely learning is incorrect. When $\omega = R$, almost surely learning is correct.*

B.2 Example 2: Partisan Bias and Unawareness

We construct this variation by adding two types to the setting considered in [Example 1](#). Types θ_1 and θ_2 are partisan types with the same signal misspecification and preferences as in [Example 1](#). Types θ_3 and θ_4 are non-partisan types that correctly interpret signals, $\hat{F}_3^\omega(s) = \hat{F}_4^\omega(s) = F^\omega(s)$; θ_3 is a social type while θ_4 is an autarkic type.¹ Both types have the same preferences as θ_1 and θ_2 , i.e. $u_i(a, \omega) = \mathbb{1}_{a=\omega}$. Assume that an equal and positive share of partisan and nonpartisan types are autarkic, $\pi(\theta_2)/(\pi(\theta_1) + \pi(\theta_2)) = \pi(\theta_4)/(\pi(\theta_3) + \pi(\theta_4)) \in (0, 1)$. Both social types have correct beliefs about the share of autarkic types, but partisan θ_1 believes all agents are partisan, $\hat{\pi}_1(\theta_1) = \pi(\theta_1) + \pi(\theta_3)$ and $\hat{\pi}_1(\theta_2) = \pi(\theta_2) + \pi(\theta_4)$, and analogously, non-partisan θ_3 believes that all agents are non-partisan. Let $q \equiv \pi(\theta_3) + \pi(\theta_4)$ denote the share of non-partisan types and $\pi_A \equiv \pi(\theta_2) + \pi(\theta_4)$ denote the share of autarkic types. To close the model, assume that the signal distributions are continuous and symmetric, $F^R(s) = 1 - F^L(1 - s)$ with support $\mathcal{S} = [0, 1]$, and $p_0 = 1/2$. Signals are aligned since partisan types order signal realizations in the same way as nonpartisan types, i.e. s^ν is increasing in s ([Assumption 1](#)).

The true action probabilities for partisan types θ_1 and θ_2 are identical to those derived in [Section 3.1](#) for [Example 1](#), as are θ_1 's subjective action probabilities for each type. A non-partisan type $\theta_i \in \{\theta_3, \theta_4\}$ who has likelihood ratio λ and observes signal realization s updates to belief $\frac{p_i(\lambda, s)}{1 - p_i(\lambda, s)} = \lambda \left(\frac{s}{1-s} \right)$. It chooses action L if this belief is less than one, which is equivalent to $s < 1/(1 + \lambda) = \bar{s}_{i,1}(\lambda)$. At likelihood ratio λ_3 , type θ_3 chooses L with probability $F^\omega(1/(1 + \lambda_3))$. Type θ_4 is autarkic. Therefore, its likelihood ratio is constant at $\lambda_4 = 1$ and it chooses action L with probability $F^\omega(.5)$. Type θ_3 has correct beliefs about the probability that θ_3 and θ_4 choose action L .

We use these subjective and true action probabilities for each type to construct $\hat{\psi}_1$, $\hat{\psi}_3$ and ψ . Partisan type θ_1 is now also misspecified about the type distribution, since it does not account for the nonpartisan types. It believes action L occurs with probability $\hat{\psi}_1(L|\omega, \boldsymbol{\lambda}) = (1 - \pi_A)F^\omega(1/(1 + \lambda_1)) + \pi_A F^\omega(.5)$. This misspecification about the type distribution leads the partisan type to underestimate the range of signal realizations for which other agents choose action L , while its signal misspecification causes it to overestimate the probability of these signal realizations. The latter effect dominates, and θ_1 overestimates the frequency of action

¹In a slight abuse of our previous notation, we maintain θ_2 as the partisan autarkic type for consistency with [Example 1](#), which violates our convention that the first k types are the social types.

L . Nonpartisan type θ_3 has a correctly specified model of the signal distribution and believes that other agents do as well, since it does not account for the partisan types. It believes action L occurs with probability $\hat{\psi}_3(L|\omega, \lambda) = (1 - \pi_A)F^\omega(1/(1 + \lambda_3)) + \pi_A F^\omega(.5)$. This type misspecification leads the nonpartisan type to believe that other agents are choosing L for a larger range of signal realizations than is actually the case, which leads it to overestimate the frequency of L actions. The true probability of action L is

$$\begin{aligned} \psi(L|\omega, \lambda) &= (1 - q)((1 - \pi_A)F^\omega((1/(1 + \lambda_1))^{1/\nu}) + \pi_A F^\omega(.5^{1/\nu})) \\ &\quad + q((1 - \pi_A)F^\omega(1/(1 + \lambda_3)) + \pi_A F^\omega(.5)). \end{aligned}$$

Although the partisan and nonpartisan social types have different models of the world, their models collapse to the same subjective probability of each action when they have the same current belief: for any λ with $\lambda_1 = \lambda_3$, $\hat{\psi}_1(L|\omega, \lambda) = \hat{\psi}_3(L|\omega, \lambda)$. Therefore, these types update their likelihood ratios in the same way following each action. For different reasons, their beliefs both move too much towards state R following R actions and too little towards state L following L actions. This implies that when there is a common prior, after any history h_t , beliefs are equal, $\lambda_{1,t} = \lambda_{3,t}$.²

Given that the two likelihood ratios move in unison, we can consider the partisan and nonpartisan social types as a single type to characterize asymptotic learning outcomes. Disagreement and mixed learning do not arise, since it is not possible to separate beliefs. Global stability immediately follows from local stability for the two agreement outcomes. Therefore, determining the set of parameters (ν, q) for which each agreement outcome is locally stable fully characterizes asymptotic learning outcomes. This leads to the following proposition.

Proposition 6 (Partisan Bias). *When $\omega = L$, there exist unique cutoffs $q_1 \in (0, 1)$ and $q_2 \in (q_1, 1)$ such that:*

- (i) *For $q < q_1$, there exist unique cutoffs $0 < \nu_1(q) < \nu_2(q) < 1$ such that if $\nu > \nu_2(q)$, then almost surely learning is correct, if $\nu \in (\nu_1(q), \nu_2(q))$, then almost surely learning is cyclical and if $\nu < \nu_1(q)$, then almost surely learning is incorrect.*
- (ii) *For $q \in (q_1, q_2)$, there exists a unique cutoff $0 < \nu_2(q) < 1$ such that if $\nu > \nu_2(q)$, then almost surely learning is correct and if $\nu < \nu_2(q)$, then almost surely learning is cyclical.*
- (iii) *For $q > q_2$, almost surely learning is correct.*

When $\omega = R$, almost surely learning is correct.

Proof. The construction of the locally stable set is similar to [Example 1](#). To simplify notation, define $\alpha_\nu \equiv F^L(.5^{1/\nu})$ as the probability that type θ_2 chooses action L in state L . Given this notation, type θ_4 chooses action L in state L with probability α_1 . As in [Example 1](#), $F^R(.5) = 1 - F^L(.5) = 1 - \alpha_1$, α_ν is strictly increasing in ν and $\alpha_1 > 1/2$. We characterize how $\Lambda(\omega)$ depends on ν and q . We write $\gamma_1(\omega, \lambda; \nu, q)$, $\gamma_3(\omega, \lambda; \nu, q)$, and $\Lambda(\omega; \nu, q)$ to make this dependence explicit. Since beliefs move in unison, $\gamma_3(\omega, \lambda; \nu, q) = \gamma_1(\omega, \lambda; \nu, q)$, and therefore, we can focus on characterizing $\gamma_1(\omega, \lambda; \nu, q)$ at $(0, 0)$ and (∞, ∞) .

²Partisan and nonpartisan types with the same likelihood ratio may choose different actions following a given signal realization s , as they have different signal cutoffs.

To determine whether $(\infty, \infty) \in \Lambda(L; \nu, q)$, we need to determine the sign of

$$\begin{aligned} \gamma_1(L, (\infty, \infty); \nu, q) &= \psi(L|L, (\infty, \infty); \nu, q) \log \frac{1 - \alpha_1}{\alpha_1} \\ &\quad + \psi(R|L, (\infty, \infty); \nu, q) \log \frac{1 - \pi_A(1 - \alpha_1)}{1 - \pi_A\alpha_1}, \end{aligned}$$

where $\psi(L|L, (\infty, \infty); \nu, q) \equiv \pi_A((1 - q)\alpha_\nu + q\alpha_1)$ and $\psi(R|L, (\infty, \infty); \nu, q) \equiv \pi_A((1 - q)(1 - \alpha_\nu) + q(1 - \alpha_1)) + 1 - \pi_A$. Since $\alpha_1 > 1/2$, the update from an L action is negative, $\log \frac{1 - \alpha_1}{\alpha_1} < 0$ and the update from an R action is positive, $\log \frac{1 - \pi_A(1 - \alpha_1)}{1 - \pi_A\alpha_1} > 0$. Note both terms are independent of ν and q . Since α_ν is strictly increasing in ν , the probability of an L action, $\psi(L|L, (\infty, \infty); \nu, q)$, is strictly increasing in ν and q , and the probability of an R action, $\psi(R|L, (\infty, \infty); \nu, q)$, is strictly decreasing in ν and q . Therefore, $\gamma_1(L, (\infty, \infty); \nu, q)$ is strictly decreasing in ν and q . At $\nu = 1$, both partisan and nonpartisan types are identical, so $\psi(L|L, (\infty, \infty); 1, q) = \pi_A\alpha_1$ and $\psi(R|L, (\infty, \infty); 1, q) = \pi_A(1 - \alpha_1) + 1 - \pi_A$. Therefore, for any $q \in [0, 1]$, $\gamma_1(L, (\infty, \infty); 1, q) < 0$ by the concavity of the log operator. Similarly, at $q = 1$, for any $\nu \in [0, 1]$, $\gamma_1(L, (\infty, \infty); \nu, 1) < 0$ by the concavity of the log operator. At $\nu = 0$, θ_2 chooses action R for all signal realizations, i.e. $\alpha_0 = 0$. Therefore, at $q = 0$, $\psi(L|L, (\infty, \infty); 0, 0) = 0$ and $\gamma_1(L, (\infty, \infty); 0, 0) = \log \frac{1 - \pi_A(1 - \alpha_1)}{1 - \pi_A\alpha_1} > 0$. This establishes that there exists a cutoff $q_1 \in (0, 1)$ such that for $q < q_1$, there exists a cutoff $\nu_1(q) \in (0, 1)$ such that for $\nu < \nu_1(q)$, $\gamma_1(L, (\infty, \infty); \nu, q) > 0$ and $(\infty, \infty) \in \Lambda(L; \nu, q)$ and for $\nu > \nu_1(q)$, $\gamma_1(L, (\infty, \infty); \nu, q) < 0$ and $(\infty, \infty) \notin \Lambda(L; \nu, q)$. For $q > q_1$, $\gamma_1(L, (\infty, \infty); \nu, q) < 0$ and $(\infty, \infty) \notin \Lambda(L; \nu, q)$.

To determine whether $(0, 0) \in \Lambda(L; \nu, q)$, we need to determine the sign of

$$\gamma_1(L, (0, 0); \nu, q) = \psi(L|L, (0, 0); \nu, q) \log \frac{1 - \pi_A\alpha_1}{\pi_A\alpha_1 + 1 - \pi_A} + \psi(R|L, (0, 0); \nu, q) \log \frac{\alpha_1}{1 - \alpha_1},$$

where $\psi(L|L, (0, 0); \nu, q) \equiv \pi_A((1 - q)\alpha_\nu + q\alpha_1) + 1 - \pi_A$ and $\psi(R|L, (0, 0); \nu, q) \equiv \pi_A((1 - q)(1 - \alpha_\nu) + q(1 - \alpha_1))$. As in the previous case, the update from an L action is negative and the probability of an L action is strictly increasing in ν and q , while the update from an R action is positive and the probability of an R action is strictly decreasing in ν and q . Therefore, $\gamma_1(L, (0, 0); \nu, q)$ is strictly decreasing in ν and q . By similar reasoning to the case of (∞, ∞) , at $\nu = 1$, $\gamma_1(L, (0, 0); 1, q) < 0$ for all $q \in [0, 1]$ and at $q = 1$, $\gamma_1(L, (0, 0); \nu, 1) < 0$ for all $\nu \in [0, 1]$ by the concavity of the log operator. At $\nu = 0$ and $q = 0$, $\psi(L|L, (0, 0); 0, 0) = 1 - \pi_A$ since $\alpha_0 = 0$. As in [Example 1](#), $\gamma_1(L, (0, 0); 0, 0) > 0$. This establishes that there exists a cutoff $q_2 \in (0, 1)$ such that for $q < q_2$, there exists a cutoff $\nu_2(q)$ such that for $\nu < \nu_2(q)$, $\gamma_1(L, (0, 0); \nu, q) > 0$ and $(0, 0) \notin \Lambda(L; \nu, q)$, and for $\nu > \nu_2(q)$, $\gamma_1(L, (0, 0); \nu, q) < 0$ and $(0, 0) \in \Lambda(L; \nu, q)$. For $q > q_2$, $\gamma_1(L, (0, 0); \nu, q) < 0$ and $(0, 0) \in \Lambda(L; \nu, q)$.

To show that $q_1 < q_2$ and $\nu_1(q) < \nu_2(q)$ for all $q < q_1$, note $\gamma_1(L, (\infty, \infty); \nu, q) - \gamma_1(L, (\infty, \infty); 1, q)$ is equal to

$$\pi_A(1 - q)(\alpha_\nu - \alpha_1) \left(\log \frac{1 - \pi_A\alpha_1}{\pi_A\alpha_1 + 1 - \pi_A} - \log \frac{\alpha_1}{1 - \alpha_1} \right)$$

and by the symmetry of the signal distributions, $\gamma_1(L, (0, 0); \nu, q) - \gamma_1(L, (0, 0); 1, q) =$

$\gamma_1(L, (\infty, \infty); \nu, q) - \gamma_1(L, (\infty, \infty); 1, q)$. Moreover, $\gamma_1(L, (0, 0); 1, q) - \gamma_1(L, (\infty, \infty); 1, q)$ is 0 at $\pi_A = 0$, 0 at $\pi_A = 1$, and concave in π_A since the second derivative is

$$\frac{\pi_A(1 - 4q + 4q^2)(2 - 2\alpha_1 - 1)}{(\pi_A(1 - \alpha_1) + 1 - \pi_A)^2(\pi_A\alpha_1 + 1 - \pi_A)^2} \leq 0.$$

Therefore, $(0, 0) \notin \Lambda(\omega; \nu, q)$ before $(\infty, \infty) \in \Lambda(\omega; \nu, q)$.

Next consider $\omega = R$. Then $\gamma(R, (\infty, \infty); 1, q) > 0$ and $\gamma(R, (0, 0); 1, q) > 0$ for all $q \in [0, 1]$, since only correct learning can occur at $\nu = 1$. The only change in the above expressions is that now the true probabilities of each action are taken with respect to state R , rather than state L . Therefore, the comparative statics are similar to the comparative statics in state L : $\gamma_1(R, (0, 0); \nu, q)$ and $\gamma_1(R, (\infty, \infty); \nu, q)$ are decreasing in ν and q . Therefore, $\gamma_1(R, (0, 0); \nu, q) > 0$ for all ν and q , which implies $(0, 0) \notin \Lambda(R; \nu, q)$ for all ν and q . Similarly, $\gamma_1(R, (\infty, \infty); \nu, q) > 0$ for all ν and q , which implies $(\infty, \infty) \in \Lambda(R; \nu, q)$ for all ν and q . Therefore, $\Lambda(R; \nu, q) = \{(\infty, \infty)\}$ for all ν and q and learning is almost surely correct. $\square$

C Proofs from Section 4

C.1 Section 4.1 (Overreaction)

Proof of Observation 1. Suppose agents observe signals directly. Modify the definition of the expected change in the log likelihood ratio to allow for an uncountable signal space (as opposed to a finite action space):

$$\tilde{\gamma}(\omega, \boldsymbol{\lambda}; \nu) \equiv \int_{s \in \mathcal{S}} \log \left(\frac{s}{1-s} \right)^\nu dF^\omega(s).$$

Then $\tilde{\gamma}(\omega, \boldsymbol{\lambda}; \nu) = \nu \tilde{\gamma}(\omega, \boldsymbol{\lambda}; 1)$ since $\int_{s \in \mathcal{S}} \log \left(\frac{s}{1-s} \right)^\nu dF^\omega(s) = \nu \int_{s \in \mathcal{S}} \log \left(\frac{s}{1-s} \right) dF^\omega(s) = \nu \tilde{\gamma}(\omega, \boldsymbol{\lambda}; 1)$, where $\tilde{\gamma}(\omega, \boldsymbol{\lambda}; 1)$ is the expected change in the log likelihood ratio in the correctly specified model. Therefore, $\tilde{\gamma}(\omega, \boldsymbol{\lambda}; \nu)$ has the same sign as $\tilde{\gamma}(\omega, \boldsymbol{\lambda}; 1)$. Since correct learning obtains almost surely when agents have a correctly specified model, $\Lambda(L; 1) = \{0\}$ and $\Lambda(R; 1) = \{\infty\}$. This implies that $\Lambda(L; \nu) = \{0\}$ and $\Lambda(R; \nu) = \{\infty\}$ for all $\nu \in [1, \infty)$. Berk (1966) shows that beliefs converges a.s. in state L to the unique element in $\Lambda(L)$. Therefore, correct learning occurs almost surely, independent of ν . $\square$

Proof of Proposition 1. Let $x \equiv \pi(\theta_1)/\pi(\theta_2)$ denote the ratio of social to autarkic types. If an agent is an autarkic type with overreaction parameter ν , then $\hat{p}^*(\nu) \equiv \frac{(p^*)^{1/\nu}}{(1-p^*)^{1/\nu} + (p^*)^{1/\nu}}$ is the signal cutoff to choose action a_1 . Note that this reduces to p^* for a correctly specified type, i.e. $\hat{p}^*(1) = p^*$.

We first construct the locally stable set. We write $\gamma_i(\omega, \boldsymbol{\lambda}; x, \nu)$ and $\Lambda(\omega; x, \nu)$ to make these expressions' dependence on parameters x and ν explicit. Define $\Gamma_0(x, \nu) \equiv \gamma_1(L, 0; x, \nu)(x+1)$ and $\Gamma_\infty(x, \nu) \equiv \gamma_1(L, \infty; x, \nu)(x+1)$. Then from the construction of $\gamma_i(\omega, \boldsymbol{\lambda}; x, \nu)$,

$$\begin{aligned} \Gamma_0(x, \nu) &\equiv (F^L(\hat{p}^*(\nu)) + x) \log \frac{F^R(p^*) + x}{F^L(p^*) + x} - F^R(\hat{p}^*(\nu)) \log \frac{F^R(p^*)}{F^L(p^*)} \\ &\quad + (F^L(1/2) - F^R(1/2) + F^R(\hat{p}^*(\nu)) - F^L(\hat{p}^*(\nu))) \log \frac{F^R(1/2) - F^R(p^*)}{F^L(1/2) - F^L(p^*)} \end{aligned}$$

$$\begin{aligned}\Gamma_\infty(x, \nu) \equiv & -(F^R(\hat{p}^*(\nu)) + x) \log \frac{F^R(p^*) + x}{F^L(p^*) + x} + F^L(\hat{p}^*(\nu)) \log \frac{F^R(p^*)}{F^L(p^*)} \\ & + (F^L(1/2) - F^R(1/2) + F^R(\hat{p}^*(\nu)) - F^L(\hat{p}^*(\nu))) \log \frac{F^R(1/2) - F^R(p^*)}{F^L(1/2) - F^L(p^*)}.\end{aligned}$$

These functions have the same sign as $\gamma_1(L, 0; x, \nu)$ or $\gamma_1(L, \infty; x, \nu)$, respectively. Therefore, the signs of $\Gamma_0(x, \nu)$ and $\Gamma_\infty(x, \nu)$ can be used to characterize the locally stable set $\Lambda(\omega; x, \nu)$. Since there is a single social type, long-run learning is fully determined by $\Lambda(\omega; x, \nu)$.

To show the desired cutoffs exist, we show (i) $\nu \mapsto \Gamma_0(x, \nu)$ crosses zero at most once for a fixed x , (ii) if $0 \notin \Lambda(L; x, \nu)$ for some x' , then $0 \notin \Lambda(L; x, \nu)$ for all $x > x'$, (iii) $0 \notin \Lambda(L; x, \nu)$ for all (x, ν) . To show (i), note that the derivative of $\Gamma_0(x, \nu)$ with respect to ν is

$$\begin{aligned}\frac{\partial \Gamma_0}{\partial \nu} = & \frac{d\hat{p}^*(\nu)}{d\nu} f^L(\hat{p}^*(\nu)) \\ & \times \left(\log \frac{F^R(p^*) + x}{F^L(p^*) + x} - \frac{\hat{p}^*(\nu)}{1 - \hat{p}^*(\nu)} \log \frac{F^R(p^*)}{F^L(p^*)} - \left(1 - \frac{\hat{p}^*(\nu)}{1 - \hat{p}^*(\nu)} \right) \log \frac{F^R(1/2) - F^R(p^*)}{F^L(1/2) - F^L(p^*)} \right),\end{aligned}$$

where we use the property that $f^R(\hat{p}^*(\nu))/f^L(\hat{p}^*(\nu)) = \hat{p}^*(\nu)/(1 - \hat{p}^*(\nu))$ which follows from the normalization that signal realizations are posterior beliefs. The sign of this derivative is the same as the sign of

$$\log \frac{F^R(p^*) + x}{F^L(p^*) + x} + \frac{\hat{p}^*(\nu)}{1 - \hat{p}^*(\nu)} \left(\log \frac{F^R(1/2) - F^R(p^*)}{F^L(1/2) - F^L(p^*)} - \log \frac{F^R(p^*)}{F^L(p^*)} \right) - \log \frac{F^R(1/2) - F^R(p^*)}{F^L(1/2) - F^L(p^*)}.$$

This expression is increasing in ν , so $\nu \mapsto \Gamma_0(x, \nu)$ is either decreasing, U-shaped or increasing. Given $\Gamma_0(x, 1) \leq 0$, $\nu \mapsto \Gamma_0(x, \nu)$ changes signs at most once. Therefore, for a fixed x , there exists a cutoff $\bar{\nu} > 1$ such that $0 \notin \Lambda(L; x, \nu)$ for all $\nu > \bar{\nu}$ and $0 \in \Lambda(L; x, \nu)$ for all $\nu < \bar{\nu}$. For (ii), note that the derivative $\partial \Gamma_0 / \partial \nu$ is strictly increasing in x . If we can show that $\Gamma_0(x, 1)$ is increasing in x , then as x increases, $\Lambda = 0$ becomes unstable at a lower value of ν . The derivative of $\Gamma_0(x, 1)$ with respect to x is

$$\frac{\partial \Gamma_0}{\partial x} = \log \frac{F^R(p^*) + x}{F^L(p^*) + x} + \frac{F^L(p^*) - F^R(p^*)}{F^R(p^*) + x}.$$

Moreover, the second derivative is

$$\frac{\partial^2 \Gamma_0}{\partial x^2} = -\frac{(F^L(p^*) - F^R(p^*))^2}{(F^L(p^*) + x)(F^R(p^*) + x)^2} < 0.$$

So $x \mapsto \Gamma_0(x, 1)$ is concave in x and $\lim_{x \rightarrow \infty} \frac{\partial \Gamma_0}{\partial x}(x, 1) = 0$. Therefore, $\frac{\partial \Gamma_0}{\partial x}(x, 1) \geq 0$ for all x . Finally, $\Gamma_0(x, \nu) \geq \Gamma_0(x', \nu)$ for $x > x'$. Therefore, as x increases, $\gamma_1(L, 0; x, \nu)$ crosses 0 at a lower ν , i.e. if $0 \notin \Lambda(L; x', \nu)$ then $0 \notin \Lambda(L; x, \nu)$. For (iii), the derivative of $\Gamma_\infty(x, \nu)$ with respect to ν is

$$\frac{\partial \Gamma_\infty}{\partial \nu} = \frac{d\hat{p}^*(\nu)}{d\nu} f^R(\hat{p}^*(\nu))$$

$$\times \left(-\log \frac{F^R(p^*) + x}{F^L(p^*) + x} + \frac{1 - \hat{p}^*(\nu)}{\hat{p}^*(\nu)} \log \frac{F^R(p^*)}{F^L(p^*)} - \left(\frac{1 - \hat{p}^*(\nu)}{\hat{p}^*(\nu)} - 1 \right) \log \frac{F^R(1/2) - F^R(p^*)}{F^L(1/2) - F^L(p^*)} \right)$$

This derivative is maximized at $x = 0$ for a fixed ν since $\log \frac{F^R(p^*) + x}{F^L(p^*) + x}$ is monotone in x . At $x = 0$, $\frac{\partial \Gamma_\infty}{\partial \nu}(0, \nu) < 0$. Therefore, $\frac{\partial \Gamma_\infty}{\partial \nu}(x, \nu) < 0$ for all (x, ν) and $\infty \notin \Lambda(L; x, \nu)$ for all (x, ν) .

The symmetric environment implies identical cutoffs in state R . Therefore, $\bar{\pi}$ and $\bar{\nu}$ exist and satisfy the desired properties. Finally

$$\lim_{x \rightarrow \infty} \lim_{\nu \rightarrow \infty} \Gamma_0(x, \nu) = F^R(p^*) - F^L(p^*) - F^R(1/2) \log \frac{F^R(p^*)}{F^L(p^*)} > 0$$

by assumption. Therefore, cyclical learning occurs for some parameters. $\square$

C.2 Section 4.2 (Naive Learning)

We first prove [Proposition 3](#) and then [Proposition 2](#), as the latter is based on the former.

Proof of Proposition 3. Let $\alpha_L \equiv F^L(1/2)$ be the probability an autarkic type plays action L in state L and $\alpha_R \equiv F^R(1/2)$ be the probability an autarkic type plays action L in state R . Note that $\alpha_L \in (0, 1)$ and $\alpha_R \in (0, 1)$, since private signals are informative. In a slight abuse of notation, let $\hat{\pi}_i$ denote $\hat{\pi}_i(\theta_A)$ and π denote $\pi(\theta_A)$ to abbreviate the following expressions.

We first construct the locally stable set. We write $\gamma_i(\omega, \boldsymbol{\lambda}; \hat{\pi}_i)$ and $\Lambda(\omega; \hat{\pi}_1, \hat{\pi}_2)$ to make these expressions' dependence on $\hat{\pi}_1$ and $\hat{\pi}_2$ explicit. The local stability of correct learning is determined by the sign of

$$\gamma_i(L, (0, 0); \hat{\pi}_i) = (\pi \alpha_L + 1 - \pi) \log \left(\frac{\hat{\pi}_i \alpha_R + 1 - \hat{\pi}_i}{\hat{\pi}_i \alpha_L + 1 - \hat{\pi}_i} \right) + \pi(1 - \alpha_L) \log \left(\frac{1 - \alpha_R}{1 - \alpha_L} \right).$$

If θ_i has a correctly specified model, $\gamma_i(L, (0, 0); \pi) < 0$. This expression is decreasing in $\hat{\pi}_i$. Therefore, $\gamma_i(L, (0, 0); \hat{\pi}_i) < 0$ for all $\hat{\pi}_i \geq \pi$. This implies that $(0, 0) \in \Lambda(L; \hat{\pi}_1, \hat{\pi}_2)$ for all $\hat{\pi}_1, \hat{\pi}_2$. Therefore, correct learning arises with positive probability at any level of heterogeneity. The local stability of incorrect learning is determined by the sign of

$$\gamma_i(L, (\infty, \infty); \hat{\pi}_i) = \pi \alpha_L \log \left(\frac{\alpha_R}{\alpha_L} \right) + (\pi(1 - \alpha_L) + 1 - \pi) \log \left(\frac{\hat{\pi}_i(1 - \alpha_R) + 1 - \hat{\pi}_i}{\hat{\pi}_i(1 - \alpha_L) + 1 - \hat{\pi}_i} \right).$$

This expression is increasing in $\hat{\pi}_i$ and is equivalent to the representative agent model at $\hat{\pi}_i = \hat{\pi}$. Therefore, if $\gamma_i(L, (\infty, \infty); \hat{\pi}) < 0$, then $\gamma_1(L, (\infty, \infty); \hat{\pi}_1) < 0$ since $\hat{\pi}_1 \leq \hat{\pi}$ by definition. This implies that if incorrect learning does not arise in the representative agent model with bias $\hat{\pi}$, i.e. $(\infty, \infty) \notin \Lambda(L; \hat{\pi}, \hat{\pi})$, then it does not arise in any corresponding heterogeneous model with average bias $\hat{\pi}$, i.e. $(\infty, \infty) \notin \Lambda(L; \hat{\pi}_1, \hat{\pi}_2)$ for all $\hat{\pi}_1, \hat{\pi}_2$ such that $(\hat{\pi}_1 + \hat{\pi}_2)/2 = \hat{\pi}$. Further, we know from [Bohren \(2016\)](#) that there exists a cutoff $\bar{\pi} \in (\pi, 1]$ such that for $\hat{\pi}_i > \bar{\pi}$, $\gamma_i(L, (\infty, \infty); \hat{\pi}_i) > 0$, with $\bar{\pi} < 1$ for small enough π . Therefore, $(\infty, \infty) \in \Lambda(L; \hat{\pi}, \hat{\pi})$ for $\hat{\pi} > \bar{\pi}$ and $(\infty, \infty) \in \Lambda(L; \hat{\pi}_1, \hat{\pi}_2)$ for $\hat{\pi}_1 > \bar{\pi}$. The local stability of

disagreement is determined by the sign of

$$\begin{aligned}\gamma_i(L, (0, \infty); \hat{\pi}_i) &= (\pi\alpha_L + (1 - \pi)/2) \log \left(\frac{\hat{\pi}_i\alpha_R + \frac{1}{2}(1 - \hat{\pi}_i)}{\hat{\pi}_i\alpha_L + \frac{1}{2}(1 - \hat{\pi}_i)} \right) \\ &\quad + (\pi(1 - \alpha_L) + (1 - \pi)/2) \log \left(\frac{\hat{\pi}_i(1 - \alpha_R) + \frac{1}{2}(1 - \hat{\pi}_i)}{\hat{\pi}_i(1 - \alpha_L) + \frac{1}{2}(1 - \hat{\pi}_i)} \right) \\ &= \pi(2\alpha_L - 1) \log \left(\frac{\hat{\pi}_i(1 - \alpha_L) + \frac{1}{2}(1 - \hat{\pi}_i)}{\hat{\pi}_i\alpha_L + \frac{1}{2}(1 - \hat{\pi}_i)} \right),\end{aligned}$$

where the second equality follows from symmetry, $\alpha_R = 1 - \alpha_L$. Given $\alpha_L > 1/2$, $(\hat{\pi}_i(1 - \alpha_L) + \frac{1}{2}(1 - \hat{\pi}_i))/(\hat{\pi}_i\alpha_L + \frac{1}{2}(1 - \hat{\pi}_i)) < 1$ and $2\alpha_L - 1 > 0$. Therefore, $\gamma_i(L, (0, \infty); \hat{\pi}_i) < 0$ for any $\hat{\pi}_i$. This implies that $(0, \infty)$ almost surely does not arise, i.e. $(0, \infty) \notin \Lambda(L; \hat{\pi}_1, \hat{\pi}_2)$. Given $\gamma_i(L, (\infty, 0); \hat{\pi}_i) = \gamma_i(L, (0, \infty); \hat{\pi}_i)$, $(\infty, 0)$ almost surely does not arise. Therefore, almost surely disagreement does not arise. The construction of $\Lambda(R; \hat{\pi}_1, \hat{\pi}_2)$ is analogous.

Next, we rule out mixed learning. Since correct learning is always locally stable, the only candidate mixed outcomes are $\lambda_1^* = \infty$ or $\lambda_2^* = \infty$. As argued above $\gamma_1(L, (0, \infty); \hat{\pi}_1) < 0$ for any $\hat{\pi}_1$ and $\gamma_2(L, (\infty, 0); \hat{\pi}_2) < 0$ for any $\hat{\pi}_2$. This implies $\Lambda_M(L) = \emptyset$. Therefore, mixed learning almost surely does not arise. The construction of $\Lambda_M(R)$ is analogous.

Given $\Lambda_M(\omega) = \emptyset$ and $\Lambda(\omega; \hat{\pi}_1, \hat{\pi}_2)$ does not contain any disagreement outcomes—and therefore, we do not need to consider maximal accessibility—by [Theorem 4](#), $\Lambda(\omega; \hat{\pi}_1, \hat{\pi}_2)$ fully characterizes the set of asymptotic learning outcomes. From the above characterization, either $\Lambda(\omega; \hat{\pi}_1, \hat{\pi}_2) = \{(0, 0)\}$ or $\Lambda(\omega; \hat{\pi}_1, \hat{\pi}_2) = \{(0, 0), (\infty, \infty)\}$. Therefore, either learning is almost surely correct, or learning is almost surely correct or incorrect with both occurring with positive probability. Further, if $\Lambda(\omega; \hat{\pi}, \hat{\pi}) = \{(0, 0)\}$, then $\Lambda(\omega; \hat{\pi}_1, \hat{\pi}_2) = \{(0, 0)\}$ for all $\hat{\pi}_1, \hat{\pi}_2$ such that $(\hat{\pi}_1 + \hat{\pi}_2)/2 = \hat{\pi}$, and if $\Lambda(\omega; \hat{\pi}_1, \hat{\pi}_2) = \{(0, 0), (\infty, \infty)\}$, then $\Lambda(\omega; \hat{\pi}, \hat{\pi}) = \{(0, 0), (\infty, \infty)\}$ at $\hat{\pi} = (\hat{\pi}_1 + \hat{\pi}_2)/2$.

Proof of Proposition 2. This result follows directly from the constructions of $\gamma_i(\omega, \boldsymbol{\lambda}; \hat{\pi}_i)$ in [Proposition 3](#). Generically, $\gamma_i(\omega, (0, 0); \hat{\pi}_i) \neq 0$ and $\gamma_i(\omega, (\infty, \infty); \hat{\pi}_i) \neq 0$ for $i = 1, 2$. Given an average bias $\hat{\pi}$, consider the case where $\gamma_i(\omega, (0, 0); \hat{\pi}) \neq 0$ and $\gamma_i(\omega, (\infty, \infty); \hat{\pi}) \neq 0$ for $i = 1, 2$. For any $\delta > 0$, there exists an ε such that for $|\hat{\pi}_1 - \hat{\pi}| < \varepsilon/2$ and $|\hat{\pi}_2 - \hat{\pi}| < \varepsilon/2$, $|\gamma_i(\omega, \boldsymbol{\lambda}; \hat{\pi}_i) - \gamma_i(\omega, \boldsymbol{\lambda}; \hat{\pi})| < \delta$ for $\boldsymbol{\lambda} \in \{(0, 0), (\infty, \infty)\}$ and $i = 1, 2$. Choosing δ small enough ensures that $\gamma_i(\omega, \boldsymbol{\lambda}; \hat{\pi}_i)$ and $\gamma_i(\omega, \boldsymbol{\lambda}; \hat{\pi})$ have the same sign. Therefore, $\Lambda(\omega; \hat{\pi}_1, \hat{\pi}_2) = \Lambda(\omega; \hat{\pi}, \hat{\pi})$ and the heterogeneous set-up has the same set of learning outcomes as the corresponding representative agent set-up. $\square$

C.3 Section 4.3 (Level-k)

Proof of Proposition 4. Let $\boldsymbol{\lambda} = (\lambda_2, \lambda_3)$ denote the vector of likelihood ratios for the social types θ_2 and θ_3 . Note $\lambda_{1,t} = 1$ for all t . When type $\theta_i \in \{\theta_1, \theta_2, \theta_3\}$ has current belief λ_i , it chooses action R iff it observes a signal realization $s \geq 1/(\lambda_i + 1)$. Given $\lambda_1 = 1$, type θ_1 chooses action L with probability $F^\omega(0.5)$ and action R with probability $1 - F^\omega(0.5)$, independent of the history. Type θ_2 's subjective probability of each L action in the history is the probability that a level-1 type chooses action L , $\hat{\psi}_2(L|\omega, \boldsymbol{\lambda}) = F^\omega(0.5)$ and its subjective probability of each R action is $\hat{\psi}_2(R|\omega, \boldsymbol{\lambda}) = 1 - F^\omega(0.5)$, independent of the history. Given belief λ_2 , level-2 chooses an L action with probability $F^\omega(1/(\lambda_2 + 1))$ and an R action with

probability $1 - F^\omega(1/(\lambda_2 + 1))$. Type θ_3 's subjective probability of each L action is the weighted average of the probability that a level-1 type and a level-2 type choose action L ,

$$\hat{\psi}_3(L|\omega, \boldsymbol{\lambda}) = (1 - \varepsilon)F^\omega(1/(\lambda_2 + 1)) + \varepsilon F^\omega(.5),$$

which does depend on the history through λ_2 . The subjective probability of an R action is analogous. Finally, the *true* probability of an L action depends on the correct distribution over types,

$$\psi(L|\omega, \boldsymbol{\lambda}) = \pi(\theta_1)F^\omega(.5) + \pi(\theta_2)F^\omega(1/(\lambda_2 + 1)) + \pi(\theta_3)F^\omega(1/(\lambda_3 + 1)).$$

To simplify the exposition, let $\alpha_L \equiv F^L(.5)$ be the probability a level-1 type plays action L in state L and $\alpha_R \equiv F^R(.5)$ be the probability a level-1 type plays action L in state R . Note that $\alpha_L \in (0, 1)$ and $\alpha_R \in (0, 1)$, since private signals are informative.

Suppose $\omega = L$. We first consider local stability for the level-3 type. At the correct learning outcome, $(0, 0)$, the level-2 type chooses action L for all signal realizations. Therefore, the level-3 type believes that L actions are approximately uninformative for small ε , $\frac{\hat{\psi}_3(L|R, (0, 0))}{\hat{\psi}_3(L|L, (0, 0))} = \frac{1 - \varepsilon + \varepsilon \alpha_R}{1 - \varepsilon + \varepsilon \alpha_L} \approx 1$ and R actions are from the level-1 type, $\frac{\hat{\psi}_3(R|R, (0, 0))}{\hat{\psi}_3(R|L, (0, 0))} = \frac{1 - \alpha_R}{1 - \alpha_L}$. Since only the level-1 type plays action R , the true probability of an R action is $\pi(\theta_1)(1 - \alpha_L)$. Therefore, for small ε , $\gamma_3(L, (0, 0)) = (\pi(\theta_1)\alpha_L + \pi(\theta_2) + \pi(\theta_3)) \log \frac{1 - \varepsilon + \varepsilon \alpha_R}{1 - \varepsilon + \varepsilon \alpha_L} + \pi(\theta_1)(1 - \alpha_L) \log \frac{1 - \alpha_R}{1 - \alpha_L} \approx \pi(\theta_1)(1 - \alpha_L) \log \frac{1 - \alpha_R}{1 - \alpha_L} > 0$ and correct learning is not locally stable for the level-3 type, $(0, 0) \notin \Lambda_3(L)$. Similarly, for small ε , $\gamma_3(L, (\infty, \infty)) \approx \pi(\theta_1)\alpha_L \log \frac{\alpha_R}{\alpha_L} < 0$ and incorrect learning is not locally stable for the level-3 type, $(\infty, \infty) \notin \Lambda_3(L)$. This establishes that correct learning and incorrect learning almost surely do not occur for small ε , as neither outcome is locally stable for level-3 types.

This leaves the disagreement outcomes as candidate learning outcomes. Consider $(0, \infty)$. As in the case of $(0, 0)$, the level-3 type believes that L actions are approximately uninformative and R actions are from the level-1 type. But now, this confirms the level-3 type's belief that the state is R , $\gamma_3(L, (0, \infty)) \approx (\pi(\theta_1)(1 - \alpha_L) + \pi(\theta_3)) \log \frac{1 - \alpha_R}{1 - \alpha_L} > 0$ and $(0, \infty) \in \Lambda_3(L)$. Similarly, $\gamma_3(L, (\infty, 0)) \approx (\pi(\theta_1)\alpha_L + \pi(\theta_3)) \log \frac{\alpha_R}{\alpha_L} < 0$ and $(\infty, 0) \in \Lambda_3(L)$. Therefore, for small ε , both disagreement outcomes are locally stable for the level-3 type, $\Lambda_3(L) = \{(0, \infty), (\infty, 0)\}$.

Next, we determine whether the disagreement outcomes are locally stable for the level-2 type. The level-2 type believes that all actions are from level-1 types. Therefore, it interprets L and R actions in the same way at both disagreement outcomes. At $(0, \infty)$, the true probability of an L action is $\pi(\theta_1)\alpha_L + \pi(\theta_2)$, while at $(\infty, 0)$, it is $\pi(\theta_1)\alpha_L + \pi(\theta_3)$. Therefore, $\gamma_2(L, (0, \infty)) = (\pi(\theta_1)\alpha_L + \pi(\theta_2)) \log \frac{\alpha_R}{\alpha_L} + (\pi(\theta_1)(1 - \alpha_L) + \pi(\theta_3)) \log \frac{1 - \alpha_R}{1 - \alpha_L}$ and $\gamma_2(L, (\infty, 0)) = (\pi(\theta_1)\alpha_L + \pi(\theta_3)) \log \frac{\alpha_R}{\alpha_L} + (\pi(\theta_1)(1 - \alpha_L) + \pi(\theta_2)) \log \frac{1 - \alpha_R}{1 - \alpha_L}$. The signs of these expressions vary with the true distribution of types. To characterize the region of the type distribution at which each disagreement outcome is locally stable, we use the inequalities (a) $\frac{\alpha_R}{\alpha_L} < 1$, (b) $\frac{1 - \alpha_R}{1 - \alpha_L} > 1$ and (c) from the correctly specified model, $\alpha_L \log \frac{\alpha_R}{\alpha_L} + (1 - \alpha_L) \log \frac{1 - \alpha_R}{1 - \alpha_L} < 0$, as well as the property that $\pi \mapsto \gamma_2(L, (0, \infty))$ and $\pi \mapsto \gamma_2(L, (\infty, 0))$ are continuous.

Case (i): As $\pi(\theta_3) \rightarrow 0$, $\gamma_2(L, (0, \infty)) \rightarrow (\pi(\theta_1)\alpha_L + 1 - \pi(\theta_1)) \log \frac{\alpha_R}{\alpha_L} + \pi(\theta_1)(1 - \alpha_L) \log \frac{1 - \alpha_R}{1 - \alpha_L} < 0$ for all $\pi(\theta_1)$, where the negative sign follows from inequalities (a) and

(c). Therefore, there exists a cutoff $c_1 > 0$ such that for $\pi(\theta_3) < c_1$, $(0, \infty) \in \Lambda_2(L)$ for all $\pi(\theta_1)$ and $\pi(\theta_2)$.

Case (ii): As $\pi(\theta_3) \rightarrow 1$, $\gamma_2(L, (0, \infty)) \rightarrow \log \frac{1-\alpha_R}{1-\alpha_L} > 0$ and $\gamma_2(L, (\infty, 0)) \rightarrow \log \frac{\alpha_R}{\alpha_L} < 0$. Therefore, there exists an interior cutoff $c_2 \in (0, 1)$ such that for $\pi(\theta_3) > c_2$, $(0, \infty) \notin \Lambda_2(L)$ and there exists a cutoff $c_3 < 1$ such that for $\pi(\theta_3) > c_3$, $(\infty, 0) \notin \Lambda_2(L)$ for all $\pi(\theta_1)$ and $\pi(\theta_2)$, where $c_2 > 0$ follows from part (i). Therefore, there exists an interior cutoff $\bar{\pi}_3 = \max\{c_2, c_3\} \in (0, 1)$ such that if $\pi(\theta_3) > \bar{\pi}_3$, neither disagreement outcome is locally stable for θ_2 . Combined with $\Lambda_3(L) = \{(0, \infty), (\infty, 0)\}$, this implies that $\Lambda(L) = \emptyset$ for $\pi(\theta_3) > \bar{\pi}_3$ and small ε .

Case (iii): As $\pi(\theta_2) \rightarrow 0$, $\gamma_2(L, (\infty, 0)) \rightarrow (\pi(\theta_1)\alpha_L + 1 - \pi(\theta_1)) \log \frac{\alpha_R}{\alpha_L} + \pi(\theta_1)(1 - \alpha_L) \log \frac{1-\alpha_R}{1-\alpha_L} < 0$ for all $\pi(\theta_1)$, where the negative sign follows from inequalities (a) and (c). Therefore, there exists a cutoff $c_4 > 0$ such that for $\pi(\theta_2) < c_4$, $(\infty, 0) \notin \Lambda_2(L)$ for all $\pi(\theta_1)$ and $\pi(\theta_3)$.

Case (iv): As $\pi(\theta_2) \rightarrow 1$, $\gamma_2(L, (0, \infty)) \rightarrow \log \frac{\alpha_R}{\alpha_L} < 0$ and $\gamma_2(L, (\infty, 0)) \rightarrow \log \frac{1-\alpha_R}{1-\alpha_L} > 0$. Therefore, there exists a cutoff $c_5 < 1$ such that for $\pi(\theta_2) > c_5$, $(0, \infty) \in \Lambda_2(L)$ and there exists an interior cutoff $c_6 \in (0, 1)$ such that for $\pi(\theta_2) > c_6$, $(\infty, 0) \in \Lambda_2(L)$ for all $\pi(\theta_1)$ and $\pi(\theta_3)$, where $c_6 > 0$ follows from case (iii). Therefore, there exists an interior cutoff $\bar{\pi}_2 = \max\{c_5, c_6\} \in (0, 1)$ such that if $\pi(\theta_2) > \bar{\pi}_2$, both disagreement outcomes are locally stable for θ_2 . Combined with $\Lambda_3(L) = \{(0, \infty), (\infty, 0)\}$, this implies that $\Lambda(L) = \{(0, \infty), (\infty, 0)\}$ for $\pi(\theta_2) > \bar{\pi}_2$ and small ε .

Case (v): As $\pi(\theta_1) \rightarrow 1$, $\gamma_2(L, (0, \infty)) \rightarrow \alpha_L \log \frac{\alpha_R}{\alpha_L} + (1 - \alpha_L) \log \frac{1-\alpha_R}{1-\alpha_L} < 0$ and $\gamma_2(L, (\infty, 0)) \rightarrow \alpha_L \log \frac{\alpha_R}{\alpha_L} + (1 - \alpha_L) \log \frac{1-\alpha_R}{1-\alpha_L} < 0$. Therefore, there exists an interior cutoff $c_7 \in (0, 1)$ such that for $\pi(\theta_1) > c_7$, $(0, \infty) \in \Lambda_2(L)$ and there exists an interior cutoff $c_8 \in (0, 1)$ such that for $\pi(\theta_1) > c_8$, $(\infty, 0) \notin \Lambda_2(L)$ for all $\pi(\theta_2)$ and $\pi(\theta_3)$, where $c_7 > 0$ and $c_8 > 0$ follow from cases (ii) and (iv). Therefore, there exists an interior cutoff $\bar{\pi}_1 = \max\{c_7, c_8\} \in (0, 1)$ such that if $\pi(\theta_1) > \bar{\pi}_1$, $(0, \infty)$ is locally stable for θ_2 and $(\infty, 0)$ is not. Combined with $\Lambda_3(L) = \{(0, \infty), (\infty, 0)\}$, this implies that $\Lambda(L) = \{(0, \infty)\}$ for $\pi(\theta_1) > \bar{\pi}_1$ and small ε .

Fixing $\pi(\theta_2)$, $\gamma_2(L, (0, \infty))$ is increasing in $\pi(\theta_3)$. Given this, we next show that the type distribution can be divided into two connected regions in the simplex such that $(0, \infty) \in \Lambda_2(L)$ or $(0, \infty) \notin \Lambda_2(L)$, and these regions are separated by the unique solution to $\gamma_2(L, (0, \infty)) = 0$. As shown above, at $\pi(\theta_2) = 0$ and $\pi(\theta_3) = 0$, $\gamma_2(L, (0, \infty)) < 0$ and at $\pi(\theta_2) = 0$ and $\pi(\theta_3) = 1$, $\gamma_2(L, (0, \infty)) > 0$. Therefore, there exists a cutoff $c_9 \in (0, 1)$ such that at $\pi(\theta_2) = 0$ and $\pi(\theta_3) = c_9$, $\gamma_2(L, (0, \infty)) = 0$. Similarly, there exists a cutoff $c_{10} \equiv \log \frac{\alpha_L}{\alpha_R} / (\log \frac{\alpha_L}{\alpha_R} - \log \frac{1-\alpha_L}{1-\alpha_R})$ such that at $\pi(\theta_1) = 0$ and $\pi(\theta_3) = c_{10}$, $\gamma_2(L, (0, \infty)) = 0$. Given $\gamma_2(L, (0, \infty))$ is linear in $\pi(\theta_2)$ and $\pi(\theta_3)$, the solution to $\gamma_2(L, (0, \infty)) = 0$ is linear in the simplex and represented by the line connecting $(1 - c_9, 0, c_9)$ and $(0, 1 - c_{10}, c_{10})$. This establishes the above statement.

Fixing $\pi(\theta_2)$, $\gamma_2(L, (\infty, 0))$ is decreasing in $\pi(\theta_3)$. Therefore, by similar reasoning, the type distribution can be divided into two connected regions such that $(\infty, 0) \in \Lambda_2(L)$ or $(\infty, 0) \notin \Lambda_2(L)$, and these regions are separated by the unique solution to $\gamma_2(L, (\infty, 0)) = 0$. Given $\gamma_2(L, (\infty, 0))$ is linear in $\pi(\theta_2)$ and $\pi(\theta_3)$, the solution to $\gamma_2(L, (\infty, 0)) = 0$ is linear in the simplex and represented by the line connecting $(1 - c_{11}, c_{11}, 0)$ and $(0, 1 - c_{12}, c_{12})$, where $c_{11} \in (0, 1)$ is the value of $\pi(\theta_2)$ such that $\gamma_2(L, (\infty, 0)) = 0$ when $\pi(\theta_3) = 0$, and

$$c_{12} \equiv \log \frac{1-\alpha_L}{1-\alpha_R} / (\log \frac{1-\alpha_L}{1-\alpha_R} - \log \frac{\alpha_L}{\alpha_R}).$$

Given the linearity of both solutions, if $c_{10} \geq c_{12}$, then the solution to $\gamma_2(L, (0, \infty)) = 0$ lies above the solution to $\gamma_2(L, (\infty, 0)) = 0$. Therefore, there are three distinct regions such that for small ε , either (i) $\Lambda(L) = \emptyset$, (ii) $\Lambda(L) = \{(0, \infty)\}$, or (iii) $\Lambda(L) = \{(0, \infty), (\infty, 0)\}$. Otherwise, if $c_{10} \leq c_{12}$, the solutions cross exactly once. Therefore, there are four distinct regions such that for small ε , either (i) $\Lambda(L) = \emptyset$, (ii) $\Lambda(L) = \{(0, \infty)\}$, (iii) $\Lambda(L) = \{(\infty, 0)\}$, or (iv) $\Lambda(L) = \{(0, \infty), (\infty, 0)\}$. Note that when the signal distributions are symmetric, $c_{10} \geq c_{12}$. The construction of $\Lambda(R)$ is analogous.

We next show that both disagreement outcomes are maximally accessible at all type distributions. Formally, we show that for any $\pi \in \Delta((\theta_1, \theta_2, \theta_3))$ and $\varepsilon \in (0, 1]$, $(0, \infty)$ and $(\infty, 0)$ are maximally accessible. At $\lambda = (0, 0)$, type θ_2 perceives L actions as stronger evidence of state L than type θ_3 ,

$$\frac{\hat{\psi}_2(L|R, (0, 0))}{\hat{\psi}_2(L|L, (0, 0))} = \frac{\alpha_R}{\alpha_L} < \frac{\varepsilon + (1 - \varepsilon)\alpha_R}{\varepsilon + (1 - \varepsilon)\alpha_L} = \frac{\hat{\psi}_3(L|R, (0, 0))}{\hat{\psi}_3(L|L, (0, 0))},$$

and both types perceive R actions in the same way,

$$\frac{\hat{\psi}_2(R|R, (0, 0))}{\hat{\psi}_2(R|L, (0, 0))} = \frac{\hat{\psi}_3(R|R, (0, 0))}{\hat{\psi}_3(R|L, (0, 0))} = \frac{1 - \alpha_R}{1 - \alpha_L}.$$

Therefore, $\theta_3 \succ_{(0,0)} \theta_2$. From [Definition 7](#), this implies that $(0, \infty)$ is maximally accessible. At $\lambda = (\infty, \infty)$, type θ_2 perceives R actions as stronger evidence of state R than type θ_3 ,

$$\frac{\hat{\psi}_2(R|R, (\infty, \infty))}{\hat{\psi}_2(R|L, (\infty, \infty))} = \frac{1 - \alpha_R}{1 - \alpha_L} > \frac{\varepsilon + (1 - \varepsilon)(1 - \alpha_R)}{\varepsilon + (1 - \varepsilon)(1 - \alpha_L)} = \frac{\hat{\psi}_3(R|R, (\infty, \infty))}{\hat{\psi}_3(R|L, (\infty, \infty))},$$

and both types perceive L actions in the same way,

$$\frac{\hat{\psi}_2(L|R, (\infty, \infty))}{\hat{\psi}_2(L|L, (\infty, \infty))} = \frac{\hat{\psi}_3(L|R, (\infty, \infty))}{\hat{\psi}_3(L|L, (\infty, \infty))} = \frac{\alpha_R}{\alpha_L}.$$

Therefore, $\theta_2 \succ_{(\infty,\infty)} \theta_3$. From [Definition 7](#), this implies that $(\infty, 0)$ is maximally accessible. Therefore, a disagreement outcome arises with positive probability if and only if it is in $\Lambda(\omega)$.

Finally, we need to rule out mixed learning outcomes. Suppose $\omega = L$ and consider the four possible mixed outcomes. Consider $(0, \theta_3)$. By the concavity of the log operator, $\alpha_L \log \frac{\alpha_R}{\alpha_L} + (1 - \alpha_L) \log \frac{1-\alpha_R}{1-\alpha_L} < 0$. Therefore, since $\frac{\alpha_R}{\alpha_L} < 0$, $\gamma_2(L, (0, 0)) = (\pi_1 \alpha_L + \pi(\theta_2) + \pi(\theta_3)) \log \frac{\alpha_R}{\alpha_L} + \pi_1(1 - \alpha_L) \log \frac{1-\alpha_R}{1-\alpha_L} < 0$. and $(0, 0) \in \Lambda_2(L)$. By the definition of $\Lambda_M(L)$, this implies that $(0, \theta_3) \notin \Lambda_M(L)$ and this mixed learning outcome almost surely does not arise. Consider (∞, θ_3) . This outcome is in $\Lambda_M(L)$ if $(\infty, \infty) \notin \Lambda_2(L)$ and $(0, \infty) \notin \Lambda_2(L)$, which is equivalent to $\gamma_2(L, (\infty, \infty)) < 0$ and $\gamma_2(L, (0, \infty)) > 0$. However, $\gamma_2(L, (\lambda_2, \infty))$ is increasing in λ_2 , so this is not possible. Therefore, $(\infty, \theta_3) \notin \Lambda_M(L)$ and this mixed learning outcome almost surely does not arise. Consider $(0, \theta_2)$. This outcome is in $\Lambda_M(L)$ if $(0, 0) \notin \Lambda_3(L)$ and $(0, \infty) \notin \Lambda_3(L)$. From the characterization of $\Lambda(L)$ above, we know that $(0, \infty) \in \Lambda_3(L)$. Therefore, $(0, \theta_2) \notin \Lambda_M(L)$ and this mixed learning outcome almost

surely does not arise. Consider (∞, θ_2) . This outcome is in $\Lambda_M(L)$ if $(\infty, 0) \notin \Lambda_3(L)$ and $(\infty, \infty) \notin \Lambda_3(L)$. From the characterization of $\Lambda(L)$ above, we know that $(\infty, 0) \in \Lambda_3(L)$. Therefore, $(\infty, \theta_2) \notin \Lambda_M(L)$ and this mixed learning outcome almost surely does not arise. Together, this establishes $\Lambda_M(L) = \emptyset$. Similar logic shows $\Lambda_M(R) = \emptyset$.

Given $\Lambda_M(\omega) = \emptyset$ and both disagreement outcomes are maximally accessible, by [Theorem 4](#), $\Lambda(\omega)$ determines the set of asymptotic learning outcomes. As $\varepsilon \rightarrow 1$, $\Lambda(\omega) \subseteq \{(0, \infty), (\infty, 0)\}$. Either $\Lambda(\omega) = \emptyset$, in which case learning is cyclical for both types, or $\Lambda(\omega) \neq \emptyset$, in which case beliefs almost surely converge to a limit random variable with support $\Lambda(\omega)$. The construction of $\Lambda(\omega)$ above establishes the cutoffs on the type distribution such that $\Lambda(\omega) = \emptyset$, $\Lambda(\omega) = \{(0, \infty)\}$, $\Lambda(\omega) = \{(\infty, 0)\}$ or $\Lambda(\omega) = \{(0, \infty), (\infty, 0)\}$. $\square$

D Learning Characterization: More than Two Social Types

This section proves analogues of the global stability of disagreement, mixed learning, and belief convergence results in [Section 3](#) and [Appendix A](#) for any finite number of social types. Together, this establishes a direct analogue of [Theorem 4](#); an analogue of [Corollary 2](#) immediately follows. These results nest the case of $k \leq 2$.

D.1 Global Stability of Disagreement

We first prove an analogue of [Theorem 7](#) to show that separability can also be used to establish the global stability of a disagreement outcome when there are more than two social types. We then extend the definition of maximal accessibility and prove that it implies the separability condition, establishing an analogue of [Theorem 3](#).

Theorem 7' (Global Stability of Disagreement ($k \geq 2$)). *Consider a learning environment that is identified at certainty and satisfies [Assumptions 1](#) to [4](#). Suppose disagreement outcome $\lambda^* \in \Lambda(\omega)$ and, starting from agreement outcome $\lambda_1^* \in \{0^k, \infty^k\}$, there exists a finite sequence of adjacent disagreement outcomes $\lambda_2^*, \dots, \lambda_L^* = \lambda^*$ such that for $l = 1, \dots, L - 1$, either (i) $(\lambda_l^*)_i = 0$, $(\lambda_{l+1}^*)_i = \infty$ and λ_l^* is separable at zero for θ_i , or (ii) $(\lambda_l^*)_i = \infty$, $(\lambda_{l+1}^*)_i = 0$ and λ_l^* is separable at infinity for θ_i . Then λ^* is globally stable in state ω .*

Proof. Given $\kappa \in \{1, \dots, k - 1\}$, consider disagreement outcome $\lambda^* = (0^\kappa, \infty^{k-\kappa})$. Suppose $\lambda^* \in \Lambda(\omega)$ and for each $l = 1, \dots, k - \kappa$, $\lambda_l^* = (0^{k-l+1}, \infty^{l-1})$ is separable at zero for type θ_{k-l+1} . Given $\lambda_l^* = (0^{k-l+1}, \infty^{l-1})$ is separable at zero for type θ_{k-l+1} , by [Lemma 5](#), $\lambda_{l+1}^* = (0^{k-l}, \infty^l)$ is adjacently accessible from λ_l^* . Since this holds for each element of the sequence starting at $\lambda_1^* = 0^k$ and ending at $\lambda_{k-\kappa+1}^* = \lambda^*$, by [Lemma 6](#), λ^* is accessible. Fix an initial belief $\lambda_1 \in (0, \infty)^k$ and choose an $\varepsilon < e^{-E}$, where E is defined in [Eq. \(11\)](#). By accessibility, there exists a finite sequence ξ of N actions that occurs with positive probability, such that following ξ , $\lambda_{N+1} \in B_\varepsilon(\lambda^*)$. Given λ^* is locally stable, this implies $Pr(\lambda_t \rightarrow \lambda^* | h = \xi) > 0$. Given $Pr(h = \xi) > 0$, from any $\lambda_1 \in (0, \infty)^k$, $Pr(\lambda_t \rightarrow \lambda^*) > 0$. This establishes that λ^* is globally stable. The case in which there is a sequence of stationary beliefs that are separable at infinity is analogous, as is the proof for other disagreement outcomes. $\square$

We next use the maximal R-order $\succ_\lambda$ to define a sufficient condition for separability, which we refer to as maximally separable. We use this condition to extend the definition of maximal accessibility to the case of more than two social types.

Definition 11 (Maximally Separable ($k \geq 2$)). Belief $\lambda^* \in \{0, \infty\}^k \setminus \infty^k$ is maximally separable at zero for type θ_i with $\lambda_i^* = 0$ if $\theta_j \succeq_{\lambda^*} \theta_i$ for all j with $\lambda_j^* = \infty$ and $\theta_i \succ_{\lambda^*} \theta_j$ for all $j \neq i$ with $\lambda_j^* = 0$. Belief $\lambda^* \in \{0, \infty\}^k \setminus 0^k$ is maximally separable at infinity for type θ_i with $\lambda_i^* = \infty$ if $\theta_j \succ_{\lambda^*} \theta_i$ for all $j \neq i$ with $\lambda_j^* = \infty$ and $\theta_i \succeq_{\lambda^*} \theta_j$ for all j with $\lambda_j^* = 0$.

Definition 7' (Maximal Accessibility ($k \geq 2$)). Disagreement outcome $\lambda^* \in \{0, \infty\}^k \setminus \{0^k, \infty^k\}$ is maximally accessible if, starting from agreement outcome $\lambda_1^* \in \{0^k, \infty^k\}$, there exists a finite sequence of adjacent disagreement outcomes $\lambda_2^*, \dots, \lambda_L^* = \lambda^*$ such that for $l = 1, \dots, L-1$, either (i) $(\lambda_l^*)_i = 0$, $(\lambda_{l+1}^*)_i = \infty$ and λ_l^* is maximally separable at zero for θ_i , or (ii) $(\lambda_l^*)_i = \infty$, $(\lambda_{l+1}^*)_i = 0$ and λ_l^* is maximally separable at infinity for θ_i .

As in the case of $k = 2$, maximal accessibility guarantees that there exists a finite sequence of a_1 and a_M actions that separates beliefs and moves them to a neighborhood of the disagreement outcome. It is straightforward to verify from the primitives of the model and is equivalent to Definition 7 when $k = 2$. Using Definition 7', the statement of Theorem 3' is identical to Theorem 3.

Theorem 3' (Global Stability of Disagreement ($k \geq 2$)). Consider a learning environment that satisfies Assumptions 1 to 4. If disagreement outcome λ^* is in $\Lambda(\omega)$ and maximally accessible, then λ^* is globally stable in state ω .

Proof. We show that Definition 7' implies the conditions for separability outlined in Theorem 7'. Given $\kappa \in \{1, \dots, k-1\}$, consider $\lambda^* = (0^\kappa, \infty^{k-\kappa})$. Suppose $\lambda^* \in \Lambda(\omega)$ and λ^* is maximally accessible, with $\lambda_l^* = (0^{k-l+1}, \infty^{l-1})$ maximally separable at zero for θ_{k-l+1} for $l = 1, \dots, k-\kappa$. For each $l = 1, \dots, k-\kappa$, $\theta_{k-l+1} \succ_{\lambda_l^*} \theta_{k-l}$ implies that the submatrix $\Psi[\theta_{k-l+1}, \theta_{k-l}; a_1, a_M](\lambda_l^*)$ defined in Eq. (14) has a positive determinant. Therefore, there exists a $c \in \mathbb{R}_+^2$ that solves

$$\Psi[\theta_{k-l+1}, \theta_{k-l}; a_1, a_M](\lambda_l^*) \cdot c = \begin{pmatrix} 1 \\ 0 \end{pmatrix}.$$

By continuity, there exists a perturbation of c to $\tilde{c} \in \mathbb{R}_+^2$ such that

$$\Psi[\theta_{k-l+1}, \theta_{k-l}; a_1, a_M](\lambda_l^*) \cdot \tilde{c} = \begin{pmatrix} G_{k-l+1} \\ G_{k-l} \end{pmatrix},$$

where $G_{k-l+1} > 0$ and $G_{k-l} < 0$. Moreover, $\Psi[\theta_j; a_1, a_M](\lambda_l^*) \cdot \tilde{c} > 0$ for any $j > k-l+1$ and $\Psi[\theta_j; a_1, a_M](\lambda_l^*) \cdot \tilde{c} < 0$ for any $j < k-l$, where $\Psi[\theta_j; a_1, a_M](\lambda)$ is the submatrix of Eq. (13) for type θ_j and actions a_1 and a_M . Therefore, by Definition 8, λ_l^* is separable at zero for θ_{k-l+1} , since the definition holds for vector $c' \in (0, \infty)^{|\mathcal{A}|}$ with $c'_1 = c_1$, $c'_M = c_2$ and $c'_i = 0$ otherwise. The case of maximal separability at infinity is analogous, as is the proof for the other disagreement outcomes. $\square$

D.2 Mixed Learning

Consider the mixed learning outcome (λ_C^*, C) in which beliefs converge to $\lambda_C^* \in \{0, \infty\}^{|\mathcal{C}|}$ for some subset of social types $C \subset \Theta_S$ with $|C| \in \{1, \dots, k-1\}$ and beliefs do not converge for the remaining social types $N \equiv \Theta_S \setminus C$, where λ_C denotes the likelihood ratio vector λ restricted to a set of types C . Without loss of generality, maintain the convention that the

first $|C|$ types are the convergent types, i.e. $C = \{\theta_1, \dots, \theta_{|C|}\}$, and the remaining types are the non-convergent types i.e. $N = \{\theta_{|C|+1}, \dots, \theta_k\}$ (it is always possible to relabel the type space so that this holds).

For example, when $k = 3$, $((0, 0), \{\theta_1, \theta_2\})$ denotes the mixed outcome where θ_1 and θ_2 's beliefs converge to zero and θ_3 's beliefs do not converge. If $(0, 0, 0) \in \Lambda_3(\omega)$ or $(0, 0, \infty) \in \Lambda_3(\omega)$, then when $\langle \lambda_{1,t}, \lambda_{2,t} \rangle \rightarrow (0, 0)$, with positive probability the beliefs of θ_3 also converge in state ω . This is a sufficient condition to establish that $((0, 0), \{\theta_1, \theta_2\})$ almost surely does not occur in state ω . Sufficient conditions to rule out mixed outcomes in which the beliefs of two or more social types do not converge are more involved, as we also need to ensure that the neighborhood of a locally stable outcome for the non-convergent types is reached with positive probability when the beliefs of the convergent types converge. For example, to rule out the mixed outcome $(0, \theta_1)$ in which θ_1 's beliefs converge to zero and θ_2 and θ_3 have cyclical learning, in addition to $(0, 0, 0) \in \Lambda_2(\omega) \cap \Lambda_3(\omega)$, we also need to show that from a neighborhood of the other stationary beliefs with $\lambda_1 = 0$, i.e. $\lambda \in \{(0, \infty, 0), (0, 0, \infty), (0, \infty, \infty)\}$, either (i) beliefs enter a neighborhood of $(0, 0, 0)$ with positive probability or (ii) $\lambda \in \Lambda_2(\omega) \cap \Lambda_3(\omega)$. The following paragraphs formalize this idea.

We first define the concept of mixed accessibility. The concept applies to pairs of stationary beliefs in which non-convergent types whose components differ between the two belief vectors agree, which we refer to as agreement adjacent beliefs.

Definition 12 (Agreement Adjacent). *Given a set of types $N \subset \Theta_S$, distinct stationary beliefs $\lambda_N \in \{0, \infty\}^{|N|}$ and $\lambda'_N \in \{0, \infty\}^{|N|}$ are agreement adjacent if $\lambda_i = \lambda_j$ for each $\theta_i, \theta_j \in N$ such that $\lambda'_i \neq \lambda_i$ and $\lambda'_j \neq \lambda_j$.*

Trivially, two stationary belief vectors that differ in only one component are agreement adjacent. Given a mixed outcome and a stationary belief for the non-convergent types, the set of stationary beliefs that are mixed accessible from this belief depends on the local stability of this belief for each non-convergent type.

Definition 13 (Mixed Accessible ($k \geq 2$)). *Given mixed outcome (λ_C^*, C) with $N \equiv \Theta_S \setminus C$, stationary belief $\lambda'_N \in \{0, \infty\}^{|N|}$ is mixed accessible from distinct stationary belief $\lambda_N \in \{0, \infty\}^{|N|}$ in state ω if λ'_N and λ_N are agreement adjacent and $(\lambda_C^*, \lambda_N) \notin \Lambda_i(\omega)$ for some $\theta_i \in N$ such that $\lambda'_i \neq \lambda_i$.*

As we will show in the proof of [Lemma 4'](#), mixed accessibility is a sufficient condition to establish that with positive probability, the likelihood ratio process either transitions between the neighborhoods of two agreement adjacent stationary beliefs or exits a neighborhood of the mixed outcome. We next define a graph to represent which stationary beliefs are mixed accessible from other stationary beliefs.

Definition 14 (Mixed Accessible Graph ($k \geq 2$)). *Given (λ_C^*, C) with $N \equiv \Theta_S \setminus C$, define the mixed accessible graph $\mathcal{G}(\lambda_C^*, C; \omega)$ with nodes $\lambda_N \in \{0, \infty\}^{|N|}$ as follows: there is a directed edge from λ_N to λ'_N if and only if λ'_N is mixed accessible from λ_N in state ω .*

We say (λ_C^*, C) is *reducible* in state ω if $\mathcal{G}(\lambda_C^*, C; \omega)$ has no cycles. We refer to a node with no edges leaving it as a *terminal node*—in other words, a node from which no other nodes are mixed accessible. It follows from the definition of mixed accessibility that λ_N is a terminal node in state ω if and only if $(\lambda_C^*, \lambda_N) \in \cap_{\theta_i \in N} \Lambda_i(\omega)$.

We use this graph to define $\Lambda_M(\omega)$ as the set of mixed outcomes that are not reducible,

$$\Lambda_M(\omega) \equiv \{(\lambda_C^*, C) \text{ a mixed outcome} \mid (\lambda_C^*, C) \text{ is not reducible in state } \omega\}, \quad (20)$$

where a mixed outcome corresponds to $\lambda_C^* \in \{0, \infty\}^{|C|}$, $C \subset \Theta_S$ and $|C| \in \{1, \dots, k-1\}$. This definition is equivalent to Eq. (5) when $k = 2$. Using Eq. (20), the statement of Lemma 4' is identical to Lemma 4.

Lemma 4' (Unstable Mixed Outcomes ($k \geq 2$)). *Consider a learning environment that is identified at certainty and satisfies Assumptions 1 to 4. If mixed outcome $(\lambda_C^*, C) \notin \Lambda_M(\omega)$, then $\Pr(\lambda_{C,t} \rightarrow \lambda_C^* \text{ and } \lambda_{N,t} \text{ does not converge}) = 0$ in state ω , where $N \equiv \Theta_S \setminus C$.*

As in the case of $k = 2$, if a mixed learning outcome arises with positive probability, then it must be in $\Lambda_M(\omega)$. Therefore, if (λ_C^*, C) is reducible, then almost surely it does not arise. Intuitively, if (λ_C^*, C) is reducible, then when the beliefs of the convergent types are in a neighborhood of λ_C^* , almost surely either the beliefs of the convergent types leave this neighborhood or the beliefs of the non-convergent types also converge. Reducibility is relatively straightforward to verify and is always satisfied in some important cases. For instance, it holds when $\gamma_i(\omega, \lambda) < 0$ for all $\lambda \in \{0, \infty\}^k$ and $\theta_i \in \Theta_S$ (this includes environments that are close to a correctly specified environment).³

Proof of Lemma 4'. Fix state ω and consider mixed outcome (λ_C^*, C) with corresponding graph $\mathcal{G}(\lambda_C^*, C; \omega)$ and non-convergent types $N \equiv \Theta_S \setminus C$. Suppose (λ_C^*, C) is reducible, i.e. $(\lambda_C^*, C) \notin \Lambda_M(\omega)$. Let $\varepsilon \in (0, e^{-E})$, where E is defined in Eq. (11), and suppose $\lambda_1 \in \text{int}(B_\varepsilon(\lambda_C^*)) \times (0, \infty)^{|N|}$. Let $\tau \equiv \min\{t \mid \lambda_t \notin B_\varepsilon(\lambda_C^*) \times (0, \infty)^{|N|}\}$ be the first time that $\langle \lambda_t \rangle$ leaves a neighborhood of the mixed outcome. We will establish the following claim: almost surely, either (i) $\tau < \infty$ or (ii) $\langle \lambda_t \rangle$ converges for all social types. By the linearity of the likelihood ratio process, this implies the same holds whenever $\lambda_t \in \text{int}(B_\varepsilon(\lambda_C^*)) \times (0, \infty)^{|N|}$, and therefore, (λ_C^*, C) almost surely does not occur.

Step 1: Show that for any terminal node $\lambda_N \in \mathcal{G}(\lambda_C^*, C; \omega)$, when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*, \lambda_N))$, then with probability uniformly bounded away from zero, either $\langle \lambda_t \rangle \rightarrow (\lambda_C^*, \lambda_N)$ or $\tau < \infty$. Given a terminal node $\lambda_N \in \mathcal{G}(\lambda_C^*, C; \omega)$, as stated above, $(\lambda_C^*, \lambda_N) \in \cap_{\theta_i \in N} \Lambda_i(\omega)$. If $(\lambda_C^*, \lambda_N) \in \cap_{\theta_i \in C} \Lambda_i(\omega)$, then (λ_C^*, λ_N) is locally stable, so by Theorem 1, when beliefs are in $B_\varepsilon(\lambda_C^*, \lambda_N)$, then $\langle \lambda_t \rangle \rightarrow (\lambda_C^*, \lambda_N)$ with probability uniformly bounded away from zero. Otherwise, if $(\lambda_C^*, \lambda_N) \notin \cap_{\theta_i \in C} \Lambda_i(\omega)$, then there exists a $\theta_i \in C$ such that when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*, \lambda_N))$, $\langle \lambda_{i,t} \rangle$ is bounded below by a process that almost surely exits $B_\varepsilon(\lambda_C^*)$ (this also follows from the proof of Theorem 1). Therefore, $\tau < \infty$ with probability uniformly bounded away from zero. Together this implies that, starting from the ε -neighborhood of any terminal node λ_N , with probability uniformly bounded away from zero either $\tau < \infty$ or $\langle \lambda_t \rangle$ converges to (λ_C^*, λ_N) .

Step 2: Show that for any non-terminal node $\lambda_N \in \mathcal{G}(\lambda_C^*, C; \omega)$, when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*, \lambda_N))$, then with probability uniformly bounded away from zero, either $\langle \lambda_t \rangle$ enters $\cup_{\lambda'_N \in \mathcal{E}(\lambda_N)} \text{int}(B_\varepsilon(\lambda_C^*, \lambda'_N))$

³To see this, consider the graph induced by any mixed outcome (λ_C^*, C) with $N \equiv \Theta_S \setminus C$. Each node where κ non-convergent types have belief $\lambda_i = \infty$ has an edge to all agreement adjacent nodes in which $\kappa' < \kappa$ non-convergent types have belief $\lambda_i = \infty$ and does not have an edge to any other nodes. Therefore, every path terminates at node $0^{|C|}$. For any mixed outcome $(0^{|C|}, C)$, this is a convergent point. For other mixed outcomes, this is a point at which some $\theta_i \in C$'s belief eventually exits a neighborhood of λ_C^* .

or $\tau < \infty$, where $\mathcal{E}(\lambda_N)$ denotes the set of nodes that λ_N has edges to. Given a non-terminal node $\lambda_N \in \mathcal{G}(\lambda_C^*, C; \omega)$, let $U(\lambda_N) \subset N$ denote the set of types such that $(\lambda_C^*, \lambda_N) \notin \Lambda_i(\omega)$ for each $\theta_i \in U(\lambda_N)$ and $(\lambda_C^*, \lambda_N) \in \Lambda_i(\omega)$ for each $\theta_i \in N \setminus U(\lambda_N)$. As stated above, $(\lambda_C^*, \lambda_N) \notin \cap_{\theta_i \in N} \Lambda_i(\omega)$ for non-terminal nodes, so $U(\lambda_N) \neq \emptyset$.

Step 2a: We first define a space $I(\lambda_N)$ adjacent to $B_\varepsilon(\lambda_N)$ and show that when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*, \lambda_N))$, then with probability uniformly bounded away from zero, either $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(\lambda_C^*) \times I(\lambda_N))$ or $\tau < \infty$. Given a set of types $u \in \mathcal{P}(U(\lambda_N))$, where $\mathcal{P}(\cdot)$ denotes the power set, let $I_{u,i}(\lambda_N) \equiv [\varepsilon, 1/\varepsilon]$ if $\theta_i \in u$ and $I_{u,i}(\lambda_N) \equiv B_\varepsilon((\lambda_N)_i)$ if $\theta_i \in N \setminus u$. Define $I_u(\lambda_N) \equiv \prod_{\theta_i \in N} I_{u,i}(\lambda_N)$ for each $u \in \mathcal{P}(U(\lambda_N))$ and $I(\lambda_N) \equiv \cup_{u \in \mathcal{P}(U(\lambda_N)) \setminus \emptyset} I_u(\lambda_N)$. In other words, $I(\lambda_N)$ is the space in which the beliefs of subsets of types in $U(\lambda_N)$ are in $[\varepsilon, 1/\varepsilon]$ and the beliefs of the remaining non-convergent types are in the ε -neighborhood of λ_N . By the proof of [Theorem 1](#), when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*, \lambda_N))$, then with probability uniformly bounded away from zero, $\langle \lambda_{i,t} \rangle$ exits $B_\varepsilon((\lambda_C^*)_i)$ for some $\theta_i \in U(\lambda_N) \cup C$. Combined with $\varepsilon < e^{-E}$, which ensures that $\langle \lambda_{i,t} \rangle$ does not enter $B_\varepsilon(\{0, \infty\} \setminus (\lambda_N)_i)$ in the same period it exits $B_\varepsilon((\lambda_N)_i)$ for any $\theta_i \in N$, this implies that with probability uniformly bounded away from zero, either $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(\lambda_C^*) \times I(\lambda_N))$ or $\tau < \infty$.

Step 2b: We next show that when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*) \times I(\lambda_N))$, then with probability uniformly bounded away from zero, either $\langle \lambda_t \rangle$ enters $\cup_{\lambda'_N \in \mathcal{E}(\lambda_N)} \text{int}(B_\varepsilon(\lambda_C^*, \lambda'_N))$ or $\tau < \infty$. First consider $u \in \mathcal{P}(U(\lambda_N))$ such that $(\lambda_N)_i = 0$ for some $\theta_i \in u$. Suppose $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*) \times I_u(\lambda_N))$. Note $I_{u,i}(\lambda_N) = [\varepsilon, 1/\varepsilon]$ for $\theta_i \in u$ and $I_{u,i}(\lambda_N) \equiv B_\varepsilon((\lambda_N)_i)$ for $\theta_i \in N \setminus u$. Let $V_\infty(\lambda_N, u) \subset \{0, \infty\}^{|N|}$ denote the set of stationary beliefs for non-convergent types in which $\theta_i \in u$ has belief ∞ , $\theta_i \in N \setminus u$ such that $(\lambda_N)_i = \infty$ has belief ∞ , and $\theta_i \in N \setminus u$ such that $(\lambda_N)_i = 0$ has belief zero or infinity. Then there exists a finite sequence of actions a_M such that, starting from any belief in $\text{int}(B_\varepsilon(\lambda_C^*) \times I_u(\lambda_N))$, $\langle \lambda_{N,t} \rangle$ enters $\cup_{\lambda'_N \in V_\infty(\lambda_N, u)} \text{int}(B_\varepsilon(\lambda'_N))$.⁴ If $\langle \lambda_{C,t} \rangle$ exits $B_\varepsilon(\lambda_C^*)$ during this sequence, then $\tau < \infty$; otherwise, $\langle \lambda_t \rangle$ is in $\cup_{\lambda'_N \in V_\infty(\lambda_N, u)} \text{int}(B_\varepsilon(\lambda_C^*, \lambda'_N))$. Each belief $\lambda'_N \in V_\infty(\lambda_N, u)$ is agreement adjacent to λ_N , as a subset of types in N have belief 0 at λ_N and ∞ at λ'_N , and all other types in N have the same belief at λ_N and λ'_N . By definition of $V_\infty(\lambda_N, u)$, for each $\lambda'_N \in V_\infty(\lambda_N, u)$, $(\lambda_N)_i \neq (\lambda'_N)_i$ for $\theta_i \in u$ such that $(\lambda_N)_i = 0$. Further, $(\lambda_C^*, \lambda_N) \notin \Lambda_i(\omega)$ for each $\theta_i \in u$. Therefore, each $\lambda'_N \in V_\infty(\lambda_N, u)$ is mixed accessible from λ_N , which implies $V_\infty(\lambda_N, u) \subset \mathcal{E}(\lambda_N)$. This establishes that, given $u \in \mathcal{P}(U(\lambda_N))$ such that $(\lambda_N)_i = 0$ for some $\theta_i \in u$, when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*) \times I_u(\lambda_N))$, then with probability uniformly bounded away from zero, either $\langle \lambda_t \rangle$ enters $\cup_{\lambda'_N \in \mathcal{E}(\lambda_N)} \text{int}(B_\varepsilon(\lambda_C^*, \lambda'_N))$ or $\tau < \infty$. Next consider $u \in \mathcal{P}(U(\lambda_N))$ such that $(\lambda_N)_i = \infty$ for some $\theta_i \in u$. Let $V_0(\lambda_N, u)$ denote the set of stationary beliefs for non-convergent types in which $\theta_i \in u$ has belief zero, $\theta_i \in N \setminus u$ such

⁴To see this, note that when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*) \times I_u(\lambda_N))$, then $\langle \lambda_{i,t} \rangle$ is in $[\varepsilon, 1/\varepsilon]$ for $\theta_i \in u$. Therefore we can construct a finite sequence sequence of a_M actions such that following this sequence, $\langle \lambda_{i,t} \rangle$ is in $\text{int}(B_\varepsilon(\infty))$ for each $\theta_i \in u$. For $\theta_i \in N \setminus u$ such that $(\lambda_N)_i = \infty$, $I_{u,i}(\lambda_N) \equiv B_\varepsilon(\infty)$, and therefore, $\langle \lambda_{i,t} \rangle$ remains in $\text{int}(B_\varepsilon(\infty))$ following this sequence. For $\theta_i \in N \setminus u$ such that $(\lambda_N)_i = 0$, $I_{u,i}(\lambda_N) \equiv B_\varepsilon(0)$. If $\langle \lambda_{i,t} \rangle$ remains in $\text{int}(B_\varepsilon(0))$ or enters $\text{int}(B_\varepsilon(\infty))$ following this sequence for each such $\theta_i \in N \setminus u$, then we are done. Otherwise, continue repeating a_M until all $\theta_i \in N \setminus u$ with $(\lambda_N)_i = 0$ have beliefs in $\text{int}(B_\varepsilon(0) \cup B_\varepsilon(\infty))$. Given that there are a finite number of such types and, for any such type, a finite number of a_M actions will move its beliefs from $[\varepsilon, 1/\varepsilon]$ to $\text{int}(B_\varepsilon(\infty))$, this will hold following a finite number of additional a_M actions. Following these additional a_M actions, $\langle \lambda_{i,t} \rangle$ remains in $\text{int}(B_\varepsilon(\infty))$ for $\theta_i \in u$, as does $\langle \lambda_{i,t} \rangle$ for $\theta_i \in N \setminus u$ such that $(\lambda_N)_i = \infty$. Therefore, following this sequence, $\langle \lambda_{N,t} \rangle$ is in $\text{int}(B_\varepsilon(\lambda'_N))$ for some $\lambda'_N \in V_\infty(\lambda_N, u)$.

that $(\lambda_N)_i = 0$ has belief zero, and $\theta_i \in N \setminus u$ such that $(\lambda_N)_i = \infty$ has belief zero or infinity. Then substituting a_1 for a_M and $V_0(\lambda_N, u)$ for $V_\infty(\lambda_N, u)$, by similar reasoning, when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*) \times I_u(\lambda_N))$, then with probability uniformly bounded away from zero, either $\langle \lambda_t \rangle$ enters $\cup_{\lambda'_N \in \mathcal{E}(\lambda_N)} \text{int}(B_\varepsilon(\lambda_C^*, \lambda'_N))$ or $\tau < \infty$. Given that one of these two cases apply to each $u \in \mathcal{P}(U(\lambda_N))$ and $I(\lambda_N) \equiv \cup_{u \in \mathcal{P}(U(\lambda_N)) \setminus \emptyset} I_u(\lambda_N)$, this establishes that when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*) \times I(\lambda_N))$, then with probability uniformly bounded away from zero, either $\langle \lambda_t \rangle$ enters $\cup_{\lambda'_N \in \mathcal{E}(\lambda_N)} \text{int}(B_\varepsilon(\lambda_C^*, \lambda'_N))$ or $\tau < \infty$.

Step 3: Show that for any non-terminal node $\lambda_N \in \mathcal{G}(\lambda_C^*, C; \omega)$, when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*, \lambda_N))$, then with probability uniformly bounded away from zero, either $\langle \lambda_t \rangle$ enters $\cup_{\lambda'_N \in \mathcal{T}} \text{int}(B_\varepsilon(\lambda_C^*, \lambda'_N))$ or $\tau < \infty$, where $\mathcal{T}$ denotes the set of terminal nodes. Given $\Lambda_M(\omega)$ is empty, (λ_C^*, C) is reducible and therefore, $\mathcal{G}(\lambda_C^*, C; \omega)$ has no cycles. Therefore, starting from any non-terminal node $\lambda_N \in \mathcal{G}(\lambda_C^*, C; \omega)$ and iterating Step 2 a finite number of times ensures that when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*, \lambda_N))$, then with probability uniformly bounded away from zero, either $\langle \lambda_t \rangle$ enters $\cup_{\lambda'_N \in \mathcal{T}} \text{int}(B_\varepsilon(\lambda_C^*, \lambda'_N))$ or $\tau < \infty$.

Step 4: Show that when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*) \times I)$, where $I \equiv (0, \infty)^{|N|} \setminus \cup_{\lambda_N \in \{0, \infty\}^{|N|}} B_\varepsilon(\lambda_N)$, then almost surely, either $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(\lambda_C^*) \times \cup_{\lambda_N \in \{0, \infty\}^{|N|}} B_\varepsilon(\lambda_N))$ or $\tau < \infty$. Consider $u \subset N$. Similar to above, let $I_{u,i} \equiv [\varepsilon, 1/\varepsilon]$ if $\theta_i \in u$, $I_{u,i} \equiv B_\varepsilon(0) \cup B_\varepsilon(\infty)$ if $\theta_i \in N \setminus u$, and $I_u \equiv \prod_{\theta_i \in N} I_{u,i}$. Then by similar reasoning to [Footnote 4](#), there exists a finite sequence of a_1 actions such that, starting from any belief in $\text{int}(B_\varepsilon(\lambda_C^*) \times I_u)$, $\langle \lambda_{N,t} \rangle$ enters $\cup_{\lambda_N \in \{0, \infty\}^{|N|}} \text{int}(B_\varepsilon(\lambda_N))$ following this sequence. If $\langle \lambda_{C,t} \rangle$ exits $B_\varepsilon(\lambda_C^*)$ during this sequence, then $\tau < \infty$; otherwise, $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*) \times \cup_{\lambda_N \in \{0, \infty\}^{|N|}} B_\varepsilon(\lambda_N))$. Given that this sequence is finite and occurs with probability uniformly bounded away from zero across $B_\varepsilon(\lambda_C^*) \times I_u$, and such a sequence exists for each $u \subset N$, when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*) \times I)$, then almost surely $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(\lambda_C^*)) \times \cup_{\lambda_N \in \{0, \infty\}^{|N|}} \text{int}(B_\varepsilon(\lambda_N))$ or $\tau < \infty$.

Taken together, Steps 2-4 establish that when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda_C^*)) \times (0, \infty)^{|N|}$, then with probability uniformly bounded away from zero, either $\langle \lambda_t \rangle$ enters $\cup_{\lambda_N \in \mathcal{T}} \text{int}(B_\varepsilon(\lambda_C^*, \lambda_N))$ or $\tau < \infty$. It follows from [Theorem 1](#) that when $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(\lambda_C^*, \lambda_N))$ for any $\lambda_N \notin \mathcal{T}$, it almost surely exits $B_\varepsilon(\lambda_C^*, \lambda_N)$, and from Step 4 that when $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(\lambda_C^*) \times I)$, it almost surely exits $B_\varepsilon(\lambda_C^*) \times I$. Therefore, when $\lambda_1 \in \text{int}(B_\varepsilon(\lambda_C^*)) \times (0, \infty)^{|N|}$,

$$Pr(\tau < \infty \text{ or } \lambda_t \in \cup_{\lambda_N \in \mathcal{T}} B_\varepsilon(\lambda_C^*, \lambda_N) \text{ i.o.}) = 1.$$

If $\langle \lambda_t \rangle$ is in $\cup_{\lambda_N \in \mathcal{T}} B_\varepsilon(\lambda_C^*, \lambda_N)$ infinitely often, then almost surely either $\langle \lambda_t \rangle$ converges to (λ_C^*, λ_N) for some $\lambda_N \in \mathcal{T}$ or $\tau < \infty$. This establishes the claim. $\square$

D.3 Learning Characterization

To establish almost sure convergence, we define an analogous graph to [Definition 14](#) for the case in which all social types are non-convergent types.

Definition 15 (Accessible Graph ($k \geq 2$)). *Define the accessible graph $\mathcal{G}(\omega)$ with nodes $\lambda \in \{0, \infty\}^k$ as follows: there is a directed edge from λ to λ' if and only if λ' is mixed accessible from λ in state ω .*

It follows from the definition of mixed accessibility that λ is a terminal node if and only if $\lambda \in \Lambda(\omega)$. Given the definitions of $\Lambda_M(\omega)$ and maximal accessibility for $k > 2$, the statement of [Lemma 7'](#) mirrors [Lemma 7](#).

Lemma 7' (Belief Convergence ($k > 2$)). Consider a learning environment that is identified at certainty and satisfies *Assumptions 1 to 4*. If $\Lambda(\omega)$ contains an agreement outcome or maximally accessible disagreement outcome, $\Lambda_M(\omega) = \emptyset$ and $\mathcal{G}(\omega)$ has no cycles, then for any initial belief $\lambda_1 \in (0, \infty)^k$, there exists a random variable λ_∞ with $\text{supp}(\lambda_\infty) = \Lambda(\omega)$ such that $\lambda_t \rightarrow \lambda_\infty$ almost surely in state ω .⁵

Proof. Fix state ω . Suppose $\Lambda_M(\omega) = \emptyset$ and $\mathcal{G}(\omega)$ has no cycles. Let $\varepsilon \in (0, e^{-E})$, where E is defined in Eq. (11). It follows from the definition of mixed accessibility that $\lambda \in \mathcal{G}(\omega)$ is a terminal node if and only if $\lambda \in \Lambda(\omega)$. Given a terminal node $\lambda \in \Lambda(\omega)$, by Theorem 1, when $\langle \lambda_t \rangle$ is in $B_\varepsilon(\lambda)$, then $\langle \lambda_t \rangle \rightarrow \lambda$ with probability uniformly bounded away from zero. By analogous reasoning to Step 2 in the proof of Lemma 4', for any non-terminal node $\lambda \in \mathcal{G}(\omega)$, when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda))$, then with probability uniformly bounded away from zero, $\langle \lambda_t \rangle$ enters $\cup_{\lambda' \in \mathcal{E}(\lambda)} \text{int}(B_\varepsilon(\lambda'))$, where $\mathcal{E}(\lambda)$ denotes the set of nodes that λ has edges to. Given $\mathcal{G}(\omega)$ has no cycles, starting from any non-terminal node $\lambda \in \mathcal{G}(\omega)$, when $\langle \lambda_t \rangle$ is in $\text{int}(B_\varepsilon(\lambda))$, then with probability uniformly bounded away from zero, $\langle \lambda_t \rangle$ enters $\cup_{\lambda' \in \Lambda(\omega)} \text{int}(B_\varepsilon(\lambda'))$. When $\langle \lambda_t \rangle$ is in $I \equiv (0, \infty)^k \setminus \cup_{\lambda \in \{0, \infty\}^k} B_\varepsilon(\lambda)$, then by similar reasoning to Step 4 in the proof of Lemma 4', almost surely $\langle \lambda_t \rangle$ enters $\cup_{\lambda \in \{0, \infty\}^k} \text{int}(B_\varepsilon(\lambda))$. Taken together, this establishes that when $\langle \lambda_t \rangle$ is in $(0, \infty)^k$, then with probability uniformly bounded away from zero, $\langle \lambda_t \rangle$ enters $\cup_{\lambda \in \Lambda(\omega)} \text{int}(B_\varepsilon(\lambda))$. Further, it follows from Theorem 1 that when $\langle \lambda_t \rangle$ enters $\text{int}(B_\varepsilon(\lambda))$ for any $\lambda \in \{0, \infty\}^k \setminus \Lambda(\omega)$, it almost surely exits $B_\varepsilon(\lambda)$. Given that $\langle \lambda_t \rangle$ also almost surely exits I , it follows that starting from any $\lambda_1 \in (0, \infty)^k$, $\Pr(\lambda_t \in \cup_{\lambda \in \Lambda(\omega)} B_\varepsilon(\lambda) \text{ i.o.}) = 1$. Therefore, almost surely $\langle \lambda_t \rangle$ converges to some $\lambda \in \Lambda(\omega)$. $\square$

The statement of the learning characterization for $k > 2$ is analogous to Theorem 4, using the generalized definitions of maximal accessibility (Definition 7') and $\Lambda_M(\omega)$ (Eq. (20)).

Theorem 4' (Learning Characterization ($k > 2$)). Consider a learning environment that is identified at certainty and satisfies *Assumptions 1 to 4*. When the state is L :

- (i) **Correct learning** occurs with positive probability if and only if $0^k \in \Lambda(L)$.
- (ii) **Incorrect learning** occurs with positive probability if and only if $\infty^k \in \Lambda(L)$.
- (iii) **Entrenched Disagreement** occurs with positive probability if $\Lambda(L)$ contains a maximally accessible disagreement outcome and almost surely does not occur if $\Lambda(L)$ contains no disagreement outcome. Each maximally accessible disagreement outcome in $\Lambda(L)$ occurs with positive probability.
- (iv) **Cyclical Learning** occurs almost surely for all social types if $\Lambda(L) = \emptyset$ and $\Lambda_M(L) = \emptyset$, occurs almost surely for at least one social type if $\Lambda(L) = \emptyset$, and almost surely does not occur if $\Lambda(L)$ contains an agreement outcome or maximally accessible disagreement outcome, $\Lambda_M(L) = \emptyset$, and $\mathcal{G}(L)$ has no cycles.

An analogous result holds in state R .

⁵An alternative condition to $\mathcal{G}(\omega)$ has no cycles is there exists a $\theta_i \in \Theta_S$ such that $\sup_{\lambda_{-i} \in [0, \infty]^{k-1}} \gamma_i(\lambda_i = 0, \lambda_{-i}) < 0$ or $\inf_{\lambda_{-i} \in [0, \infty]^{k-1}} \gamma_i(\lambda_i = \infty, \lambda_{-i}) > 0$. This follows from the observation in Lemma 4' that either beliefs converge or visit each mixed outcome with $|C| = 1$ infinitely often. If the latter occurs with positive probability, then almost surely $\langle \lambda_{i,t} \rangle \rightarrow \lambda_\infty \subset \{0, \infty\}$, which contradicts $\Lambda_M(\omega) = \emptyset$.

The proof mirrors the case of two social types: it directly follows from [Lemma 3](#), [Theorems 1 and 2](#), [Theorem 3'](#), and [Lemmas 4'](#) and [7'](#).

E Belief-Dependent Models of Inference

In this section, we extend our framework to allow a type's model of inference to vary with its belief about the state. We show that with this extension, our framework nests [Rabin and Schrag \(1999\)](#) and [Epstein et al. \(2010\)](#).

E.1 Framework

Modify a type's model of inference as follows. Given likelihood ratio $\lambda \in [0, \infty]^k$, type θ_i has subjective private signal distribution $\hat{F}_i^\omega(s; \lambda)$ in state ω and subjective type distribution $\hat{\pi}_i(\theta; \lambda)$. An agent uses likelihood ratio λ_t to interpret signal $\tilde{s}_t$ or action $\tilde{a}_t$ at time t . Maintain the assumption from [Section 2](#) that $\hat{F}_i^L(\cdot, \lambda)$ and $\hat{F}_i^R(\cdot, \lambda)$ are mutually absolutely continuous with full support on $\mathcal{S}$ for each $\lambda \in [0, \infty]^k$. Further, social and autarkic types believe that the signal is uniformly informative. When signals are aligned, this can be written as follows: for all $s \in [0, 1]$, either $\hat{F}_i^L(s; \lambda) = \hat{F}_i^R(s; \lambda) = 0$ for all $\lambda \in [0, \infty]^k$, $\hat{F}_i^L(s; \lambda) = \hat{F}_i^R(s; \lambda) = 1$ for all $\lambda \in [0, \infty]^k$, or $\inf_{\lambda \in [0, \infty]^k} \hat{F}_i^L(s; \lambda) - \hat{F}_i^R(s; \lambda) > 0$, with the final case holding for some $s \in [0, 1]$.

As in [Section 2](#), we focus on settings in which social types believe that actions are informative. We need to modify [Assumption 3](#) so that this holds uniformly across $[0, \infty]^k$.

Assumption 3' (Informative Actions). *For actions $a \in \{a_1, a_M\}$, there exists an autarkic type $\theta_j \in \Theta_A$ with $\pi(\theta_j) > 0$ that plays a with probability uniformly bounded away from zero across $[0, \infty]^k$, and each social type $\theta_i \in \Theta_S$ believes that such an autarkic type exists with probability uniformly bounded away from zero, $\inf_{\lambda \in [0, \infty]^k} \hat{\pi}_i(\theta_j; \lambda) > 0$.*

For technical reasons, we also make the following continuity assumption.

Assumption 5 (Continuity). *For each $\theta_i \in \Theta$, the mapping $\lambda \mapsto (\hat{F}_i^L, \hat{F}_i^R, \hat{\pi}_i)$ is continuous under the total variation norm except at at most a finite number of interior likelihood ratios $\lambda \in (0, \infty)^k$ and $\lambda \rightarrow 1/(1 + d\hat{F}_i^L/d\hat{F}_i^R(s; \lambda))$ is continuous at $\lambda \in \{0, \infty\}^k$.*

Substituting [Assumption 3'](#) for [Assumption 3](#) and adding [Assumption 5](#), the modified version of [Eq. \(1\)](#) is:

$$\hat{\psi}_i(a_m | \omega, \lambda) \equiv \sum_{j=1}^n \hat{\pi}_i(\theta_j; \lambda) (\hat{F}_i^\omega(\bar{s}_{j,m}(\lambda_j; \lambda); \lambda) - \hat{F}_i^\omega(\bar{s}_{j,m-1}(\lambda_j; \lambda); \lambda)), \quad (21)$$

where $\bar{s}_{j,m}(\lambda_j; \lambda)$ denotes the signal cutoff for θ_j when it has belief λ_j and social types have belief λ . Note $\bar{s}_{j,m}$ depends on $(\hat{F}_j^L, \hat{F}_j^R)$, and hence, when these distributions depend on λ , so does $\bar{s}_{j,m}$. The proof of [Lemma 2](#) continues to hold for [Eq. \(21\)](#) with minor modifications.⁶ [Theorems 1 to 6](#) follow.

⁶Aside from minor changes to notion and a straightforward application of the continuity assumed in [Assumption 5](#), there are two main changes. To establish the uniform bound for $a \in \{a_1, a_M\}$ and bounded informativeness for $a \in \mathcal{A}$, it is necessary to account for the subjective type and signal distributions' dependence on λ . Let $\theta_j \in \Theta_A$ be an autarkic type that $\theta_i \in \Theta_S$ believes satisfies [Assumption 3'](#) for action a_1 and

E.2 Nesting Under- and Overreaction in Epstein et al. (2010)

Epstein et al. (2010) consider an individual learning model where an agent under- or overreacts to signals. They parameterize this bias with the following updating rule: an agent with prior $p \in [0, 1]$ who observes signal realization $s \in \mathcal{S}$ updates her posterior to

$$Pr(\omega = R|s, p) = (1 - \alpha) \underbrace{\left(\frac{ps}{ps + (1-p)(1-s)} \right)}_{\text{correct posterior}} + \alpha p \quad (22)$$

for some $\alpha \leq 1$. Underreaction corresponds to $\alpha > 0$, overreaction to corresponds to $\alpha < 0$, and the correctly specified model corresponds to $\alpha = 0$. This parametric form of under- and overreaction can be represented in the individual learning version of our extended framework as follows. Eq. (22) uniquely maps to a type in our framework that forms subjective posterior

$$\hat{s}(s, \lambda) = \frac{(1 - \alpha) \left(\frac{s}{s\lambda + (1-s)} \right) + \alpha \left(\frac{1}{1+\lambda} \right)}{1 + (1 - \alpha) \left(\frac{(1-\lambda)s}{s\lambda + (1-s)} \right) + \alpha \left(\frac{1-\lambda}{1+\lambda} \right)} \quad (23)$$

following signal realization $s \in \mathcal{S}$ when it has belief $\lambda \in (0, \infty)$. Eq. (22) does not map into a unique $\hat{s}(s, \lambda)$ at $\lambda \in \{0, \infty\}$, since the prior and the posterior are the same regardless of the signal realization. Since our learning characterization utilizes the limit of $\hat{s}(s, \lambda)$ as $\lambda \rightarrow \{0, \infty\}$, we need to specify how the signal is interpreted at certainty to close the model. At $\lambda = 0$, we use Eq. (23) evaluated at $\lambda = 0$. Eq. (23) is not well-defined at $\lambda = \infty$, so we define

$$\hat{s}(s, \infty) \equiv \lim_{\lambda \rightarrow \infty} \hat{s}(s, \lambda) = \frac{s}{(1 - \alpha)(1 - s) + (1 + \alpha)s}.$$

This is the unique subjective posterior that satisfies the continuity property required by Lemma 2.⁷ This set-up satisfies the properties in Lemma 2, so our learning characterization applies.

$\bar{s}_{j,1}^* \equiv \inf_{\lambda \in [0, \infty]^k} \bar{s}_{j,1} \left(\frac{p_0}{1-p_0}; \lambda \right)$. Then the analogue of Eq. (6) is:

$$\begin{aligned} \frac{\hat{\psi}_i(a_1|R, \lambda)}{\hat{\psi}_i(a_1|L, \lambda)} &\leq \frac{\hat{\pi}_i(\theta_j; \lambda) \hat{F}_i^R(\bar{s}_{j,1}^*; \lambda) + \hat{\pi}_i(\Theta_S \cup \Theta_A \setminus \{\theta_j\}; \lambda)}{\hat{\pi}_i(\theta_j; \lambda) \hat{F}_i^L(\bar{s}_{j,1}^*; \lambda) + \hat{\pi}_i(\Theta_S \cup \Theta_A \setminus \{\theta_j\}; \lambda)} \\ &\leq \sup_{\lambda \in [0, \infty]^k} \frac{(\inf_{\lambda' \in [0, \infty]^k} \hat{\pi}_i(\theta_j; \lambda')) \hat{F}_i^R(\bar{s}_{j,1}^*; \lambda) + 1 - \inf_{\lambda' \in [0, \infty]^k} \hat{\pi}_i(\theta_j; \lambda')}{(\inf_{\lambda' \in [0, \infty]^k} \hat{\pi}_i(\theta_j; \lambda')) \hat{F}_i^L(\bar{s}_{j,1}^*; \lambda) + 1 - \inf_{\lambda' \in [0, \infty]^k} \hat{\pi}_i(\theta_j; \lambda')} < 1, \end{aligned}$$

where the last line follows from Assumption 3', which ensures that $\inf_{\lambda \in [0, \infty]^k} \hat{\pi}_i(\theta_j; \lambda) > 0$ and the uniform informativeness of the subjective signal distributions, which ensures that $\inf_{\lambda \in [0, \infty]^k} (\hat{F}_i^L(\bar{s}_{j,1}^*; \lambda) - \hat{F}_i^R(\bar{s}_{j,1}^*; \lambda)) > 0$. Similar logic establishes that $\hat{\psi}_i(a|\omega, \lambda)$ is uniformly bounded away from 0 for all $\lambda \in [0, \infty]^k$, $a \in \mathcal{A}$ and $\omega \in \{L, R\}$, and therefore, a is boundedly informative.

⁷In an individual learning setting, any pair of subjective signal distributions that induce the same $\hat{s}$ must satisfy $\frac{\hat{\psi}(s|R, \lambda)}{\hat{\psi}(s|L, \lambda)} = \frac{\hat{s}(s, \lambda)}{1 - \hat{s}(s, \lambda)}$, so $\hat{s}$ determines the properties required by Lemma 2. A consequence of this is that any misspecified distribution that rationalizes $\hat{s}$ will lead to the same behavior. In Bohren and Hauser (2021b) we show that there exist subjective distributions $\hat{F}^L$ and $\hat{F}^R$ that rationalize this $\hat{s}$ and satisfy Assumption 1, Assumption 3' and Assumption 5.

E.3 Nesting Confirmation Bias in Rabin and Schrag (1999)

Rabin and Schrag (1999) consider an individual learning model where an agent exhibits confirmation bias. The agent observes a signal that takes one of two possible values, s_L or s_R , where s_ω is more likely in state ω than state ω' . Confirmation bias takes the following form: if the agent observes s_ω when she believes ω' is more likely, then with probability $q \in (0, 1)$ she misinterprets the signal realization as $s_{\omega'}$. To represent this model in the individual learning version of our extended framework, we make one additional change to allow multiple signal realizations to induce the same posterior belief. This allows $\hat{s}$ to map two signal realizations that induce the same true posterior to different subjective posteriors. Given this minor extension, this form of confirmation bias can be represented as follows. Suppose $\mathcal{S} = \{l_1, l_2, r_1, r_2\}$. Assume $Pr(l_1 \text{ or } l_2 | \omega = L) = Pr(r_1 \text{ or } r_2 | \omega = R) = s > 1/2$, conditional on observing l_1 or l_2 , l_2 is realized with probability q , and similarly for r_2 . Signal realizations l_1 and l_2 induce the same true posterior, as do r_1 and r_2 . When $\lambda > 1$, the agent interprets the signal as if $\hat{\psi}(l_1 | L, \lambda) = s$, $\hat{\psi}(l_1 | R, \lambda) = 1 - s$, $\hat{\psi}(l_2 | L, \lambda) = \hat{\psi}(r_1 | L, \lambda) = \hat{\psi}(r_2 | L, \lambda) = (1 - s)/3$ and $\hat{\psi}(l_2 | R, \lambda) = \hat{\psi}(r_1 | R, \lambda) = \hat{\psi}(r_2 | R, \lambda) = s/3$. Similarly if $\lambda \leq 1$, the agent interprets the signal as if $\hat{\psi}(r_1 | R, \lambda) = s$, $\hat{\psi}(r_1 | L, \lambda) = 1 - s$, $\hat{\psi}(l_1 | L, \lambda) = \hat{\psi}(l_2 | L, \lambda) = \hat{\psi}(r_2 | L, \lambda) = s/3$ and $\hat{\psi}(l_1 | R, \lambda) = \hat{\psi}(l_2 | R, \lambda) = \hat{\psi}(r_2 | R, \lambda) = (1 - s)/3$. This set-up satisfies the properties in Lemma 2, so our learning characterization applies.