

Mislearning from censored data: The gambler's fallacy and other correlational mistakes in optimal-stopping problems

KEVIN HE

Department of Economics, University of Pennsylvania

I study endogenous learning dynamics for people who misperceive intertemporal correlations in random sequences. Biased agents face an optimal-stopping problem. They are uncertain about the underlying distribution and learn its parameters from predecessors. Agents stop when early draws are “good enough,” so predecessors’ experiences contain negative streaks but not positive streaks. When agents wrongly expect systematic reversals (the “gambler’s fallacy”), they understate the likelihood of consecutive below-average draws, converge to over-pessimistic beliefs about the distribution’s mean, and stop too early. Agents uncertain about the distribution’s variance overestimate it to an extent that depends on predecessors’ stopping thresholds. I also analyze how other misperceptions of intertemporal correlation interact with endogenous data censoring.

KEYWORDS. Misspecified learning, gambler’s fallacy, Berk–Nash equilibrium, endogenous data censoring, fictitious variation.

JEL CLASSIFICATION. D83, D91.

1. INTRODUCTION

When a fair coin lands on tails three times in a row, many people wrongly expect the same coin to have an increased chance of landing on heads on the next toss to “balance things out.” This mistaken belief stems from a widespread statistical bias called the *gambler’s fallacy*, where people expect too much reversal from sequential realizations of independent random events. Studies have documented the gambler’s fallacy in settings where it is strictly costly, such as lotteries with parimutuel payouts (Terrell (1994), Suetens, Galbo-Jørgensen, and Tyran (2016)) and incentivized lab experiments (Benjamin, Moore, and Rabin (2017)). The same bias also affects experienced decision-makers in high-stakes environments, including immigration judges

Kevin He: hesichao@gmail.com

I am indebted to Drew Fudenberg, Matthew Rabin, Tomasz Strzalecki, and Ben Golub for their guidance and support. I thank the anonymous referees, Isaiah Andrews, Ruiqing Cao, In-Koo Cho, Martin Cripps, Krishna Dasaratha, Jetlir Duraj, Ben Enke, Ignacio Esponda, Jiacheng Feng, Mira Frick, Tristan Gagnon-Bartsch, Ashvin Gandhi, Oliver Hart, Johannes Hörner, Alice Hsiaw, Ryota Iijima, Yuhta Ishii, Lawrence Jin, Yizhou Jin, Michihiro Kandori, Max Kasy, Shengwu Li, Jonathan Libgober, Matthew Lilley, George Mailath, Eric Maskin, Weicheng Min, Xiaosheng Mu, Andy Newman, Harry Pei, Joshua Schwartzstein, Roberto Serrano, Philipp Strack, Elie Tamer, Omer Tamuz, Michael Thaler, Linh T. Tô, Maria Voronina, Yuichi Yamamoto, and seminar participants for their insightful comments. I thank the California Institute of Technology for hospitality where some of the work on this paper was completed.

(Chen, Moskowitz, and Shue (2016)) and Masters in Business Administration (MBA) admissions interviewers (Simonsohn and Gino (2013)).

The gambler's fallacy affects people's behavior and beliefs in optimal-stopping problems, an important class of economic environments where agents act on sequential signal realizations. For instance, Mueller, Spinnewijn, and Topa (2021) use survey data to document beliefs consistent with the gambler's fallacy in job search, finding that job seekers' perceived probability of becoming employed within the next few months *increases* over the course of the unemployment spell. In settings like this, how does the bias affect society's long-run beliefs about the economic fundamentals (e.g., the labor market conditions) and how does it influence agents' behavior? These questions are challenging because the biased agents do not passively observe an exogenous data stream, but take stopping actions that censor the observation of future signal realizations. The stopping decisions, in turn, depend on the agents' (possibly mistaken) beliefs about the fundamentals.

In this paper, I study novel implications of the gambler's fallacy and other correlational mistakes in optimal-stopping problems when a society of biased agents learn about the underlying distributions. Agents take turns playing the same stage game: an optimal-stopping problem with draws in different periods generated from fixed but unknown distributions. Agents learn about the means of the distributions from experience, but start with a dogmatic and wrong belief about the correlation between the draws. For instance, when the draws are objectively independent but agents expect the draws to exhibit reversals conditional on the means, they suffer from the gambler's fallacy. I show the non-self-confirming steady state of misspecified Bayesian learning in this environment involves distorted beliefs about the marginal distributions and suboptimal-stopping behavior, and the directions of these errors depend on details of the correlational mistake. I derive further results about how changes in the stage game affect long-run learning outcomes and how additional uncertainty about the variance of the distributions interacts with stopping incentives.

To illustrate the main mechanism behind these results, consider as a running example human resources (HR) managers who suffer from the gambler's fallacy. Each manager sequentially interviews candidates for a single job opening and exaggerates how unlikely it is to get consecutive above-average or consecutive below-average applicants (relative to the labor pool mean). This error stems from the same psychology that leads people to exaggerate how unlikely it is to get consecutive heads or consecutive tails when tossing a fair coin. Evidence from MBA admissions suggests this bias can have a sizable effect on sequential interviews: following applicants who are 1 standard deviation worse than usual, interviewers expect the next candidate to exceed average quality by the equivalent of 2 years of work experience (Simonsohn and Gino (2013)).

Suppose the managers are initially uncertain about the labor pool quality and collectively learn about this fundamental over time. Every manager is responsible for hiring in a different year. Each junior manager consults with senior managers and adopts their beliefs about the labor pool based on their recruiting experience for similar positions in the past. The junior manager then implements a stopping strategy for her own recruiting problem, updates her belief at the end of the hiring season, and shares this

new belief with her successors.¹ How does the gambler's fallacy influence the managers' beliefs and behavior in the long run?

In this example, agents tend to stop when early draws are deemed “good enough,” causing an asymmetric truncation of experience. When a manager discovers a sufficiently strong candidate early in the hiring cycle, she stops her recruitment efforts and does not observe what additional candidates would have been found for the same job opening with a longer search. This endogenous *censoring effect* on histories interacts with the gambler's fallacy bias and generates pessimistic inference about the labor pool. Managers continue searching only when their early candidates are below average. They misinterpret subsequent above-average candidates as the expected positive reversal after bad initial outcomes, not as strong signals about the labor pool. On the other hand, they are surprised by subsequent below-average candidates since their bias leads them to understate the likelihood of bad streaks, misreading consecutive bad draws as very strong negative signals about the pool. That is, after bad early draws, managers under-infer from subsequent good draws but over-infer from subsequent bad draws. On average, they communicate an overpessimistic impression of the labor pool to future junior managers. This pessimism informs the junior managers' stopping strategy, and affects the kind of censored history they observe and the new beliefs they pass down to their own successors.

The key mechanism behind my results is the *interaction* between psychological bias and data censoring in stopping problems. Neither is dispensable. Agents who do not suffer from correlational mistakes learn the fundamentals correctly even from censored histories. Conversely, in an environment without censoring where agents observe ex post what would have been drawn in each period of the optimal-stopping problem, even biased agents learn the fundamentals correctly. In particular, the gambler's fallacy is a “symmetric” bias; the “asymmetric” learning outcome of over-pessimism only obtains when the bias interacts with an (endogenous) asymmetric censoring mechanism that tends to produce data containing negative streaks but not positive streaks. More broadly, the selective censoring of sequential signals represents a natural source of data endogeneity whose impact on different biases remains understudied.

The misinference mechanism central to this paper implies novel comparative statics predictions about how the economic environment affects learning outcomes under the gambler's fallacy. Returning to [Mueller, Spinnewijn, and Topa's \(2021\)](#) context of job seekers, my results suggest that government policies subsidizing longer search, such as extended unemployment insurance, help mitigate belief distortions for job seekers who commit the gambler's fallacy. This is because such policies lead agents to use higher acceptance thresholds and generate less censored histories, which in turn induce less pessimistic beliefs for their successors. Comparative statics of this sort are unique to a setting where biased agents learn from endogenously censored histories: changing the stage game has no effect on the long-run learning outcomes if data are exogenous or if agents are correctly specified.

¹This environment where managers pass down their beliefs is equivalent to biased managers updating their beliefs using all past managers' hiring experience.

Finally, I extend the analysis for the case of the gambler's fallacy by considering uncertainty about both the means and variances of the distributions. In this joint estimation, agents misinfer means by the same amounts as in the baseline model and exaggerate variances. The idea is that agents attribute streaks of good or bad draws to "noise." The degree of belief in this *fictitious variation* both depends on the severity of history censoring (as the amount of noise inferred depends on the kind of data) and influences the agents' stopping strategy (as higher variance encourages continuing in search problems due to option value). To illustrate how this belief in fictitious variation interacts with endogenous learning, I show that a society where agents are uncertain about the variances end up with a less distorted long-run belief about the means than another society where agents know the correct variances. This is despite the fact that agents in both societies would make the same (mis)inference about the means when given the same data.

The rest of the paper is organized as follows. Section 2 presents the model and discusses the modeling assumptions. The model is general enough to capture various misperceptions of intertemporal correlation, with the gambler's fallacy as a special case. Section 3 analyzes the steady state of learning and contains the main results of the paper. Section 4 proves the convergence of misspecified learning dynamics to the steady state. Section 5 discusses related theoretical literature. Section 6 concludes.

2. MODEL

2.1 *The objective environment*

The stage game is a two-period optimal-stopping problem. In the first period, the agent draws $x_1 \in \mathbb{R}$ and decides whether to stop. If she stops, her payoff is $u_1(x_1) = x_1$ and the stage game ends. If she continues, she incurs a cost $\kappa \in \mathbb{R}$, enters the second period, and then draws $x_2 \in \mathbb{R}$. (This κ may also be negative, a subsidy for continuing.) There is probability $0 \leq q < 1$ that the first draw can be recalled in the second period and the agent can pick the best of the two draws, but with complementary probability the first draw is no longer available. So the agent's expected payoff from continuing, conditional on the draws, is $u_2(x_1, x_2) = q \cdot \max(x_1, x_2) + (1 - q)x_2 - \kappa$. Both q and κ are known parameters.

This stage game fits a number of economic situations:

- Many industries have an annual hiring cycle. Consider a firm in such an industry and an HR manager who must fill a job opening during this year's cycle. In the early phase of the hiring cycle, she finds a candidate with quality x_1 . She must decide between hiring this candidate immediately or waiting. Waiting lets her continue searching in the late phase of the cycle, but carries the risk that the early candidate accepts an offer from a different firm in the interim.
- A homeowner lists his house for sale and receives an offer in each period. The homeowner must decide whether to accept the first offer he gets and take his house off the market or to wait for the second offer, incurring a waiting cost and risking the first buyer leaving the market.

- An unemployed worker searches for jobs. While unemployed, she receives a job offer in each period and decides whether to continue her job search. Once she becomes employed, she stops searching and no longer receives further offers.

The draws x_1 and x_2 are the realizations of two possibly correlated Gaussian random variables X_1 and X_2 , with unconditional means $\mu_1^\bullet, \mu_2^\bullet \in \mathbb{R}$. We have $X_1 = \mu_1^\bullet + \epsilon_1$ and $X_2 = \mu_2^\bullet + \epsilon_2$, where $\epsilon_1 \sim \mathcal{N}(0, \sigma^2)$ and $(\epsilon_2 | \epsilon_1) \sim \mathcal{N}(-r\epsilon_1, \sigma^2)$ for some fixed value of $r \in \mathbb{R}$. The parameters $\mu_1^\bullet, \mu_2^\bullet \in \mathbb{R}$ are the *true fundamentals* that stand for the average qualities of the two pools in the two periods. (In general we may have $\mu_1^\bullet \neq \mu_2^\bullet$. For instance, this might happen due to dynamic adverse selection in the labor pool over time in the example of the HR manager.) The ϵ_1 and ϵ_2 terms represent the idiosyncratic factors that determine how the agent's actual draws deviate from the average qualities of the respective pools, with r the *true reversal parameter*. When $r > 0$, the idiosyncratic factors that lead to an unusually good first draw relative to the early pool quality also portend a below-average second draw. (Such reversals may happen, for instance, if the agent is exhausting a small pool.) Note that X_1 and X_2 are independent when $r = 0$, negatively correlated when $r > 0$, and positively correlated when $r < 0$.

2.2 Gambler's fallacy and other correlational mistakes

I introduce a general model of misperceptions of intertemporal correlation, with the gambler's fallacy as a special case. Section 3 will both analyze how different kinds of correlational mistakes interact with endogenous data censoring and present more in-depth results that focus on the gambler's fallacy.

Agents are uncertain about both the fundamentals and the reversal parameter. They believe that if the average qualities of the pools are $\mu_1, \mu_2 \in \mathbb{R}$, then the draws are generated by $X_1 = \mu_1 + \epsilon_1$ and $X_2 = \mu_2 + \epsilon_2$ with $\epsilon_1 \sim \mathcal{N}(0, \sigma^2)$ and $(\epsilon_2 | \epsilon_1) \sim \mathcal{N}(-\gamma\epsilon_1, \sigma^2)$ for some unknown $\gamma \in [\gamma_l, \gamma_h]$. If $0 = r < \gamma_l$, then the agents suffer from the gambler's fallacy. This may represent a superstitious belief in an environment where the two draws are objectively independent that if someone gets lucky on the first draw, then bad luck is "due" to befall them in the near future. More generally, when $r < \gamma_l$ (but r may not be 0), agents exaggerate the amount of reversal in the idiosyncratic factors across the draws. On the other hand, we may also have $r > \gamma_h$, in which case agents dogmatically underestimate the amount of reversal. This might be called a form of "hot-hand fallacy," where following a "lucky" first draw agents systematically overestimate the chance of another good draw (and symmetrically for bad draws).²

Denote by $\phi(\cdot | \mu)$ the Gaussian density with mean μ and variance σ^2 , and let $\Psi(\mu_1, \mu_2; \gamma)$ refer to the joint distribution $X_1 = \mu_1 + \epsilon_1$ and $X_2 = \mu_2 + \epsilon_2$ with $\epsilon_1 \sim \phi(\cdot | 0)$ and $(\epsilon_2 | \epsilon_1) \sim \phi(\cdot | -\gamma\epsilon_1)$. Agents believe the joint distribution of (X_1, X_2) is described by one of the *feasible models*, $\{\Psi(\mu_1, \mu_2; \gamma) : (\mu_1, \mu_2) \in \mathbb{R}^2, \gamma \in [\gamma_l, \gamma_h]\}$. If $r \notin [\gamma_l, \gamma_h]$, then the set of feasible models excludes the true model, $\Psi^\bullet := \Psi(\mu_1^\bullet, \mu_2^\bullet; r)$,

²Rabin and Vayanos (2010) propose a different mechanism for the hot-hand fallacy: agents expect reversals (not streaks) conditional on the fundamentals, but misinfer fundamentals. This also leads agents to predict that streaks will continue.

so Bayesian updating within the class of feasible models amounts to misspecified learning. I use misspecification as a tool to represent and study the gambler's fallacy and other correlational mistakes.

Throughout, I maintain the assumption that $r, \gamma_l, \gamma_h \neq -1$. It turns out that for the model $\Psi(\mu_1, \mu_2; -1)$ with any μ_1, μ_2 , all stopping strategies are optimal. So I rule out this knife-edge case by assuming that neither the true reversal parameter nor one of the end points of $[\gamma_l, \gamma_h]$ is exactly equal to -1 . I still allow the case that the interval of subjectively feasible reversal parameters contains -1 in its interior. Finally, denote $\gamma_n := \arg \min_{\gamma \in [\gamma_l, \gamma_h]} |\gamma - r|$ as the nearest point in the interval $[\gamma_l, \gamma_h]$ to r . Note that if $r \in [\gamma_l, \gamma_h]$, then the nearest point is $\gamma_n = r$ itself. Otherwise, γ_n is one of the end points, γ_l or γ_h .

2.3 The steady state

Suppose a sequence of agents arrive one per round ($t = 1, 2, 3, \dots$) and take turns playing the stage game. All agents have the same set of reversal parameters $[\gamma_l, \gamma_h]$ that they find plausible. They face the same but unknown objective pool qualities $(\mu_1^\bullet, \mu_2^\bullet)$ and true reversal parameter r . At the end of each round t , the t th agent updates her belief about qualities and about the reversal parameter using her experience, then communicates her updated belief to her successor. The successor acts based on the inherited belief, then passes down an updated belief at the end of the round to his own successor, and so forth. I now define the steady state of this learning system.

Roughly speaking, a *steady state* of the system consists of a strategy $S^\infty : \mathbb{R} \rightarrow \{\text{Stop}, \text{Continue}\}$ that maps the realization of the first draw $X_1 = x_1$ into a stopping decision, and point-mass beliefs about the pool qualities and the reversal parameter, $(\mu_1^\infty, \mu_2^\infty, \gamma^\infty) \in \mathbb{R}^2 \times [\gamma_l, \gamma_h]$, so that (i) agents find it optimal to follow strategy S^∞ given beliefs $(\mu_1^\infty, \mu_2^\infty, \gamma^\infty)$, and (ii) $(\mu_1^\infty, \mu_2^\infty, \gamma^\infty)$ are the “best-fitting” beliefs about the pool qualities and the reversal parameter given data generated from the strategy S^∞ . The steady state corresponds to [Esponda and Pouzo's \(2016\)](#) Berk–Nash equilibrium adapted to the current setting.

To make precise the meaning of “best-fitting” beliefs for misspecified learners, the *history* of the stage game is an element $h \in \mathbb{H} := \mathbb{R} \times (\mathbb{R} \cup \{\emptyset\})$. If an agent decides to stop after $X_1 = x_1$, her history is $(x_1, \emptyset)$. If an agent continues after $X_1 = x_1$ and gets a second draw $X_2 = x_2$, her history is (x_1, x_2) . The symbol $\emptyset$ is a *censoring indicator*, emphasizing if the agent stops, then the counterfactual second draw that she would have found had she continued remains unobserved.

Consider the strategy S and the parameters (μ_1, μ_2, γ) . The agent's subjective likelihood of the history $h = (x_1, x_2)$ with $S(x_1) = \text{Continue}$ is $\phi(x_1 | \mu_1) \cdot \phi(x_2 | \mu_2 - \gamma(x_1 - \mu_1))$, while that of the history $h = (x_1, \emptyset)$ with $S(x_1) = \text{Stop}$ is $\phi(x_1 | \mu_1)$. Let $(\mu_1^*(S), \mu_2^*(S), \gamma^*(S)) \in \mathbb{R}^2 \times [\gamma_l, \gamma_h]$ be the *pseudo-true parameters* with respect to S that maximize the expected log likelihood of the agent's history, with the expectation taken over the true distribution of histories generated by S . Intuitively speaking, these correspond to the long-run inferences about the fundamentals and the reversal parameter when a large sample of histories is generated using the stopping strategy S .

Equivalently, the pseudo-true parameters minimize the *KL divergence* between the expected and the objective distributions over histories. Let $\mathcal{H}(\Psi(\mu_1, \mu_2; \gamma); S)$ refer to the distribution of histories when the draws have the joint distribution $\Psi(\mu_1, \mu_2; \gamma)$ and histories are censored according to the strategy S . The true distribution of histories given strategy S is $\mathcal{H}(\Psi(\mu_1^\bullet, \mu_2^\bullet; r); S)$, which I abbreviate as $\mathcal{H}^\bullet(S)$. To avoid trivialities, I will focus on steady states where agents continue with positive probability (otherwise their beliefs are not disciplined by the observation of any second-period draws), that is to say strategies S where $S(x_1) = \text{Continue}$ for a positive Lebesgue measure of $x_1 \in \mathbb{R}$. For such an S , the *Kullback–Leibler (KL) divergence* from $\mathcal{H}^\bullet(S)$ to $\mathcal{H}(\Psi(\mu_1, \mu_2; \gamma); S)$, denoted by $D_{\text{KL}}(\mathcal{H}^\bullet(S) \parallel \mathcal{H}(\Psi(\mu_1, \mu_2; \gamma); S))$, is

$$\begin{aligned} & \int_{x_1 \in S^{-1}(\text{Stop})} \phi(x_1 \mid \mu_1^\bullet) \cdot \ln \left(\frac{\phi(x_1 \mid \mu_1^\bullet)}{\phi(x_1 \mid \mu_1)} \right) dx_1 \\ & + \int_{x_1 \in S^{-1}(\text{Cont.})} \left\{ \int_{-\infty}^{\infty} \phi(x_1 \mid \mu_1^\bullet) \cdot \phi(x_2 \mid \mu_2^\bullet - r(x_1 - \mu_1)) \right. \\ & \cdot \ln \left[\frac{\phi(x_1 \mid \mu_1^\bullet) \cdot \phi(x_2 \mid \mu_2^\bullet - r(x_1 - \mu_1))}{\phi(x_1 \mid \mu_1) \cdot \phi(x_2 \mid \mu_2 - \gamma(x_1 - \mu_1))} \right] dx_2 \Big\} dx_1. \end{aligned} \quad (1)$$

So the KL divergence in (1) is the expected log-likelihood ratio of the history under the true process versus under the model $\Psi(\mu_1, \mu_2; \gamma)$, where expectation over histories is taken under the true process. In general, this optimization objective depends on the stopping strategy S . It is simple to see that the minimizers of KL divergence are the same as the maximizers of expected log likelihood of the history.

I formalize the definition of a steady state.

DEFINITION 1. A *steady state* consists of $\mu_1^\infty, \mu_2^\infty \in \mathbb{R}$, $\gamma^\infty \in [\gamma_l, \gamma_h]$, and a strategy S^∞ such that (i) S^∞ continues with positive probability and is optimal among all stopping strategies for the model $\Psi(\mu_1^\infty, \mu_2^\infty; \gamma^\infty)$ and (ii) $\mu_1^\infty = \mu_1^*(S^\infty)$, $\mu_2^\infty = \mu_2^*(S^\infty)$, $\gamma^\infty = \gamma^*(S^\infty)$.

The steady state is not a self-confirming equilibrium. There is positive KL divergence between the true data distribution in the steady state and the data distribution under $\Psi(\mu_1^\infty, \mu_2^\infty; \gamma^\infty)$, so even the best-fitting beliefs do not perfectly explain the data. To see this, consider the special case of $r = 0$, $\gamma_n > 0$. Objectively, the conditional distribution $X_2 \mid (X_1 = x_1)$ has a mean of $\mu_2^\bullet$ for every $x_1 \in \mathbb{R}$. In the steady state, the biased agents believe the same conditional distribution has a mean of $\mu_2^\infty - \gamma_n(x_1 - \mu_1^\infty)$, which only equals $\mu_2^\bullet$ for one value of x_1 . The histories cannot be fully explained by $\Psi(\mu_1^\infty, \mu_2^\infty; \gamma^\infty)$, as the predicted conditional distribution $X_2 \mid (X_1 = x_1)$ does not match what is in the data for almost all x_1 values where the steady-state strategy chooses to continue.

We may view the steady state as a stand-alone equilibrium concept that captures the optimality of behavior given beliefs and the constrained optimality of inferences given behavior, in the sense of minimizing KL divergence. Alternatively, Section 4 provides a

Bayesian-learning foundation for the steady state, in an environment where agents are not actually solving the KL divergence minimization problem given in (1) and do not observe any history of the stage game other than the history they personally experience. In that setting, (1) is involved in characterizing the steady state when a sequence of agents each play the stage game once and pass down their updated Bayesian beliefs to their successors.

2.4 Discussion of behavioral assumptions

In this paper, the agents' correlational mistake stems from their dogmatic belief in the interval $[\gamma_l, \gamma_h]$, which may exclude the true reversal parameter r . One story about how the agents erroneously think $\gamma_l, \gamma_h > 0$ in an environment with $r = 0$ (that is, suffer from the gambler's fallacy) relates to [Kahneman and Tversky's \(1972\)](#) representativeness heuristic in judging the likelihoods of random sequences. Objectively, the idiosyncratic factors ϵ_i (e.g., luck) that govern how draws in different periods deviate from their respective pool averages are sampled independently and identically distributed (i.i.d.) from a mean-zero distribution. The representativeness heuristic states that people know certain "essential characteristics" of the parent population generating these idiosyncratic factors (perhaps by observing their luck in other settings where the fundamentals are known), but exaggerate the extent to which small samples typically represent these characteristics. Agents who expect a sample of size two (ϵ_1, ϵ_2) to approximate the mean-zero property of the parent population of idiosyncratic factors should believe in a reversal of luck, that is, $\gamma_l, \gamma_h > 0$.

This is not a fully detailed and satisfactory microfoundation for the gambler's fallacy bias, and unfortunately there is limited work on the origin and persistence of biases in learning contexts. This literature typically studies the implications of a dogmatically wrong belief about one parameter on the Bayesian inference about a different parameter (e.g., [Heidhues, Kőszegi, and Strack \(2018, 2019\)](#)). Better understanding why mistakes persist is an important next step.

My setup corresponds to the model of the gambler's fallacy introduced in [Rabin and Vayanos \(2010\)](#), but applied to a different fundamental process. [Rabin and Vayanos \(2010\)](#) study a setting where a signal $s_t = \theta_t + \epsilon_t$ is generated each period t around the fundamental θ_t . Objectively $\epsilon_t \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma_\epsilon^2)$, but agents believe $\epsilon_t = \omega_t - \alpha\delta\epsilon_{t-1} - \alpha\delta^2\epsilon_{t-2} - \dots$ for $\omega_t \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma_\omega^2)$ and some $\alpha > 0, \delta \in (0, 1)$. This specializes to my model with $r = 0$ when there are two periods $t = 1, 2$, the fundamental process is $\theta_t = \mu_t$ for deterministic but unknown μ_1, μ_2 , agents know the variance $\sigma_\omega^2 = \sigma_\epsilon^2$, and $\gamma_l = \gamma_h = \alpha\delta$. For [Rabin and Vayanos \(2010\)](#), the fundamentals (θ_t) follow an (autoregressive) AR(1) process instead of being deterministic, and they study agents who exogenously observe all signals and estimate the long-run mean and persistence of the fundamental process. I study a different environment with endogenous data where agents' stopping decisions censor the observation of future signals.

3. STEADY-STATE RESULTS

3.1 Inference about parameters from censored data

A *cutoff strategy* is a strategy S whose stopping region $S^{-1}(\text{Stop})$ is either $[c, \infty)$ for some $c \in \mathbb{R} \cup \{\infty\}$ or $(-\infty, c]$ for some $c \in \mathbb{R} \cup \{-\infty\}$. The next proposition provides a closed-form expression for the pseudo-true parameters as a function of the cutoff threshold c in a cutoff strategy S . This result can be thought of as a one-sided benchmark of how biased learners misinfer the fundamentals and the reversal parameter using data censored at an exogenously given threshold. The subsequent steady-state analysis considers stopping strategies that best respond to the beliefs they induce. All proofs appear in the Appendix.

PROPOSITION 1. *For any strategy S that continues with positive probability, $\mu_1^*(S) = \mu_1^\bullet$, $\gamma^*(S) = \gamma_n$. If S is a cutoff strategy that stops when $x_1 \geq c$ for some $c \in \mathbb{R} \cup \{\infty\}$, then $\mu_2^*(c) = \mu_2^\bullet + (r - \gamma_n) \cdot (\mu_1^\bullet - \mathbb{E}[X_1 | X_1 \leq c])$. If S is a cutoff strategy that stops when $x_1 \leq c$ for some $c \in \mathbb{R} \cup \{-\infty\}$, then $\mu_2^*(c) = \mu_2^\bullet + (r - \gamma_n) \cdot (\mu_1^\bullet - \mathbb{E}[X_1 | X_1 \geq c])$.*

Proposition 1 shows that the misinference phenomenon requires both data censoring and the correlational mistake. Even biased agents with $r \notin [\gamma_l, \gamma_h]$ correctly estimate the fundamentals in the absence of censoring (i.e., under the strategy S that never stops). Conversely, agents whose prior belief does not contain a dogmatic correlational mistake (i.e., when $r \in [\gamma_l, \gamma_h]$) end up with correct beliefs about the fundamentals for any level of censoring.

Whether biased agents with $r \notin [\gamma_l, \gamma_h]$ will hold overpessimistic or overoptimistic beliefs about the fundamentals depends on the direction of their correlational mistake and the direction of data censoring. When $r - \gamma_n < 0$ and the strategy stops for high values of X_1 , and when $r - \gamma_n > 0$ and the strategy stops for low values of X_1 , agents have overpessimistic beliefs. When $r - \gamma_n < 0$ and the strategy stops for low values of X_1 , and when $r - \gamma_n > 0$ and the strategy stops for high values of X_1 , agents have overoptimistic beliefs. In all cases, more severe censoring (i.e., a cutoff strategy that stops for more realizations of X_1) exacerbates the belief distortion. Details of the intertemporal correlation misperception interact with the region of selective censoring to determine agents' long-run beliefs.

Turning to our main application, when agents exaggerate reversals $r - \gamma_n < 0$ and observe data generated from a cutoff rule that stops for high X_1 (e.g., stop searching if and only if the early candidate's quality is higher than some c), they have overpessimistic beliefs about μ_2 and their beliefs decrease without bound as the stopping threshold c decreases. I will use this application to explain why directional data censoring leads to belief distortions for biased learners.

Suppose $r = 0$ and $\gamma_n > 0$. Under the gambler's fallacy, the expected realization of X_2 depends on two factors: the second-period pool quality μ_2 and a reversal effect based on the realization of X_1 . The society of biased agents who stop for low values of X_1 cannot end up with a correct or overoptimistic belief about μ_2 , else they would be systematically disappointed by the realizations of X_2 in their own histories in an environment

where X_1 and X_2 are objectively independent. This is because the second draw is only observed when the first draw's quality is low enough, a contingency that leads biased agents to expect positive reversal on average. The long-run beliefs of the agents thus feature two mistakes partially canceling each other out to better fit the data, as their pessimism about the quality of the late-phase pool counteracts their false expectation of positive reversals when the first draw is bad enough to be rejected.

The severity of the biased agents' pessimism increases with the severity of censoring. The intuition is that the bias leads agents to infer a lower μ_2^* to better match X_2 s in histories that start with bad X_1 s, but doing so carries the cost of a worse model fit for histories that start with intermediate X_1 s. More severe censoring—generated by a strategy that stops not only after the very good early early draws, but also after the intermediate ones—alleviates this cost, as histories that start with intermediate X_1 s no longer contain their associated X_2 s. The extra censoring thus decreases the optimal inference μ_2^* .

The agents jointly estimate the reversal parameter r and the fundamentals $\mu_1^\bullet$ and $\mu_2^\bullet$. Proposition 1 says that agents always end up believing the nearest feasible parameter γ_n to the true reversal parameter r . To gain some geometric intuition for this result, view the agents' inference problem as using a scatter plot of (x_1, x_2) data points to estimate a conditional expectation, $\mathbb{E}[X_2 | X_1 = x_1]$. This conditional expectation is a linear function in x_1 with a slope of $-\gamma$ and an intercept determined by μ_2 . The conditional expectation in the true data-generating process has the slope $-r$. The agent is free to infer any intercept, but must pick a slope such that $\gamma \in [\gamma_l, \gamma_h]$. Geometrically speaking, the best-fitting regression line will have the slope $-\gamma_n$. A line with a slope as close as possible to the data-generating slope and the best-fitting intercept given this slope will better describe the data points than a line with any other feasible slope and any other intercept.

Proposition 1 also tells us that the quality of the early pool is always correctly estimated with any stopping strategy. This is because the first draw's quality X_1 is always observed, and $\mu_1^* = \mu_1^\bullet$ provides the best fit for the first-period data. The agents cannot improve the fit of second-period data by distorting their inference about the early pool: for any reversal parameter γ , fundamentals (μ_1', μ_2) and $(\mu_1^\bullet, \mu_2 - \gamma(\mu_1^\bullet - \mu_1'))$ generate the same conditional distributions of $X_2 | (X_1 = x_1)$ for any realization x_1 . Any distortion of the inference about early pool from $\mu_1^\bullet$ to μ_1' to better explain X_2 data can be equivalently done by keeping $\mu_1^* = \mu_1^\bullet$ and shifting μ_2^* by $-\gamma(\mu_1^\bullet - \mu_1')$. There is no trade-off between fitting X_1 and fitting X_2 , so the agents correctly infer $\mu_1^\bullet$ to provide the best fit for the early-pool mean.

Mueller, Spinnewijn, and Topa (2021) report in their Figure 3 that very recently unemployed workers underestimate their probability of finding a job in the next three months. This is consistent with Proposition 1's prediction of ex ante pessimistic beliefs at the start of the search, in a world where people suffer from the gambler's fallacy and accept early draws (i.e., job offers) that are sufficiently good.

3.2 Steady-state stopping behavior

In this section, I turn to behavior in the steady state. In the main application of the gambler's fallacy ($r = 0$, $\gamma_n > 0$), we know from Proposition 1 that agents end up with

overpessimistic beliefs about μ_2 if they infer from histories that are censored when $X_1 \geq c$ for any threshold $c \in \mathbb{R}$. But this pessimistic belief does not by itself imply that the misspecified agents must stop too often compared to a rational agent who knows the true fundamentals and r . Outside of the steady state, there is an intuition that an agent with the gambler's fallacy may stop less often than a rational one, even if the biased agent is overpessimistic about μ_2 . Consider an environment with $r = 0$, $\gamma_n > 0$, and suppose the stopping problem satisfies $\kappa = 0$, $q = 0$, so there is no cost of continuing but also no probability of recall. Suppose the true fundamentals are $\mu_1^\bullet \gg \mu_2^\bullet$. If a biased agent has the correct beliefs about the fundamentals, she perceives a greater continuation value after $X_1 = \mu_2^\bullet$ than a rational agent with the same correct beliefs, since the former holds a false expectation of positive reversals after a bad (relative to $\mu_1^\bullet$) early draw. The rational stopping cutoff is $c^\bullet = \mu_2^\bullet$ and the rational agent is willing to stop after $X_1 = \mu_2^\bullet$, but the biased agent strictly prefers to continue after such an early draw and has an indifference threshold strictly above $c^\bullet$. By continuity, the biased agent's cutoff threshold remains strictly above $c^\bullet$ even under slightly pessimistic beliefs about μ_2 .

Such ambiguity about behavior disappears in the steady state. The main result of this section, Proposition 3, compares the *steady-state* stopping behavior of the biased learners to the objectively optimal thresholds. Toward this result, I begin with a lemma that characterizes the optimal behavior for an agent who believes in the model $\Psi(\mu_1, \mu_2; \gamma)$, and a sufficient condition about the existence and uniqueness of the steady state.

LEMMA 1. *Consider the model $\Psi(\mu_1, \mu_2; \gamma)$ for any $\mu_1, \mu_2, \gamma \in \mathbb{R}$. When $\gamma \neq -1$, there is a unique cutoff $C(\mu_1, \mu_2; \gamma)$ so that the agent is indifferent between continuing and stopping after $X_1 = C(\mu_1, \mu_2; \gamma)$. When $\gamma > -1$, the optimal strategy is to stop when $X_1 \geq C(\mu_1, \mu_2; \gamma)$, and $\mu_2 \mapsto C(\mu_1, \mu_2; \gamma)$ is strictly increasing. When $\gamma < -1$, the optimal strategy is to stop when $X_1 \leq C(\mu_1, \mu_2; \gamma)$, and $\mu_2 \mapsto C(\mu_1, \mu_2; \gamma)$ is strictly decreasing.*

Lemma 1 says the optimal behavior under the model $\Psi(\mu_1, \mu_2; \gamma)$ is a cutoff strategy, and whether the agent stops after high enough or low enough values of X_1 depends on whether $\gamma > -1$ or $\gamma < -1$. To understand why, note that if the agent thinks X_1 and X_2 are independent ($\gamma = 0$), then she will choose to stop when the realization of X_1 is so large that the known payoff from stopping exceeds the expectation of the uncertain payoff from continuing and drawing an independent X_2 . But if the agent thinks X_1 and X_2 are sufficiently positively correlated ($\gamma < -1$), then larger realizations of X_1 make it even more attractive to continue. In this case, it is bad realizations of X_1 that cause the agent to stop, for the positive correlation makes the agent pessimistic about X_2 after a bad X_1 .

Suppose $\gamma_n > -1$, and consider a simplified setting where the agents know $\mu_1 = \mu_1^\bullet$ and always believe in $\gamma = \gamma_n$. For agents who exaggerate reversals ($r - \gamma_n < 0$), there is a positive feedback loop between distorted beliefs and distorted strategies: a more pessimistic belief about the second-period pool leads to a lower stopping cutoff by Lemma 1, and a lower stopping cutoff leads to more pessimistic beliefs by Proposition 1. On the other hand, for agents who suffer from the opposite correlational mistake

($r - \gamma_n > 0$), there is instead a negative feedback loop: a more pessimistic belief about μ_2 still leads to a lower stopping cutoff, but a lower stopping cutoff leads to more *optimistic* beliefs by Proposition 1. Heidhues, Kőszegi, and Strack (2018) show that overconfidence and underconfidence biases in a static effort-choice problem also lead to positive and negative feedback loops, respectively. In both environments, reversing the direction of the bias changes the nature of the feedback cycle between distorted actions and distorted beliefs.

The next result gives a sufficient condition for the existence and uniqueness of the steady state.

PROPOSITION 2. *There exists a unique steady state if $|(r - \gamma_n)/(1 + \gamma_n)| < 1$.*

When $r = 0$, so the draws are objectively independent, Proposition 2 says a unique steady state exists under any amount of the gambler's fallacy ($\gamma_n > 0$), and also under a moderate amount of the opposite correlational mistake ($-1/2 < \gamma_n < 0$). In general, a steady state may fail to exist when Proposition 2's condition is violated, as the following example shows.

EXAMPLE 1. Suppose $\kappa = 0$ and $q = 0$ (no cost of continuing and no probability of recall), and let $\gamma_l = \gamma_h = 0$, $r = -2$, and $\mu_1^\bullet = \mu_2^\bullet = 0$. No steady state exists in this setting. This is because by Lemma 1, steady-state behavior must involve stopping for $X_1 \geq c$ for some $c \in \mathbb{R}$. In fact, since the agent believes X_1 and X_2 are independent, she is indifferent between continuing and stopping if the early draw equals μ_2 , her belief about the mean of the second-period draw. Proposition 1 implies her belief μ_2 is related to c by $\mu_2^*(c) = 2 \cdot \mathbb{E}[X_1 \mid X_1 \leq c] < 0$. We need to find a $c < 0$ such that $c = 2 \cdot \mathbb{E}[X_1 \mid X_1 \leq c]$, which is impossible. Intuitively, the feedback cycle between more pessimistic beliefs and lower cutoff thresholds is expansionary and tends to $-\infty$. $\diamond$

As Example 1 hints at, the condition $|(r - \gamma_n)/(1 + \gamma_n)| < 1$ in Proposition 2 ensures that the feedback between beliefs and behavior is a contraction map.

Under the condition $|(r - \gamma_n)/(1 + \gamma_n)| < 1$, the next result compares the (unique) steady-state cutoff threshold c^∞ with the objectively optimal one, $c^\bullet$. Of course, by Lemma 1, if r and γ_n are on the opposite sides of -1 , then the comparison of thresholds is meaningless as the steady-state behavior will have the “opposite” kind of stopping region relative to the optimal behavior. When they are on the same side of -1 , Proposition 3 shows that whether $c^\infty < c^\bullet$ or $c^\infty > c^\bullet$ depends on the direction of the correlational mistake.

PROPOSITION 3. *Suppose $|(r - \gamma_n)/(1 + \gamma_n)| < 1$, and suppose either both $r, \gamma_n > -1$ or both $r, \gamma_n < -1$. Let c^∞ be the cutoff where the steady-state strategy switches between continuing and stopping, and let $c^\bullet$ be switching cutoff of the objectively optimal strategy. If $r - \gamma_n < 0$, then $c^\infty < c^\bullet$. If $r - \gamma_n > 0$, then $c^\infty > c^\bullet$.*

Combined with Proposition 1 and Lemma 1, Proposition 3 gives us the following conclusions when $|(r - \gamma_n)/(1 + \gamma_n)| < 1$: if $r, \gamma_n > -1$ (so that steady-state and optimal

strategies stop after good first-period draws), then $\gamma_n > r$ implies that the agent stops too often and underestimates μ_2 , while $\gamma_n < r$ implies that the agent stops too rarely and overestimates μ_2 . By contrast, if $r, \gamma_n < -1$ (so that steady-state and optimal strategies stop after bad first-period draws), then the implications of these two biases are reversed.

In particular, when $r = 0$ and $\gamma_n > 0$, Proposition 3's early-stopping conclusion strengthens Proposition 1's overpessimism result. In the steady state, agents must be *sufficiently* pessimistic as to overcome the opposite intuition about late stopping under the gambler's fallacy discussed earlier. To understand the intuition, note biased agents believe in different conditional distributions of X_2 following different realizations of X_1 , with more pessimistic beliefs after higher realizations. In a steady state $((\mu_1^\infty, \mu_2^\infty, \gamma_n), c^\infty)$, the agents' subjective distribution of X_2 following $X_1 = c^\infty$ must be a leftward shift of the true distribution $\phi(\cdot | \mu_2^\bullet)$. Else, their subjective distributions of X_2 would stochastically dominate the true distribution following *all* x_1 values in the continuation region, so heuristically they could improve the fit of their model by lowering their belief about μ_2 . The biased agents' indifference at c^∞ is thus based on an overly pessimistic belief about the continuation value, so we must have $c^\infty < c^\bullet$.

3.3 Gambler's fallacy with independent draws

In this section, I derive additional steady-state results for the main application of agents who suffer from the gambler's fallacy in an environment with independent X_1 and X_2 : that is, $r = 0$ and $\gamma_n > 0$.

3.3.1 Comparative statics in the stage game's parameters How do steady-state beliefs react to changes in the stage game's parameters, q and κ ? In general, when learners infer from exogenous data, their decision problem does not influence learning outcomes. This observation holds independently of whether learners are misspecified. On the other hand, correctly specified learners in my setting always end up with correct beliefs in the long run, so the game parameters are again irrelevant. With misspecified learners in an endogenous-data setting, however, changes in the stage game carry long-run consequences on society's beliefs about the fundamentals.

PROPOSITION 4. *Suppose $r = 0$ and $\gamma_n > 0$. Let $((\mu_1^{(q,\kappa)}, \mu_2^{(q,\kappa)}, \gamma_n), c^{(q,\kappa)})$ denote the unique steady-state beliefs and cutoff under parameters $q \in [0, 1)$, $\kappa \in \mathbb{R}$. The steady-state belief $\mu_2^{(q,\kappa)}$ is strictly increasing in q and strictly decreasing in κ , but always satisfies $\mu_2^{(q,\kappa)} < \mu_2^\bullet$. The steady-state cutoff threshold $c^{(q,\kappa)}$ is strictly increasing in q and strictly decreasing in κ .*

Proposition 4 provides novel predictions about how the economic environment affects biased inference under the gambler's fallacy. It says when agents are more patient (i.e., suffer a lower waiting cost or receive a higher subsidy for continuing) or when they have a higher chance of recalling previous draws, then they will end up with less distorted beliefs about the pool in the long run. These changes in environmental parameters partially correct society's long-run beliefs by incentivizing longer search and mitigating the censoring effect.

3.3.2 Fictitious variation and censoring So far, I have assumed agents hold dogmatic and correct beliefs about the variance of X_1 and the conditional variance of $X_2 | (X_1 = x_1)$. Now consider agents who are uncertain about these variances and jointly estimate them together with the means of the pools. I show that agents end up exaggerating the variances in a way that depends on the severity of data censoring.

For $\mu_1, \mu_2 \in \mathbb{R}$, $\sigma_1^2, \sigma_2^2 \geq 0$, and $\gamma \in \mathbb{R}$, let $\Psi(\mu_1, \mu_2, \sigma_1^2, \sigma_2^2; \gamma)$ refer to the joint distribution $X_1 = \mu_1 + \epsilon_1$, $X_2 = \mu_2 + \epsilon_2$ with $\epsilon_1 \sim \mathcal{N}(0, \sigma_1^2)$, and $(\epsilon_2 | \epsilon_1) \sim \mathcal{N}(-\gamma\epsilon_1, \sigma_2^2)$. In this section, “fundamentals” refer to the four parameters μ_1 , μ_2 , σ_1^2 , and σ_2^2 , and I assume for simplicity $\gamma_l = \gamma_h = \gamma > 0$. Objectively, X_1 and X_2 are independent Gaussian random variables each with a variance of $(\sigma^\bullet)^2 > 0$, so the true joint distribution of (X_1, X_2) is $\Psi^\bullet := \Psi(\mu_1^\bullet, \mu_2^\bullet, (\sigma^\bullet)^2, (\sigma^\bullet)^2; 0)$.

Following (1), write $D_{\text{KL}}(\mathcal{H}^\bullet(c) \parallel \mathcal{H}(\Psi(\mu_1, \mu_2, \sigma_1^2, \sigma_2^2; \gamma); c))$ to denote the KL divergence between the true distribution of histories with X_2 censored whenever $X_1 > c$ and the implied history distribution under the fundamentals μ_1 , μ_2 , σ_1^2 , and σ_2^2 . This divergence is given by

$$\begin{aligned} & \int_c^\infty \phi(x_1 | \mu_1^\bullet, (\sigma^\bullet)^2) \cdot \ln \left(\frac{\phi(x_1 | \mu_1^\bullet, (\sigma^\bullet)^2)}{\phi(x_1 | \mu_1, \sigma_1^2)} \right) dx_1 \\ & + \int_{-\infty}^c \left\{ \int_{-\infty}^\infty \phi(x_1 | \mu_1^\bullet, (\sigma^\bullet)^2) \cdot \phi(x_2 | \mu_2^\bullet, (\sigma^\bullet)^2) \right. \\ & \quad \cdot \ln \left[\frac{\phi(x_1 | \mu_1^\bullet, (\sigma^\bullet)^2) \cdot \phi(x_2 | \mu_2^\bullet, (\sigma^\bullet)^2)}{\phi(x_1 | \mu_1, \sigma_1^2) \cdot \phi(x_2 | \mu_2 - \gamma(x_1 - \mu_1), \sigma_2^2)} \right] dx_2 \Big\} dx_1, \end{aligned} \quad (2)$$

where $\phi(x | \mu, \sigma^2)$ is the Gaussian density with mean μ and variance σ^2 , evaluated at x .

The next proposition gives closed-form expressions for the pseudo-true fundamentals μ_1^* , μ_2^* , $(\sigma_1^*)^2$, and $(\sigma_2^*)^2$ that minimize (2).

PROPOSITION 5. *Suppose $r = 0$. The solutions of*

$$\min_{\mu_1, \mu_2 \in \mathbb{R}, \sigma_1^2, \sigma_2^2 \geq 0} D_{\text{KL}}(\mathcal{H}^\bullet(c) \parallel \mathcal{H}(\Psi(\mu_1, \mu_2, \sigma_1^2, \sigma_2^2; \gamma); c))$$

are $\mu_1^(c) = \mu_1^\bullet$, $\mu_2^*(c) = \mu_2^\bullet - \gamma(\mu_1^\bullet - \mathbb{E}[X_1 | X_1 \leq c])$, $(\sigma_1^*)^2(c) = (\sigma^\bullet)^2$, and $(\sigma_2^*)^2(c) = (\sigma^\bullet)^2 + \gamma^2 \text{Var}[X_1 | X_1 \leq c]$. So $(\sigma_2^*)^2(c)$ strictly increases in c .*

Comparing Proposition 5 with the expressions for $\mu_1^*(c)$ and $\mu_2^*(c)$ in Proposition 1 (for the special case of $r = 0$, $\gamma_n = \gamma > 0$, and a strategy that stops when $X_1 \geq c$) shows that agents misinfer the means in the same way regardless of whether they know the variances. Biased agents correctly estimate the first-period variance, $(\sigma_1^*)^2 = (\sigma^\bullet)^2$, but overestimate second-period variance. They exaggerate the variation in quality among the late-phase draws. This phenomenon relates to findings in Rabin (2002) and Rabin and Vayanos (2010), who refer to exaggeration of variance under the gambler’s fallacy as *fictitious variation*. The key innovation of Proposition 5 is to show, in an endogenous-

data setting, how the degree of fictitious variation depends on the severity of censoring.

The magnitude of this distortion increases in the severity of the gambler's fallacy but decreases with the severity of the censoring, as $\text{Var}[X_1 | X_1 \leq c]$ increases in c for X_1 Gaussian. Here is the intuition. Whereas the objective conditional distribution of $X_2 | (X_1 = x_1)$ is independent of x_1 , the biased agents entertain different beliefs about this distribution for different x_1 s. The agents' best-fitting inference about μ_2 ensures their belief about $X_2 | (X_1 = x_1)$ fits the data well following "typical" realizations of x_1 in the continuation region $(-\infty, c]$, but they are still surprised when they experience a streak of bad draws in their own stage game. Agents who observe such surprising streaks attribute the unexpectedly low realizations of X_2 to "noise," and thus pass down beliefs that estimate a higher conditional variance of $X_2 | (X_1 = x_1)$. A larger fraction of the agents attribute their data to "noise" when $\text{Var}[X_1 | X_1 \leq c]$ is larger, because the frequency of the surprising streaks depends on how much X_1 tends to deviate from its typical value of $\mathbb{E}[X_1 | X_1 \leq c]$ conditional on the event $\{X_1 \leq c\}$.

The next result demonstrates the interplay between fictitious variation and endogenous censoring in the steady state. Consider two societies of agents who have the same bias, play the same stage game, and face the same true fundamentals. Agents in society A know the true variances and only infer about (μ_1, μ_2) , while those in society B do not know the variances and infer about $(\mu_1, \mu_2, \sigma_1^2, \sigma_2^2)$.

PROPOSITION 6. *Suppose $r = 0$, $\gamma_l = \gamma_h = \gamma$, and the probability of recall is interior, $0 < q < 1$. Let (μ_1^A, μ_2^A, c^A) and $(\mu_1^B, \mu_2^B, (\sigma_1^B)^2, (\sigma_2^B)^2, c^B)$ be the steady-state beliefs about the fundamentals and the steady-state cutoffs in the two societies. Then $\mu_2^B > \mu_2^A$ and $c^B > c^A$. Also, $\sigma_2^B > \sigma_2^*(c^A)$.*

The endogenous-data setting leads to two novel implications of fictitious variation relative to [Rabin and Vayanos's \(2010\)](#) exogenous-data world. First, even though Proposition 5 implies that the two societies would make the same inferences about the pool means if they were given the same data, in steady state society B holds more optimistic (i.e., more correct) beliefs about μ_2 and uses a higher cutoff than society A. Allowing uncertainty on one dimension (variance) ends up affecting society's long-run inference in another dimension (mean), because a belief in fictitious variation increases the agents' perceived option value of continuing and thus changes their behavior and the kind of data they observe in the steady state. Second, fictitious variation has a "multiplier effect," as formalized by the final statement of Proposition 6. Society B's steady-state belief about σ_2 is higher than what it would have been had they simply inferred using data generated from society A's steady-state cutoff c^A . Allowing for uncertainty about the pool variances leads to fictitious variation that increases society B's cutoff above c^A . This is because when the agent can recall the first draw with an interior probability, the option value of waiting for the second draw is larger when the second labor pool has a larger variance in quality. This higher cutoff further heightens society B's belief in fictitious variation, since Proposition 5 implies $\sigma_2^*(c)$ is strictly increasing, and so forth.

4. CONVERGENCE TO THE STEADY STATE

This section shows the steady state defined and studied earlier corresponds to the long-run learning outcome for a society of biased agents acting one by one.

Time is discrete and partitioned into rounds $t = 1, 2, 3, \dots$. One short-lived agent arrives per round. For simplicity, in analyzing convergence I focus on learning about the fundamentals μ_1 and μ_2 , and suppose agents have a degenerate belief about the reversal parameter, $\gamma_l = \gamma_h > -1$. Agent 1 starts with a prior belief M_0 given by a continuously differentiable prior density $m_0 : [\underline{\mu}_1, \bar{\mu}_1] \times [\underline{\mu}_2, \bar{\mu}_2] \rightarrow \mathbb{R}_{>0}$, while each agent $t \geq 2$ adopts the final belief $\tilde{M}_{t-1}$ of agent $t-1$ as her prior belief. Since all agents commit the same statistical bias, each agent's inherited belief aggregates all the information in all predecessors' histories. The same learning dynamics obtain in an environment where every agent starts with the common prior belief M_0 and observes the stage-game histories of all predecessors.

In each round t , agent t chooses a cutoff threshold $\tilde{C}_t$ to maximize her expected payoff based on her prior belief.³ She observes the outcome of her game and updates her belief from $\tilde{M}_{t-1}$ to $\tilde{M}_t$ by applying Bayes' rule to her stage-game history, $\tilde{H}_t \in \mathbb{H}$. She then passes down $\tilde{M}_t$ as the prior belief of agent $t+1$.

By Proposition 2, there exists a unique steady state $((\mu_1^\infty, \mu_2^\infty, \gamma_n), c^\infty)$ when $|(r - \gamma_n)/(1 + \gamma_n)| < 1$. Proposition 7 shows that almost surely behavior and belief converge to this steady state for any prior density m_0 , provided the support $[\underline{\mu}_1, \bar{\mu}_1] \times [\underline{\mu}_2, \bar{\mu}_2]$ includes the steady-state beliefs $(\mu_1^\infty, \mu_2^\infty)$. To state this convergence result formally, I need to develop the probability space underlying the learning system.

The sequences $(\tilde{M}_t)$, $(\tilde{C}_t)$, and $(\tilde{H}_t)$ are stochastic processes whose randomness stems from randomness of the stage-game draws in different rounds. The convergence result is about the almost-sure convergence of the processes $(\tilde{M}_t)$ and $(\tilde{C}_t)$. Consider the $\mathbb{R}^2$ -valued stochastic process $(X_t)_{t \geq 1} = (X_{1,t}, X_{2,t})_{t \geq 1}$, where X_t and $X_{t'}$ are independent for $t \neq t'$. Within each t , $X_{1,t} \sim \phi(\cdot | \mu_1^\bullet)$ and $X_{2,t} | (X_{1,t} = x_{1,t}) \sim \phi(\cdot | \mu_2^\bullet - r(x_{1,t} - \mu_1^\bullet))$ are jointly Gaussian. Interpret X_t as the pair of potential draws in the t th round of the stage game. Clearly, there exists a probability space $(\Omega, \mathcal{A}, \mathbb{P})$, with sample space $\Omega = (\mathbb{R}^2)^\infty$ interpreted as paths of the process just described, $\mathcal{A}$ the Borel σ -algebra on Ω , and $\mathbb{P}$ the measure on sample paths so that the process $X_t(\omega) = \omega_t$ has the desired distribution. The term “almost surely” means with probability 1 with respect to the realization of the infinite sequence of all (potential) draws, i.e., $\mathbb{P}$ -almost surely. The processes $(\tilde{M}_t)$, $(\tilde{C}_t)$, and $(\tilde{H}_t)$ are defined on this probability space and adapted to the filtration $(\mathcal{F}_t)_{t \geq 1}$, where $\mathcal{F}_t$ is the sub- σ -algebra generated by draws up to round t , $\mathcal{F}_t = \sigma((X_s)_{s=1}^t)$. Write $(\tilde{\mu}_{1,t}, \tilde{\mu}_{2,t})$ for the random element in $[\underline{\mu}_1, \bar{\mu}_1] \times [\underline{\mu}_2, \bar{\mu}_2]$ given by the belief $\tilde{M}_t$.

PROPOSITION 7. *Suppose $|(r - \gamma_n)/(1 + \gamma_n)| < 1$, $r \neq \gamma_n$, and $\gamma_n > -1$. Provided $\underline{\mu}_1 < \mu_1^\bullet < \bar{\mu}_1$ and $\underline{\mu}_2 < \mu_2^\infty < \bar{\mu}_2$, almost surely $\lim_{t \rightarrow \infty} \tilde{C}_t = c^\infty$ and $(\tilde{\mu}_{1,t}, \tilde{\mu}_{2,t})_{t \geq 1}$ converges in L^1 to $(\mu_1^\bullet, \mu_2^\infty)$, where $((\mu_1^\bullet, \mu_2^\infty, \gamma_n), c^\infty)$ is the unique steady state.*

³I focus on learning across different iterations of the stage game and assume agents do not update beliefs within the stage game.

4.1 Proof outline for Proposition 7

The argument for Proposition 7 adapts techniques from [Heidhues, Kőszegi, and Strack \(2018\)](#); in particular, a law of large numbers for martingale increments. I discuss the novelties specific to my environment below.

4.1.1 When μ_1 is known First consider a simpler situation where agents dogmatically know that $\mu_1 = \mu_1^\bullet$ and only entertain uncertainty about μ_2 in some bounded interval $[\underline{\mu}_2, \bar{\mu}_2]$ that includes μ_2^∞ . I use a statistical tool from [Heidhues, Kőszegi, and Strack \(2018\)](#): a version of the law of large numbers for martingales whose quadratic variation grows linearly.

PROPOSITION 10 FROM [HEIDHUES, KŐSZEGI, AND STRACK \(2018\)](#). *Let $(y_t)_t$ be a martingale that satisfies a.s. $|y_t| \leq vt$ for some constant $v \geq 0$. We have that a.s. $\lim_{t \rightarrow \infty} \frac{y_t}{t} = 0$.*

After simplifying the problem with this result, I establish a pair of mutual bounds on asymptotic behavior and asymptotic beliefs. If cutoff thresholds are asymptotically bounded between c^l and c^h , $c^l < c^h$, then beliefs about μ_2 must be asymptotically supported on the interval $[\mu_2^*(c^l), \mu_2^*(c^h)]$ when $r - \gamma_n < 0$ and asymptotically supported on the interval $[\mu_2^*(c^h), \mu_2^*(c^l)]$ when $r - \gamma_n > 0$. Conversely, if belief is asymptotically supported on the subinterval $[\mu_2^l, \mu_2^h] \subseteq [\underline{\mu}_2, \bar{\mu}_2]$, then cutoff thresholds must be asymptotically bounded between $C(\mu_1^\bullet, \mu_2^l; \gamma_n)$ and $C(\mu_1^\bullet, \mu_2^h; \gamma_n)$.

Applying this pair of lemmas to $[\underline{\mu}_2, \bar{\mu}_2]$, I conclude that asymptotically $\tilde{M}_t$ must be supported on the subinterval with the end points $\mathcal{I}(\underline{\mu}_2)$ and $\mathcal{I}(\bar{\mu}_2)$, where $\mathcal{I}$ is the composition $\mathcal{I}(\mu_2) := \mu_2^*(C(\mu_1^\bullet, \mu_2; \gamma))$. The proof of Proposition 2 implies that $\mathcal{I}$ is a contraction map whose iterates converge to μ_2^∞ . Therefore by repeatedly applying the pair of lemmas, the bound on asymptotic beliefs gets refined down to the singleton $\{\mu_2^\infty\}$, showing the almost-sure convergence of beliefs and behavior.

4.1.2 Uncertainty about μ_1 In the hypothesis of Proposition 7, both μ_1 and μ_2 are unknown, so there is two-dimensional uncertainty about the fundamentals. This complication prevents a direct application of [Heidhues, Kőszegi, and Strack's \(2018\)](#) statistical tools, as their tools are only designed to work with a one-dimensional fundamental. But the structure of the inference problem is such that I can separately bound the agents' asymptotic beliefs in two "directions," thus reducing the task of proving a two-dimensional belief bound into a pair of tasks involving one-dimensional belief bounds.

Consider a pair of fundamentals, (μ_1, μ_2) and $(\mu'_1, \mu'_2) = (\mu_1 + d, \mu_2 - \gamma d)$ for some $d > 0$, satisfying $\mu_1, \mu'_1 \leq \mu_1^\bullet$. That is, (μ_1, μ_2) and (μ'_1, μ'_2) lie on the same line with slope $-\gamma$. For any uncensored history $(x_1, x_2) \in \mathbb{R}^2$, the likelihood of second-period draw x_2 is the same under both pairs of fundamentals, $\phi(x_2 | \mu_2 - \gamma(x_1 - \mu_1)) = \phi(x_2 | \mu'_2 - \gamma(x_1 - \mu'_1))$. So both pairs of fundamentals (μ_1, μ_2) and (μ'_1, μ'_2) explain X_2 data equally well in *all* uncensored histories. At the same time, (μ'_1, μ'_2) provides a strictly better fit for X_1 data on average than (μ_1, μ_2) , since $\mu_1 < \mu'_1 \leq \mu_1^\bullet$. This means in the long run,

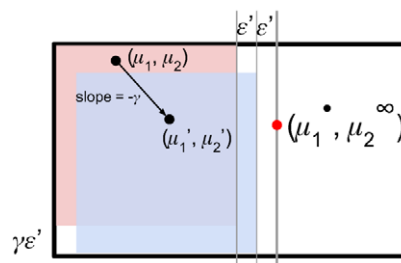


FIGURE 1. Directional derivative for data log likelihood along the vector $(1, -\gamma)$ in the space of fundamentals implies the upper shaded region receives zero posterior probability asymptotically.

fundamentals (μ_1, μ_2) should receive much less posterior probability than (μ_1', μ_2') , as the latter better rationalize the data overall.

To formalize this, I compute the directional derivative for data log likelihood along the vector $(1, -\gamma)$ in the space of fundamentals. I establish an (almost-sure) positive lower bound on this directional derivative at all points at least $2\epsilon'$ to the left of μ_1^* , and an analogous negative upper bound to the right of μ_1^* . Figure 1 is an illustration for the case of $\gamma > 0$. This allows me to show the upper shaded region receives 0 posterior probability asymptotically, by comparing each point in this upper shaded region with a corresponding point in the lower shaded region along a line of slope $-\gamma$.

By repeating this argument for small values of ϵ' (and applying the symmetric bound to the right of μ_1^*), I show that belief is asymptotically concentrated either along a small vertical strip containing the steady-state beliefs, (μ_1^*, μ_2^∞) or along an edge of belief's support, as shown in Figure 2. The latter possibility requires belief in an extreme value of $\mu_2 \in \{\underline{\mu}_2, \bar{\mu}_2\}$ in the support of the prior m_0 and can be ruled out by showing that, within these regions, slightly increasing or decreasing belief in μ_2 leads to better fit.

Having restricted the long-run belief to a thin vertical strip, the first “direction” of the belief bounds is complete and the dimensionality of uncertainty is effectively reduced back to one. The rest of the argument proceeds similarly to the case where agents know μ_1^* discussed above.

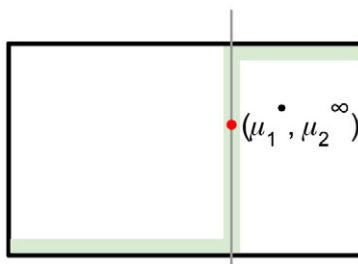


FIGURE 2. Belief must be asymptotically concentrated in the shaded area.

5. RELATED THEORETICAL LITERATURE

A strand of behavioral economics literature has focused on a different cognitive error when agents learn from partial data: selection neglect. Theory papers in this area have studied agents who observe a selective sample in different settings: good's quality in a bilateral trade game (Esponda (2008)), investment outcomes by past entrepreneurs (Jehiel (2018)), government policy effectiveness (Esponda and Pouzo (2017, 2019)), and outcomes of recent experiments (Chen (2021)). In all of these settings, the sample selection depends on some unobserved private information of other players. Biased agents fail to account for the informational content of selection,⁴ thus make wrong inferences. While I also consider a setting where agents learn from partial data, I focus on the implications of a different bias in such environments: the gambler's fallacy. Selection neglect and the gambler's fallacy can be conceptually unified under the broader category of correlational mistakes. As Spiegel (2016, 2017) points out, many examples of selection neglect can be viewed as biases stemming from incorrect conditional-independence assumptions. I emphasize that the biased agents in my world do not additionally suffer from selection neglect. Agents derive different inferences from histories censored at different thresholds purely as a result of misperceiving the reversal parameter that relates different draws; this conclusion does not come from the combination of multiple behavioral biases.

Rabin (2002) and Rabin and Vayanos (2010) are the first to study the inferential mistakes implied by the gambler's fallacy. Like these papers, I consider agents who believe in reversals conditional on the underlying fundamentals and mislearn some parameters of the world as a result. Except for an example in Rabin (2002), all such investigations focus on passive inference, whereby learners observe an exogenous signal process. By contrast, this paper examines an endogenous learning setting where actions affect observables. Section 7 of Rabin (2002) discusses an example of endogenous learning with a finite-urn model of the gambler's fallacy. The nature of Rabin's (2002) endogenous data, however, is unrelated to the censoring effect central to the current paper.⁵

This work joins a strand of literature on the implications of misspecified Bayesian learning when the learner's actions affect the data she observes. The earliest example is Nyarko (1991). Esponda and Pouzo (2016) propose an equilibrium concept for such settings: the Berk–Nash equilibrium. Subsequently, a number of papers have studied the properties of Berk–Nash equilibria in different applied contexts (Fudenberg, Romanayuk, and Strack (2017), Heidhues, Kőszegi, and Strack (2018), Frick, Iijima, and Ishii (2020)) and the persistence and stability of misspecifications (Frick, Iijima, and Ishii (2021b), Fudenberg and Lanzani (2021), He and Libgober (2021)). In addition to using this framework to explore the gambler's fallacy, I also highlight a new source of data

⁴Some recent experiments have demonstrated selection neglect in laboratory subjects: Enke (2020), Baron, Huck, and Jehiel (2019), and Araujo, Wang, and Wilson (2021).

⁵In Rabin's (2002) example, biased agents (correctly) believe that the part of the data that is always observable is independent of the part of the data that is sometimes missing. However, what I term the censoring effect is about misinference resulting from agents wrongly believing in negative correlation between the early draws that are always observed and the later draws that may be censored, depending on the realizations of the early draws. I discuss this further in the Online Appendix of an earlier version of this paper: <https://arxiv.org/pdf/1803.08170v5.pdf>.

endogeneity relative to the existing papers: the censoring effect in an optimal-stopping problem. Agents' stopping decisions determine how many signals they observe about the fundamentals. Other recent papers (Esponda, Pouzo, and Yamamoto (2021), Fudenberg, Lanzani, and Strack (2021), Frick, Iijima, and Ishii (2021a), Heidhues, Kőszegi, and Strack (2021)) prove general theorems about the convergence of misspecified learning in different settings. Though not the primary contribution of this work, the convergence result in Proposition 7 deals with a setting that is not covered by these papers: a multi-dimensional inference problem with a continuum of states, signals, and actions.

Although Section 4 considers a learning system with a sequence of short-lived agents, the "social learning" aspect of the framework is not central to the results. In fact, the environment where a sequence of short-lived agents act one at a time is equivalent to an environment where a single long-lived agent plays the stage game repeatedly, myopically maximizing her expected payoff in each iteration of the stage game. In the growing literature on social learning with misspecified Bayesians (e.g., Eyster and Rabin (2010), Gaurino and Jehiel (2013), Bohren (2016), Bohren and Hauser (2021), Dasaratha and He (2020), Frick, Iijima, and Ishii (2020), Bushong and Gagnon-Bartsch (2022)), agents observe their predecessors' actions but make errors when inverting these actions to deduce said predecessors' information. This kind of action inversion does not take place here: later agents inherit all the information that their predecessors have seen by adopting their beliefs, so predecessors' actions are uninformative.

The econometrics literature has also studied data-generating processes with censoring, for example, the Tobit model and models of competing risks.⁶ This literature has primarily focused on the issue of model identification from censored data (Cox (1962), Tsiatis (1975), Heckman and Honoré (1989)). In my setting, there is no identification problem for correctly specified agents. Instead, I study how agents make wrong parameter estimates from censored data when they infer using a family of misspecified models. Another contrast is that the econometrics literature has focused on exogenous data-censoring mechanisms, but censoring is endogenous in this paper and depends on the beliefs of previous agents.

6. CONCLUDING REMARKS

This paper studies endogenous learning dynamics of misspecified agents. The general framework allows different correlational mistakes, and shows that the interaction between the statistical bias and data censoring in optimal-stopping problems distorts beliefs and behavior. When agents suffer from the gambler's fallacy, they hold overly pessimistic beliefs about the fundamentals and stop too frequently in the steady state. Lower continuation costs, as well as initial uncertainty about the distribution's variance, partially correct asymptotic beliefs about the distribution's mean.

An earlier version of this paper⁷ shows that the steady-state results (about overpessimistic inference and early stopping) and the convergence result continue to hold for

⁶References can be found in Amemiya (1985) and Crowder (2001).

⁷Available at <https://arxiv.org/pdf/1803.08170v5.pdf>.

a larger class of stage games and any symmetric, log-concave distributions. That earlier version also contains an extension with any finite number $L \geq 2$ of periods instead of two periods.

In line with previous work on the gambler's fallacy, I take the behavioral error as given and do not try to explain the origin of the bias. Endogenizing the gambler's fallacy and other common statistical errors is an interesting open question.

I have studied a particular environment where censoring happens (histories in optimal-stopping problems). The key mechanism I highlight—the interaction between data censoring and bias—applies more broadly and delivers different predictions in different contexts. Environments that feature different censoring patterns would produce different predictions, but again through the same basic mechanism—interaction between censoring and bias. More broadly, other kinds of “symmetric” behavioral biases may lead to “asymmetric” predictions in environments that feature directional data censoring. I am leaving open the interaction of other kinds of behavioral learning with other censoring mechanisms to future work.

APPENDIX: PROOFS

A.1 Proof of Proposition 1

In the true model, $X_2|(X_1 = x_1) \sim \mathcal{N}(\mu_2^\bullet - r(x_1 - \mu_1^\bullet), \sigma^2)$, while the agents' feasible model $\Psi(\mu_1, \mu_2; \gamma)$ has $X_2|(X_1 = x_1) \sim \mathcal{N}(\mu_2 - \gamma(x_1 - \mu_1), \sigma^2)$. Suppose histories are generated with a stopping rule that continues in the positive Lebesgue measure set $K \subseteq \mathbb{R}$. The objective in (1) is

$$\begin{aligned} & \int_{x_1 \notin K} \phi(x_1 | \mu_1^\bullet) \cdot \ln \left(\frac{\phi(x_1 | \mu_1^\bullet)}{\phi(x_1 | \mu_1)} \right) dx_1 \\ & + \int_{x_1 \in K} \phi(x_1 | \mu_1^\bullet) \cdot \left\{ \int_{-\infty}^{\infty} \phi(x_2 | \mu_2^\bullet - r(x_1 - \mu_1^\bullet)) \right. \\ & \quad \cdot \ln \left[\frac{\phi(x_1 | \mu_1^\bullet) \cdot \phi(x_2 | \mu_2^\bullet - r(x_1 - \mu_1^\bullet))}{\phi(x_1 | \mu_1) \cdot \phi(x_2 | \mu_2 - \gamma(x_1 - \mu_1))} \right] dx_2 \Big\} dx_1. \end{aligned}$$

This can be rewritten as

$$\begin{aligned} & \int_{x_1 \notin K} \phi(x_1 | \mu_1^\bullet) \ln \left(\frac{\phi(x_1 | \mu_1^\bullet)}{\phi(x_1 | \mu_1)} \right) dx_1 \\ & + \int_{x_1 \in K} \phi(x_1 | \mu_1^\bullet) \left\{ \int_{-\infty}^{\infty} \phi(x_2 | \mu_2^\bullet - r(x_1 - \mu_1^\bullet)) \ln \left[\frac{\phi(x_1 | \mu_1^\bullet)}{\phi(x_1 | \mu_1)} \right] dx_2 \right\} dx_1 \\ & + \int_{x_1 \in K} \phi(x_1 | \mu_1^\bullet) \cdot \left\{ \int_{-\infty}^{\infty} \phi(x_2 | \mu_2^\bullet - r(x_1 - \mu_1^\bullet)) \right. \\ & \quad \cdot \ln \left[\frac{\phi(x_2 | \mu_2^\bullet - r(x_1 - \mu_1^\bullet))}{\phi(x_2 | \mu_2 - \gamma(x_1 - \mu_1))} \right] dx_2 \Big\} dx_1, \end{aligned}$$

which is

$$\begin{aligned} & \int_{-\infty}^{\infty} \phi(x_1 | \mu_1^\bullet) \cdot \ln\left(\frac{\phi(x_1 | \mu_1^\bullet)}{\phi(x_1 | \mu_1)}\right) dx_1 \\ & + \int_{x_1 \in K} \phi(x_1 | \mu_1^\bullet) \cdot \int_{-\infty}^{\infty} \phi(x_2 | \mu_2^\bullet - r(x_1 - \mu_1^\bullet)) \cdot \ln\left[\frac{\phi(x_2 | \mu_2^\bullet - r(x_1 - \mu_1^\bullet))}{\phi(x_2 | \mu_2 - \gamma(x_1 - \mu_1))}\right] dx_2 dx_1. \end{aligned}$$

The KL divergence between $\mathcal{N}(\mu_{\text{true}}, \sigma_{\text{true}}^2)$ and $\mathcal{N}(\mu_{\text{model}}, \sigma_{\text{model}}^2)$ is $\ln(\sigma_{\text{model}}/\sigma_{\text{true}}) + [(\sigma_{\text{true}}^2 + (\mu_{\text{true}} - \mu_{\text{model}})^2)/(2\sigma_{\text{model}}^2)] - \frac{1}{2}$, so we may simplify the first term and the inner integral of the second term:

$$\frac{(\mu_1 - \mu_1^\bullet)^2}{2\sigma^2} + \int_{x_1 \in K} \phi(x_1 | \mu_1^\bullet) \cdot \frac{(\mu_2 - \gamma(x_1 - \mu_1) - \mu_2^\bullet + r(x_1 - \mu_1^\bullet))^2}{2\sigma^2} dx_1.$$

Multiplying through by σ^2 , we get a simplified objective with the same minimizers:

$$\xi(\mu_1, \mu_2, \gamma) = \frac{(\mu_1 - \mu_1^\bullet)^2}{2} + \int_{x_1 \in K} \phi(x_1 | \mu_1^\bullet) \cdot \frac{1}{2} \cdot [\mu_2 - \gamma(x_1 - \mu_1) - \mu_2^\bullet + r(x_1 - \mu_1^\bullet)]^2 dx_1.$$

We have the partial derivatives by differentiating under the integral sign:

$$\begin{aligned} \frac{\partial \xi}{\partial \mu_2} &= \int_{x_1 \in K} \phi(x_1 | \mu_1^\bullet) \cdot [\mu_2 - \gamma(x_1 - \mu_1) - \mu_2^\bullet + r(x_1 - \mu_1^\bullet)] dx_1 \\ \frac{\partial \xi}{\partial \mu_1} &= (\mu_1 - \mu_1^\bullet) + \gamma \int_{x_1 \in K} \phi(x_1 | \mu_1^\bullet) \cdot [\mu_2 - \gamma(x_1 - \mu_1) - \mu_2^\bullet + r(x_1 - \mu_1^\bullet)] dx_1 \\ &= (\mu_1 - \mu_1^\bullet) + \gamma \frac{\partial \xi}{\partial \mu_2} \\ \frac{\partial \xi}{\partial \gamma} &= - \int_{x_1 \in K} \phi(x_1 | \mu_1^\bullet) \cdot [x_1 - \mu_1] \cdot [\mu_2 - \gamma(x_1 - \mu_1) - \mu_2^\bullet + r(x_1 - \mu_1^\bullet)] dx_1. \end{aligned}$$

Suppose $(\mu_1^*, \mu_2^*, \gamma^*)$ is the minimum. By the first-order conditions for μ_1 and μ_2 , we have

$$\frac{\partial \xi}{\partial \mu_1}(\mu_1^*, \mu_2^*, \gamma^*) = \frac{\partial \xi}{\partial \mu_2}(\mu_1^*, \mu_2^*, \gamma^*) = 0 \quad \Rightarrow \quad \mu_1^* = \mu_1^\bullet.$$

Substituting this into the first-order condition for μ_2 ,

$$\frac{\partial \xi}{\partial \mu_2}(\mu_1^\bullet, \mu_2^*, \gamma^*) = 0 \quad \Rightarrow \quad \mu_2^* = \mu_2^\bullet + (r - \gamma^*) \cdot (\mu_1^\bullet - \mathbb{E}[X_1 | X_1 \in K]).$$

It remains to find γ^* . We have

$$\frac{\partial \xi}{\partial \gamma}(\mu_1^*, \mu_2^*, \gamma^*) = -\mathbb{P}[X_1 \in K] \cdot \mathbb{E}[(X_1 - \mu_1^*) \cdot (\mu_2^* - \gamma^*(X_1 - \mu_1^*) - \mu_2^\bullet + r(X_1 - \mu_1^\bullet)) | X_1 \in K].$$

We rearrange the expectation term as

$$\begin{aligned} & \mathbb{E}[(X_1 - \mu_1^*) \cdot (\mu_2^* - \gamma^*(X_1 - \mu_1^*) - \mu_2^\bullet + r(X_1 - \mu_1^\bullet)) | X_1 \in K] \\ &= \mathbb{E}[(X_1 - \mu_1^*) | X_1 \in K] \cdot \mathbb{E}[(\mu_2^* - \gamma^*(X_1 - \mu_1^*) - \mu_2^\bullet + r(X_1 - \mu_1^\bullet)) | X_1 \in K] \\ &+ \text{Cov}[X_1 - \mu_1^*, \mu_2^* - \gamma^*(X_1 - \mu_1^*) - \mu_2^\bullet + r(X_1 - \mu_1^\bullet) | X_1 \in K]. \end{aligned}$$

The first-order condition (FOC) for μ_2 implies $\mathbb{E}[(\mu_2^* - \gamma^*(X_1 - \mu_1^*) - \mu_2^\bullet + r(X_1 - \mu_1^\bullet)) | X_1 \in K] = 0$ at the optimum $(\mu_1^*, \mu_2^*, \gamma^*)$. Also, we may drop terms without X_1 in the conditional covariance operator, and we get

$$\frac{\partial \xi}{\partial \gamma}(\mu_1^*, \mu_2^*, \gamma^*) = \mathbb{P}[X_1 \in K] \cdot (\gamma^* - r) \cdot \text{Cov}(X_1, X_1 | X_1 \in K).$$

We have $\mathbb{P}[X_1 \in K] > 0$ and $\text{Cov}(X_1, X_1 | X_1 \in K) > 0$; hence we conclude

$$\frac{\partial \xi}{\partial \gamma}(\mu_1^*, \mu_2^*, \gamma^*) \begin{cases} > 0 & \text{for } \gamma^* > r \\ = 0 & \text{for } \gamma^* = r \\ < 0 & \text{for } \gamma^* < r. \end{cases}$$

When $r \in [\gamma_l, \gamma_h]$, $(\mu_1^*, \mu_2^*, \gamma^*)$ cannot minimize ξ if $\gamma^* \neq r$ at either end point where FOC in γ does not hold, ξ can be strictly reduced by changing γ slightly. In case that $\gamma_l > r$, at the optimum we must have $\frac{\partial \xi}{\partial \gamma}(\mu_1^*, \mu_2^*, \gamma^*) > 0$. By the Karush–Kuhn–Tucker condition, this means the minimizer is $\gamma^* = \gamma_l$. Conversely, when $\gamma_h < r$, at the optimum we must have $\frac{\partial \xi}{\partial \gamma}(\mu_1^*, \mu_2^*, \gamma^*) < 0$. In that case, the minimizer is $\gamma^* = \gamma_h$. So in both cases, $\gamma^* = \gamma_n$ as desired.

Finally, by using $\mu_2^* = \mu_2^\bullet + (r - \gamma_n) \cdot (\mu_1^\bullet - \mathbb{E}[X_1 | X_1 \in K])$ and specializing to the case where the continuation region K is either $(-\infty, c]$ or $[c, \infty)$, we get the closed-form expression of $\mu_2^*(c)$.

A.2 Proving Lemma 1

I state and prove a stronger result, which will be used in some of the later proofs.

LEMMA A.1. *Consider the model $\Psi(\mu_1, \mu_2; \gamma)$ for any $\mu_1, \mu_2, \gamma \in \mathbb{R}$. Let $D(x_1)$ be the difference between the expected payoff in stopping and continuing after $X_1 = x_1$ in the model. If $\gamma = -1$, then $D(x_1)$ is constant in x_1 . If $\gamma > -1$, then $D(x_1)$ is continuous and strictly increasing in x_1 with $\lim_{x_1 \rightarrow \pm\infty} D(x_1) = \pm\infty$. If $\gamma < -1$, then $D(x_1)$ is continuous and strictly decreasing in x_1 with $\lim_{x_1 \rightarrow \pm\infty} D(x_1) = \mp\infty$. When $\gamma \neq -1$, there is a unique $C(\mu_1, \mu_2; \gamma)$ so that the agent is indifferent between continuing and stopping after $x_1 = C(\mu_1, \mu_2; \gamma)$. For fixed $\mu_1 \in \mathbb{R}$, $\gamma \neq -1$, the function $\mu_2 \mapsto C(\mu_1, \mu_2; \gamma)$ is linear with a slope of $1/(\gamma + 1)$.*

Using Lemma A.1, agents stop after high values of X_1 when $\gamma > -1$ and stop after low values of X_1 when $\gamma < -1$, because D is strictly increasing when $\gamma > -1$ and strictly decreasing when $\gamma < -1$. Also, since $\mu_2 \mapsto C(\mu_1, \mu_2; \gamma)$ has a slope of $1/(\gamma + 1)$, it is strictly increasing if $\gamma > -1$ and strictly decreasing if $\gamma < -1$.

PROOF OF LEMMA A.1. In the model $\Psi(\mu_1, \mu_2; \gamma)$, the expected difference between stopping and continuing after $X_1 = x_1$ is

$$D(x_1) = x_1 - q\mathbb{E}[\max(x_1, [X_2 | x_1])] - (1 - q)\mathbb{E}[X_2 | x_1] + \kappa,$$

where $[X_2 | x_1] \sim \mathcal{N}(\mu_2 - \gamma(x_1 - \mu_1), \sigma^2)$. This is clearly continuous in x_1 . When $\gamma = -1$, D is constant because we have for every $a \in \mathbb{R}$, $\delta > 0$,

$$D(a + \delta) - D(a) = \delta - q\{\mathbb{E}[\max(a + \delta, [X_2 | a + \delta])] - \mathbb{E}[\max(a, [X_2 | a])]\} - \delta(1 - q).$$

In comparing $\max(a + \delta, [X_2 | a + \delta])$ and $\max(a, [X_2 | a])$, note the distribution $[X_2 | a + \delta]$ is $[X_2 | a]$ shifted to the right by δ , so the distribution $\max(a + \delta, [X_2 | a + \delta])$ is just $\max(a, [X_2 | a])$ shifted to the right by δ . Thus, $\mathbb{E}[\max(a + \delta, [X_2 | a + \delta])] - \mathbb{E}[\max(a, [X_2 | a])] = \delta$. So overall, $D(a + \delta) - D(a) = 0$.

When $\gamma > -1$, $[X_2 | a + \delta]$ is strictly stochastically dominated by $\delta + [X_2 | a]$; therefore, $\mathbb{E}[\max(a + \delta, [X_2 | a + \delta])] < \delta + \mathbb{E}[\max(a, [X_2 | a])]$. Also, we have $\mathbb{E}[X_2 | a + \delta] - \mathbb{E}[X_2 | a] = -\gamma\delta$, so we get $D(a + \delta) - D(a) > (1 - q)(1 + \gamma)\delta > 0$. This shows D is strictly increasing at a rate of at least $(1 - q)(1 + \gamma)$ at every point in the domain; therefore, $\lim_{x_1 \rightarrow \pm\infty} D(x_1) = \pm\infty$.

When $\gamma < -1$, $[X_2 | a + \delta]$ strictly stochastically dominates $\delta + [X_2 | a]$; therefore, $\mathbb{E}[\max(a + \delta, [X_2 | a + \delta])] > \delta + \mathbb{E}[\max(a, [X_2 | a])]$. Also, we have $\mathbb{E}[X_2 | a + \delta] - \mathbb{E}[X_2 | a] = -\gamma\delta$, so we get $D(a + \delta) - D(a) < (1 - q)(1 + \gamma)\delta < 0$. This shows D is strictly decreasing at a rate of at least $(1 - q)(1 + \gamma)$ at every point in the domain; therefore, $\lim_{x_1 \rightarrow \pm\infty} D(x_1) = \mp\infty$.

When $\gamma \neq -1$, the existence and uniqueness of $C(\mu_1, \mu_2; \gamma)$ come from the fact that D is strictly monotonic and takes on both positive and negative values, so it must cross 0 at a unique point.

In fact, $C(\mu_1, \mu_2; \gamma)$ is linear in μ_2 with a coefficient of $1/(\gamma + 1)$. To see this, fix μ_1 and γ , and consider the difference $x_1 - q\mathbb{E}[\max(x_1, [X_2 | x_1])] - (1 - q)\mathbb{E}[X_2 | x_1] + \kappa$ as a function $G(x_1, \mu_2)$ of x_1 and μ_2 . For every $\delta > 0$, we have $G(x_1 + \delta/(\gamma + 1), \mu_2 + \delta) = G(x_1, \mu_2)$. This is because

$$\begin{aligned} \mathcal{N}\left((\mu_2 + \delta) - \gamma\left(x_1 + \frac{\delta}{\gamma + 1}\right) - \mu_1, \sigma^2\right) &= \mathcal{N}(\mu_2 - \gamma(x_1 - \mu_1), \sigma^2) + \delta - \delta\frac{\gamma}{\gamma + 1} \\ &= \mathcal{N}(\mu_2 - \gamma(x_1 - \mu_1), \sigma^2) + \delta\frac{1}{\gamma + 1}; \end{aligned}$$

therefore, $q\mathbb{E}_{\mu_2 + \delta}[\max(x_1 + \delta/(\gamma + 1), [X_2 | x_1])] = q(\mathbb{E}_{\mu_2}[\max(x_1, [X_2 | x_1])] + \delta/(\gamma + 1))$. Also $(1 - q)\mathbb{E}_{\mu_2 + \delta}[X_2 | x_1] = (1 - q)(\mathbb{E}_{\mu_2}[X_2 | x_1] + \delta/(\gamma + 1))$. Using these two facts,

$$G\left(x_1 + \frac{\delta}{\gamma + 1}, \mu_2 + \delta\right) - G(x_1, \mu_2) = \frac{\delta}{\gamma + 1} - q\left(\delta\frac{1}{\gamma + 1}\right) - (1 - q)\left(\delta\frac{1}{\gamma + 1}\right) = 0.$$

That is, increasing belief about μ_2 by δ and also increasing the realization of the early draw by $\delta/(\gamma + 1)$ cancel each other out in terms of the difference between the expected payoffs in stopping and continuing. Therefore, we must have $C(\mu_1, \mu_2 + \delta; \gamma) = C(\mu_1, \mu_2; \gamma) + \delta/(\gamma + 1)$. $\square$

A.3 Proof of Proposition 2

Consider the map $\mathcal{I} : \mathbb{R} \rightarrow \mathbb{R}$ defined by $\mathcal{I}(\mu_2) := \mu_2^*(C(\mu_1^\bullet, \mu_2; \gamma_n))$, where we define $\mu_2^*(c) = \mu_2^\bullet + (r - \gamma_n) \cdot (\mu_1^\bullet - \mathbb{E}[X_1 | X_1 \leq c])$ if $\gamma_n > -1$ and $\mu_2^*(c) = \mu_2^\bullet + (r - \gamma_n) \cdot (\mu_1^\bullet - \mathbb{E}[X_1 | X_1 \geq c])$ if $\gamma_n < -1$. Lemma A.1 shows $\mu_2 \mapsto C(\mu_1^\bullet, \mu_2; \gamma_n)$ is linear with a slope of $1/(\gamma_n + 1)$. Also, by a property of the Gaussian distribution, both $c \mapsto \mathbb{E}[X_1 | X_1 \leq c]$ and $c \mapsto \mathbb{E}[X_1 | X_1 \geq c]$ are Lipschitz continuous with a Lipschitz constant of 1. Therefore, the composition $\mathcal{I}$ is Lipschitz continuous with a Lipschitz constant of $|(r - \gamma_n)/(1 + \gamma_n)| < 1$, hence a contraction map. By a property of contraction maps, $\mathcal{I}$ has a unique fixed point, which we denote μ_2^∞ . When $\gamma_n > -1$, the beliefs $(\mu_1^\infty, \mu_2^\infty, \gamma_n)$ together with the cutoff strategy that stops when $c \geq C(\mu_1^\infty, \mu_2^\infty; \gamma_n)$ make up a steady state by Proposition 1 and Lemma 1. When $\gamma_n < -1$, the beliefs $(\mu_1^\infty, \mu_2^\infty, \gamma_n)$ together with the cutoff strategy that stops when $c \leq C(\mu_1^\infty, \mu_2^\infty; \gamma_n)$ make up a steady state for the same reason. Also, this steady state is unique. By Proposition 1, in any steady-state beliefs $(\mu'_1, \mu'_2, \gamma')$, we must have $\mu'_1 = \mu_1^\bullet$, $\gamma' = \gamma_n$. This implies μ'_2 must be a fixed point of $\mathcal{I}$ by the optimality of behavior and the KL divergence minimization of beliefs, yet μ_2^∞ is the unique fixed point of $\mathcal{I}$.

A.4 Proof of Proposition 3

Under the condition $|(r - \gamma_n)/(1 + \gamma_n)| < 1$, by Proposition 2 there exists a unique steady state where $\gamma^\infty = \gamma_n$ and the agent uses a cutoff strategy with some threshold c^∞ . The agent stops when $X_1 \geq c^\infty$ if $\gamma_n > -1$ and stops when $X_1 \leq c^\infty$ if $\gamma_n < -1$.

Suppose $r, \gamma_n > -1$. Then by Proposition 1, $\mu_2^\infty = \mu_2^\bullet + (r - \gamma_n) \cdot (\mu_1^\bullet - \mathbb{E}[X_1 | X_1 \leq c^\infty])$. Since $\mathbb{E}[X_1 | X_1 \leq c^\infty] < c^\infty$, we get $\mu_2^\infty < \mu_2^\bullet + (r - \gamma_n) \cdot (\mu_1^\bullet - c^\infty) \iff \mu_2^\infty - \gamma_n(c^\infty - \mu_1^\bullet) < \mu_2^\bullet - r(c^\infty - \mu_1^\bullet)$ if $r - \gamma_n < 0$ and, symmetrically, $\mu_2^\infty - \gamma_n(c^\infty - \mu_1^\bullet) > \mu_2^\bullet - r(c^\infty - \mu_1^\bullet)$ if $r - \gamma_n > 0$. In the $r - \gamma_n < 0$ case, it shows the agent's belief about the second-period mean of X_2 conditional on $X_1 = c^\infty$ is strictly lower than the truth. As the agent who believes in the model $\Psi(\mu_1^\bullet, \mu_2^\infty; \gamma_n)$ is indifferent between continuing and stopping after $X_1 = c^\infty$, an agent who believes in the model $\Psi(\mu_1^\bullet, \mu_2^\bullet; r)$ finds it strictly better to continue after $X_1 = c^\infty$. Under the model $\Psi(\mu_1^\bullet, \mu_2^\bullet; r)$ with $r > -1$, by Lemma A.1 the agent strictly prefers continuing only at those c with $c < C(\mu_1^\bullet, \mu_2^\bullet; r) = c^\bullet$, which shows $c^\infty < c^\bullet$. The $r - \gamma_n > 0$ case symmetrically leads to the conclusion that $c^\infty > c^\bullet$.

Suppose both $r, \gamma_n < -1$. Then by Proposition 1, $\mu_2^\infty = \mu_2^\bullet + (r - \gamma_n) \cdot (\mu_1^\bullet - \mathbb{E}[X_1 | X_1 \geq c^\infty])$. Since $\mathbb{E}[X_1 | X_1 \geq c^\infty] > c^\infty$, we get $\mu_2^\infty > \mu_2^\bullet + (r - \gamma_n) \cdot (\mu_1^\bullet - c^\infty) \iff \mu_2^\infty - \gamma_n(c^\infty - \mu_1^\bullet) > \mu_2^\bullet - r(c^\infty - \mu_1^\bullet)$ if $r - \gamma_n < 0$ and, symmetrically, $\mu_2^\infty - \gamma_n(c^\infty - \mu_1^\bullet) < \mu_2^\bullet - r(c^\infty - \mu_1^\bullet)$ if $r - \gamma_n > 0$. In the $r - \gamma_n < 0$ case, it shows the agent's belief about the second-period mean of X_2 conditional on $X_1 = c^\infty$ is strictly higher than the truth. As the agent who believes in the model $\Psi(\mu_1^\bullet, \mu_2^\infty; \gamma_n)$ is indifferent between continuing and stopping after $X_1 = c^\infty$, an agent who believes in the model $\Psi(\mu_1^\bullet, \mu_2^\bullet; r)$ finds it strictly better to stop after $X_1 = c^\infty$. Under the model $\Psi(\mu_1^\bullet, \mu_2^\bullet; r)$ with $r < -1$, by Lemma A.1 the agent strictly prefers stopping only at those c with $c < C(\mu_1^\bullet, \mu_2^\bullet; r) = c^\bullet$, which shows $c^\infty < c^\bullet$. The $r - \gamma_n > 0$ case symmetrically leads to the conclusion that $c^\infty > c^\bullet$.

A.5 Proof of Proposition 4

I will show a stronger statement. Given a pair of second-period payoff functions u_2^H and u_2^L , say u_2^H *payoff dominates* u_2^L (abbreviated $u_2^H \succ u_2^L$) if for every $x_1 \in \mathbb{R}$, $u_2^H(x_1, x_2) \geq u_2^L(x_1, x_2)$ for every $x_2 \in \mathbb{R}$, and also $u_2^H(x_1, x_2) > u_2^L(x_1, x_2)$ for a positive-measure set of x_2 in $\mathbb{R}$. It is clear that increasing q or decreasing κ in the statement of Proposition 4 leads to a payoff dominating game. There is a unique steady state for any (q, κ) by Proposition 2 since $r = 0$ and $\gamma_n > 0$. The next part of Proposition 4 is implied by the following proposition.

PROPOSITION A.1. *Let $r = 0$ and $\gamma_n > 0$. Suppose both (u_1, u_2^H) and (u_1, u_2^L) correspond to stage games with some (q, κ) , and that $u_2^H \succ u_2^L$. The steady state of (u_1, u_2^H) features strictly more optimistic belief about the second-period fundamental and a strictly higher cutoff threshold than the steady state of (u_1, u_2^L) .*

I require an auxiliary lemma for the proof.

LEMMA A.2. *Suppose both (u_1, u_2^H) and (u_1, u_2^L) correspond to stage games with some (q, κ) , and that $u_2^H \succ u_2^L$. For all $\mu_1, \mu_2 \in \mathbb{R}$, $\gamma > 0$, $C_{u_1, u_2^H}(\mu_1, \mu_2; \gamma) > C_{u_1, u_2^L}(\mu_1, \mu_2; \gamma)$.*

PROOF. Indifference $c^L = C_{u_1, u_2^L}(\mu_1, \mu_2; \gamma)$ implies $u_1(c^L) = \mathbb{E}_{\tilde{X}_2 \sim \phi(\cdot | \mu_2 - \gamma(c^L - \mu_1))} [u_2^L(c^L, \tilde{X}_2)]$. Since $u_2^H(c^L, x_2) \geq u_2^L(c^L, x_2)$ for all $x_2 \in \mathbb{R}$, with strict inequality on a positive-measure set, this shows $u_1(c^L) < \mathbb{E}_{\tilde{X}_2 \sim \phi(\cdot | \mu_2 - \gamma(c^L - \mu_1))} [u_2^H(c^L, \tilde{X}_2)]$. The best stopping strategy in the model $\Psi(\mu_1, \mu_2; \gamma)$ with the utility functions (u_1, u_2^H) has a cutoff form by Lemma A.1. This shows $C_{u_1, u_2^H}(\mu_1, \mu_2; \gamma)$ is strictly above c^L . $\square$

PROOF OF PROPOSITION A.1. Now to the proof of Proposition A.1. Say the unique steady states under (u_1, u_2^H) and (u_1, u_2^L) are $((\mu_1^\bullet, \mu_{2,H}^\infty, \gamma_n), c_H^\infty)$ and $((\mu_1^\bullet, \mu_{2,L}^\infty, \gamma_n), c_L^\infty)$, respectively. Let $\mathcal{I}_H$ and $\mathcal{I}_L$ be the iteration maps corresponding to these two stage games, that is to say,

$$\mathcal{I}_H(\mu_2) := \mu_2^*(C_{u_1, u_2^H}(\mu_1^\bullet, \mu_2; \gamma_n))$$

$$\mathcal{I}_L(\mu_2) := \mu_2^*(C_{u_1, u_2^L}(\mu_1^\bullet, \mu_2; \gamma_n)).$$

From the proof of Proposition 2, both $\mathcal{I}_H$ and $\mathcal{I}_L$ are contraction maps. Consider their iterates with a starting value of 0. That is, put $\mu_{2,H}^{[0]} = 0$, $\mu_{2,L}^{[0]} = 0$, and let $\mu_{2,H}^{[t]} = \mathcal{I}_H(\mu_{2,H}^{[t-1]})$, $\mu_{2,L}^{[t]} = \mathcal{I}_L(\mu_{2,L}^{[t-1]})$ for $t \geq 1$. By a property of contraction maps and since the fixed points of the iteration maps are the steady-state beliefs, $\mu_{2,H}^{[t]} \rightarrow \mu_{2,H}^\infty$ and $\mu_{2,L}^{[t]} \rightarrow \mu_{2,L}^\infty$.

By induction, I will show $\mu_{2,L}^{[t]} \leq \mu_{2,H}^{[t]}$ for every $t \geq 0$. The base case of $t = 0$ is true by definition. If $\mu_{2,L}^{[T]} \leq \mu_{2,H}^{[T]}$, then

$$C_{u_1, u_2^L}(\mu_1^\bullet, \mu_{2,L}^{[T]}; \gamma) \leq C_{u_1, u_2^L}(\mu_1^\bullet, \mu_{2,H}^{[T]}; \gamma) < C_{u_1, u_2^H}(\mu_1^\bullet, \mu_{2,H}^{[T]}; \gamma).$$

The first inequality comes from C being increasing in the second argument and the inductive hypothesis, while the second inequality is due to Lemma A.2. Therefore,

$\mathcal{I}_L(\mu_{2,L}^{[T]}) \leq \mathcal{I}_H(\mu_{2,H}^{[T]})$ using the fact that μ_2^* is increasing by Proposition 1, so $\mu_{2,L}^{[T+1]} \leq \mu_{2,H}^{[T+1]}$.

Since weak inequalities are preserved by limits, we have $\mu_{2,H}^\infty \geq \mu_{2,L}^\infty$. It is impossible to have $\mu_{2,H}^\infty = \mu_{2,L}^\infty$, because this would lead to $c_H^\infty > c_L^\infty$ by Lemma A.2, which in turn implies $\mu_{2,H}^\infty = \mu_2^*(c_H^\infty) > \mu_2^*(c_L^\infty) = \mu_{2,L}^\infty$. This inequality contradicts $\mu_{2,H}^\infty = \mu_{2,L}^\infty$. Therefore, we in fact have $\mu_{2,H}^\infty > \mu_{2,L}^\infty$. The conclusion that $c_H^\infty > c_L^\infty$ follows from Lemma A.2 and the fact that C increases in its second argument. $\square$

A.6 Proof of Proposition 5

Rewrite (2) as

$$\begin{aligned} & \int_{-\infty}^{\infty} \phi(x_1 | \mu_1^\bullet, (\sigma^\bullet)^2) \cdot \ln \left(\frac{\phi(x_1 | \mu_1^\bullet, (\sigma^\bullet)^2)}{\phi(x_1 | \mu_1, \sigma_1^2)} \right) dx_1 \\ & + \int_{-\infty}^c \phi(x_1 | \mu_1^\bullet, (\sigma^\bullet)^2) \\ & \cdot \int_{-\infty}^{\infty} \phi(x_2 | \mu_2^\bullet, (\sigma^\bullet)^2) \ln \left[\frac{\phi(x_2 | \mu_2^\bullet, (\sigma^\bullet)^2)}{\phi(x_2 | \mu_2 - \gamma(x_1 - \mu_1), \sigma_2^2)} \right] dx_2 dx_1. \end{aligned}$$

KL divergence between $\mathcal{N}(\mu_{\text{true}}, \sigma_{\text{true}}^2)$ and $\mathcal{N}(\mu_{\text{model}}, \sigma_{\text{model}}^2)$ is $\ln(\sigma_{\text{model}}/\sigma_{\text{true}}) + [(\sigma_{\text{true}}^2 + (\mu_{\text{true}} - \mu_{\text{model}})^2)/(2\sigma_{\text{model}}^2)] - \frac{1}{2}$, so we may simplify the first term and the inner integral of the second term:

$$\begin{aligned} & \ln \frac{\sigma_1}{\sigma^\bullet} + \frac{(\mu_1 - \mu_1^\bullet)^2}{2\sigma_1^2} + \frac{(\sigma^\bullet)^2}{2\sigma_1^2} - \frac{1}{2} \\ & + \int_{-\infty}^c \phi(x_1 | \mu_1^\bullet, \sigma^\bullet) \cdot \left[\ln \frac{\sigma_2}{\sigma^\bullet} + \frac{(\sigma^\bullet)^2 + (\mu_2 - \gamma(x_1 - \mu_1) - \mu_2^\bullet)^2}{2\sigma_2^2} - \frac{1}{2} \right] dx_1. \end{aligned}$$

Dropping terms not dependent on any of the four variables gives a simplified version of the objective:

$$\begin{aligned} & \xi(\mu_1, \mu_2, \sigma_1, \sigma_2) \\ & := \ln \frac{\sigma_1}{\sigma^\bullet} + \frac{(\mu_1 - \mu_1^\bullet)^2}{2\sigma_1^2} + \frac{(\sigma^\bullet)^2}{2\sigma_1^2} \\ & + \int_{-\infty}^c \phi(x_1 | \mu_1^\bullet, (\sigma^\bullet)^2) \cdot \left[\ln \frac{\sigma_2}{\sigma^\bullet} + \frac{(\sigma^\bullet)^2 + (\mu_2 - \gamma(x_1 - \mu_1) - \mu_2^\bullet)^2}{2\sigma_2^2} \right] dx_1. \end{aligned}$$

Differentiating under the integral sign,

$$\frac{\partial \xi}{\partial \mu_2} = \int_{-\infty}^c \phi(x_1 | \mu_1^\bullet, (\sigma^\bullet)^2) \cdot \left[\frac{(\mu_2 - \gamma(x_1 - \mu_1) - \mu_2^\bullet)}{\sigma_2^2} \right] dx_1$$

$$\begin{aligned}\frac{\partial \xi}{\partial \mu_1} &= \frac{(\mu_1 - \mu_1^*)}{\sigma_1^2} + \gamma \int_{-\infty}^c \phi(x_1 | \mu_1^*, (\sigma^*)^2) \cdot \left[\frac{(\mu_2 - \gamma(x_1 - \mu_1) - \mu_2^*)}{\sigma_2^2} \right] dx_1 \\ &= \frac{(\mu_1 - \mu_1^*)}{\sigma_1^2} + \gamma \frac{\partial \xi}{\partial \mu_2}.\end{aligned}$$

At FOC $(\mu_1^*, \mu_2^*, \sigma_1^*, \sigma_2^*)$, we have $\frac{\partial \xi}{\partial \mu_2}(\mu_1^*, \mu_2^*, \sigma_1^*, \sigma_2^*) = 0$, hence $\mu_1^* = \mu_1^*$. Similar arguments as before then establish $\mu_2^* = \mu_2^* - \gamma(\mu_1^* - \mathbb{E}[X_1 | X_1 \leq c])$, where expectation is taken with respect to the true distribution of X_1 (with the true variance $(\sigma^*)^2$). Then $\frac{\partial \xi}{\partial \sigma_1}(\mu_1^*, \mu_2^*, \sigma_1^*, \sigma_2^*) = (1/\sigma_1^*) - ((\sigma^*)^2/(\sigma_1^*)^3) = 0$, which gives $\sigma_1^* = \sigma^*$ (since $\sigma_1^* \geq 0$).

Finally, from the FOC for σ_2 ,

$$\int_{-\infty}^c \phi(x_1; \mu_1^*, (\sigma^*)^2) \cdot \left[\frac{1}{\sigma_2^*} - \frac{(\sigma^*)^2 + (\mu_2^* - \gamma(x_1 - \mu_1^*) - \mu_2^*)^2}{(\sigma_2^*)^3} \right] dx_1 = 0.$$

Substituting in values of μ_1^* and μ_2^* already solved for,

$$\begin{aligned}(\sigma_2^*)^2 &= (\sigma^*)^2 + \mathbb{E}[(\mu_2^* - \gamma(X_1 - \mu_1^*) - \mu_2^*)^2 | X_1 \leq c] \\ &= (\sigma^*)^2 + \mathbb{E}[(\mu_2^* - \gamma(\mu_1^* - \mathbb{E}[X_1 | X_1 \leq c]) - \gamma(X_1 - \mu_1^*) - \mu_2^*)^2 | X_1 \leq c] \\ &= (\sigma^*)^2 + \gamma^2 \mathbb{E}[(X_1 - \mu_1^*) - (\mathbb{E}[X_1 | X_1 \leq c] - \mu_1^*)]^2 | X_1 \leq c] \\ &= (\sigma^*)^2 + \gamma^2 \text{Var}[X_1 | X_1 \leq c]\end{aligned}$$

as desired. Finally, $\sigma_2^*(c)$ is an increasing function of c because $\text{Var}[X_1 | X_1 \leq c]$ increases in c for X_1 Gaussian (Mailhot (1985)).

A.7 Proving Proposition 6

I start with a lemma that says if the decision problem is convex, a stronger belief in fictitious variation increases the subjectively optimal cutoff threshold.

LEMMA A.3. *Suppose that under the feasible model $\Psi(\mu_1, \mu_2, \sigma_1^2, \sigma_2^2; \gamma)$, the agent is indifferent between stopping at c and continuing. Suppose $\hat{\sigma}_2^2 > \sigma_2^2$. Then if $x_2 \mapsto u_2(c, x_2)$ is convex with strict convexity for x_2 in a positive-measure set, then under the feasible model $\Psi(\mu_1, \mu_2, \sigma_1^2, \hat{\sigma}_2^2; \gamma)$, the agent strictly prefers continuing at c .*

PROOF. Indifference at $x_1 = c$ under $\Psi(\mu_1, \mu_2, \sigma_1^2, \sigma_2^2; \gamma)$ implies

$$u_1(c) = \mathbb{E}_{X_2 \sim \mathcal{N}(\mu_2 - \gamma(x_1 - \mu_1), \sigma_2^2)}[u_2(c, X_2)].$$

When hypothesis is satisfied,

$$\mathbb{E}_{X_2 \sim \mathcal{N}(\mu_2 - \gamma(x_1 - \mu_1), \sigma_2^2)}[u_2(c, X_2)] < \mathbb{E}_{X_2 \sim \mathcal{N}(\mu_2 - \gamma(x_1 - \mu_1), \hat{\sigma}_2^2)}[u_2(c, X_2)]$$

since $\hat{\sigma}_2^2 > \sigma_2^2$ implies that $\mathcal{N}(\mu_2 - \gamma(x_1 - \mu_1), \hat{\sigma}_2^2)$ is a strict mean-preserving spread of $\mathcal{N}(\mu_2 - \gamma(x_1 - \mu_1), \sigma_2^2)$. The right-hand side is the expected continuation payoff under model $\Psi(\mu_1, \mu_2, \sigma_1^2, \hat{\sigma}_2^2; \gamma)$, so the agent strictly prefers continuing when $X_1 = c$. $\square$

Now I give the proof of Proposition 6.

PROOF OF PROPOSITION 6. By the proof of Proposition 2, $\mathcal{I}(\mu_2; \gamma) := \mu_2^*(C(\mu_1^\bullet, \mu_2, \gamma))$ for society A is a contraction map in μ_2 . By way of contradiction, suppose $c^B \leq c^A$. Then $\mu_2^B \leq \mu_2^A$ by Proposition 5. In society A, $C(\mu_1^\bullet, \mu_2^B; \gamma) < c^B$ by Lemma A.3, as there is strictly positive probability of recall. This shows $\mathcal{I}(\mu_2^B; \gamma) < \mu_2^B$. In fact, for the t -times iteration, we have $\mathcal{I}^{(t)}(\mu_2^B; \gamma) \leq \mathcal{I}(\mu_2^B; \gamma) < \mu_2^B$, which means $\mathcal{I}$ has a fixed point strictly smaller than μ_2^A . This contradicts μ_2^A being the only fixed point of $\mathcal{I}$. Hence we must have $\mu_2^B > \mu_2^A$ and $c^B > c^A$. We have $\sigma_2^B = \sigma_2^*(c^B)$, which is larger than $\sigma_2^*(c^A)$ by combining $c^B > c^A$ with Proposition 5. $\square$

A.8 Proving Proposition 7

I introduce some new notation. Abbreviate $\square := [\underline{\mu}_1, \bar{\mu}_1] \times [\underline{\mu}_2, \bar{\mu}_2]$. Let $\gamma = \gamma_n$ and let $\text{li}(\mu_2)$ be the line in $\mathbb{R}^2$ with slope $-\gamma$ that passes through the point $(\mu_1^\bullet, \mu_2)$. There are some minimal and maximal $\underline{\mu}_2^\circ$ and $\bar{\mu}_2^\circ$ so that $\text{li}(\underline{\mu}_2^\circ) \cap \square \neq \emptyset$ and $\text{li}(\bar{\mu}_2^\circ) \cap \square \neq \emptyset$. Finally, for $\mu_2^l < \mu_2^h$, let $\diamond[\mu_2^l, \mu_2^h] := \bigcup_{\mu_2 \in [\mu_2^l, \mu_2^h]} \text{li}(\mu_2)$. So we have $\square \subseteq \diamond[\underline{\mu}_2^\circ, \bar{\mu}_2^\circ]$. Similarly the half-open versions $\diamond[\mu_2^l, \mu_2^h)$ and $\diamond(\mu_2^l, \mu_2^h]$ are defined as the unions $\bigcup_{\mu_2 \in [\mu_2^l, \mu_2^h)}$ and $\bigcup_{\mu_2 \in (\mu_2^l, \mu_2^h]}$ $\text{li}(\mu_2)$. Figure 3 illustrates a case with $\gamma > 0$.

A.8.1 Preliminary results First I consider how the predicted second-period payoff after $X_1 = x_1$ depends on the parameters of the feasible model $\Psi(\mu_1, \mu_2; \gamma)$.

LEMMA A.4. For every $\mu_1, \mu_2, x_1 \in \mathbb{R}$, the conditional distribution $X_2|X_1 = x_1$ is the same under $\Psi(\mu_1^\bullet, \mu_2 + \gamma(\mu_1 - \mu_1^\bullet); \gamma)$ and $\Psi(\mu_1, \mu_2; \gamma)$. So in particular, $C(\mu_1, \mu_2; \gamma) = C(\mu_1^\bullet, \mu_2 + \gamma(\mu_1 - \mu_1^\bullet); \gamma)$.

PROOF. Under the feasible model $\Psi(\mu_1^\bullet, \mu_2 + \gamma(\mu_1 - \mu_1^\bullet); \gamma)$, the conditional density of X_2 given $X_1 = x_1$ is $\phi(\cdot | \mu_2 + \gamma(\mu_1 - \mu_1^\bullet) - \gamma(x_1 - \mu_1^\bullet))$, which simplifies to $\phi(\cdot | \mu_2 - \gamma(x_1 - \mu_1))$. It is easy to see that this is also the expression for the same conditional density under $\Psi(\mu_1, \mu_2; \gamma)$.

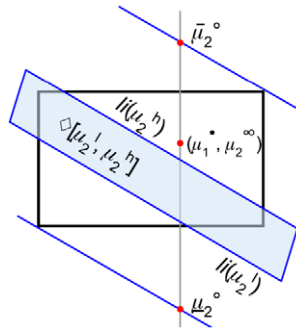


FIGURE 3. Notation for the proof of Proposition 7.

Suppose $C(\mu_1, \mu_2; \gamma) = c$. This implies $u_1(c) = \mathbb{E}_{\Psi(\mu_1, \mu_2; \gamma)}[u_2(c, X_2) | X_1 = c]$. But by the equivalence of conditional distribution given above,

$$u_1(c) = \mathbb{E}_{\Psi(\mu_1^\bullet, \mu_2 + \gamma(\mu_1 - \mu_1^\bullet); \gamma)}[u_2(c, X_2) | X_1 = c].$$

This means c is also the indifference threshold for the model $\Psi(\mu_1^\bullet, \mu_2 + \gamma(\mu_1 - \mu_1^\bullet); \gamma)$. $\square$

As a corollary, this lemma shows the restriction to cutoff strategies is without loss and that $\tilde{C}_l$ is well defined. That is, for any belief given by a density on $\square$, there exists a cutoff strategy that is weakly optimal among the class of all stopping strategies and, further, this cutoff strategy is strictly optimal among the class of cutoff strategies. This is because for any $x_1 \in \mathbb{R}$ and any density $\tilde{m}$ on $\square$,

$$\begin{aligned} & \int_{\square} \mathbb{E}_{\Psi(\mu_1, \mu_2; \gamma)}[u_2(x_1, X_2) | X_1 = x_1] \cdot \tilde{m}(\mu_1, \mu_2) d(\mu_1, \mu_2) \\ &= \int_{\underline{\mu}_2^\circ}^{\bar{\mu}_2^\circ} \mathbb{E}_{\Psi(\mu_1^\bullet, \mu_2; \gamma)}[u_2(x_1, X_2) | X_1 = x_1] \cdot \tilde{m}^V(\mu_2) d\mu_2, \end{aligned}$$

where $\tilde{m}^V(\mu_2)$ is the integral of $\tilde{m}(\mu_1, \mu_2)$ over $\text{li}(\mu_2)$. This equality holds because by Lemma A.4, all fundamentals on $\text{li}(\mu_2)$ imply the same continuation payoff after $X_1 = x_1$ as the fundamentals $(\mu_1^\bullet, \mu_2)$.

LEMMA A.5. *If $\gamma > -1$, then the function*

$$x_1 \mapsto u_1(x_1) - \int_{\underline{\mu}_2^\circ}^{\bar{\mu}_2^\circ} \mathbb{E}_{\Psi(\mu_1^\bullet, \mu_2; \gamma)}[u_2(x_1, X_2) | X_1 = x_1] \tilde{m}^V(\mu_2) d\mu_2$$

is strictly increasing, continuous, and crosses 0.

PROOF. Let $d\nu(\mu_2) = \tilde{m}^V(\mu_2) d\mu_2$. Consider the payoff difference between accepting x_1 and continuing under belief ν :

$$D(x_1; \nu) := u_1(x_1) - \int \mathbb{E}_{X_2 \sim \phi(\cdot | \mu_2 - \gamma(x_1 - \mu_1^\bullet))} [u_2(x_1, X_2)] d\nu(\mu_2).$$

Note that $D(x_1, \nu) = \int D(x_1; \mu_1^\bullet, \mu_2, \gamma) d\nu(\mu_2)$. When $\gamma > -1$, Lemma A.1 shows that for every $\mu_2 \in \mathbb{R}$, $D(x_1; \mu_1^\bullet, \mu_2, \gamma)$ is strictly increasing in x_1 . Hence, the same must hold for $D(x_1, \nu)$.

Lemma A.1 shows there exists some $x'_1 \in \mathbb{R}$ so that $D(x'_1; \mu_1^\bullet, \underline{\mu}_2, \gamma) < 0$ and that there exists some $x''_1 \in \mathbb{R}$ satisfying $D(x''_1; \mu_1^\bullet, \bar{\mu}_2, \gamma) > 0$. Since u_2 increases in its second argument, we also get $D(x'_1; \mu_1^\bullet, \mu_2, \gamma) < 0$ and $D(x''_1; \mu_1^\bullet, \mu_2, \gamma) > 0$ for all $\mu_2 \in [\underline{\mu}_2, \bar{\mu}_2]$. This implies $D(x'_1; \nu) < 0$ and $D(x''_1; \nu) > 0$, as ν is supported on (a subset of) $[\underline{\mu}_2, \bar{\mu}_2]$.

To show $D(x_1; \nu)$ is continuous in x_1 , fix some $\bar{x}_1$. Let $\pi(\mu_2)$ represent the expectation of the absolute value of a normal random variable with mean $\mu_2 - \gamma((\bar{x}_1 - 1) - \mu_1^\bullet)$

and variance σ^2 . Here $\pi(\mu_2)$ is bounded by a constant plus a linear function of μ_2 as we vary μ_2 . For $|x_1 - \bar{x}_1| \leq 1$,

$$|\mathbb{E}_{X_2 \sim \phi(\cdot | \mu_2 - \gamma(x_1 - \mu_1^\bullet))}[u_2(x_1, X_2)]| \leq q(|\bar{x}_1| + 1 + \pi(\mu_2)) + (1 - q)\pi(\mu_2) + |\kappa|,$$

and the right-hand side is a positive and integrable function with respect to $d\nu(\mu_2)$. For a sequence $x_1^{(n)} \rightarrow \bar{x}_1$, the integrand in $D(x_1^{(n)}; \nu)$ is dominated by $q(|\bar{x}_1| + 1 + \pi(\mu_2)) + (1 - q)\pi(\mu_2) + |\kappa|$ for all large enough n , so by the dominated convergence theorem, $D(x_1^{(n)}; \nu) \rightarrow D(\bar{x}_1; \nu)$. So $D(\cdot; \nu)$ is continuous. $\square$

Now the key step is to separate the two-dimensional inference problem into a pair of one-dimensional problems.

A.8.2 Learning $\mu_1^\bullet$ I define the stochastic process of data log likelihood (for a given fundamental). For each $\mu_1, \mu_2 \in \text{supp}(m_0)$, let $\ell_t(\mu_1, \mu_2)(\omega)$ be the log likelihood that the fundamentals are (μ_1, μ_2) and histories $(\tilde{H}_s)_{s \leq t}(\omega)$ are generated by the end of round t . It is given by

$$\ell_t(\mu_1, \mu_2)(\omega) := \ln(m_0(\mu_1, \mu_2)) + \sum_{s=1}^t \ln(\text{lik}(\tilde{H}_s(\omega); \mu_1, \mu_2)),$$

where $\text{lik}(x_1, \emptyset; \mu_1, \mu_2) := \phi(x_1 | \mu_1)$ and $\text{lik}(x_1, x_2; \mu_1, \mu_2) := \phi(x_1 | \mu_1) \cdot \phi(x_2 | \mu_2 - \gamma(x_1 - \mu_1))$. Let $g_1(\cdot)$ and $g_2(\cdot | x_1)$ be the true densities for the distributions of X_1 and $X_2 | (X_1 = x_1)$, incorporating the true parameters $\mu_1^\bullet, \mu_2^\bullet$, and r . Let $f_2(z)$ be the Gaussian distribution with the mean $\mu_2^\bullet$ and variance σ^2 evaluated at z . By simple algebra, we may expand

$$\begin{aligned} \ell_t(\mu_1, \mu_2)(\omega) &= \ln(m_0(\mu_1, \mu_2)) + \sum_{s=1}^t \ln[g_1(X_{1,s}(\omega) - \mu_1 + \mu_1^\bullet)] \\ &\quad + \sum_{s=1}^t \mathbf{1}\{X_{1,s}(\omega) \leq \tilde{C}_s(\omega)\} \cdot \ln[f_2(X_{2,s}(\omega) - \mu_2 + \mu_2^\bullet + \gamma(X_{1,s}(\omega) - \mu_1))]. \end{aligned}$$

I first establish that, without knowing anything about the process $(\tilde{C}_t)$, we can conclude agents either learn $\mu_1^\bullet$ arbitrarily well or they believe in a boundary value of μ_2 ; that is, either $\mu_2 = \underline{\mu}_2$ or $\mu_2 = \bar{\mu}_2$. (We can later rule out these boundary beliefs of μ_2 .)

LEMMA A.6. *Let $\epsilon > 0$ be given. If $\gamma > 0$, then*

$$\begin{aligned} \lim_{t \rightarrow \infty} \tilde{M}_t \{ \square \cap (([\mu_1^\bullet - \epsilon, \mu_1^\bullet + \epsilon] \times \mathbb{R}) \cup ([\underline{\mu}_1, \mu_1^\bullet] \times [\underline{\mu}_2, \underline{\mu}_2 + \epsilon]) \\ \cup ([\mu_1^\bullet, \bar{\mu}_1] \times [\bar{\mu}_2 - \epsilon, \bar{\mu}_2])) \} = 1. \end{aligned}$$

If $\gamma \leq 0$, then

$$\begin{aligned} \lim_{t \rightarrow \infty} \tilde{M}_t \{ \square \cap (([\mu_1^\bullet - \epsilon, \mu_1^\bullet + \epsilon] \times \mathbb{R}) \cup ([\underline{\mu}_1, \mu_1^\bullet] \times [\bar{\mu}_2 - \epsilon, \bar{\mu}_2]) \\ \cup ([\mu_1^\bullet, \bar{\mu}_1] \times [\underline{\mu}_2, \underline{\mu}_2 + \epsilon])) \} = 1. \end{aligned}$$

PROOF. First calculate the directional derivative $\nabla_v \frac{1}{t} \ell_t(\mu_1, \mu_2)$, where $v = (1/\sqrt{1+\gamma^2}, -\gamma/\sqrt{1+\gamma^2})'$ is the unit vector with slope $-\gamma$. We have

$$\begin{aligned} \frac{\partial(\ell_t/t)}{\partial\mu_1}(\mu_1, \mu_2) &= \frac{1}{t} \frac{D_1 m_0(\mu_1, \mu_2)}{m_0(\mu_1, \mu_2)} - \frac{1}{t} \sum_{s=1}^t \frac{g'_1(X_{1,s} - \mu_1 + \mu_1^\bullet)}{g_1(X_{1,s} - \mu_1 + \mu_1^\bullet)} \\ &\quad - \frac{\gamma}{t} \sum_{s=1}^t \mathbf{1}\{X_{1,s} \leq \tilde{C}_s\} \cdot \lambda(X_{2,s} - \mu_2 + \mu_2^\bullet + \gamma(X_{1,s} - \mu_1)) \\ \frac{\partial(\ell_t/t)}{\partial\mu_2}(\mu_1, \mu_2) &= \frac{1}{t} \frac{D_2 m_0(\mu_1, \mu_2)}{m_0(\mu_1, \mu_2)} \\ &\quad - \frac{1}{t} \sum_{s=1}^t \mathbf{1}\{X_{1,s} \leq \tilde{C}_s\} \cdot \lambda(X_{2,s} - \mu_2 + \mu_2^\bullet + \gamma(X_{1,s} - \mu_1)), \end{aligned}$$

where $D_1 m_0$ and $D_2 m_0$ are the two partial derivatives of m_0 , and $\lambda(\cdot) := f'_2(\cdot)/f_2(\cdot)$. At every ω and every (μ_1, μ_2) , note the last summand in $\frac{\partial(\ell_t/t)}{\partial\mu_1}$ is γ times the last summand in $\frac{\partial(\ell_t/t)}{\partial\mu_2}$. Therefore,

$$\begin{aligned} \nabla_v \frac{1}{t} \ell_t(\mu_1, \mu_2) &= \frac{-1}{\sigma^2 \sqrt{1+\gamma^2}} \left(\frac{1}{t} \sum_{s=1}^t \frac{g'_1(X_{1,s} - \mu_1 + \mu_1^\bullet)}{g_1(X_{1,s} - \mu_1 + \mu_1^\bullet)} \right) + \frac{1}{t \sqrt{1+\gamma^2}} \frac{1}{t} \frac{D_1 m_0(\mu_1, \mu_2)}{m_0(\mu_1, \mu_2)} \\ &\quad - \frac{\gamma}{t \sqrt{1+\gamma^2}} \frac{D_2 m_0(\mu_1, \mu_2)}{m_0(\mu_1, \mu_2)}. \end{aligned}$$

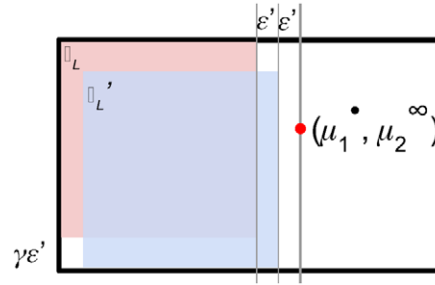
Since m_0 , $D_1 m_0$, and $D_2 m_0$ are continuous on the compact set $\square$, there exists some $0 < B < \infty$ so that $|(D_1 m_0(\mu_1, \mu_2))/(m_0(\mu_1, \mu_2))| < B$ and $|(D_2 m_0(\mu_1, \mu_2))/(m_0(\mu_1, \mu_2))| < B$ for all $(\mu_1, \mu_2) \in \square$. Pick any $\epsilon' > 0$. We have that for every ω ,

$$\inf_{(\mu_1, \mu_2) \in \square_L} \left[\left(\nabla_v \frac{1}{t} \ell_t(\mu_1, \mu_2) \right) + \frac{1}{\sigma^2 \sqrt{1+\gamma^2}} \left(\frac{1}{t} \sum_{s=1}^t \frac{g'_1(X_{1,s} - \mu_1 + \mu_1^\bullet)}{g_1(X_{1,s} - \mu_1 + \mu_1^\bullet)} \right) \right] \geq -\frac{1}{t} \frac{(1+\gamma)}{\sqrt{1+\gamma^2}} B,$$

where $\square_L := [\underline{\mu}_1, \mu_1^\bullet - 2\epsilon'] \times [\underline{\mu}_2 + \gamma\epsilon', \bar{\mu}_2]$ when $\gamma > 0$ and $\square_L := [\underline{\mu}_1, \mu_1^\bullet - 2\epsilon'] \times [\underline{\mu}_2, \bar{\mu}_2 + \gamma\epsilon']$ when $\gamma \leq 0$ is a subrectangle to the left of $\mu_1^\bullet - \epsilon'$. By the law of large numbers applied to the independent and identically distributed (i.i.d.) sequence $(g'_1(X_{1,s} - (\mu_1^\bullet - \epsilon) + \mu_1^\bullet) - \epsilon) + \mu_1^\bullet)/g_1(X_{1,s} - (\mu_1^\bullet - \epsilon) + \mu_1^\bullet)_{s \geq 1}$, almost surely

$$\frac{1}{t} \sum_{s=1}^t \frac{g'_1(X_{1,s} - (\mu_1^\bullet - \epsilon) + \mu_1^\bullet)}{g_1(X_{1,s} - (\mu_1^\bullet - \epsilon) + \mu_1^\bullet)} \rightarrow \mathbb{E}_{X \sim g_1} \left[\frac{g'_1(X_1 + \epsilon)}{g_1(X_1 + \epsilon)} \right].$$

Since $\mathbb{E}_{X \sim g_1} [g'_1(X_1)/g_1(X_1)] = 0$ and since $z \mapsto g'_1(z)/g_1(z) = \frac{d}{dz}(\ln(g_1(z)))$ is strictly decreasing by log concavity of the normal distribution, there is some $\delta > 0$ so that $\mathbb{E}_{X \sim g_1} [g'_1(X_1 + \epsilon')/g_1(X_1 + \epsilon')] = -\delta$. Furthermore, for any $\mu_1 \leq \mu_1^\bullet - \epsilon'$, then for any $x_1 \in \mathbb{R}$, $g'_1(x_1 - \mu_1 + \mu_1^\bullet)/g_1(x_1 - \mu_1 + \mu_1^\bullet) \leq g'_1(x_1 + \epsilon')/g_1(x_1 - \epsilon')$. Along any ω where

FIGURE 4. Eventually posterior belief assigns zero probability to $\square_L$.

$\frac{1}{t} \sum_{s=1}^t [g'_1(X_{1,s} - (\mu_1^\bullet - \epsilon') + \mu_1^\bullet) / g_1(X_{1,s} - (\mu_1^\bullet - \epsilon') + \mu_1^\bullet)] \rightarrow -\delta$, we therefore also have

$$\limsup_{t \rightarrow \infty} \sup_{\mu_1 \geq \mu_1^\bullet - \epsilon'} \frac{1}{t} \sum_{s=1}^t \frac{g'_1(X_{1,s} - \mu_1 + \mu_1^\bullet)}{g_1(X_{1,s} - \mu_1 + \mu_1^\bullet)} \leq -\delta.$$

Therefore, almost surely

$$\liminf_{t \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in \square_L} \left(\nabla_v \frac{1}{t} \ell_t(\mu_1, \mu_2) \right) \geq \frac{\delta}{\sigma^2 \sqrt{1 + \gamma^2}}.$$

Let $\square'_L$ be $\square_L$ shifted by the vector $(\epsilon', -\gamma\epsilon')$, so it remains in $\square$ and at least ϵ' to the left of $\mu_1^\bullet$. That is, $\square'_L := [\underline{\mu}_1 + \epsilon', \mu_1^\bullet - \epsilon'] \times [\underline{\mu}_2, \bar{\mu}_2 - \gamma\epsilon']$ if $\gamma > 0$ (illustrated in Figure 4) and $\square'_L := [\underline{\mu}_1 + \epsilon', \mu_1^\bullet - \epsilon'] \times [\underline{\mu}_2 - \gamma\epsilon', \bar{\mu}_2]$ if $\gamma \leq 0$. I will show that $\lim_{t \rightarrow \infty} \tilde{M}_t(\square_L) = 0$ almost surely. The idea is we can map every point in $\square_L$ to another point in $\square'_L$ in the direction of v (see Figure 4). For every point, its image under the map will have much higher posterior probability, since we have a uniform, strictly positive lower bound on the directional derivative of log likelihood ℓ_t in the direction of v :

$$\begin{aligned} \tilde{M}_t(\square_L) &= \int_{\square_L} \tilde{m}_t(\mu_1, \mu_2) d\mu \\ &= \int_{\square'_L} \tilde{m}_t(\mu_1, \mu_2) \cdot \frac{\tilde{m}_t(\mu_1 - \epsilon', \mu_2 - \gamma\epsilon')}{\tilde{m}_t(\mu_1, \mu_2)} d\mu \\ &= \int_{\square'_L} \tilde{m}_t(\mu_1, \mu_2) \exp(\ell_t(\mu_1 - \epsilon', \mu_2 - \gamma\epsilon') - \ell_t(\mu_1, \mu_2)) d\mu \\ &= \int_{\square'_L} \tilde{m}_t(\mu_1, \mu_2) \exp\left(-\int_0^{\epsilon'} \nabla_v \ell_t(\mu_1 - \epsilon' + z, \mu_2 - \gamma\epsilon' + \gamma z) dz\right) d\mu. \end{aligned}$$

Almost surely,

$$\liminf_{t \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in \square'_L, z \in [0, \epsilon']} (\nabla_v \ell_t(\mu_1 - \epsilon' + z, \mu_2 - \gamma\epsilon' + \gamma z)) \geq \frac{t\delta}{\sigma^2 \sqrt{1 + \gamma^2}},$$

so almost surely

$$\limsup_{t \rightarrow \infty} \tilde{M}_t(\square_L) \leq \limsup_{t \rightarrow \infty} \exp\left(-\frac{\epsilon' t \delta}{\sigma^2 \sqrt{1 + \gamma^2}}\right) \cdot \int_{\square'_L} \tilde{m}_t(\mu_1, \mu_2) d\mu.$$

But for every ω and t , the right-hand side is bounded above by $\exp(-(\epsilon' t \delta)/(\sigma^2 \sqrt{1 + \gamma^2}))$, which tends to 0 as $t \rightarrow \infty$ since $\epsilon', \delta > 0$. So in fact $\tilde{M}_t(\square_L) \rightarrow 0$ almost surely.

Since the choice of $\epsilon' > 0$ was arbitrary, this shows for every $\epsilon > 0$, almost surely $\lim_{t \rightarrow \infty} \tilde{M}_t([\underline{\mu}_1, \mu_1^\bullet - \epsilon] \times [\underline{\mu}_2 + \epsilon, \bar{\mu}_2]) = 0$ when $\gamma > 0$ and $\lim_{t \rightarrow \infty} \tilde{M}_t([\underline{\mu}_1, \mu_1^\bullet - \epsilon] \times [\underline{\mu}_2, \bar{\mu}_2 - \epsilon]) = 0$ when $\gamma \leq 0$. And by a symmetric argument, $\lim_{t \rightarrow \infty} \tilde{M}_t([\mu_1^\bullet + \epsilon, \bar{\mu}_1] \times [\underline{\mu}_2, \bar{\mu}_2 - \epsilon]) = 0$ when $\gamma > 0$ and $\lim_{t \rightarrow \infty} \tilde{M}_t([\mu_1^\bullet + \epsilon, \bar{\mu}_1] \times [\underline{\mu}_2 + \epsilon, \bar{\mu}_2]) = 0$ when $\gamma \leq 0$. Taking the complement of these sets that get assigned probability 0 in the limit establishes the result. $\square$

A.8.3 Decomposing partial derivative of log likelihood with respect to μ_2 I record a decomposition of $\frac{\partial \ell}{\partial \mu_2}(\mu_1, \mu_2)$, the partial derivative of the log-likelihood process with respect to its second argument.

Define two stochastic processes,

$$\begin{aligned} \varphi_s(\mu_1, \mu_2) &:= -\lambda(X_{2,s} - \mu_2 + \mu_2^\bullet + \gamma(X_{1,s} - \mu_1)) \cdot 1\{X_{1,s} \leq \tilde{C}_s\} \\ \bar{\varphi}_s(\mu_1, \mu_2) &:= \frac{\partial}{\partial \mu_2} \bar{L}(\mu_2 + \gamma(\mu_1 - \mu_1^\bullet) \mid \tilde{C}_s), \end{aligned}$$

where $\bar{L}(\mu_2 \mid c) := \int_{-\infty}^c g_1(x_1) \cdot \int_{-\infty}^{\infty} g_2(x_2 \mid x_1) \cdot \ln(\phi(x_2 \mid \mu_2 - \gamma(x_1 - \mu_1^\bullet))) dx_2 dx_1$. Note that $\bar{\varphi}_s(\mu_1, \mu_2)$ is measurable with respect to $\mathcal{F}_{s-1}$, since $(\tilde{C}_t)$ is a predictable process. Write $\xi_s(\mu_1, \mu_2) := \varphi_s(\mu_1, \mu_2) - \bar{\varphi}_s(\mu_1, \mu_2)$ and $y_t(\mu_1, \mu_2) := \sum_{s=1}^t \xi_s(\mu_1, \mu_2)$. Write $z_t(\mu_1, \mu_2) := \sum_{s=1}^t \bar{\varphi}_s(\mu_1, \mu_2)$.

LEMMA A.7. *We have $\frac{\partial \ell_t}{\partial \mu_2}(\mu_1, \mu_2) = \frac{D_2 m_0(\mu_1, \mu_2)}{m_0(\mu_1, \mu_2)} + y_t(\mu_1, \mu_2) + z_t(\mu_1, \mu_2)$.*

PROOF. This comes from expanding $\ell_t(\mu_1, \mu_2)$ and taking its derivative as in the proof of Lemma A.6. $\square$

Now I derive a result about the $\xi_t(\mu_1, \mu_2)$ processes for different pairs (μ_1, μ_2) .

LEMMA A.8. *There exists $\kappa_\xi < \infty$ so that for every $(\mu_1, \mu_2) \in \square$ and for every $t \geq 1$, $\omega \in \Omega$, $\mathbb{E}[\xi_t^2(\mu_1, \mu_2) \mid \mathcal{F}_{t-1}](\omega) \leq \kappa_\xi$.*

PROOF. Note that $\bar{\varphi}_t(\mu_1, \mu_2)$ is measurable with respect to $\mathcal{F}_{t-1}$. Also, $\varphi_t(\mu_1, \mu_2) \mid \mathcal{F}_{t-1} = \varphi_t(\mu_1, \mu_2) \mid \tilde{C}_t$, because by independence of X_t from $(X_s)_{s=1}^{t-1}$, the only information that $\mathcal{F}_{t-1}$ contains about $\varphi_t(\mu_1, \mu_2)$ is in determining the cutoff threshold $\tilde{C}_t$.

At a sample path ω so that $\tilde{C}_t(\omega) = c \in \mathbb{R}$,

$$\begin{aligned}
 & \mathbb{E}[\varphi_s(\mu_1, \mu_2) | \mathcal{F}_{t-1}](\omega) \\
 &= \mathbb{E}[-\lambda(X_{2,s} - \mu_2 + \mu_2^\bullet + \gamma(X_{1,s} - \mu_1)) \cdot \mathbf{1}\{X_1 \leq c\}] \\
 &= \frac{\partial}{\partial \mu_2} \int_{-\infty}^c g_1(x_1) \cdot \int_{-\infty}^{\infty} g_2(x_2 | x_1) \cdot \ln(\phi(x_2 | \mu_2 - \gamma(x_1 - \mu_1))) dx_2 dx_1 \\
 &= \frac{\partial}{\partial \mu_2} \int_{-\infty}^c g_1(x_1) \cdot \int_{-\infty}^{\infty} g_2(x_2 | x_1) \cdot \ln(\phi(x_2 | [\mu_2 + \gamma(\mu_1 - \mu_1^\bullet)] - \gamma(x_1 - \mu_1^\bullet))) dx_2 dx_1 \\
 &= \frac{\partial}{\partial \mu_2} \bar{L}(\mu_2 + \gamma(\mu_1 - \mu_1^\bullet) | c).
 \end{aligned}$$

This shows that $\mathbb{E}[\varphi_s(\mu_1, \mu_2) | \mathcal{F}_{t-1}](\omega) = \bar{\varphi}_s(\mu_1, \mu_2)(\omega)$. Since this holds regardless of c , we get that $\mathbb{E}[\varphi_s(\mu_1, \mu_2) | \mathcal{F}_{t-1}] = \bar{\varphi}_t(\mu_1, \mu_2)$ for all ω , that is to say,

$$\begin{aligned}
 \mathbb{E}[\xi_t^2(\mu_1, \mu_2) | \mathcal{F}_{t-1}] &= \text{Var}[\varphi_t(\mu_1, \mu_2) | \mathcal{F}_{t-1}] \leq \mathbb{E}[\varphi_t^2(\mu_1, \mu_2) | \mathcal{F}_{t-1}] \\
 &\leq \mathbb{E}[(\lambda(X_{2,s} - \mu_2 + \mu_2^\bullet + \gamma(X_{1,s} - \mu_1)))^2].
 \end{aligned}$$

It suffices to show $\mathbb{E}[(\lambda(X_2 - \mu_2 + \mu_2^\bullet + \gamma(X_1 - \mu_1)))^2]$ exists for all $\mu_1, \mu_2 \in \mathbb{R}$ and is continuous. The (finite) maximum value this expectation takes on the compact set $\square$ can be taken as κ_ξ .

Since the second derivative of the log of the normal density is uniformly bounded, there exists some $\kappa_{f_2} < \infty$ so that for all $z \in \mathbb{R}$, $-\kappa_{f_2} < \lambda'(z) < 0$. So $\lambda(z)$ is Lipschitz continuous with constant κ_{f_2} . Let $b_0 := \lambda(-\mu_2 + \mu_2^\bullet - \gamma\mu_1)$.

For any $x_1, x_2 \in \mathbb{R}$,

$$\begin{aligned}
 & (\lambda(x_2 - \mu_2 + \mu_2^\bullet + \gamma(x_1 - \mu_1)))^2 \\
 &= b_0^2 + (\lambda(x_2 - \mu_2 + \mu_2^\bullet + \gamma(x_1 - \mu_1)))^2 - (\lambda(-\mu_2 + \mu_2^\bullet - \gamma\mu_1))^2 \\
 &\leq b_0^2 + |\lambda(x_2 - \mu_2 + \mu_2^\bullet + \gamma(x_1 - \mu_1)) - \lambda(-\mu_2 + \mu_2^\bullet - \gamma\mu_1)| \\
 &\quad \times |\lambda(x_2 - \mu_2 + \gamma(x_1 + \mu_2^\bullet - \mu_1)) + \lambda(-\mu_2 + \mu_2^\bullet - \gamma\mu_1)| \\
 &\leq b_0^2 + (\kappa_{f_2} \cdot (|x_2| + \gamma|x_1|)) \cdot (2b_0 + (\kappa_{f_2} \cdot (|x_2| + \gamma|x_1|))).
 \end{aligned}$$

Note the bound is a second-order polynomial in $|x_1|$ and $|x_2|$. We have

$$\begin{aligned}
 & \mathbb{E}[(\lambda(X_2 - \mu_2 + \mu_2^\bullet + \gamma(X_1 - \mu_1)))^2] \\
 &\leq \mathbb{E}[b_0^2 + (\kappa_{f_2} \cdot (|X_2| + \gamma|X_1|)) \cdot (2b_0 + (\kappa_{f_2} \cdot (|X_2| + \gamma|X_1|)))] < \infty,
 \end{aligned}$$

where the last inequality is due to the fact that X_1 and X_2 have finite second moments. $\square$

A.8.4 A law of large numbers for martingale increments I use a statistical result from [Heidhues, Kőszegi, and Strack \(2018\)](#) to show that the y_t/t term in the decomposition of $\frac{1}{t} \frac{\partial \ell_t}{\partial \mu_2}$ almost surely converges to 0 in the long run and, furthermore, this convergence is

uniform on $\square$. This lets me focus on terms of the form $\bar{\varphi}_s(\mu_1, \mu_2)$, which can be interpreted as the *expected* contribution to the log-likelihood derivative from round s data. This lends tractability to the problem as $\bar{\varphi}_s(\mu_1, \mu_2)$ only depends on $\tilde{C}_s$, but not on $X_{1,s}$ or $X_{2,s}$.

LEMMA A.9. *For every $(\mu_1, \mu_2) \in \square$, $\lim_{t \rightarrow \infty} |y_t(\mu_1, \mu_2)/t| = 0$ almost surely.*

PROOF. [Heidhues, Kőszegi, and Strack's \(2018\)](#) Proposition 10 shows that if (y_t) is a martingale such that there exists some constant $v \geq 0$ satisfying $[y]_t \leq vt$ almost surely, where $[y]_t$ is the quadratic variation of (y_t) , then almost surely $\lim_{t \rightarrow \infty} (y_t/t) = 0$.

Consider the process $y_t(\mu_1, \mu_2)$ for a fixed $(\mu_1, \mu_2) \in \square$. By definition $y_t = \sum_{s=1}^t \varphi_s(\mu_1, \mu_2) - \bar{\varphi}_s(\mu_1, \mu_2)$. As established in the proof of Lemma A.8, for every s , $\bar{\varphi}_s(\mu_1, \mu_2) = \mathbb{E}[\varphi_s(\mu_1, \mu_2) | \mathcal{F}_{s-1}]$. So for $t' < t$,

$$\begin{aligned} \mathbb{E}[y_t(\mu_1, \mu_2) | \mathcal{F}_{t'}] &= \sum_{s=1}^{t'} \varphi_s(\mu_1, \mu_2) - \bar{\varphi}_s(\mu_1, \mu_2) + \mathbb{E}\left[\sum_{s=t'+1}^t \varphi_s(\mu_1, \mu_2) - \bar{\varphi}_s(\mu_1, \mu_2) \middle| \mathcal{F}_{t'}\right] \\ &= \sum_{s=1}^{t'} \varphi_s(\mu_1, \mu_2) - \bar{\varphi}_s(\mu_1, \mu_2) + \sum_{s=t'+1}^t \mathbb{E}[\mathbb{E}[\varphi_s(\mu_1, \mu_2) - \bar{\varphi}_s(\mu_1, \mu_2) | \mathcal{F}_{s-1}] | \mathcal{F}_{t'}] \\ &= \sum_{s=1}^{t'} \varphi_s(\mu_1, \mu_2) - \bar{\varphi}_s(\mu_1, \mu_2) + 0 = y_{t'}(\mu_1, \mu_2). \end{aligned}$$

This shows $(y_t(\mu_1, \mu_2))_t$ is a martingale. Also,

$$[y(\mu_1, \mu_2)]_t = \sum_{s=1}^{t-1} \mathbb{E}[(y_s(\mu_1, \mu_2) - y_{s-1}(\mu_1, \mu_2))^2 | \mathcal{F}_{s-1}] = \sum_{s=1}^{t-1} \mathbb{E}[\xi_s^2(\mu_1, \mu_2) | \mathcal{F}_{s-1}] \leq \kappa_\xi \cdot t$$

by Lemma A.8. Therefore, [Heidhues, Kőszegi, and Strack \(2018\)](#) Proposition 10 applies. $\square$

LEMMA A.10. *$\lim_{t \rightarrow \infty} \sup_{(\mu_1, \mu_2) \in \square} |y_t(\mu_1, \mu_2)/t| = 0$ almost surely.*

PROOF. This argument is similar to Lemma 11 in [Heidhues, Kőszegi, and Strack \(2018\)](#). I apply Lemma 2 of [Andrews \(1992\)](#), which says to prove this result I just need to check conditions BD, P-SSLN, and S-LIP from [Andrews \(1992\)](#). BD holds because $\square$ is a bounded subset of $\mathbb{R}^2$. P-SLLN holds by Lemma A.9, which shows for all $(\mu_1, \mu_2) \in \square$, $\lim_{t \rightarrow \infty} |y_t(\mu_1, \mu_2)/t| = 0$ almost surely.

Condition S-LIP is essentially a Lipschitz continuity condition. It requires finding a sequence of random variables B_t such that $|\xi_t(\mu_1, \mu_2) - \xi_t(\mu'_1, \mu'_2)| \leq B_t \cdot (|\mu_1 - \mu'_1| + |\mu_2 - \mu'_2|)$ almost surely, such that these random variables satisfy $\sup_{t \geq 1} \frac{1}{t} \sum_{s=1}^t \mathbb{E}[B_s] < \infty$, and $\lim_{t \rightarrow \infty} \frac{1}{t} \sum_{s=1}^t (B_s - \mathbb{E}[B_s]) = 0$ almost surely.

But for every ω , $\varphi_s(\mu_1, \mu_2) := -\lambda(X_{2,s} - \mu_2 + \mu_2^\bullet + \gamma(X_{1,s} - \mu_1)) \cdot 1\{X_{1,s} \leq \tilde{C}_s\}$,

$$\begin{aligned} & |\varphi_s(\mu_1, \mu_2) - \varphi_s(\mu'_1, \mu'_2)| \\ & \leq |\lambda(X_{2,s} - \mu_2 + \mu_2^\bullet + \gamma(X_{1,s} - \mu_1)) - \lambda(X_{2,s} - \mu'_2 + \mu_2^\bullet + \gamma(X_{1,s} - \mu'_1))|. \end{aligned}$$

As $\ln(f_2(\cdot))$ has a bounded second derivative, the right-hand side is bounded by $\kappa_{f_2} \cdot (|\mu_2 - \mu'_2| + \gamma \cdot |\mu_1 - \mu'_1|)$.

Now that we know $|\varphi_s(\mu_1, \mu_2) - \varphi_s(\mu'_1, \mu'_2)|(\omega) \leq \kappa_{f_2} \cdot (|\mu_2 - \mu'_2| + \gamma \cdot |\mu_1 - \mu'_1|)$ for all ω , we must also have $|\bar{\varphi}_s(\mu_1, \mu_2) - \bar{\varphi}_s(\mu'_1, \mu'_2)|(\omega) \leq \kappa_{f_2} \cdot (|\mu_2 - \mu'_2| + \gamma \cdot |\mu_1 - \mu'_1|)$ for all ω since $\bar{\varphi}_s(\mu_1, \mu_2) = \mathbb{E}[\varphi_s(\mu_1, \mu_2) | \mathcal{F}_{s-1}]$.

Setting B_s as the constant $2\kappa_{f_2}$ for every s satisfies S-LIP. $\square$

A.8.5 Bounds on asymptotic beliefs and asymptotic cutoffs Recall that Lemma A.4 implies that for any μ_2 , all pairs of fundamentals on the line $\text{li}(\mu_2)$ have the same optimal cutoff threshold. Then against any feasible model $\Psi(\mu_1, \mu_2; \gamma)$ with $(\mu_1, \mu_2) \in \square$, the best cutoff strategy is between $C(\mu_1^\bullet, \underline{\mu}_2^\circ; \gamma)$ and $C(\mu_1^\bullet, \bar{\mu}_2^\circ; \gamma)$. Define these cutoffs as $\underline{c}^\circ$ and $\bar{c}^\circ$, respectively.

LEMMA A.11. *Let $\underline{c}^\circ \leq c \leq \bar{c}^\circ$. If $r - \gamma < 0$, then $\liminf_{t \rightarrow \infty} \tilde{C}_t \geq c$ almost surely implies $\lim_{t \rightarrow \infty} \tilde{M}_t(\diamond[\underline{\mu}_2^\circ, \mu_2^*(c)]) = 0$ almost surely and $\limsup_{t \rightarrow \infty} \tilde{C}_t \leq c$ almost surely implies $\lim_{t \rightarrow \infty} \tilde{M}_t(\diamond[\mu_2^*(c), \bar{\mu}_2^\circ]) = 0$ almost surely. If $r - \gamma > 0$, then $\liminf_{t \rightarrow \infty} \tilde{C}_t \geq c$ almost surely implies $\lim_{t \rightarrow \infty} \tilde{M}_t(\diamond[\mu_2^*(c), \bar{\mu}_2^\circ]) = 0$ almost surely and $\limsup_{t \rightarrow \infty} \tilde{C}_t \leq c$ almost surely implies $\lim_{t \rightarrow \infty} \tilde{M}_t(\diamond[\underline{\mu}_2^\circ, \mu_2^*(c)]) = 0$ almost surely.*

PROOF. We prove the liminf statement for the case of $r - \gamma < 0$ and briefly discuss the argument for the limsup statement for the case of $r - \gamma > 0$; the arguments for the other two statements are very similar.

Consider the first statement when $r - \gamma < 0$, fixing some $\underline{c}$ with $\underline{c}^\circ \leq \underline{c} \leq \bar{c}^\circ$. We show that for all $\epsilon > 0$, there exists $\delta > 0$ such that almost surely,

$$\liminf_{t \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in \square \cap \diamond[\underline{\mu}_2^\circ, \mu_2^*(\underline{c}) - \epsilon]} \frac{1}{t} \frac{\partial \ell_t}{\partial \mu_2}(\mu_1, \mu_2) \geq \delta.$$

From Lemma A.7, we may rewrite the left-hand side as

$$\liminf_{t \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in \square \cap \diamond[\underline{\mu}_2^\circ, \mu_2^*(\underline{c}) - \epsilon]} \left[\frac{1}{t} \frac{D_2 m_0(\mu_1, \mu_2)}{m_0(\mu_1, \mu_2)} + \frac{y_t(\mu_1, \mu_2)}{t} + \frac{z_t(\mu_1, \mu_2)}{t} \right],$$

which is no smaller than taking the inf separately across the three terms in the bracket:

$$\begin{aligned} & \liminf_{t \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in \square \cap \diamond[\underline{\mu}_2^\circ, \mu_2^*(\underline{c}) - \epsilon]} \frac{1}{t} \frac{D_2 m_0(\mu_1, \mu_2)}{m_0(\mu_1, \mu_2)} + \liminf_{t \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in \diamond[\underline{\mu}_2^\circ, \mu_2^*(\underline{c}) - \epsilon]} \frac{y_t(\mu_1, \mu_2)}{t} \\ & + \liminf_{t \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in \square \cap \diamond[\underline{\mu}_2^\circ, \mu_2^*(\underline{c}) - \epsilon]} \frac{z_t(\mu_1, \mu_2)}{t}. \end{aligned}$$

Since $D_2 m_0 / m_0$ is bounded on $\square$ as $D_2 m_0$ is continuous, and m_0 is continuous and strictly positive on the compact set $\square$, the first term is 0 for every ω . To deal with the second term,

$$\begin{aligned} \liminf_{t \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in \square \cap \diamond[\underline{\mu}_2^\circ, \mu_2^*(\underline{c}) - \epsilon]} \frac{y_t(\mu_1, \mu_2)}{t} &\geq \liminf_{t \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in \square} \left| \frac{y_t(\mu_1, \mu_2)}{t} \right| \\ &= \liminf_{t \rightarrow \infty} \left\{ -1 \cdot \sup_{(\mu_1, \mu_2) \in \square} \left| \frac{y_t(\mu_1, \mu_2)}{t} \right| \right\}. \end{aligned}$$

Lemma A.10 gives $\lim_{t \rightarrow \infty} \sup_{(\mu_1, \mu_2) \in \square} |y_t(\mu_1, \mu_2)/t| = 0$ almost surely. Hence, we conclude that, almost surely,

$$\liminf_{t \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in \square \cap \diamond[\underline{\mu}_2^\circ, \mu_2^*(\underline{c}) - \epsilon]} \frac{y_t(\mu_1, \mu_2)}{t} \geq 0.$$

It suffices then to find $\delta > 0$ and show $\liminf_{t \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in \square \cap \diamond[\underline{\mu}_2^\circ, \mu_2^*(\underline{c}) - \epsilon]} (z_t(\mu_1, \mu_2)/t) \geq \delta$ almost surely. To do this, I first show $\bar{\varphi}_s(\mu_1, \mu_2)(\omega) \geq \delta$ whenever $\tilde{C}_s(\omega) \geq \underline{c}$ and $\mu_2 \leq \mu_2^*(\underline{c}) - \epsilon$. At every $\underline{c}^\circ \leq c' \leq \bar{c}^\circ$, we get

$$\frac{\partial}{\partial \mu_2} \bar{L}(\mu_2 | c') = \int_{-\infty}^{c'} g_1(x_1) \cdot \int_{-\infty}^{\infty} (-1) \cdot g_2(x_2 | x_1) \cdot \lambda(x_2 - \mu_2 + \mu_2^\bullet + \gamma(x_1 - \mu_1^\bullet)) dx_2 dx_1.$$

The first-order condition implies that $\frac{\partial}{\partial \mu_2} \bar{L}(\mu_2^*(c') | c') = 0$. Since λ is strictly decreasing, we also get $\frac{\partial}{\partial \mu_2} \bar{L}(\mu_2 | c') > 0$ for any $\mu_2 < \mu_2^*(c')$. Since we have $r - \gamma < 0$, $\mu_2^*(\cdot)$ is strictly increasing, which means $\frac{\partial}{\partial \mu_2} \bar{L}(\mu_2^*(\underline{c}) - \epsilon | c') > 0$ for any $\underline{c} \leq c' \leq \bar{c}^\circ$. Let $\delta > 0$ satisfy $\min_{c' \in [\underline{c}, \bar{c}^\circ]} \frac{\partial}{\partial \mu_2} \bar{L}(\mu_2^*(\underline{c}) - \epsilon | c') > \delta$, which exists because $c' \mapsto \frac{\partial}{\partial \mu_2} \bar{L}(\mu_2^*(\underline{c}) - \epsilon | c')$ is continuous on the compact domain $[\underline{c}, \bar{c}^\circ]$. When $\tilde{C}_s(\omega) = c' \in [\underline{c}, \bar{c}^\circ]$ and for any $(\mu_1, \mu_2) \in \square \cap \diamond[\underline{\mu}_2^\circ, \mu_2^*(\underline{c}) - \epsilon]$, we have $\bar{\varphi}_s(\mu_1, \mu_2)(\omega) = \frac{\partial}{\partial \mu_2} \bar{L}(\mu_2 | c') \geq \frac{\partial}{\partial \mu_2} \bar{L}(\mu_2^*(\underline{c}) - \epsilon | c') > \delta$.

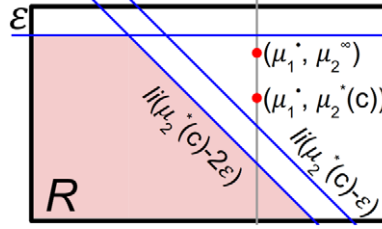
Along any ω where $\liminf_{t \rightarrow \infty} \tilde{C}_t \geq \underline{c}$, we therefore have

$$\liminf_{s \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in \square \cap \diamond[\underline{\mu}_2^\circ, \mu_2^*(\underline{c}) - \epsilon]} \bar{\varphi}_s(\mu_1, \mu_2) \geq \delta$$

and, thus,

$$\begin{aligned} \liminf_{t \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in \square \cap \diamond[\underline{\mu}_2^\circ, \mu_2^*(\underline{c}) - \epsilon]} \frac{z_t(\mu_1, \mu_2)}{t} \\ = \liminf_{t \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in \square \cap \diamond[\underline{\mu}_2^\circ, \mu_2^*(\underline{c}) - \epsilon]} \frac{1}{t} \left[\sum_{s=1}^t \bar{\varphi}_s(\mu_1, \mu_2) \right] \geq \delta. \end{aligned}$$

Let $R := [\underline{\mu}_1, \bar{\mu}_1] \times [\underline{\mu}_2, \bar{\mu}_2 - \epsilon] \cap \diamond[\underline{\mu}_2^\circ, \mu_2^*(\underline{c}) - 2\epsilon]$ and let $R' := R + (0, \epsilon)'$ be R -shifted upward by ϵ . We have both $R, R' \subseteq \square \cap \diamond[\underline{\mu}_2^\circ, \mu_2^*(\underline{c}) - \epsilon]$. Figure 5 illustrates the case of $\gamma > 0$.

FIGURE 5. Bounding the posterior belief assigned to the region R .

So using the same argument as in the proof of Lemma A.6,

$$\begin{aligned}\tilde{M}_t(R) &= \int_{R'} \tilde{m}_t(\mu_1, \mu_2) \cdot \frac{\tilde{m}_t(\mu_1, \mu_2 - \epsilon)}{\tilde{m}_t(\mu_1, \mu_2)} d\mu \\ &= \int_{R'} \tilde{m}_t(\mu_1, \mu_2) \exp(\ell_t(\mu_1, \mu_2 - \epsilon) - \ell_t(\mu_1, \mu_2)) d\mu \\ &= \int_{R'} \tilde{m}_t(\mu_1, \mu_2) \exp\left(-\int_0^\epsilon \frac{\partial \ell_t}{\partial \mu_2}(\mu_1, \mu_2 - \epsilon + z) dz\right) d\mu.\end{aligned}$$

Almost surely,

$$\liminf_{t \rightarrow \infty} \inf_{(\mu_1, \mu_2) \in R', z \in [0, \epsilon]} \left(\frac{\partial \ell_t}{\partial \mu_2}(\mu_1, \mu_2 - \epsilon + z) \right) \geq t\delta,$$

so almost surely

$$\limsup_{t \rightarrow \infty} \tilde{M}_t(R) \leq \limsup_{t \rightarrow \infty} \exp(-t\epsilon\delta) \cdot \int_{R'} \tilde{m}_t(\mu_1, \mu_2) d\mu = 0.$$

Letting $\epsilon \rightarrow 0$ and noting that $\text{li}(\mu_2^*(\underline{c}))$ crosses the top edge of $\square$ to the left of μ_1^* when $\gamma > 0$, we get $\lim_{t \rightarrow \infty} \tilde{M}_t(\diamond[\mu_2^*(\underline{c}), \bar{\mu}_2^*] \cup [\underline{\mu}_1, \mu_1^*] \times \{\bar{\mu}_2\}) = 1$ almost surely. But from Lemma A.6, the set $[\underline{\mu}_1, \mu_1^*] \times \{\bar{\mu}_2\}$ must receive no weight in the limit; hence $\lim_{t \rightarrow \infty} \tilde{M}_t(\diamond[\mu_2^*(\underline{c}), \mu_2^*(\underline{c})]) = 0$ almost surely as desired. (The case of $\gamma < 0$ is analogous.)

Now consider any $\underline{c}^\circ \leq \bar{c} \leq \bar{c}^\circ$. I briefly discuss why $\limsup_{t \rightarrow \infty} \tilde{C}_t \leq \bar{c}$ almost surely implies $\lim_{t \rightarrow \infty} \tilde{M}_t(\diamond[\mu_2^*(\underline{c}^\circ), \mu_2^*(\bar{c})]) = 0$ almost surely when $r - \gamma > 0$. As in the argument before, the key is to find some $\delta > 0$ such that $\frac{\partial}{\partial \mu_2} \bar{L}(\mu_2 | c') > \delta$ whenever $c' \in [\underline{c}^\circ, \bar{c}]$ and $\mu_2 \leq \mu_2^*(\bar{c})$. For each $c \in [\underline{c}^\circ, \bar{c}^\circ]$, the FOC implies $\frac{\partial}{\partial \mu_2} \bar{L}(\mu_2^*(c') | c') = 0$. Since λ is strictly decreasing, we also get $\frac{\partial}{\partial \mu_2} \bar{L}(\mu_2 | c') > 0$ for any $\mu_2 < \mu_2^*(c')$. Since we now consider $r - \gamma > 0$, $\mu_2^*(c)$ is strictly decreasing in c , and this shows $\frac{\partial}{\partial \mu_2} \bar{L}(\mu_2^*(\bar{c}) - \epsilon | c') > 0$ for any $\underline{c}^\circ \leq c' \leq \bar{c}$. We can find $\delta > 0$ such that $\frac{\partial}{\partial \mu_2} \bar{L}(\mu_2^*(\bar{c}) - \epsilon | c') > \delta$ for every $\underline{c}^\circ \leq c' \leq \bar{c}$ by continuity, so we also get $\frac{\partial}{\partial \mu_2} \bar{L}(\mu_2 | c') > \delta$ for any $\mu_2 \leq \mu_2^*(\bar{c}) - \epsilon$. $\square$

Now I use a bound on agents' asymptotic beliefs about μ_2 to deduce asymptotic restrictions on their cutoffs.

LEMMA A.12. Suppose that there are $\underline{\mu}_2^\circ \leq \mu_2^l < \mu_2^h \leq \bar{\mu}_2^\circ$ such that $\lim_{t \rightarrow \infty} \tilde{M}_t(\diamond[\mu_2^l, \mu_2^h]) = 1$ almost surely. Then $\liminf_{t \rightarrow \infty} \tilde{C}_t \geq C(\mu_1^\bullet, \mu_2^l; \gamma)$ and $\limsup_{t \rightarrow \infty} \tilde{C}_t \leq C(\mu_1^\bullet, \mu_2^h; \gamma)$ almost surely.

PROOF. I show $\liminf_{t \rightarrow \infty} \tilde{C}_t \geq C(\mu_1^\bullet, \mu_2^l; \gamma)$ almost surely. The argument establishing $\limsup_{t \rightarrow \infty} \tilde{C}_t \leq C(\mu_1^\bullet, \mu_2^h; \gamma)$ is symmetric.

Let $c^l = C(\mu_1^\bullet, \mu_2^l; \gamma)$, and recall before we defined $\underline{c}^\circ := C(\mu_1^\bullet, \underline{\mu}_2^\circ; \gamma)$ and $\bar{c}^\circ := C(\mu_1^\bullet, \bar{\mu}_2^\circ; \gamma)$.

Let $U(c; \mu_1, \mu_2)$ be the expected payoff of using the stopping strategy S_c when $(X_1, X_2) \sim \Psi(\mu_1, \mu_2; \gamma)$. I first show $c \mapsto U(c; \mu_1, \mu_2)$ is single peaked: it is strictly increasing up to $c = c^*$, the subjectively optimal cutoff under $\Psi(\mu_1, \mu_2; \gamma)$, then strictly decreasing afterward. Recall (from the proof of Lemma A.1 when $\gamma \geq -1$) the cutoff form of the best stopping strategy comes from the fact that $u_1(x_1) < \mathbb{E}_{\Psi(\mu_1, \mu_2; \gamma)}[u_2(x_1, X_2) | X_1 = x_1]$ for $x_1 < c^*$, but $u_1(x_1) > \mathbb{E}_{\Psi(\mu_1, \mu_2; \gamma)}[u_2(x_1, X_2) | X_1 = x_1]$ for $x_1 > c^*$. For two cutoffs $c_1 < c_2 < c^*$, the two stopping strategies S_{c_1} and S_{c_2} only differ in how they treat first-period draws in the interval $[c_1, c_2]$, so we can write the difference in their expected payoffs as $\int_{c_1}^{c_2} (\mathbb{E}_{\Psi(\mu_1, \mu_2; \gamma)}[u_2(x_1, X_2) | X_1 = x_1] - u_1(x_1)) \phi(x_1 | \mu_1) dx_1$. The integrand is strictly positive on $[c_1, c_2]$; therefore, $U(c_1; \mu_1, \mu_2) < U(c_2; \mu_1, \mu_2)$. This shows $U(\cdot; \mu_1, \mu_2)$ is strictly increasing up until c^* ; a symmetric argument shows it is strictly decreasing after c^* .

By Lemma A.4, $C(\mu_1', \mu_2'; \gamma) = C(\mu_1^\bullet, \mu_2; \gamma)$ for all $(\mu_1', \mu_2') \in \text{li}(\mu_2)$. Since $c \mapsto U(c; \mu_1, \mu_2)$ is single peaked for every (μ_1, μ_2) and since $c^l \leq C(\mu_1^\bullet, \mu_2; \gamma)$ for all $\mu_2 \in [\mu_2^l, \mu_2^h]$, we also get $c^l \leq C(\mu_1', \mu_2'; \gamma)$ for every $(\mu_1', \mu_2') \in \diamond[\mu_2^l, \mu_2^h]$, since $\diamond[\mu_2^l, \mu_2^h]$ is the union of the line segments, $\diamond[\mu_2^l, \mu_2^h] = \bigcup_{\mu_2 \in [\mu_2^l, \mu_2^h]} \text{li}(\mu_2)$.

Fix some $\epsilon > 0$. We get $U(c^l; \mu_1, \mu_2) - U(c^l - \epsilon; \mu_1, \mu_2) > 0$ for every $(\mu_1, \mu_2) \in \diamond[\mu_2^l, \mu_2^h]$. As $(\mu_1, \mu_2) \mapsto (U(c^l; \mu_1, \mu_2) - U(c^l - \epsilon; \mu_1, \mu_2))$ is continuous, there exists some $\kappa^* > 0$ so that $U(c^l; \mu_1, \mu_2) - U(c^l - \epsilon; \mu_1, \mu_2) > \kappa^*$ for all $(\mu_1, \mu_2) \in \diamond[\mu_2^l, \mu_2^h]$. In particular, if $\nu \in \Delta(\diamond[\mu_2^l, \mu_2^h])$ is a belief about fundamentals, then $\int U(c^l; \mu_1, \mu_2) - U(c^l - \epsilon; \mu_1, \mu_2) d\nu(\mu) > \kappa^*$.

Now let $\bar{\kappa} := \sup_{c \in [\underline{c}^\circ, \bar{c}^\circ]} \sup_{(\mu_1, \mu_2) \in \square} U(c; \mu_1, \mu_2)$ and $\underline{\kappa} := \inf_{c \in [\underline{c}^\circ, \bar{c}^\circ]} \inf_{(\mu_1, \mu_2) \in \square} U(c; \mu_1, \mu_2)$. Find $p \in (0, 1)$ so that $p\kappa^* - (1-p)(\bar{\kappa} - \underline{\kappa}) = 0$. At any belief $\hat{\nu} \in \Delta(\square)$ that assigns more than probability p to the parallelogram $\diamond[\mu_2^l, \mu_2^h]$, the optimal cutoff is larger than $c^l - \epsilon$. To see this, take any $\hat{c} \leq c^l - \epsilon$ and I will show $\hat{c}$ is suboptimal. If $\hat{c} < \underline{c}$, then it is suboptimal after any belief on $\diamond$. If $\underline{c} \leq \hat{c} \leq c^l - \epsilon$, I show that $\int U(c^l; \mu_1, \mu_2) - U(\hat{c}; \mu_1, \mu_2) d\hat{\nu}(\mu) > 0$. To see this, we may decompose $\hat{\nu}$ as the mixture of a probability measure ν on $\diamond[\mu_2^l, \mu_2^h]$ and another probability measure ν^c on $\square \setminus \diamond[\mu_2^l, \mu_2^h]$. Let $\hat{p} > p$ be the probability that ν assigns to $\diamond[\mu_2^l, \mu_2^h]$. The above integral is equal to

$$\begin{aligned} & \hat{p} \int_{\diamond[\mu_2^l, \mu_2^h]} U(c^l; \mu_1, \mu_2) - U(\hat{c}; \mu_1, \mu_2) d\nu(\mu) \\ & + (1 - \hat{p}) \int_{\square \setminus \diamond[\mu_2^l, \mu_2^h]} U(c^l; \mu_1, \mu_2) - U(\hat{c}; \mu_1, \mu_2) d\nu^c(\mu). \end{aligned}$$

Since c^l is to the left of the optimal cutoff for all $(\mu_1, \mu_2) \in \diamond[\mu_2^l, \mu_2^h]$ and $\hat{c} \leq c^l - \epsilon$, then $U(\hat{c}; \mu_1, \mu_2) \leq U(c^l - \epsilon; \mu_1, \mu_2)$ for all $(\mu_1, \mu_2) \in \diamond[\mu_2^l, \mu_2^h]$. The first summand

is no less than $\hat{p} \int_{\Diamond[\mu_2^l, \mu_2^h]} U(c^l; \mu_1, \mu_2) - U(c^l - \epsilon; \mu_1, \mu_2) d\nu(\mu) \geq \hat{p}\kappa^*$. Also, the integrand in the second summand is no smaller than $-(\bar{\kappa} - \underline{\kappa})$; therefore, $\int U(c^l; \mu_1, \mu_2) - U(c^l - \epsilon; \mu_1, \mu_2) d\hat{\nu}(\mu) \geq \hat{p}\kappa^* - (1 - \hat{p})(\bar{\kappa} - \underline{\kappa})$. Since $\hat{p} > p$, we get $\hat{p}\kappa^* - (1 - \hat{p})(\bar{\kappa} - \underline{\kappa}) > 0$.

Along any sample path ω where $\lim_{t \rightarrow \infty} \tilde{M}_t(\Diamond[\mu_2^l, \mu_2^h])(\omega) = 1$, eventually $\tilde{M}_t(\Diamond[\mu_2^l, \mu_2^h])(\omega) > p$ for all large enough t , meaning $\liminf_{t \rightarrow \infty} \tilde{C}_t(\omega) \geq c^l - \epsilon$. Since $\lim_{t \rightarrow \infty} \tilde{M}_t(\Diamond[\mu_2^l, \mu_2^h]) = 1$ almost surely, this shows $\liminf_{t \rightarrow \infty} \tilde{C}_t \geq C(\mu_1^\bullet, \mu_2^l; \gamma) - \epsilon$ almost surely. As the choice of $\epsilon > 0$ was arbitrary, we conclude $\liminf_{t \rightarrow \infty} \tilde{C}_t \geq C(\mu_1^\bullet, \mu_2^l; \gamma)$ almost surely. $\square$

A.8.6 The contraction map I now combine the results established so far to prove the convergence statement in Proposition 7.

PROOF OF PROPOSITION 7 (CONVERGENCE). Let $\mu_{2,[1]}^A := \underline{\mu}_2^\circ$ and $\mu_{2,[1]}^B := \bar{\mu}_2^\circ$. For $k = 2, 3, \dots$, iteratively define $\mu_{2,[k]}^A := \mathcal{I}(\mu_{2,[k-1]}^A; \gamma)$ and $\mu_{2,[k]}^B := \mathcal{I}(\mu_{2,[k-1]}^B; \gamma)$. Let $\mu_{2,[k]}^l := \min(\mu_{2,[k]}^A, \mu_{2,[k]}^B)$ and $\mu_{2,[k]}^h := \max(\mu_{2,[k]}^A, \mu_{2,[k]}^B)$. I show by induction that for every k , $\lim_{t \rightarrow \infty} \tilde{M}_t(\Diamond[\mu_{2,[k]}^l, \mu_{2,[k]}^h]) = 1$ almost surely. (The base case of $k = 1$ holds by the support of the prior belief.)

Inductive step when $r - \gamma < 0$. From Lemma A.12, if $\lim_{t \rightarrow \infty} \tilde{M}_t(\Diamond[\mu_{2,[k]}^l, \mu_{2,[k]}^h]) = 1$ almost surely, then $\liminf_{t \rightarrow \infty} \tilde{C}_t \geq C(\mu_1^\bullet, \mu_{2,[k]}^l; \gamma)$ and $\limsup_{t \rightarrow \infty} \tilde{C}_t \leq C(\mu_1^\bullet, \mu_{2,[k]}^h; \gamma)$ almost surely. Using these conclusions in Lemma A.11, we deduce that almost surely,

$$\lim_{t \rightarrow \infty} \tilde{M}_t(\Diamond[\mu_2^*(C(\mu_1^\bullet, \mu_{2,[k]}^l; \gamma)), \mu_2^*(C(\mu_1^\bullet, \mu_{2,[k]}^h; \gamma))]) = 1.$$

Both $C(\mu_1^\bullet, \cdot; \gamma)$ and $\mu_2^*(\cdot)$ are strictly increasing, so $\lim_{t \rightarrow \infty} \tilde{M}_t(\Diamond[\mu_{2,[k+1]}^l, \mu_{2,[k+1]}^h]) = 1$ almost surely.

Inductive step when $r - \gamma > 0$. Now $C(\mu_1^\bullet, \cdot; \gamma)$ is strictly increasing but $\mu_2^*(\cdot)$ is strictly decreasing. From Lemma A.12, if $\lim_{t \rightarrow \infty} \tilde{M}_t(\Diamond[\mu_{2,[k]}^l, \mu_{2,[k]}^h]) = 1$ almost surely, then $\liminf_{t \rightarrow \infty} \tilde{C}_t \geq C(\mu_1^\bullet, \mu_{2,[k]}^l; \gamma)$ and $\limsup_{t \rightarrow \infty} \tilde{C}_t \leq C(\mu_1^\bullet, \mu_{2,[k]}^h; \gamma)$ almost surely. But using these conclusions in Lemma A.11, for the case of $r - \gamma > 0$, we further deduce that

$$\lim_{t \rightarrow \infty} \tilde{M}_t(\Diamond[\mu_2^*(C(\mu_1^\bullet, \mu_{2,[k]}^h; \gamma)), \mu_2^*(C(\mu_1^\bullet, \mu_{2,[k]}^l; \gamma))]) = 1.$$

So now we have $\mu_{2,[k+1]}^l = \mu_2^*(C(\mu_1^\bullet, \mu_{2,[k]}^h; \gamma))$ and $\mu_{2,[k+1]}^h = \mu_2^*(C(\mu_1^\bullet, \mu_{2,[k]}^l; \gamma))$, but still conclude $\lim_{t \rightarrow \infty} \tilde{M}_t(\Diamond[\mu_{2,[k+1]}^l, \mu_{2,[k+1]}^h]) = 1$ almost surely.

The iterates $(\mu_{2,[k]}^A)_{k \geq 1}$ and $(\mu_{2,[k]}^B)_{k \geq 1}$ are the iterates of a contraction map, so $\lim_{k \rightarrow \infty} \mu_{2,[k]}^A = \mu_2^* = \lim_{k \rightarrow \infty} \mu_{2,[k]}^B$. Thus, the agent's posterior converges in L^1 to $\text{li}(\mu_2^\infty)$ almost surely (since the support of the prior is bounded). In addition, the sequences of bounds on asymptotic actions also converge by continuity: $\lim_{k \rightarrow \infty} C(\mu_1^\bullet, \mu_{2,[k]}^A; \gamma) = c^\infty = \lim_{k \rightarrow \infty} C(\mu_1^\bullet, \mu_{2,[k]}^B; \gamma)$. This implies $\lim_{t \rightarrow \infty} \tilde{C}_t = c^\infty$ almost surely. Finally, combining the asymptotic belief result with Lemma A.6, we see that in fact $\tilde{M}_t$ converges in L^1 to the point $(\mu_1^\bullet, \mu_2^\infty)$ almost surely. $\square$

REFERENCES

- Amemiya, Takeshi (1985), *Advanced Econometrics*. Harvard University Press. [1288]
- Andrews, W. K. Donald (1992), “Generic uniform convergence.” *Econometric Theory*, 8, 241–257. [1304]
- Araujo, Felipe A., Stephanie W. Wang, and Alistair J. Wilson (2021), “The times they are A-changing: Experimenting with dynamic adverse selection.” *American Economic Journal: Microeconomics*, 13, 1–22. [1287]
- Barron, Kai, Steffen Huck, and Philippe Jehiel (2019), “Everyday econometricians: Selection neglect and overoptimism when learning from others.” Working Paper. [1287]
- Benjamin, Daniel J., Don A. Moore, and Matthew Rabin (2017), “Biased beliefs about random samples: Evidence from two integrated experiments.” Working Paper. [1269]
- Bohren, J. Aislinn (2016), “Informational herding with model misspecification.” *Journal of Economic Theory*, 163, 222–247. [1288]
- Bohren, J. Aislinn, and Daniel N. Hauser (2021), “Learning with heterogeneous misspecified models: Characterization and robustness.” *Econometrica*, 89, 3025–3077. [1288]
- Bushong, Benjamin and Tristan Gagnon-Bartsch (2022), “Learning with misattribution of reference dependence.” *Journal of Economic Theory*. [1288]
- Chen, Daniel L., Tobias J. Moskowitz, and Kelly Shue (2016), “Decision making under the gambler’s fallacy: Evidence from asylum judges, loan officers, and baseball umpires.” *Quarterly Journal of Economics*, 131, 1181–1242. [1270]
- Chen, Wanyi (2021), “Dynamic survival bias in optimal stopping problems.” *Journal of Economic Theory*, 196, 105286. [1287]
- Cox, David R. (1962), *Renewal Theory*. Methuen. [1288]
- Crowder, Martin J. (2001), *Classical Competing Risks*. Chapman and Hall/CRC. [1288]
- Dasaratha, Krishna and Kevin He (2020), “Network structure and naive sequential learning.” *Theoretical Economics*, 15, 415–444. [1288]
- Enke, Benjamin (2020), “What you see is all there is.” *Quarterly Journal of Economics*, 135, 1363–1398. [1287]
- Esponda, Ignacio (2008), “Behavioral equilibrium in economies with adverse selection.” *American Economic Review*, 98, 1269–1291. [1287]
- Esponda, Ignacio and Demian Pouzo (2016), “Berk–Nash equilibrium: A framework for modeling agents with misspecified models.” *Econometrica*, 84, 1093–1130. [1274, 1287]
- Esponda, Ignacio and Demian Pouzo (2017), “Conditional retrospective voting in large elections.” *American Economic Journal: Microeconomics*, 9, 54–75. [1287]
- Esponda, Ignacio and Demian Pouzo (2019), “Retrospective voting and party polarization.” *International Economic Review*, 60, 157–186. [1287]

- Esponda, Ignacio, Demian Pouzo, and Yuichi Yamamoto (2021), “Asymptotic behavior of Bayesian learners with misspecified models.” *Journal of Economic Theory*, 195, 105260. [1288]
- Eyster, Erik and Matthew Rabin (2010), “Naive herding in rich-information settings.” *American Economic Journal: Microeconomics*, 2, 221–243. [1288]
- Frick, Mira, Ryota Iijima, and Yuhta Ishii (2020), “Misinterpreting others and the fragility of social learning.” *Econometrica*, 88, 2281–2328. [1287, 1288]
- Frick, Mira, Ryota Iijima, and Yuhta Ishii (2021a), “Belief convergence under misspecified learning: A martingale approach.” Working Paper. [1288]
- Frick, Mira, Ryota Iijima, and Yuhta Ishii (2021b), “Welfare comparisons for biased learning.” Working Paper. [1287]
- Fudenberg, Drew and Giacomo Lanzani (2021), “Which misperceptions persist?” Working Paper. [1287]
- Fudenberg, Drew, Giacomo Lanzani, and Philipp Strack (2021), “Limit points of endogenous misspecified learning.” *Econometrica*, 89, 1065–1098. [1288]
- Fudenberg, Drew, Gleb Romanyuk, and Philipp Strack (2017), “Active learning with a misspecified prior.” *Theoretical Economics*, 12, 1155–1189. [1287]
- Gaurino, Antonio and Philippe Jehiel (2013), “Social learning with coarse inference.” *American Economic Journal: Microeconomics*, 5, 147–174. [1288]
- He, Kevin and Jonathan Libgober (2021), “Evolutionarily stable (mis)specifications: Theory and applications.” Working Paper. [1287]
- Heckman, James J. and Bo E. Honoré (1989), “The identifiability of the competing risks model.” *Biometrika*, 76, 325–330. [1288]
- Heidhues, Paul, Botond Köszegi, and Philipp Strack (2018), “Unrealistic expectations and misguided learning.” *Econometrica*, 86, 1159–1214. [1276, 1280, 1285, 1287, 1303, 1304]
- Heidhues, Paul, Botond Köszegi, and Philipp Strack (2019), “Overconfidence and prejudice.” Working Paper. [1276]
- Heidhues, Paul, Botond Köszegi, and Philipp Strack (2021), “Convergence in models of misspecified learning.” *Theoretical Economics*, 16, 73–99. [1288]
- Jehiel, Philippe (2018), “Investment strategy and selection bias: An equilibrium perspective on overoptimism.” *American Economic Review*, 108, 1582–1597. [1287]
- Kahneman, Daniel and Amos Tversky (1972), “Subjective probability: A judgment of representativeness.” *Cognitive psychology*, 3, 430–454. [1276]
- Mailhot, Louis (1985), “Une propriété de la variance de certaines lois de probabilité réelles tronquées.” *Comptes rendus de l'Académie des sciences. Série 1, Mathématique*, 301, 241–244. [1296]

Mueller, Andreas I., Johannes Spinnewijn, and Giorgio Topa (2021), “Job seekers’ perceptions and employment prospects: Heterogeneity, duration dependence, and bias.” *American Economic Review*, 111, 324–363. [1270, 1271, 1278]

Nyarko, Yaw (1991), “Learning in mis-specified models and the possibility of cycles.” *Journal of Economic Theory*, 55, 416–427. [1287]

Rabin, Matthew (2002), “Inference by believers in the law of small numbers.” *Quarterly Journal of Economics*, 117, 775–816. [1282, 1287]

Rabin, Matthew and Dimitri Vayanos (2010), “The gambler’s and hot-hand fallacies: Theory and applications.” *Review of Economic Studies*, 77, 730–778. [1273, 1276, 1282, 1283, 1287]

Simonsohn, Uri and Francesca Gino (2013), “Daily horizons: Evidence of narrow bracketing in judgment from 10 years of MBA admissions interviews.” *Psychological Science*, 24, 219–224. [1270]

Spiegler, Ran (2016), “Bayesian networks and boundedly rational expectations.” *Quarterly Journal of Economics*, 131, 1243–1290. [1287]

Spiegler, Ran (2017), ““Data monkeys”: A procedural model of extrapolation from partial statistics.” *Review of Economic Studies*, 84, 1818–1841. [1287]

Suetens, Sigrid, Claus B. Galbo-Jørgensen, and Jean-Robert Tyran (2016), “Predicting lotto numbers: A natural experiment on the gambler’s fallacy and the hot-hand fallacy.” *Journal of the European Economic Association*, 14, 584–607. [1269]

Terrell, Dek (1994), “A test of the gambler’s fallacy: Evidence from pari-mutuel games.” *Journal of Risk and Uncertainty*, 8, 309–317. [1269]

Tsiatis, Anastasios (1975), “A nonidentifiability aspect of the problem of competing risks.” *Proceedings of the National Academy of Sciences*, 72, 20–22. [1288]

Co-editor Ran Spiegler handled this manuscript.

Manuscript received 10 December, 2020; final version accepted 19 August, 2021; available online 23 August, 2021.