

Pathwise concentration bounds for Bayesian beliefs

DREW FUDENBERG

Department of Economics, MIT

GIACOMO LANZANI

Department of Economics, MIT

PHILIPP STRACK

Department of Economics, Yale University

We show that Bayesian posteriors concentrate on the outcome distributions that approximately minimize the Kullback–Leibler divergence from the empirical distribution, uniformly over sample paths, even when the prior does not have full support. This generalizes Diaconis and Freedman’s (1990) uniform convergence result to, e.g., priors that have finite support, are constrained by independence assumptions, or have a parametric form that cannot match some probability distributions. The concentration result lets us provide a rate of convergence for Berk’s (1966) result on the limiting behavior of posterior beliefs when the prior is misspecified. We provide a bound on approximation errors in “anticipated-utility” models, and extend our analysis to outcomes that are perceived to follow a Markov process.

KEYWORDS. Misspecified learning, Bayesian consistency.

JEL CLASSIFICATION. C11, D81.

1. INTRODUCTION

Learning from repeated observations is a key feature of many economic settings, and almost all economic studies of learning model it as Bayesian inference. To understand the medium and long-run implications of Bayesian learning, it is useful to know how quickly beliefs concentrate around the data generating processes that best explain the observations. Our main result, Theorem 1, shows that the probability the posterior assigns to distributions that do not approximately maximize the likelihood assigned to the data vanishes exponentially quickly. Importantly, we identify conditions for this to hold not only with high probability, but for every possible realization of the data. More specifically, Theorem 1 establishes that for every $\varepsilon > 0$ the posterior probability of the

Drew Fudenberg: drew.fudenberg@gmail.com

Giacomo Lanzani: lanzani@mit.edu

Philipp Strack: philipp.strack@yale.edu

We thank Simone Cerreia-Vioglio, Xiaohong Chen, Roberto Corrao, Jana Gieselmann, Paul Heidhues, Yuhta Ishii, Mats Köster, Pooya Molavi, and Todd Sarver for helpful comments and conversations, and National Science Foundation grant SES 1951056, the Guido Cazzavillan Scholarship, and the Sloan Foundation for financial support.

distributions that do not ε -minimize the Kullback–Leibler (KL) divergence vanishes at an exponential rate. In contrast to earlier pathwise concentration bounds, our result holds even if the agent’s prior does not have full support, or satisfies parametric restrictions, and regardless of the true data generating process.

Our results generalize [Diaconis and Freedman \(1990\)](#), which showed that a ϕ -positivity condition implies that Bayesian posteriors converge to the empirical distribution at a uniform exponential rate. This condition requires that the support of the agent’s prior includes every distribution over outcomes, and thus rules out many settings of economic interest in which the set of outcome distributions is naturally restricted. For example, it does not apply to agents whose prior has finite support (which is natural in settings such as urn problems with a finite number of balls), agents who each period observe a set of Bernoulli trials that they think are i.i.d., or agents who believe (mistakenly or not) that some variables are positively correlated. In addition, ϕ -positivity rules out all cases where the support of the agent’s prior does not contain the true data generating process, so that the agent is misspecified.

Theorem 1 guarantees that beliefs concentrate on the approximate KL minimizers for the empirical frequency. We show that this is equivalent to concentration on a ball around the exact KL minimizers when priors have full support, but not in general. Moreover, since the KL minimizer is not unique, the theorem does not imply that beliefs converge.

We use Theorem 1 to prove Theorem 2, which provides a rate of convergence for [Berk’s \(1966\)](#) result that posterior beliefs concentrate around the Kullback–Leibler minimizers with respect to the true data generating process. Berk’s result, like our paper, is stated for an exogenous data generating process. It was extended to learning from endogenous data by [Esponda and Pouzo \(2016\)](#), which led to a renewed interest in misspecified learning in the economics literature.¹ A key step in the analysis of such models is often to establish that Bayesian beliefs concentrate quickly around the KL minimizers, and as our Theorem 1 holds pathwise it immediately implies such a result.² In [Fudenberg, Lanzani, and Strack \(2021b\)](#), we use the concentration result to characterize the long-run beliefs of a correctly specified agent who has an imperfect and selective memory.

Theorem 3 extends Theorem 1 to beliefs that result from observing a Markov process whose transition probabilities are unknown. This complements recent work by [Molavi \(2019\)](#) and [Esponda and Pouzo \(2021\)](#), which studied analogs of [Berk \(1966\)](#) and [Esponda and Pouzo \(2016\)](#) for Markovian environments.

¹Subsequent papers include [Fudenberg, Romanyuk, and Strack \(2017\)](#), [Molavi \(2019\)](#), [Bohren and Hauser \(2021\)](#), [Fudenberg and Lanzani \(2023\)](#), [He and Libgober \(2021\)](#), [Esponda, Pouzo, and Yamamoto \(2021\)](#), [Heidhues, Köszegi, and Strack \(2021\)](#), [Levy, Moreno de Barreda, and Razin \(2021\)](#), [He \(2022\)](#), and [Frick, Iijima, and Ishii \(2023\)](#). Before this, [Arrow and Green \(1973\)](#) gave the first general framework for this problem, and [Nyarko \(1991\)](#) pointed out that the combination of misspecification and endogenous observations can lead to cycles. In a setting with finitely many states, [Frick, Iijima, and Ishii \(2022\)](#) provide a convergence-in-probability result on the relative speed at which beliefs converge to the truth for agents with different likelihood functions; [Frick, Iijima, and Ishii \(2021\)](#) extend this to endogenous data.

²For example, Theorem 1 provides a shorter way to prove Proposition 1 of [Fudenberg, Lanzani, and Strack \(2021a\)](#).

Most of the paper assumes that the set of outcomes is finite, as it is in [Diaconis and Freedman \(1990\)](#). Section 6 discusses the extent to which the results extend to settings with infinitely many outcomes. It also uses our results to show that the play of a Bayesian agent converges to that predicted by the anticipated utility model (e.g., [Kreps \(1998\)](#), [Preston \(2005\)](#), [Eusepi and Preston \(2018\)](#)), and quantifies its rate of convergence. This clarifies when the anticipated utility model is a good approximation of rational play, and complements numerical studies by [Cogley, Colacito, and Sargent \(2007\)](#), [Cogley and Sargent \(2008\)](#), and [Cogley, Colacito, Hansen, and Sargent \(2008\)](#).

1.1 *The importance of pathwise concentration*

The statistical literature has many concentration results for beliefs; see, e.g., [Shen and Wasserman \(2001\)](#). Our results differ in two important ways. First, ours hold for every sample realization, and thus even when the true data generating process is time-varying and endogenous, while the existing statistics results show that beliefs concentrate with probability converging to 1 with respect to a fixed-data generating process. Second, the statistics results are for concentration around the parameters or distributions that minimize the KL divergence from the true-data generating process, while our result are for concentration around the KL-minimizers with respect to an arbitrary empirical distribution.³

Pathwise concentration has played an important role in a number of economic applications, starting with the analysis of nonequilibrium learning in games in [Fudenberg and Levine \(1993\)](#).⁴ The result has also been used to analyze selective attention ([Schwartzstein \(2014\)](#)), the merging of opinions ([Acemoglu, Chernozhukov, and Yildiz \(2016\)](#)), recursive utility functions ([Al-Najjar and Shmaya \(2019\)](#)), and persuasion ([Schwartzstein and Sunderam \(2021\)](#)).

To help motivate our analysis, we describe why pathwise concentration (rather than concentration in probability) is needed in four papers on very different problems. [Fudenberg and Levine \(1993\)](#) study the steady states of a model of nonequilibrium learning. The uniform concentration result implies that agents play myopically at any information set h that has been reached many times. Because which information sets are reached is endogenous, the proof of the main theorem uses the pathwise concentration to rule out the possibility that posteriors only concentrate conditional to histories that induce the player to make choices that prevent reaching h many times. This is not ruled out by concentration in probability, because the sets of histories under which the information set is reached less than N times could have probability approaching 1 as N grows.

[Al-Najjar and Shmaya \(2019\)](#) provide a representation result for Epstein–Zin preferences over stochastic consumption streams for patient agents. To do so, they bound

³Most of these papers also assume that the prior is correctly specified, so that the unique KL minimizer is the true distribution, but see [Kleijn and Van der Vaart \(2012\)](#) and the references therein for generalizations to misspecified priors.

⁴Subsequent applications to learning in games are [Fudenberg and Levine \(2006\)](#), [Fudenberg and He \(2018\)](#), [Gonçalves \(2020\)](#), and [Clark and Fudenberg \(2021\)](#). [Clark, Fudenberg, and He \(2022\)](#) apply our generalization of the result.

the distance between the certain equivalents of period- t consumption with period $t - 1$ and period 0 information. The (relative) impact of small-probability events on utility increases in the high-patience limit, so the representation result needs the posterior consumption variance to vanish uniformly over all sample paths, which follows from the uniform concentration of the posterior.

Gonçalves (2020) introduces an equilibrium concept for games that allows the agents to sample from the opponents' strategies at a cost before playing. To show existence of the equilibrium, the maximization problem of the agent is transformed into an optimal stopping problem. There the uniform concentration result guarantees that the stopping time is uniformly bounded by a deterministic horizon, thus transforming an infinite-horizon problem into a finite-horizon one that is then solved by backward induction.

Finally, Theorem 1 can be used to study the limit points of misspecified learning when the distribution of outcomes depends on the action played by an agent, and that action depends on the agent's beliefs. For example, the agent might be a customer who wants to learn which of two products she prefers, decides every period which one to buy, and receives a signal about the product they bought. To understand if the action a can be played in the long run, we need to understand whether the resulting process of beliefs makes it optimal to play a . Fudenberg, Lanzani, and Strack (2021a) showed that a limit action must be a best reply to *all* of the associated KL minimizers when the prior has subexponential decay. An earlier version of this paper, Fudenberg, Lanzani, and Strack (2022), use the results here to give a simpler and more transparent proof of this result.

1.2 The importance of relaxing full support

The following are examples of commonly studied situations where uniform concentration results that require a full support prior are not applicable, but our results apply.

Finite support priors In some problems, the reasonable priors have finite support as, e.g., if outcomes correspond to the color of balls drawn with replacement from an urn with known size but with unknown composition.

Correlation restrictions When the outcome space Y has a product structure, and the agent's prior imposes a qualitative restriction on how the components are correlated (e.g., that they are positively correlated), the Diaconis and Freedman (1990) result does not apply, while ours does. This naturally arises in economic problems as, e.g., in Spiegel (2020), where the agent neglects the mediating role of expectations in the Phillips curve and is mistakenly convinced that money supply and output are positively correlated. Similarly, our model can be used to study situations where the agent mistakenly believes outcomes are independent as in, e.g., Enke and Zimmermann (2019).

Moreover, whenever the agent observes the outcome of a game and believes (correctly or not) that the players have coordinated on a correlated equilibrium, the correlation structure is naturally restricted, as the possible joint distributions must satisfy the obedience constraint.⁵

⁵Formally, let Y_i be the set of actions available to player i , let the outcome space be $Y = \times_{i=1}^n Y_i$, and let $(u_i)_{i=1}^n$ be the payoff functions. If the agent is certain that the outcome corresponds to a correlated

Markov models Another context where support restrictions arise naturally is in the study of Markov models. For example, if an analyst assumes that beliefs follow Bayes rule or that a stock price process is a martingale (Bachelier (1900), Fama (1965)), the techniques of Diaconis and Freedman (1990) cannot be applied, as they would require full support over the set of transition matrices. However, it is easy to extend our analysis to analyze belief concentration in Markov models, as we do in Section 5. And our Markov model can also be used to study mistaken beliefs about the correlation between signals and outcomes, as in Esponda (2008).

Extending the applications of Section 1.1 Our results can be used to extend some past applications of the pathwise concentration results. They permit an extension of Al-Najjar and Shmaya's (2019) representation theorem to beliefs about the consumption process that do not have full support, such as its illustrative example (which is not covered by the paper's result), and also to Markovian consumption processes. For the experimentation in games considered by Gonçalves (2020), our extension allows, e.g., beliefs concentrated on the pure strategies for the opponents, or certainty that the opponent does not play a dominated strategy.

2. SETUP

We study Bayesian beliefs induced by a sequence of subjectively i.i.d. data. Let Y be a finite set of possible outcomes, and let $P = \Delta(Y)$ be the set of probability measures over Y endowed with $\|\cdot\|$, the total variation distance of probability measures.⁶

Let $\mu_0 \in \Delta(P) = \Delta(\Delta(Y))$ denote a prior belief over distributions of outcomes, and $\Theta = \text{supp } \mu_0$ denote its support.⁷ A data set $y^t = (y_1, y_2, \dots, y_t) \in Y^t$ is a vector of outcomes. For every data set y^t , we let μ_t be the *posterior belief*, which is required to satisfy Bayes rule whenever the denominator is different from 0:

$$\mu_t(C) = \frac{\int_C \prod_{\tau=1}^t p(y_\tau) d\mu_0(p)}{\int_P \prod_{\tau=1}^t p(y_\tau) d\mu_0(p)}. \quad (\text{Bayes rule})$$

The *empirical distribution* $f_t \in P$ is

$$f_t(z) = \frac{1}{t} \sum_{\tau=1}^t \mathbb{I}_{\{z\}}(y_\tau).$$

Our main result is that along any path of realized outcomes, the probability the posterior belief assigns to the outcome distributions that do not best approximate the empirical distribution converges to zero at a uniform and exponential rate. To state this

equilibrium, then every $p \in \Theta$ must satisfy $\sum_{y_{-i} \in Y_{-i}} u_i(y_i, y_{-i}) p(y_{-i}|y_i) \geq \sum_{y_{-i} \in Y_{-i}} u_i(y'_i, y_{-i}) p(y_{-i}|y_i)$ for all $i \in I$, $y_i, y'_i \in Y_i$ with $p(y_i) > 0$.

⁶Formally, for every $p, q \in P$, $\|p - q\| = \sup_{Y' \subseteq Y} |p(Y') - q(Y')|$.

⁷For every $X \subseteq \mathbb{R}^k$, we let $\Delta(X)$ denote the set of Borel probability distributions on X .

conclusion formally, we adopt the convention that $0/0 = 0$ and $0 \log 0 = 0$, and define $H : P \times P \rightarrow \mathbb{R}$ to be the (possibly infinite) Kullback–Leibler divergence of q with respect to p :

$$H(q, p) = \sum_{z \in Y} q(z) \log \left(\frac{q(z)}{p(z)} \right).$$

Let $M : P \rightrightarrows P$ be the correspondence that maps a distribution q to the set of minimizers of the Kullback–Leibler divergence over the support of the prior:

$$M(q) = \operatorname{argmin}_{p \in \Theta} H(q, p).$$

The log-likelihood assigned to the data set y^t under outcome distribution p is

$$\log \left(\prod_{\tau=1}^t p(y_\tau) \right) = \sum_{z \in Y} t f_t(z) \log p(z) = -t H(f_t, p) + t \sum_{z \in Y} f_t(z) \log f_t(z). \quad (1)$$

Minimizing the Kullback–Leibler divergence relative to the empirical distribution is hence the same as maximizing the log-likelihood assigned to the data set, so the *KL minimizers* $M(f_t)$ at time t correspond to the outcome distributions that maximize the likelihood of y^t . Throughout, $B_\varepsilon(D)$ denotes the ball of radius ε around a set $D \subseteq P$ in total variation distance, and denote by $M_\varepsilon : P \rightrightarrows P$ the correspondence that maps a distribution q to the distributions that come within ε of the minimum KL divergence:

$$M_\varepsilon(q) = \left\{ p' \in \Theta : H(q, p') \leq \min_{p \in \Theta} H(q, p) + \varepsilon \right\}.$$

3. THE RATE OF CONVERGENCE OF BAYESIAN BELIEFS

To show that Bayesian beliefs concentrate around the empirical distribution at a uniform rate, [Diaconis and Freedman \(1990\)](#) used the following condition.

DEFINITION 1 (ϕ positivity). The prior μ_0 is ϕ positive if for $\phi : \mathbb{R}_{++} \rightarrow \mathbb{R}_{++}$, $\mu_0(B_\varepsilon(p)) \geq \phi(\varepsilon)$ for every $p \in P$ and $\varepsilon > 0$.

Since ϕ positivity requires the prior to assign strictly positive probability to every ε ball, it requires the prior to have full support.

THEOREM A ([Diaconis and Freedman \(1990\)](#)). For every $\phi : \mathbb{R}_{++} \rightarrow \mathbb{R}_{++}$ and every $\varepsilon \in (0, 1)$ there are $\tilde{A}(\varepsilon) \in \mathbb{R}_{++}$ and $g(\varepsilon) \in \mathbb{R}_{++}$ such that

$$\frac{\mu_t(B_\varepsilon(f_t))}{1 - \mu_t(B_\varepsilon(f_t))} \geq \tilde{A}(\varepsilon) \exp(g(\varepsilon)t),$$

for all ϕ positive μ_0 , $t \in \mathbb{N}$, and $f_t \in \Delta(Y)$.

Theorem A shows that for ϕ positive priors, the probability that Bayesian beliefs assign to distributions that are more than ε away from the empirical distribution vanishes exponentially quickly, so it quantifies the speed at which a Bayesian with full support prior becomes more certain when observing i.i.d. data. The strength of this theorem is that it holds not only in probability, but for every realization of outcomes.

Clearly, ϕ positivity plays a crucial role in Theorem A, as if the prior is not ϕ positive the empirical distribution need not be in its support, so beliefs cannot concentrate around it. However, requiring the prior to satisfy ϕ positivity rules out several practically relevant cases. For example, ϕ positivity cannot be satisfied if the prior has finite support, reduces the dimensionality of the problem, or is supported only on unimodal distributions.

Moreover, models of misspecified learning suppose that the true data generating process is not in the support of the prior, which rules out ϕ positivity. We extend Theorem A to cases where ϕ positivity fails. Loosely, we require that either the prior gives all neighborhoods of a distribution sufficient weight or the prior gives zero weight to a small neighborhood of the distribution.

DEFINITION 2 (ϕ positivity on Θ). The prior μ_0 is *ϕ positive on Θ* if for $\phi : \mathbb{R}_{++} \rightarrow \mathbb{R}_{++}$, $\mu_0(B_\varepsilon(p)) \geq \phi(\varepsilon)$ for every $p \in \Theta$ and $\varepsilon > 0$.

Note that ϕ positivity on Θ reduces to ϕ positivity when $\Theta = P$, i.e., the prior has full support. In Diaconis and Freedman (1990), Bayes rule is well-defined everywhere, but this is not true when the prior does not have full support. We define $\Delta^\Theta(Y)$ to be the (compact) set of empirical frequencies for which Bayesian updating is well-defined for a prior with support Θ .⁸ Theorem 1 below establishes that if beliefs are ϕ positive on Θ , for every $\varepsilon \in (0, 1)$, the posterior concentrates on $M_\varepsilon(f_t)$.

THEOREM 1. For every $\phi : \mathbb{R}_{++} \rightarrow \mathbb{R}_{++}$, $\alpha \in (0, 1)$, and $\varepsilon \in (0, 1)$, there is $A(\varepsilon) > 0$ such that

$$\frac{\mu_t(M_\varepsilon(f_t))}{1 - \mu_t(M_\varepsilon(f_t))} \geq A(\varepsilon) \exp(\alpha \varepsilon t)$$

for all $t \in \mathbb{N}$, $f_t \in \Delta^\Theta(Y)$, and μ_0 that is ϕ positive on Θ .

Moreover, if $\underline{q} := \inf_{q \in \Theta} \min_{z \in \text{supp } q} q(z) > 0$, then we can set

$$A(\varepsilon) = \phi(\min\{\underline{q}/2, (1 - \alpha)\varepsilon\}\underline{q}/2).$$

Theorem 1 only requires ϕ positivity on Θ , in contrast to Theorem A, which assumes ϕ positivity on the whole space of distributions. When the prior is not ϕ positive on all of $\Delta(Y)$, beliefs need not concentrate around the empirical frequency, because this frequency might not be in the prior's support. This is why Theorem 1 bounds the probability assigned to the distributions in $M_\varepsilon(f_t)$, which are the ε minimizers of the KL divergence, while Theorem A bounds the probability assigned to $B_\varepsilon(f_t)$, the ε ball around

⁸That is, $\Delta^\Theta(Y) = \{q \in \Delta(Y) : \exists p \in \Theta, \text{supp } q \subseteq \text{supp } p\}$.

the empirical distribution. Moreover, as Example 1 below shows, the theorem does not apply to the ε ball $B_\varepsilon(M(f_t))$ around the exact minimizers $M(f_t)$, because when Θ is not convex points far from the minimizers can attain almost the same divergence.

The theorem implies that the probability assigned to all distributions that do not ε best explain the empirical frequency f_t vanishes at the exponential rate $\alpha\varepsilon$:

$$\mu_t(\Theta \setminus M_\varepsilon(f_t)) \leq \frac{1}{A(\varepsilon)} \exp(-\alpha\varepsilon t).$$

Notice that the multiplicative constant $A(\varepsilon)$ depends on the prior μ_0 only through Θ and the function ϕ . The second part of the statement guarantees that in the widely-studied case of finite support priors, there is an explicit formula to compute the rate of convergence as a function of Θ and ϕ , with the intuitive comparative statics that the rate of convergence improves when ϕ is higher and when the support is smaller.

The next example shows why Theorem 1 does not apply to the ε ball $B_\varepsilon(M(f_t))$ around the exact minimizers $M(f_t)$.

EXAMPLE 1. Let $Y = \{0, 1\}$, identify each $p \in \Delta(Y)$ with the probability of $y = 1$, and let $\mu_0(\{1/4\}) = \mu_0(\{3/4\}) = 1/2$. Consider the sequence of outcomes $(y_t)_{t=1}^\infty$ where $y_t = 1$ if t is odd and $y_t = 0$ if t is even. In the even periods $2t$, the data is uninformative about the state, and both $1/4$ and $3/4$ are minimizers. At every odd period $2t + 1$, for every $p \in \Theta$,

$$H(f_{2t+1}, p) = K_t - \frac{t}{2t+1} \log(1-p) - \frac{t+1}{2t+1} \log(p),$$

where the term K_t does not depend on p . Thus in the odd periods $M(f_{2t+1}) = \{3/4\}$, so for $\varepsilon < 1/2$, $B_\varepsilon(M(f_{2t+1})) = \{3/4\}$. However,

$$\frac{\mu_{2t+1}(B_\varepsilon(M(f_{2t+1})))}{1 - \mu_{2t+1}(B_\varepsilon(M(f_{2t+1})))} = \frac{\mu_{2t+1}(\{3/4\})}{\mu_{2t+1}(\{1/4\})} = \frac{\mu_0(\{3/4\})(1/4)^t(3/4)^{t+1}}{\mu_0(\{1/4\})(1/4)^{t+1}(3/4)^t} = 3$$

so beliefs do not concentrate on the neighborhood of the KL minimizer.⁹ The concentration result fails because the difference between the KL-divergences is $(\log(3/4) - \log(1/4))/(2t+1)$, which converges to 0. Thus even a very large data set provides only weak evidence in favor of $p = 3/4$. $\diamond$

3.1 Proof sketch of Theorem 1

The proofs of all our results are in the Appendix. The proof of Theorem 1 has three steps. Step 1 proves a local Lipschitz property of the KL divergence, step 2 gives an explicit rate of concentration for each realized empirical frequency, while step 3 concludes by turning this explicit local rate of convergence into an exponential (but with possibly implicit constant) global rate of convergence.

⁹In this example, Θ is not connected. Example 4 in the Appendix shows that the same problem can arise when it is. The example there adds a third outcome to this one, and specifies a Θ that connects $1/4$ and $3/4$ via distributions that are not relevant under the specified outcome sequence.

Specifically, Lemma 3 shows that Kullback–Leibler divergence $H(q, \cdot)$ is locally Lipschitz continuous in its second argument:

$$|H(q, p) - H(q, \tilde{p})| \leq 2\|p - \tilde{p}\| \max_{z \in Y} \max \left\{ \frac{q(z)}{p(z)}, \frac{q(z)}{\tilde{p}(z)} \right\}. \quad (2)$$

With this, we are able to prove the next lemma, which is at the heart of our results. The lemma uses the following bound on the likelihood ratio between the empirical distribution and the elements of the agent's prior: For every $f_t \in \Delta^\Theta(Y)$, $\bar{q} \in \Theta$, and $\kappa \in \mathbb{R}_+$, let

$$R(f_t, \kappa, \bar{q}) := \max_{q \in \Theta \cap B_\kappa(\bar{q})} \max_{z \in Y} \frac{f_t(z)}{q(z)}.$$

LEMMA 1. *If μ_0 is ϕ positive on Θ , then for every $\varepsilon, \varepsilon', \kappa \in \mathbb{R}_+$, $t \in \mathbb{N}$, $f_t \in \Delta^\Theta(Y)$, and $\bar{q} \in M_{\varepsilon'}(f_t)$ with $\varepsilon' + \kappa \leq \varepsilon$ and $R(f_t, \kappa, \bar{q}) < \infty$, we have*

$$\frac{\mu_t(M_\varepsilon(f_t))}{1 - \mu_t(M_\varepsilon(f_t))} \geq \phi(\kappa/2R(f_t, \kappa, \bar{q})) \exp((\varepsilon - \kappa - \varepsilon')t). \quad (3)$$

To prove the lemma, we use the Lipschitz condition (2) to establish that the Kullback–Leibler divergence from f_t is at most $\min_{p \in \Theta} H(f_t, p) + \varepsilon' + \kappa$ in a ball of radius $\kappa/2R(f_t, \kappa, \bar{q})$ around the ε' -minimizer $\bar{q}$.¹⁰ Therefore, $M_{\varepsilon'+\kappa}(f_t)$ contains a ball of radius $\kappa/2R(f_t, \kappa, \bar{q})$ around $\bar{q} \in \Theta$. As μ_0 is ϕ positive on Θ , that ball has prior probability at least $\phi(\kappa/2R(f_t, \kappa, \bar{q}))$, so the ratio between the prior probabilities of $M_{\varepsilon'+\kappa}(f_t) \supseteq B_{\kappa/2R(f_t, \kappa, \bar{q})}(\bar{q})$ and $\Theta \setminus M_\varepsilon(f_t)$ is at least

$$\frac{\mu_0(M_{\varepsilon'+\kappa}(f_t))}{1 - \mu_0(M_\varepsilon(f_t))} \geq \frac{\mu_0(B_{\kappa/2R(f_t, \kappa, \bar{q})}(\bar{q}))}{1} \geq \phi(\kappa/2R(f_t, \kappa, \bar{q})),$$

which delivers the multiplicative constant in the right-hand side of the lemma. The exponential term follows from the fact that the the posterior probability ratio of $D \subseteq P$ over $P \setminus D$ grows exponentially in the difference between the KL divergence from f_t of the distribution inside and outside D :

$$\frac{\mu_t(D)}{1 - \mu_t(D)} = \frac{\int_D \exp(-H(f_t, p)t) d\mu_0(p)}{\int_{P \setminus D} \exp(-H(f_t, p)t) d\mu_0(p)},$$

and that by definition distributions outside $M_\varepsilon(f_t)$ have a KL-divergence from f_t of at least $\min_{p \in \Theta} H(f_t, p) + \varepsilon$. Thus at time t the posterior probability ratio is larger than the lower bound on the prior probability ratio $\phi(\kappa/2R(f_t, \kappa, \bar{q}))$ multiplied by the exponential of t times the difference in divergence between the two sets $\varepsilon - \kappa - \varepsilon'$.

¹⁰For some triplets $(f_t, \kappa, \bar{q})$, $R(f_t, \kappa, \bar{q})$ can be infinite. However, $f_t \in \Delta^\Theta(Y)$ implies there is at least one p in Θ with finite KL divergence from f_t . Thus the ε' -minimizer $\bar{q}$ also has a finite divergence from f_t , and so has $\bar{q}(z) > 0$ for all $z \in \text{supp } f_t$. The same then holds for the elements of $B_\kappa(\bar{q})$ for sufficiently small κ , so R is finite, which is enough to derive the theorem.

To derive Theorem 1 from Lemma 1, we bound the multiplicative constant away from 0 on $\Delta^\Theta(Y)$. We do this by contradiction, using the compactness of $\Delta^\Theta(Y)$ and the lower semicontinuity of H to show there are $c, \kappa > 0$ such that for every $f_t \in \Delta^\Theta(Y)$ we can pick $\hat{q}_{f_t} \in M_{(1-\alpha)\varepsilon/2}$ such that $R(f_t, \kappa, \hat{q}_{f_t}) < c$. This $\hat{q}_{f_t}$ may not be an exact KL minimizer for f_t , since a minimizer q' that is close to the simplex boundary may have a high value of the ratio $f_t(z)/q'(z)$, so that the KL-divergence changes quickly around q' . Loosely speaking, moving away from the boundary decreases this ratio, and since the minimizers assign very low probability only to outcomes with very low probability under f_t , this does not have much effect on the KL fit, i.e., there is a $(1 - \alpha)\varepsilon/2$ -minimizer sufficiently far from the boundary. Finally, to show that the concentration speed scales linearly in ε (i.e., $\exp(\alpha\varepsilon t)$ for some α), we use the fact that $\hat{q}_{f_t}$ can be chosen in $M_{(1-\alpha)\varepsilon/2}(f_t)$.¹¹

3.2 Implications of Theorem 1

The most direct implication of Theorem 1 is Theorem A, which is the special case where $\Theta = \Delta(Y)$. Here, we use Pinsker's inequality (which gives a lower bound on the KL divergence of q from f_t as a function of $\|q - f_t\|$) and the fact that for a full support prior, $M(f_t) = \{f_t\}$, i.e., the unconstrained minimizer of the Kullback–Leibler divergence is the distribution itself.

PROOF OF THEOREM A. Consider $\varepsilon \in (0, 1)$ and a ϕ positive prior μ_0 . As $H(f_t, f_t) = 0$, all $p \in M_\varepsilon(f_t)$ satisfy $H(p, f_t) \leq \varepsilon$. Pinsker's inequality (Lemma 5) implies that $M_\varepsilon(f_t) \subseteq B_{\sqrt{\varepsilon/2}}(f_t)$. Defining $\tilde{\varepsilon} = \sqrt{\varepsilon/2}$, by Theorem 1 there exists $A(\varepsilon)$ such that

$$\frac{\mu_t(B_{\tilde{\varepsilon}}(f_t))}{1 - \mu_t(B_{\tilde{\varepsilon}}(f_t))} \geq \frac{\mu_t(M_\varepsilon(f_t))}{1 - \mu_t(M_\varepsilon(f_t))} \geq A(\varepsilon) \exp\left(\frac{\varepsilon}{2}t\right) = A(2\tilde{\varepsilon}^2) \exp(\tilde{\varepsilon}^2 t).$$

The result follows by letting $\tilde{A}(\varepsilon) = A(2\varepsilon^2)$ and $g(\varepsilon) = \varepsilon^2$. $\square$

Our result is also closely related to the seminal work by Berk (1966) on long-run beliefs in a misspecified model. The paper showed that when the objective data generating process is i.i.d., beliefs almost surely concentrate on every ε ball around the set of KL minimizers relative to the true outcome distribution p^* .

THEOREM B (Berk (1966)). Let $\mathbb{P}$ be the probability measure induced by i.i.d. draws from p^* . For all $\varepsilon \in (0, 1)$,

$$\lim_{t \rightarrow \infty} \mu_t(B_\varepsilon(M(p^*))) = 1 \quad \mathbb{P}\text{-a.s.}$$

Theorem 1 lets us use the assumption of ϕ positivity to add a rate of convergence to Theorem B when the number of outcomes is finite, as we have assumed so far (Section 6

¹¹Although ε enters linearly in the exponential term of Lemma 1, different values of ε may need different values of ε' and κ , so the overall effect of ε on the concentration rate is not linear.

discusses the case of infinitely many outcomes). The rate of convergence has important implications when the beliefs of an agent are used to solve a decision problem, as it lets us bound the probability of choosing actions that are not optimal with respect to the KL-minimizers as a function of how many outcomes have been observed.

THEOREM 2. *Let $\mathbb{P}$ be the probability measure induced by i.i.d. draws from p^* . If μ_0 is ϕ positive on Θ , then for every $\varepsilon \in (0, 1)$ there is a $K \in \mathbb{R}_{++}$ such that*

$$\mathbb{P}[\mu_t(B_\varepsilon(M(p^*))) < 1 - \exp(-Kt)] = O(\exp(-Kt)).$$

The idea is that if μ_0 is ϕ positive on Θ , then for all $\alpha \in (0, 1)$ there is a function $A: \mathbb{R}_{++} \rightarrow \mathbb{R}_{++}$ such that for all $\varepsilon \in (0, 1)$, $t \in \mathbb{N}$, and $f_t \in \Delta^\Theta(Y)$,

$$\mu_t(M_\varepsilon(f_t)) \geq 1 - \frac{1}{A(\varepsilon)} \exp(-\alpha \varepsilon t).$$

We then show that there is $\hat{\varepsilon} \in \mathbb{R}_{++}$ such that if the empirical frequency is in an $\hat{\varepsilon}$ ball around the objective distribution, i.e., $f_t \in B_{\hat{\varepsilon}}(p^*)$, then $M_{\hat{\varepsilon}}(f_t) \subseteq B_\varepsilon(M(p^*))$, so we can use Theorem 1 and Sanov's theorem to obtain the stated conclusion.

4. GENERAL PARAMETRIC MODELS

In our setting with a finite number of outcomes, we can view the probability distributions themselves as the parameters. However, when the size of the outcome space is large, so is the dimension of $\Delta(Y)$, and people tend to use lower-dimensional parametric models to make the distribution of outcomes easier to think about and analyze. Here, the Diaconis and Freedman (1990) result does not apply, but the extension in this section guarantees that beliefs concentrate exponentially fast around the KL-minimizing parameters themselves, rather than on the ε minimizers, whenever the realized frequency is such that KL divergence is “sufficiently convex” in the parameters.

Consider a parametric model where $\Pi = \{p_\theta: \theta \in \Theta\}$ with p_θ (Gateaux) differentiable in θ , and $\Theta \subset \mathbb{R}^k$ closed and convex, and define $\tilde{H}(f, \theta) = H(f, p_\theta)$.

DEFINITION 3. Let $m > 0$. $\tilde{H}$ is *uniformly strongly m -convex* if for all $f \in \Delta(Y)$,

$$(\nabla_\theta \tilde{H}(f, \theta) - \nabla_\theta \tilde{H}(f, \theta'))^T (\theta - \theta') \geq m \|\theta - \theta'\|_2^2 \quad \forall \theta, \theta' \in \Theta.$$

Intuitively, strong m -convexity ensures that a change in the parameter θ has an effect on the KL divergence that is at least proportional to the square of the change. In the single-dimensional case, strong m -convexity requires that the second derivative of $\tilde{H}$ in θ is bounded away from zero. In the multidimensional case, strong m -convexity is equivalent to the smallest eigenvalue of the Hessian being greater than m . Uniform strong m -convexity extends this property to parametric models. Given a uniformly strongly m -convex $\tilde{H}$, let $\theta^*(f) = \operatorname{argmin}_{\theta \in \Theta} \tilde{H}(f, \theta)$ be the (unique) parameter that minimize the KL divergence between p_θ and f .

Let $B_\varepsilon(\theta) = \{p_\eta: \|\eta - \theta\|_2 \leq \varepsilon\}$ be the set of all distributions whose parameter are at most ε away from θ . The following result establishes the concentration of beliefs about the parameter.

PROPOSITION 1. *If $\tilde{H}$ is uniformly strongly m -convex then for every $\phi : \mathbb{R}_{++} \rightarrow \mathbb{R}_{++}$ and every $\alpha \in (0, 1)$,*

$$\frac{\mu_t(B_\varepsilon(\theta^*(f_t)))}{1 - \mu_t(B_\varepsilon(\theta^*(f_t)))} \geq A\left(\frac{m\varepsilon^2}{2}\right) \exp\left(\alpha \frac{m\varepsilon^2}{2} t\right)$$

for all μ_0 that are ϕ positive on Θ , $\varepsilon \in (0, 1)$, $t \in \mathbb{N}$, and $f_t \in \Delta(Y)$, where A is the function whose existence is guaranteed by Theorem 1.

Intuitively, uniform strong m -convexity ensures that a parameter that is far from the log-likelihood maximizer is assigned a low log-likelihood. Without this assumption, parameters arbitrarily far away from the maximizer could be assigned a likelihood that is arbitrarily close to that of the log-likelihood maximizer, which precludes uniform concentration results.

To see why the proposition is true, note that because Θ is convex, when $\tilde{H}$ is strongly m -convex the KL minimizer $\theta^*(f)$ is a singleton. In addition, the convexity of Θ guarantees that small movements away from the minimizer to other parameters in Θ increase the KL divergence. Uniform strong convexity provides a lower bound on this increase, so we can conclude that parameters outside of $B_\varepsilon(\theta^*(f_t))$ are not $m\varepsilon^2/2$ minimizers. The proposition then follows from Theorem 1.

In the next example, the support restriction comes from the assumption that successive trials are i.i.d., which is a way of simplifying a complex environment.

EXAMPLE 2. (*Bernoulli trials*) Suppose the outcome $y \in Y = \{1, \dots, n\}$ corresponds to the number of Bernoulli trials needed to get one success, with $y = n - 1$ denoting the maximum number of allowed trials.¹² If the agent believes the trials are independent, their subjective distribution for outcome y is a truncated geometric distribution,

$$p_\theta(y) = \begin{cases} \theta(1 - \theta)^{y-1} & \text{for } y < n, \\ (1 - \theta)^{y-1} & \text{for } y = n, \end{cases}$$

and the support of the agent's prior only includes these distributions, i.e., $\Pi \subseteq \{p_\theta : \theta \in [0, 1]\}$. No prior with this support can satisfy the ϕ -positivity condition of Diaconis and Freedman (1990), but our results apply if the prior μ_0 has a density that is bounded away from zero on $[0, 1]$ or is Beta.¹³ We have

$$\tilde{H}(f, \theta) = -\log(1 - \theta) \left[\sum_{z=1}^n (z-1)f(z) \right] - (1 - f(n)) \log(\theta) + \sum_{z=1}^n f(z) \log(f(z)).$$

¹²So when $y = n$, no success occurred in the allowed $n - 1$ trials.

¹³Of course, the particular function ϕ will change. If the density is bounded, ϕ can be chosen linear in ε , while for Beta priors ϕ can be chosen to be a power function of ε .

This function is uniformly strongly 1-convex on $\Delta(Y)$, and the first-order condition shows that the unique KL minimizing parameter is given by

$$\theta^*(f) = \frac{1 - f(n)}{\sum_{z \in Y} z f(z) - f(n)}.$$

Thus from Proposition 1 beliefs about θ concentrate on any $B_\varepsilon(\theta^*(f))$ exponentially fast. Moreover, if the data has no realizations of n , the KL minimizer is the reciprocal of the average outcome $\sum_{z \in Y} z f(z)$. This is intuitive, as the expectation of a geometric distribution is the reciprocal of the parameter θ , i.e., $\lim_{n \rightarrow \infty} \sum_{z \in Y} p_\theta(z) z = 1/\theta$. $\diamond$

5. SUBJECTIVELY MARKOVIAN ENVIRONMENTS

We now show how to generalize Theorem 1 to the case of beliefs that the signals y are generated by a Markov process, which is a key environment in macroeconomics. For example, this is the setting where the approximation properties of the anticipated utility model have been analyzed.

In the Markov setting, the agent is learning about n different outcome distributions; let $\mathcal{P} = \Delta(Y)^Y$ be the set of transition matrices over Y endowed with the total variation distance.¹⁴ Let $\mu_0 \in \Delta(\mathcal{P}) = \Delta(\Delta(Y)^Y)$ denote a prior distribution over transition matrices and $\Theta = \text{supp } \mu_0$ its support.¹⁵

To initialize the process, we fix an observed period 0 outcome y_0 . For every data set y^t , we let μ_t be the posterior belief, which is required to satisfy Bayes rule whenever the data set has positive prior probability:

$$\mu_t(C) = \frac{\int_C \prod_{\tau=1}^t \pi(y_\tau | y_{\tau-1}) d\mu_0(\pi)}{\int_{\mathcal{P}} \prod_{\tau=1}^t \pi(y_\tau | y_{\tau-1}) d\mu_0(\pi)}. \quad (\text{Bayes Rule})$$

The *empirical transition distribution* $f_t \in \Delta(Y \times Y)$ is

$$f_t(z, z') = \frac{1}{t} \sum_{\tau=1}^t \mathbb{I}_{\{z', z\}}(y_\tau, y_{\tau-1}).$$

We define $\mathcal{H} : \Delta(Y \times Y) \times \mathcal{P} \rightarrow \bar{\mathbb{R}}$ as

$$\mathcal{H}(f, \pi) = - \sum_{(z, z') \in Y \times Y} f(z, z') \log(\pi(z' | z)).$$

¹⁴That is, for all $\chi \in (\mathbb{R}^Y)^Y$, $\|\chi\| = \sum_{z, z' \in Y} |\chi(z' | z)|/2$.

¹⁵Note that the μ_0 need not be a product measure, and that this reduces to the subjectively i.i.d. environment of the previous sections if for every $\pi \in \Theta$ and every $z, z' \in Y$, $\pi(\cdot | z) = \pi(\cdot | z')$.

The function $\mathcal{H}$ generalizes H to the non-i.i.d. case, as $\mathcal{H}(f, \pi)$ measures the log-likelihood assigned to the empirical transitions distribution f given the transition probability π :

$$\log \left(\prod_{\tau=1}^t \pi(y_\tau | y_{\tau-1}) \right) = t \sum_{(z, z') \in Y \times Y} f_t(z, z') \log \pi(z' | z) = -t \mathcal{H}(f_t, \pi).$$

Let $\mathcal{M} : \Delta(Y \times Y) \rightrightarrows \mathcal{P}$ be the correspondence that maps an empirical transition distribution f to the minimizers of $\mathcal{H}$ over the support of the prior: $\mathcal{M}(f) = \operatorname{argmin}_{\pi \in \Theta} \mathcal{H}(f, \pi)$. We let $\mathcal{M}_\varepsilon(f)$ be the set of distributions that come within ε of the minimum of $\mathcal{H}$:

$$\mathcal{M}_\varepsilon(f) = \left\{ \pi' \in \Theta : \mathcal{H}(f, \pi') \leq \min_{\pi \in \Theta} \mathcal{H}(f, \pi) + \varepsilon \right\},$$

and let $\Delta^\Theta(Y \times Y)$ denote the set of empirical transition distributions for which Bayes rule is well-defined.¹⁶

THEOREM 3. *Suppose that for all $\pi, \pi' \in \Theta$, $z, z' \in Y$, $\pi(z' | z) > 0$ if and only if $\pi'(z' | z) > 0$ and that μ_0 is ϕ positive on Θ . Then for all $\alpha \in (0, 1)$ and $\varepsilon \in (0, 1)$ there is $A(\varepsilon) > 0$ such that*

$$\frac{\mu_t(\mathcal{M}_\varepsilon(f_t))}{1 - \mu_t(\mathcal{M}_\varepsilon(f_t))} \geq A(\varepsilon) \exp(\alpha \varepsilon t),$$

for all $t \in \mathbb{N}$ and $f_t \in \Delta^\Theta(Y \times Y)$.

The proof of this result is similar in spirit to that of Theorem 1, because we can consider the data set to be a sequence of pairs (y_t, y_{t+1}) in place of a sequence of y_t .

EXAMPLE 3. In Markov settings, priors without full support arise very naturally. For example, when $Y \subseteq \mathbb{R}$, the agent may believe that the data generating process is a martingale, so that the support of the beliefs consist of all distributions $\pi \in \mathcal{P}$ for which

$$\sum_{z' \in Y} \pi(z' | z) z' = z \quad \forall z \in Y.$$

Alternatively, they may believe that the outcome process is unlikely to make large jumps between consequent periods: for some $\alpha_z \in (0, 1)$,

$$\pi(z' | z) = \frac{\alpha_z^{|z' - z|}}{\sum_{z' \in Y} \alpha_z^{|z' - z|}} \quad \forall z \in Y.$$

◇

¹⁶That is, $f \in \Delta^\Theta(Y \times Y)$ if there is a $\pi \in \Theta$ such that for all $(z, z') \in \operatorname{supp} f$, $\pi(z' | z) > 0$.

6. INFINITE OUTCOME SPACES

This section discusses pathwise concentration in the case of an infinite outcome space.

Noncompact prior support The most common approach when dealing with infinitely many outcomes is to use a parametric description of the data generating process, as was done in Section 4 for a finite outcome space. However, if the set of parameters Θ indexing the data generating process is not compact, it will be typically not be possible to satisfy ϕ -positivity. For example, if the prior is supported over the normal distributions with some fixed variance σ^2 and unknown mean $\theta \in \mathbb{R}$, and all values of the mean are considered possible, the prior cannot be ϕ -positive, as for any $\varepsilon > 0$, $\mu(B_\varepsilon(\theta)) \geq \phi(\varepsilon) > 0$ for all $\theta \in \Theta$ would imply that $\mu(\mathbb{R}) = \infty$.

Similarly, pathwise concentration fails: Pathwise concentration requires that a finite number of observations can outweigh the prior, but no fixed finite number of observations can outweigh the prior if the prior probability of the ε -minimizing set can be arbitrarily low. For this reason, pathwise concentration can be obtained only for priors with compact support.

Divergence vs. likelihood The empirical distribution is always discrete, but the KL divergence from a discrete distribution to a nonatomic one is infinite. To handle this, we shift from concentration around the KL-minimizer to concentration around the maximizer of the empirical log-likelihood. Since the likelihood is the negative of the divergence plus a constant, these coincide in the case of a finite Y , but only the empirical log-likelihood maximizers are always well-defined for a continuum of outcomes. With this change, we now extend Lemma 1 to the case of infinitely many outcomes.

Let Y be a metric space, and suppose that there exists a σ -finite measure ξ on Y such that for every $\theta \in \Theta$, the probability measure associated with θ is absolutely continuous with respect to ξ with Radon–Nykodim derivative $p_\theta \in \mathbb{R}^Y$. Let $\Pi = \{p_\theta : \theta \in \Theta\}$ and P_s be the set of simple (finite support) distributions over Y . Balls in Π are taken with respect to the supnorm. For every $\theta \in \Theta$ and $q \in P_s$, let

$$L(q||p_\theta) = \sum_{y \in Y} q(y) \log p_\theta(y)$$

be the empirical log-likelihood of the empirical distribution q under $p_\theta \in \Pi$. Also, let

$$M_\varepsilon(q) = \left\{ p_{\theta'} \in \Pi : L(q||p_{\theta'}) + \varepsilon \geq \max_{\theta \in \Theta} L(q||p_\theta) \right\}$$

be the set of ε maximizers of the empirical log-likelihood. Recall that

$$\Delta^\Theta(Y) := \{f \in P_s : \exists \theta \in \Theta, \forall y \in \text{supp } f, p_\theta(y) > 0\}$$

is the set of empirical frequencies for which Bayesian updating is well-defined.

LEMMA 2. For every $\varepsilon, \varepsilon', \kappa \in \mathbb{R}_+$, $t \in \mathbb{N}$, $f_t \in \Delta^\Theta(Y)$, $\bar{q} \in M_{\varepsilon'}(f_t)$, with $\varepsilon' + \kappa \leq \varepsilon$,

$$\frac{\mu_t(M_\varepsilon(f_t))}{1 - \mu_t(M_\varepsilon(f_t))} \geq \mu_0(B_{\kappa/R(f_t, \kappa, \bar{q})}(\bar{q})) \exp((\varepsilon - \kappa - \varepsilon')t),$$

where

$$R(f_t, \kappa, \bar{q}) = \max \left\{ \max_{\substack{q \in \Pi \cap B_\kappa(\bar{q}) \\ z \in \text{supp } f_t}} \frac{1}{q(z)}, 1 \right\}.$$

As in the finite case, m -convexity is useful for guaranteeing belief concentration, but it is harder to satisfy m -convexity when Y is infinite. For this reason, we generalize the m -convexity to only hold on a given set of empirical frequencies.

DEFINITION 4. Let $m > 0$. L is *uniformly strongly m -concave* on $\mathcal{F} \subseteq P_s$ if for all $f \in \mathcal{F}$,

$$(\nabla_\theta L(f||p_\theta) - \nabla_\theta L(f||p_{\theta'}))^T (\theta - \theta') \leq -m \|\theta - \theta'\|_2^2$$

for all $\theta, \theta' \in \Theta$.

Let $\theta^*(f_t)$ denote the empirical likelihood maximizer. In the case of a single-dimensional parameter, uniform strong m -concavity on $\mathcal{F}$ is still enough to prove that posteriors concentrate on a neighborhood of the unique maximizer at a rate that is uniform over paths with empirical frequency in $\mathcal{F}$. Example 5 in the Appendix shows that convergence need not be uniform over frequencies that do not make the likelihood function m -uniformly concave.

The main difficulty is that we have little information about which ε movements from $\theta^*(f_t)$ least decrease the empirical likelihood. The proof uses the fact that when the parameters are unidimensional there are at most two candidates for a best fitting parameter outside $B_\varepsilon(\theta^*(f_t))$ (either $\theta^*(f_t) - \varepsilon$ or $\theta^*(f_t) + \varepsilon$) to overcome this difficulty.

PROPOSITION 2. If L is uniformly strongly m -concave on $\mathcal{F}$, then for every $\phi : \mathbb{R}_{++} \rightarrow \mathbb{R}_{++}$,

$$\frac{\mu_t(B_\varepsilon(\theta^*(f_t)))}{1 - \mu_t(B_\varepsilon(\theta^*(f_t)))} \geq \phi\left(\frac{\varepsilon}{4}\right) \exp\left\{t \frac{\varepsilon^2 m}{2}\right\}$$

for all μ_0 that are ϕ positive on $\Theta \subseteq \mathbb{R}$, $\varepsilon \in (0, 1)$, $t \in \mathbb{N}$, and $f_t \in \Delta^\Theta(Y) \cap \mathcal{F}$.

As an immediate corollary, there is pathwise belief concentration for an agent who believes the data are generated by a normal distribution with known variance and unknown mean.

COROLLARY 1. Let $\sigma^2 \in \mathbb{R}_{++}$ and

$$p_\theta(y) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(y - \theta)^2}{2\sigma^2}\right).$$

For every $\phi : \mathbb{R}_{++} \rightarrow \mathbb{R}_{++}$, and every $\alpha \in (0, 1)$,

$$\frac{\mu_t(B_\varepsilon(\theta^*(f_t)))}{1 - \mu_t(B_\varepsilon(\theta^*(f_t)))} \geq \phi\left(\frac{\varepsilon}{4}\right) \exp\left\{t \frac{\varepsilon^2}{2\sigma}\right\}$$

for all μ_0 that are ϕ positive on $\Theta \subseteq \mathbb{R}$, $\varepsilon \in (0, 1)$, $t \in \mathbb{N}$, and $f_t \in \Delta^\Theta(Y)$.

No concentration on the ε maximizers Without additional assumptions, Proposition 2 cannot be strengthened to obtain pathwise concentration on the ε maximizers, as Example 6 in Appendix A.11 shows. Intuitively, with infinitely many signals, the informativeness of a single signal may be unbounded, so that the set of ε maximizers after a single signal can be arbitrarily small. If the prior probability assigned to these sets vanishes at a sufficiently high exponential rate, their good match to the data does not guarantee that the posterior concentrates on them. More precisely, for some priors the conclusion of Theorem 1 does not even hold for $t = 1$. That is, it is not possible to have a concentration that holds uniformly over all the same-length realizations, let alone concentration rate that is uniform over same-length realizations and grows exponentially in the sample size.¹⁷ We leave for future work the challenge of determining just what sorts of restrictions on the prior would allow a uniform concentration result.

Anticipated utility Much of the macroeconomics literature assumes that the data agents observe can take infinitely many different values. This is true in particular for the literature on “anticipated utility” (Kreps (1998)), which assumes that agents in the economy choose actions that maximize their payoff under a point estimate that maximizes the likelihood of their sample, ignoring uncertainty about the state. This is a simpler problem than the maximization of expected utility, and the reduction in complexity and dimension makes anticipated utility models more tractable and easier to analyze. However, it has not been clear how much error the approximation induces. For example, Cogley and Sargent (2008) wrote:

Macroeconomists might justify anticipated-utility models as an approximation to a correctly formulated Bayesian decision problem... (the models) would be more compelling if one could also show that anticipated-utility decisions well approximate Bayesian decisions. As far as we know, no one has assessed the quality of the approximation...

There is also a small literature that addresses this question using numerical simulations (Cogley, Colacito, and Sargent (2007), Cogley and Sargent (2008), Cogley et al. (2008)). Our result on Bayesian updating can be used to derive analytical results that complement these numerical studies. In particular, they imply that the long-run behavior under anticipated utility models converges to that of an expected utility maximizer. This provides a formal justification for the use of anticipated as an approximation of expected utility models in studies of long-run behavior.

To develop this link, suppose that in each period $t \in \{1, 2, 3, \dots\}$ the agent chooses an action from A . We assume that A is a convex set, endowed with a metric d that makes it a compact set. The action does not affect the outcome distribution but influences the agent’s utility function $u : A \times Y \rightarrow \mathbb{R}$, which is strictly concave in a . Let $A^*(\nu)$ denote the (unique) optimal action given belief ν , i.e.,

$$A^*(\nu) = \operatorname{argmax}_{a \in A} \int_{\Theta} \mathbb{E}_{p_{\theta}}[u(a, y)] d\nu(\theta),$$

¹⁷Fudenberg, He, and Imhof (2017) and Fudenberg, Lanzani, and Strack (2021a) point out other odd implications of priors that decay exponentially quickly.

and suppose A^* is uniformly continuous when $\Delta(\Theta)$ is endowed with the topology of weak convergence of measures.¹⁸

Let $A^*(M(f_t))$ denote the action that is optimal for a point belief in the likelihood maximizer $M(f_t)$.

PROPOSITION 3. *Suppose that $\Theta \subseteq \mathbb{R}$ is convex and that μ_0 is ϕ positive on Θ . If L is uniformly strongly m -concave on $\mathcal{F}$, then for all $\varepsilon > 0$ there is a $T \in \mathbb{N}$ such that $d(A^*(\mu_t), A^*(M(f_t))) \leq \varepsilon$ for every $t > T$ and every $f_t \in \mathcal{F}$.*

7. CONCLUSION

We have shown that for every realization of the data, Bayesian beliefs concentrate exponentially quickly on the models that best explain the empirical frequency of outcomes. One implication of this concentration result is that optimal actions can be determined directly from the empirical frequency without computing beliefs. More precisely, once the sample is sufficiently large, neither the exact sample size nor calendar time is needed to compute the optimal action; the empirical frequency is sufficient. As the dynamics and distribution of the empirical frequency are well understood, this insight can greatly simplify the analysis of the long-run behavior of Bayesian agents.

In addition to the applications developed in this paper, Theorem 1 may allow generalizations of other results about misspecified Bayesian agents who learn from endogenous data. One recurrent theme in this literature is the possibility that when actions are endogenous, misspecified beliefs can lead to cycles in setting that would not occur with correctly specified beliefs, because repeated play of an action generates evidence in favor of another action.¹⁹ In such situations, our concentration result may be used to bound the number of periods spent in each phase of the cycles. This would complement [Esponda, Pouzo, and Yamamoto \(2021\)](#), which characterized the asymptotic frequencies of these cycles when the space of beliefs can be partitioned into a finite number of attracting sets and the support of the prior is one-dimensional. Our uniform speed of convergence result might be useful in extending this to more general settings. In addition, as we provide a concentration bound for every finite time, our result can be used to characterize behavior in the “medium-run” before the asymptotic results apply.

Mazumdar, Pacchiano, Ma, Bartlett, and Jordan (2020) prove that with high probability, the posteriors of a correctly-specified Bayesian concentrate around the true parameter at rate $\sqrt{n}$, and use this result to study the long-run properties of Thompson sampling. The paper allows for infinitely many outcomes, but imposes additional strong conditions such as log-concavity of the true data generating process, and a prior density that is bounded away from 0. Our results enable extensions to Thompson sampling with less restricted priors in the finite outcome case.

In settings where multiple agents choose their actions based on the same observables, our concentration results can be used to quantify the minimal extent of the differences in their prior beliefs needed to rationalize different choices. For example, Olea,

¹⁸A sufficient condition for this is that Θ is compact and A^* is continuous.

¹⁹See, e.g., [Nyarko \(1991\)](#), [Fudenberg, Romanyuk, and Strack \(2017\)](#), [Levy, Razin, and Young \(2020\)](#), and [Lanzani \(2022\)](#).

Luis, Ortoleva, Pai, and Prat (2021) showed that when observing signals of an object's value, misspecified agents with lower-dimensional models have a higher willingness to pay after the first few observations, while correctly specified agents have a higher willingness to pay in the long-run; our result on the speed of convergence may help to better identify the switching time.

The learning in games literature has assumed correctly specified beliefs to appeal to Diaconis and Freedman (1990). Our generalization will facilitate the extension of the results from this literature to cases where the agents in the learning model have misspecified beliefs about the extensive form of the game. It will also enable extensions to incorrect beliefs about a complex network structure in Bowen, Dmitriev, and Galperti (forthcoming), and to overconfident agents as in Heidhues, Kőszegi, and Strack (2018).

APPENDIX

A.1 Properties of the KL divergence

LEMMA 3. For all $p, \tilde{p}, q \in P$,

$$|H(q, p) - H(q, \tilde{p})| \leq 2 \max_{z \in Y} \max \left\{ \frac{q(z)}{p(z)}, \frac{q(z)}{\tilde{p}(z)} \right\} \|p - \tilde{p}\|.$$

PROOF. Let

$$R := \max_{z \in Y} \max \left\{ \frac{q(z)}{p(z)}, \frac{q(z)}{\tilde{p}(z)} \right\},$$

$\hat{Y} = \{z : p(z) > \tilde{p}(z)\}$, and suppose without loss of generality that $p(\text{supp } q) \geq \tilde{p}(\text{supp } q)$. Then

$$\begin{aligned} & |H(q, p) - H(q, \tilde{p})| \\ &= \left| \sum_{z \in \text{supp } q} \left(\log \left(\frac{p(z)}{q(z)} \right) - \log \left(\frac{\tilde{p}(z)}{q(z)} \right) \right) q(z) \right| \\ &= \left| \sum_{z \in \text{supp } q} \int_{\tilde{p}(z)/q(z)}^{p(z)/q(z)} \frac{1}{r} dr q(z) \right| \\ &\leq \sum_{z \in \text{supp } q} \max \left\{ \frac{q(z)}{\tilde{p}(z)}, \frac{q(z)}{p(z)} \right\} \left| \frac{p(z)}{q(z)} - \frac{\tilde{p}(z)}{q(z)} \right| q(z) \\ &\leq R \sum_{z \in \text{supp } q} \left| \frac{p(z)}{q(z)} - \frac{\tilde{p}(z)}{q(z)} \right| q(z) \\ &= R \sum_{z \in \text{supp } q} (2\mathbb{I}_{\hat{Y}}(z) - 1) \left(\frac{p(z)}{q(z)} - \frac{\tilde{p}(z)}{q(z)} \right) q(z) \\ &= R \sum_{z \in \text{supp } q} (2\mathbb{I}_{\hat{Y}}(z) - 1) p(z) - R \sum_{z \in \text{supp } q} (2\mathbb{I}_{\hat{Y}}(z) - 1) \tilde{p}(z) \end{aligned}$$

$$= 2R \left[\sum_{z \in \text{supp } q} \mathbb{I}_{\hat{Y}}(z) p(z) - \sum_{z \in \text{supp } q} \mathbb{I}_{\hat{Y}}(z) \tilde{p}(z) \right] + R(\tilde{p}(\text{supp } q) - p(\text{supp } q)).$$

As $p(\text{supp } q) \geq \tilde{p}(\text{supp } q)$, the above term is bounded by

$$\begin{aligned} &\leq 2R \left[\sum_{z \in \text{supp } q} \mathbb{I}_{\hat{Y}}(z) p(z) - \sum_{z \in \text{supp } q} \mathbb{I}_{\hat{Y}}(z) \tilde{p}(z) \right] \leq 2R \left[\sum_{z \in \hat{Y}} p(z) - \sum_{z \in \hat{Y}} \tilde{p}(z) \right] \\ &= 2R \|p - \tilde{p}\|, \end{aligned}$$

where the last inequality follows from the definition of $\hat{Y}$ and the last equality by the definition of the total variation distance. $\square$

Recall that a probability distribution $p \in \Delta(Y)$ is absolutely continuous with respect to $q \in \Delta(Y)$, denoted as $p \ll q$, if $\text{supp } p \subseteq \text{supp } q$.

LEMMA 4. *Let $\varepsilon \in \mathbb{R}_+$. Then $M_\varepsilon(\cdot) = \{q' \in \Theta : H(\cdot, q') \leq \min_{q \in \Theta} H(\cdot, q) + \varepsilon\}$ is nonempty-valued and compact-valued. Moreover, for all $p \in P$, $M_\varepsilon(\cdot)$ is upper hemicontinuous on $B_{\inf_{z \in \text{supp } p} p(z)/2}(p) \cap \{q : q \ll p\}$.*

PROOF. If $H(p, q) = \infty$ for all $q \in \Theta$, $M_\varepsilon(p) = \Theta$ is nonempty and compact. If there is $\hat{q}$ such that $H(p, \hat{q}) = K < \infty$, the set $\Theta' = \{q \in \Theta : H(p, q) \leq K + \varepsilon\}$ is compact by the continuity of $H(p, \cdot)$, and $M_\varepsilon(p) \subseteq \Theta'$. So, the continuous and real-valued restriction of $H(p, \cdot)$ to Θ' has compact lower contour sets, and it attains a minimum. Thus $M_\varepsilon(p) \supseteq M(p) \neq \emptyset$ is nonempty and compact.

For the second part of the statement, observe that if $p \notin \Delta^\Theta(Y)$, then $M_\varepsilon(p) = \Theta$ and, therefore, $M_\varepsilon(p)$ is trivially upper hemicontinuous at p since by definition $M_\varepsilon(p') \subseteq \Theta$ for all $p' \in \Delta(Y)$. If instead $p \in \Delta^\Theta(Y)$, there exist $\hat{q} \in \Theta$ and $K \in \mathbb{R}_+$ with $H(p, \hat{q}) = K$. Moreover, the finiteness of $H(p, \hat{q}) = K$ implies that $p \ll \hat{q}$. So, there exists $K' > 0$ such that

$$H(p', \hat{q}) \leq K' \quad \forall p' \in B_{\inf_{z \in \text{supp } p} p(z)/2}(p) \cap \{q : q \ll p\}.$$

We use this equation to show that there exists C such that $r \in M_\varepsilon(p')$, $p' \in B_{\inf_{z \in \text{supp } p} p(z)/2}(p) \cap \{q : q \ll p\}$ implies $r(y) \geq C$ for all $y \in \text{supp } p$. Suppose by contradiction that this is not the case. Then there exist a convergent sequence $(r_n, p_n)_{n \in \mathbb{N}} \in (\Theta \times (B_{\inf_{z \in \text{supp } p} p(z)/2}(p) \cap \{q : q \ll p\}))^{\mathbb{N}}$ and an $\hat{y} \in \text{supp } p$ with $r_n \in M_\varepsilon(p_n)$ for all $n \in \mathbb{N}$ and $\lim_{n \rightarrow \infty} r_n(\hat{y}) = 0$. But we have

$$H(p_n, r_n) \geq \sum_{y \in Y} p_n(y) \log p_n(y) - p_n(\hat{y}) \log r_n(\hat{y}) \geq \sum_{y \in Y} p_n(y) \log p_n(y) - \frac{p}{2} \log r_n(\hat{y})$$

and the right-hand side is diverging to ∞ . So, eventually $H(p_n, r_n) \geq \varepsilon + K' \geq \varepsilon + H(p_n, \hat{q})$, a contradiction with $r_n \in M_\varepsilon(p_n)$.

This shows that for all $p' \in B_{\inf_{z \in \text{supp } p} p(z)/2}(p) \cap \{q : q \ll p\}$,

$$\begin{aligned} M_\varepsilon(p') &= \left\{ r \in \Theta : H(p', r) \leq \min_{\bar{r} \in \Theta} H(p', \bar{r}) + \varepsilon \right\} \\ &= \left\{ r \in \Theta : H(p', r) \leq \min_{\bar{r} \in \Theta} H(p', \bar{r}) + \varepsilon, r(y) \geq C, \forall y \in \text{supp } p \right\}. \end{aligned} \quad (4)$$

Also, observe that the function H is continuous on the set

$$(B_{\inf_{z \in \text{supp } p} p(z)/2}(p) \cap \{q : q \ll p\}) \times \{r \in \Theta : r(y) \geq C, \forall y \in \text{supp } p\}.$$

Therefore, if we define $G : (B_{\inf_{z \in \text{supp } p} p(z)/2}(p) \cap \{q : q \ll p\}) \rightarrow \mathbb{R}$ as

$$G(p') = \min_{\{r \in \Theta : r(y) \geq C, \forall y \in \text{supp } p\}} H(p', r)$$

G is continuous by the maximum theorem. Moreover, by equation (4) $G(p') = \min_{r \in \Theta} H(p', r)$ for all $B_{\inf_{z \in \text{supp } p} p(z)/2}(p) \cap \{q : q \ll p\}$, showing that $\min_{r \in \Theta} H(\cdot, r)$ is a continuous function when restricted on $B_{\inf_{z \in \text{supp } p} p(z)/2}(p) \cap \{q : q \ll p\}$. To conclude, we show that $M_\varepsilon(\cdot)$ is upper hemicontinuous on $B_{\inf_{z \in \text{supp } p} p(z)/2}(p) \cap \{q : q \ll p\}$ by showing that it has a closed graph. Indeed, let $(p_n, r_n) \in B_{\inf_{z \in \text{supp } p} p(z)/2}(p) \cap \{q : q \ll p\} \times \Theta$ be such that $r_n \in M_\varepsilon(p_n)$ for all $n \in \mathbb{N}$ and $\lim_{n \rightarrow \infty} (r_n, p_n) = (\hat{r}, \hat{p})$. By equation (4), for all $n \in \mathbb{N}$, we have $r_n \in \{r \in \Theta : r(y) \geq C, \forall y \in \text{supp } p\}$ and since this last set is close $\hat{r} \in \{r \in \Theta : r(y) \geq C, \forall y \in \text{supp } p\}$. By the continuity of H on $(B_{\inf_{z \in \text{supp } p} p(z)/2}(p) \cap \{q : q \ll p\}) \times \{r \in \Theta : r(y) \geq C, \forall y \in \text{supp } p\}$, and of G on $B_{\inf_{z \in \text{supp } p} p(z)/2}(p) \cap \{q : q \ll p\}$, we have

$$\min_{r \in \Theta} H(\hat{p}, r) - H(\hat{p}, \hat{r}) = G(\hat{p}) - H(\hat{p}, \hat{r}) = \lim_{n \rightarrow \infty} (G(p_n) - H(p_n, r_n)) \leq \varepsilon$$

proving that $\hat{r} \in M_\varepsilon(\hat{p})$. □

LEMMA 5 (Pinsker's inequality). *For every $p, q \in \Delta(Y)$,*

$$\|p - q\| \leq \sqrt{\frac{H(p, q)}{2}}.$$

A.2 Proof of Lemma 1

LEMMA 1. *If μ_0 is ϕ positive on Θ , then for every $\varepsilon, \varepsilon', \kappa \in \mathbb{R}_+$, $t \in \mathbb{N}$, $f_t \in \Delta^\Theta(Y)$, and $\bar{q} \in M_{\varepsilon'}(f_t)$ with $\varepsilon' + \kappa \leq \varepsilon$ and $R(f_t, \kappa, \bar{q}) < \infty$, we have*

$$\frac{\mu_t(M_\varepsilon(f_t))}{1 - \mu_t(M_\varepsilon(f_t))} \geq \phi(\kappa/2R(f_t, \kappa, \bar{q})) \exp((\varepsilon - \kappa - \varepsilon')t). \quad (5)$$

PROOF. The proof uses the following bound on how much the Kullback–Leibler divergence can increase when moving from an ε' minimizer to a nearby distribution.

CLAIM 1. *For every $p' \in \Theta$, $f \in \Delta^\Theta(Y)$, $\varepsilon', \kappa \in \mathbb{R}_+$, and $\bar{q} \in M_{\varepsilon'}(f)$,*

$$p' \in B_{\kappa/2R(f, \kappa, \bar{q})}(\bar{q}) \implies H(f, p') \leq \min_{p \in \Theta} H(f, p) + \varepsilon' + \kappa.$$

PROOF. For every two distributions $f, q \in \Delta(Y)$, there is at least one outcome that is weakly more likely under f than under q , so

$$R(f, \kappa, \bar{q}) = \max_{q \in \Theta \cap B_\kappa(\bar{q})} \max_{z \in Y} \frac{f(z)}{q(z)}$$

is bounded below by 1. Thus $p' \in B_{\kappa/2R(f, \kappa, \bar{q})}(\bar{q})$ implies $p' \in B_\kappa(\bar{q})$. Therefore, both $p', \bar{q}$ are in $\Theta \cap B_\kappa(\bar{q})$, so from the definition of R ,

$$\max_{z \in Y} \max \left\{ \frac{f(z)}{p'(z)}, \frac{f(z)}{\bar{q}(z)} \right\} \leq R(f, \kappa, \bar{q}).$$

Moreover, $p' \in B_{\kappa/2R(f, \kappa, \bar{q})}(\bar{q})$ implies that $\bar{q} \in B_{\kappa/2R(f, \kappa, \bar{q})}(p') \cap M_{\varepsilon'}(f)$, so by Lemma 3,

$$H(f, p') - H(f, \bar{q}) \leq \frac{\kappa}{R(f, \kappa, \bar{q})} \max_{z \in Y} \max \left\{ \frac{f(z)}{p'(z)}, \frac{f(z)}{\bar{q}(z)} \right\} \leq \frac{\kappa}{R(f, \kappa, \bar{q})} R(f, \kappa, \bar{q}) = \kappa,$$

and hence $H(f, p') \leq H(f, \bar{q}) + \kappa \leq \min_{p \in \Theta} H(f, p) + \varepsilon' + \kappa$. $\square$

We use Claim 1 to provide a lower bound on the probability of the ε minimizers given the empirical frequency f_t . Observe that

$$\begin{aligned} \frac{\mu_t(M_\varepsilon(f_t))}{1 - \mu_t(M_\varepsilon(f_t))} &= \frac{\int_{M_\varepsilon(f_t)} \exp(-H(p, f_t)t) d\mu_0(dp)}{\int_{\Theta \setminus M_\varepsilon(f_t)} \exp(-H(p, f_t)t) d\mu_0(dp)} \\ &\geq \frac{\int_{M_{\kappa+\varepsilon'}(f_t)} \exp(-H(p, f_t)t) d\mu_0(dp)}{\int_{\Theta \setminus M_\varepsilon(f_t)} \exp(-H(p, f_t)t) d\mu_0(dp)} \\ &\geq \frac{\exp\left(-\left(\min_{p \in \Theta} H(f_t, p) + \kappa + \varepsilon'\right)t\right)}{\exp\left(-\left(\min_{p \in \Theta} H(f_t, p) + \varepsilon\right)t\right)} \frac{\mu_0(M_{\kappa+\varepsilon'}(f_t))}{\mu_0(\Theta \setminus M_\varepsilon(f_t))} \\ &= \exp((\varepsilon - \kappa - \varepsilon')t) \frac{\mu_0(M_{\kappa+\varepsilon'}(f_t))}{\mu_0(\Theta \setminus M_\varepsilon(f_t))} \\ &\geq \exp((\varepsilon - \kappa - \varepsilon')t) \mu_0(B_{\kappa/2R(f_t, \kappa, \bar{q})}(\bar{q})) \\ &\geq \exp((\varepsilon - \kappa - \varepsilon')t) \phi(\kappa/2R(f_t, \kappa, \bar{q})). \end{aligned}$$

The first equality follows from equation (1). The first inequality follows from $\varepsilon' + \kappa \leq \varepsilon$, the second inequality from pointwise bounding the integrands and the definition of M_ε , the third inequality from Claim 1, and the fourth because μ_0 is ϕ positive on Θ . $\square$

A.3 Proof of Theorem 1

THEOREM 1. For every $\phi : \mathbb{R}_{++} \rightarrow \mathbb{R}_{++}$, $\alpha \in (0, 1)$, and $\varepsilon \in (0, 1)$ there is $A(\varepsilon)$ such that

$$\frac{\mu_t(M_\varepsilon(f_t))}{1 - \mu_t(M_\varepsilon(f_t))} \geq A(\varepsilon) \exp(\alpha \varepsilon t)$$

for all $t \in \mathbb{N}$, $f_t \in \Delta^\Theta(Y)$, and μ_0 that is ϕ positive on Θ . Moreover, if $\underline{q} := \inf_{q \in \Theta} \min_{z \in \text{supp } q} q(z) > 0$, then we can set

$$A(\varepsilon) = \phi(\min\{\underline{q}/2, (1 - \alpha)\varepsilon\}\underline{q}/2).$$

We prove the theorem for ϕ nondecreasing. This is without loss of generality, as if μ_0 is ϕ positive on Θ , it is also $\hat{\phi}$ positive on Θ where

$$\hat{\phi}(\varepsilon) = \sup_{\varepsilon' \leq \varepsilon} \phi(\varepsilon') \quad \forall \varepsilon \in \mathbb{R}_{++}.$$

Clearly, $\hat{\phi}$ is nondecreasing and only depends on ϕ . Moreover, it also pointwise weakly dominates ϕ , so when $\underline{q} > 0$, if we prove the statement for $\hat{\phi}$ we have

$$\begin{aligned} \frac{\mu_t(M_\varepsilon(f_t))}{1 - \mu_t(M_\varepsilon(f_t))} &\geq \hat{\phi}(\min\{\underline{q}/2, (1 - \alpha)\varepsilon\}\underline{q}/2) \exp(\alpha \varepsilon t) \\ &\geq \phi(\min\{\underline{q}/2, (1 - \alpha)\varepsilon\}\underline{q}/2) \exp(\alpha \varepsilon t) \end{aligned}$$

proving the statement for ϕ as well.

We first show that if $\underline{q} := \inf_{q \in \Theta} \min_{z \in \text{supp } q} q(z) > 0$, Lemma 1 yields the desired uniform rate of convergence. If $(1 - \alpha)\varepsilon < \underline{q}$, then for all $f_t \in \Delta^\Theta(Y)$ and $\bar{q} \in M(f_t)$, if $p \in \Theta \cap B_{(1-\alpha)\varepsilon}(\bar{q})$, then $\text{supp } p = \text{supp } \bar{q} \subseteq \text{supp } f_t$, so

$$\frac{1}{R(f_t, (1 - \alpha)\varepsilon, \bar{q})} = \left(\max_{p \in \Theta \cap B_{(1-\alpha)\varepsilon}(\bar{q})} \max_{z \in Y} \frac{f_t(z)}{p(z)} \right)^{-1} \geq \frac{1}{(1/\underline{q})} = \underline{q}.$$

If instead $(1 - \alpha)\varepsilon \geq \underline{q}$, it is enough to observe that $1/R(f_t, \underline{q}/2, \bar{q}) \geq \underline{q}$ for all $f_t \in \Delta^\Theta(Y)$, $\bar{q} \in M(f_t)$.

Now we move to the proof for the general case where some outcomes might have an arbitrarily low probability under data-generating processes in the support of the prior, i.e., $\underline{q}$ might equal 0. Recall that $\Delta^\Theta(Y) = \{q \in \Delta(Y) : \exists p \in \Theta, \text{supp } q \subseteq \text{supp } p\}$ is the set of distributions for which Bayes rule is well-defined, and that Theorem 1 applies only to empirical distributions $f \in \Delta^\Theta(Y)$. To provide an upper bound on R , we show that the likelihood ratio f/q (which determines the value of R) can be uniformly bounded for all probability distributions q that are sufficiently close to an $(1 - \alpha)\varepsilon/2$ minimizer of the Kullback–Leibler divergence. Intuitively, as $f \in \Delta^\Theta(Y)$ some distribution assigns nonvanishing probability to every outcome which has positive probability under f , and thus a distribution that assigns vanishing probability to some of these outcomes leads to an excessively low log-likelihood (and thus a high Kullback–Leibler divergence).

CLAIM 2. For every $\varepsilon \in (0, 1)$, there exist $\bar{\kappa}_\alpha(\varepsilon) \in (0, (1 - \alpha)\varepsilon/2]$ and $c \geq 1$ such that for all $\kappa \leq \bar{\kappa}_\alpha$ and $f \in \Delta^\Theta(Y)$, there is $\bar{q} \in M_{(1-\alpha)\varepsilon/2}(f)$ such that

$$\max_{y \in Y, q \in B_\kappa(\bar{q})} \frac{f(y)}{q(y)} \leq c.$$

PROOF. If not, then since $\Delta^\Theta(Y)$ and Θ are compact, there is a sequence $(f'_n, q_n) \in \Delta^\Theta(Y) \times \Theta$ with $q_n \in M(f'_n)$ that converges to $(\hat{f}, \hat{q})$, and such that

$$\inf_{\bar{q} \in M_{\frac{(1-\alpha)\varepsilon}{2}}(f'_n)} \left(\max_{y \in Y, q \in B_{1/n}(\bar{q})} \frac{f'_n(y)}{q(y)} \right) \geq n. \quad (6)$$

Since Y is finite, so is the set of possible supports, so there is a subsequence $(f_n)_{n \in \mathbb{N}}$ such that each element of the sequence has common support, with $(f_n(y))_n$ weakly decreasing for all $y \in Y \setminus \text{supp } \hat{f}$. Moreover, since

$$\sum_{z \in Y} f_n(z) \log f_n(z) \in \left[\log \left(\frac{1}{|Y|} \right), 0 \right] \quad \forall n \in \mathbb{N}$$

the subsequence can also be taken such that $\sum_{z \in Y} f_n(z) \log f_n(z)$ converges.

Since all the f_n are in $\Delta^\Theta(Y)$ and have common support, $q_n \in M(f_n)$, and $\log(f_n(z)) \leq 0$ for all $z \in Y$, we have

$$\begin{aligned} H(f_n, q_n) &\leq H(f_n, q_1) = \sum_{z \in \text{supp } f_1} f_n(z) \log f_n(z) - \sum_{z \in \text{supp } f_1} f_n(z) \log q_1(z) \\ &\leq - \sum_{z \in \text{supp } f_1} f_n(z) \log q_1(z) \leq - \min_{z \in \text{supp } f_1} \log q_1(z) < \infty, \end{aligned}$$

so $(H(f_n, q_n))_{n \in \mathbb{N}}$ is bounded. Moreover, if there exist $z^* \in Y$ and $l \in \mathbb{R}_{++}$ with $\lim_{n \rightarrow \infty} q_n(z^*) = 0$ and $\lim_{n \rightarrow \infty} f_n(z^*) = l$, then

$$\begin{aligned} \limsup_{n \rightarrow \infty} H(f_n, q_n) &= \limsup_{n \rightarrow \infty} \sum_{y \in Y} f_n(y) (\log f_n(y) - \log q_n(y)) \\ &= \sum_{y \in Y} \hat{f}(y) \log \hat{f}(y) + \limsup_{n \rightarrow \infty} - \sum_{y \in Y} f_n(y) \log q_n(y) \\ &\geq \sum_{y \in Y} \hat{f}(y) \log \hat{f}(y) + \limsup_{n \rightarrow \infty} -f_n(z^*) \log q_n(z^*) \\ &= \sum_{y \in Y} \hat{f}(y) \log \hat{f}(y) - l \log \left(\lim_{n \rightarrow \infty} q_n(z^*) \right) = \infty, \end{aligned}$$

which contradicts $(H(f_n, q_n))_{n \in \mathbb{N}}$ being bounded. So, for all $z \in Y$,

$$\lim_{n \rightarrow \infty} q_n(z) = 0 \implies \lim_{n \rightarrow \infty} f_n(z) = 0 \implies z \notin \text{supp } \hat{f}.$$

Thus $G := \inf_{n \in \mathbb{N}, y \in \text{supp } \hat{f}} \log q_n(y) > -\infty$.

Since $(f_n)_{n \in \mathbb{N}}$ converges, there is $N \in \mathbb{N}$ such that for all n and m larger than N ,

$$\|f_n - f_m\| \leq \frac{1}{|G| \cdot |\text{supp } f_1|} \frac{(1-\alpha)\varepsilon}{16}. \quad (7)$$

Because $\hat{H} := \liminf H(f_n, q_n) < \infty$, there exists $N' \geq N$ such that

$$H(f_{N'}, q_{N'}) \leq \hat{H} + \frac{(1-\alpha)\varepsilon}{8} \quad (8)$$

and

$$\left| \sum_{z \in Y} f_n(z) \log f_n(z) - \sum_{z \in Y} f_m(z) \log f_m(z) \right| \leq \frac{(1-\alpha)\varepsilon}{8} \quad \forall n, m \geq N'. \quad (9)$$

Moreover, there exists $N'' > N'$ such that for all $n \geq N''$,

$$H(f_n, q_n) \geq \hat{H} - \frac{(1-\alpha)\varepsilon}{8}. \quad (10)$$

Thus for every $n \geq N''$, we have

$$\begin{aligned} & H(f_n, q_{N'}) - H(f_n, q_n) \\ & \leq H(f_n, q_{N'}) - \hat{H} + \frac{(1-\alpha)\varepsilon}{8} \\ & = H(f_n, q_{N'}) - H(f_{N'}, q_{N'}) + H(f_{N'}, q_{N'}) - \hat{H} + \frac{(1-\alpha)\varepsilon}{8} \\ & = \sum_{z \in Y} f_n(z) \log f_n(z) - \sum_{z \in Y} f_{N'}(z) \log f_{N'}(z) \\ & \quad + \sum_{z \in Y} (f_n(z) - f_{N'}(z))(-\log q_{N'}(z)) + H(f_{N'}, q_{N'}) - \hat{H} + \frac{(1-\alpha)\varepsilon}{8} \\ & \leq \frac{(1-\alpha)\varepsilon}{8} + \sum_{z \in \text{supp } \hat{f}} (f_n(z) - f_{N'}(z))(-\log q_{N'}(z)) \\ & \quad + \sum_{z \in Y \setminus \text{supp } \hat{f}} (f_n(z) - f_{N'}(z))(-\log q_{N'}(z)) + H(f_{N'}, q_{N'}) - \hat{H} + \frac{(1-\alpha)\varepsilon}{8} \\ & \leq \frac{(1-\alpha)\varepsilon}{4} + \sum_{z \in \text{supp } \hat{f}} (f_n(z) - f_{N'}(z))(-\log q_{N'}(z)) + H(f_{N'}, q_{N'}) - \hat{H} \\ & \leq \frac{(1-\alpha)\varepsilon}{4} + \frac{1}{|G|} \frac{(1-\alpha)\varepsilon}{16} \cdot 2|G| + H(f_{N'}, q_{N'}) - \hat{H} \\ & \leq \frac{(1-\alpha)\varepsilon}{4} + \frac{(1-\alpha)\varepsilon}{8} + \hat{H} + \frac{(1-\alpha)\varepsilon}{8} - \hat{H} = \frac{(1-\alpha)\varepsilon}{2}, \end{aligned}$$

where the first inequality follows by equation (10), the second by equation (9), the third because $f_m(z)$ is decreasing for the outcomes outside $\text{supp } \hat{f}$, the fourth by equation (7)

and the definition of G , and the fifth by equation (8). Therefore,

$$q_{N'} \in M_{\frac{(1-\alpha)\varepsilon}{2}}(f_n)$$

for all $n \geq N''$. But this is a contradiction with equation (6), as

$$\lim_{n \rightarrow \infty} \max_{y \in Y, q \in B_{1/n}(q_{N'})} \frac{f_n(y)}{q(y)} \leq \lim_{n \rightarrow \infty} \max_{y \in \text{supp } q_{N'}} \frac{1}{q_{N'}(y)/2} < \infty. \quad \square$$

Now for every $f_t \in \Delta^\Theta(Y)$ let c , $\bar{\kappa}_\alpha(\varepsilon)$, and $\bar{q} \in M_{(1-\alpha)\varepsilon/2}(f_t)$ be the values whose existence is established by Claim 2. Then

$$\begin{aligned} \frac{\mu_t(M_\varepsilon(f_t))}{1 - \mu_t(M_\varepsilon(f_t))} &\geq \phi(\bar{\kappa}_\alpha(\varepsilon)/2R(f_t, \bar{\kappa}_\alpha(\varepsilon), \bar{q})) \exp\left(\left(\varepsilon - \bar{\kappa}_\alpha(\varepsilon) - \frac{(1-\alpha)\varepsilon}{2}\right)t\right) \\ &\geq \phi(\bar{\kappa}_\alpha(\varepsilon)/2c) \exp\left(\left(\varepsilon - \frac{(1-\alpha)\varepsilon}{2} - \frac{(1-\alpha)\varepsilon}{2}\right)t\right) \\ &\geq \phi(\bar{\kappa}_\alpha(\varepsilon)/2c) \exp(\alpha \varepsilon t), \end{aligned}$$

where the first inequality follows from applying Lemma 1 with $\varepsilon' = (1-\alpha)\varepsilon/2$ and $\kappa = \bar{\kappa}_\alpha(\varepsilon)$, the second inequality follows because Claim 2 shows that

$$c \geq \max_{y \in Y, q \in B_{\bar{\kappa}_\alpha(\varepsilon)}(\bar{q})} \frac{f_t(y)}{q(y)} \geq R(f_t, \bar{\kappa}_\alpha(\varepsilon), \bar{q})$$

and $\bar{\kappa}_\alpha(\varepsilon) \leq (1-\alpha)\varepsilon/2$, and the third inequality is algebra. Theorem 1 then follows by letting

$$A(\varepsilon) = \phi\left(\frac{\bar{\kappa}_\alpha(\varepsilon)}{2c}\right).$$

A.4 Proof of Theorem 2

THEOREM 4 (Sanov's theorem, Sanov (1961)). *Let $\mathbb{P}$ be the probability measure induced by i.i.d. draws from p^* . Then for all $A \subseteq \Delta(Y)$ and for $t \in \mathbb{N}$,*

$$\mathbb{P}[f_t \in A] \leq (t+1)^{|Y|} 2^{-t \min_{p \in A} H(p, p^*)}.$$

PROOF. See, e.g., Dupuis and Ellis (2011). □

THEOREM 2. *Let $\mathbb{P}$ be the probability measure induced by i.i.d. draws from p^* . If μ_0 is ϕ positive on Θ , then for every $\varepsilon \in (0, 1)$ there is a $K \in \mathbb{R}_{++}$ such that*

$$\mathbb{P}[\mu_t(B_\varepsilon(M(p^*))) < 1 - \exp(-Kt)] = O(\exp(-Kt)).$$

CLAIM 3. *For every $\varepsilon > 0$ and $p^* \in P$, there exists $\varepsilon' > 0$ such that*

$$M_{\varepsilon'}(p^*) \subseteq B_\varepsilon(M(p^*)).$$

PROOF. Assume the claim is false, so for every $n \in \mathbb{N}$, there exists $q_n \in \Theta \setminus B_\varepsilon(M(p^*))$ such that

$$H(p^*, q_n) - \min_{p \in \Theta} H(p^*, p) \leq \frac{1}{n}.$$

Since Θ is compact, $(q_n)_{n \in \mathbb{N}}$ admits a convergent subsequence with limit $q^* \in \Theta$. Since $H(p^*, \cdot)$ is lower semicontinuous, $q^* \in M(p^*)$. But this would imply that the subsequence is eventually in $B_\varepsilon(M(p^*))$, a contradiction. $\square$

By Claim 3, $M_{\varepsilon'}(p^*) \subseteq B_\varepsilon(M(p^*))$ for some $\varepsilon' > 0$. Since $M_{\varepsilon'}(\cdot)$ is upper hemi-continuous by Lemma 4, there exists $\hat{\varepsilon}$ such that if $\hat{q} \in B_{\hat{\varepsilon}}(p^*) \cap \{q : q \ll p^*\}$, then $M_{\varepsilon'/2}(\hat{q}) \subseteq M_{\varepsilon'}(p^*)$. Because the actual data generating process is p^* and Y is finite, $\mathbb{P}[f_t \ll p^*] = 1$ for all $t \in \mathbb{N}$. By Sanov's theorem (Theorem 4) and Pinsker's inequality (Lemma 5), for all t large enough to have $(t+1)^{|Y|} \leq 2^{\hat{\varepsilon}^2 t}$,

$$\mathbb{P}[f_t \notin B_{\hat{\varepsilon}}(p^*) \cap \{q : q \ll p^*\}] \leq \mathbb{P}[H(f_t, p^*) \geq 2\hat{\varepsilon}^2] \leq (t+1)^{|Y|} 2^{-2\hat{\varepsilon}^2 t} \leq 2^{-\hat{\varepsilon}^2 t}.$$

So,

$$\mathbb{P}[\mu_t(B_\varepsilon(M(p^*))) < 1 - \bar{K} \exp(-\hat{K}t)] = O(\exp(-K't))$$

follows from Theorem 1 by letting $\hat{K} = \alpha\varepsilon'/2$ for $\alpha \in (0, 1)$, $\bar{K} = 1/A(\varepsilon'/2)$, and $K' = \hat{\varepsilon}^2 \log 2$. Since $\lim_{t \rightarrow \infty} \frac{\bar{K} \exp(-\hat{K}t)}{\exp(-Ct)} = 0$ for all $C < \hat{K}$, the result follows by letting $K = \min\{\hat{K}/2, K'\}$.

A.5 Proof of Proposition 1

PROPOSITION 1. *If $\tilde{H}$ is uniformly strongly m -convex then for every $\phi : \mathbb{R}_{++} \rightarrow \mathbb{R}_{++}$ and every $\alpha \in (0, 1)$,*

$$\frac{\mu_t(B_\varepsilon(\theta^*(f_t)))}{1 - \mu_t(B_\varepsilon(\theta^*(f_t)))} \geq A\left(\frac{m\varepsilon^2}{2}\right) \exp\left(\alpha \frac{m\varepsilon^2}{2} t\right)$$

for all μ_0 that are ϕ positive on Θ , $\varepsilon \in (0, 1)$, $t \in \mathbb{N}$, and $f_t \in \Delta(Y)$, where A is the function whose existence is guaranteed by Theorem 1.

PROOF. We claim first that for every $\theta \in \Theta$ and $f \in \Delta^\Theta(Y)$, $\nabla_\theta \tilde{H}(f, \theta^*(f))^T (\theta - \theta^*(f)) \geq 0$. If not,

$$0 > \nabla_\theta \tilde{H}(f, \theta^*(f))^T (\theta - \theta^*(f)) = \lim_{k \rightarrow 0} \frac{\tilde{H}(f, \theta^*(f) + k(\theta - \theta^*(f))) - \tilde{H}(f, \theta^*(f))}{k}.$$

But this means that there is $\hat{k} \in (0, 1)$ such that $\tilde{H}(f, \theta^*(f) + \hat{k}(\theta - \theta^*(f))) - \tilde{H}(f, \theta^*(f)) < 0$ or $\tilde{H}(f, (1 - \hat{k})\theta^*(f) + \hat{k}\theta) < \tilde{H}(f, \theta^*(f))$. As Θ is convex, $(1 - \hat{k})\theta^*(f) + \hat{k}\theta$ belongs to Θ , but then $\theta^*(f)$ would not be a KL-minimizer.

Next, for every $\theta \in \Theta$,

$$\begin{aligned}\tilde{H}(f, \theta) - \tilde{H}(f, \theta^*(f)) &= \int_0^1 \nabla_{\theta} \tilde{H}(f, \lambda \theta + (1 - \lambda) \theta^*(f))^T (\theta - \theta^*(f)) d\lambda \\ &\geq \nabla_{\theta} \tilde{H}(f, \theta^*(f))^T (\theta - \theta^*(f)) + \frac{m}{2} \|\theta - \theta^*(f)\|_2^2 \geq \frac{m}{2} \|\theta - \theta^*(f)\|_2^2,\end{aligned}$$

where the second inequality comes from the uniform strong m -convexity of $\tilde{H}$. As a consequence,

$$M_{\frac{m\varepsilon^2}{2}}(f) \subseteq B_{\varepsilon}(\theta^*(f))$$

and the result follows from Theorem 1. $\square$

A.6 Proof of Theorem 3

THEOREM 3. *Suppose that for all $\pi, \pi' \in \Theta$, $z, z' \in Y$, $\pi(z'|z) > 0$ if and only if $\pi'(z'|z) > 0$ and that μ_0 is ϕ positive on Θ . Then for all $\alpha \in (0, 1)$ and $\varepsilon \in (0, 1)$ there is an $A(\varepsilon) > 0$ such that*

$$\frac{\mu_t(\mathcal{M}_{\varepsilon}(f_t))}{1 - \mu_t(\mathcal{M}_{\varepsilon}(f_t))} \geq A(\varepsilon) \exp(\alpha \varepsilon t),$$

for all $t \in \mathbb{N}$ and $f_t \in \Delta^{\Theta}(Y \times Y)$.

We begin with a continuity result that extends Lemma 3 to the Markov setting.

CLAIM 4. *Let $R := \max_{\pi \in \Theta, z \in Y, z' \in \text{supp } \pi(\cdot|z)} (1/\pi(z'|z))$. For all $\pi, \tilde{\pi} \in \Theta$ and $f \in \Delta^{\Theta}(Y \times Y)$,*

$$|\mathcal{H}(f, \pi) - \mathcal{H}(f, \tilde{\pi})| \leq 2R \|\pi - \tilde{\pi}\|.$$

PROOF.

$$\begin{aligned}|\mathcal{H}(f, \pi) - \mathcal{H}(f, \tilde{\pi})| &= \left| \sum_{(z, z') \in \text{supp } f} (\log(\pi(z'|z)) - \log(\tilde{\pi}(z'|z))) f(z, z') \right| \\ &= \left| \sum_{(z, z') \in \text{supp } f} \int_{\tilde{\pi}(z'|z)}^{\pi(z'|z)} \frac{1}{r} f(z, z') dr \right| \\ &\leq \sum_{(z, z') \in \text{supp } f} \max \left\{ \frac{1}{\tilde{\pi}(z'|z)}, \frac{1}{\pi(z'|z)} \right\} |\pi(z'|z) - \tilde{\pi}(z'|z)| f(z, z') \\ &\leq R \sum_{(z, z') \in \text{supp } f} |\pi(z'|z) - \tilde{\pi}(z'|z)| f(z, z') \leq 2R \|\pi - \tilde{\pi}\|.\end{aligned}$$

Here, the first inequality follows from pointwise bounding the integrand, and the second inequality follows from the fact that for all $\pi, \pi' \in \Theta$, $z, z' \in Y$, $\pi(z'|z) > 0$ if and only

$\pi'(z'|z) > 0$ and $f \in \Delta^\Theta(Y \times Y)$. The last inequality follows from the definition of the total variation distance. $\square$

We now use Claim 4 to establish the theorem. Fix $\varepsilon \in \mathbb{R}_{++}$. Rewrite the likelihood ratio for distributions inside and outside of $\mathcal{M}_\varepsilon(f_t)$ as follows:

$$\begin{aligned} \frac{\mu_t(\mathcal{M}_\varepsilon(f_t))}{1 - \mu_t(\mathcal{M}_\varepsilon(f_t))} &= \frac{\int_{\mathcal{M}_\varepsilon(f_t)} \prod_{\tau=1}^t \pi(y_\tau | y_{\tau-1}) d\mu_0(\pi)}{\int_{\Theta \setminus \mathcal{M}_\varepsilon(f_t)} \prod_{\tau=1}^t \pi(y_\tau | y_{\tau-1}) d\mu_0(\pi)} \\ &= \frac{\int_{\mathcal{M}_\varepsilon(f_t)} \exp\left(\sum_{z, z' \in Y} f_t(z, z') \log(\pi(z|z'))\right) t d\mu_0(\pi)}{\int_{\Theta \setminus \mathcal{M}_\varepsilon(f_t)} \exp\left(\sum_{z, z' \in Y} f_t(z, z') \log(\pi(z|z'))\right) t d\mu_0(\pi)} \\ &= \frac{\int_{\mathcal{M}_\varepsilon(f_t)} \exp(-\mathcal{H}(f_t, \pi)t) d\mu_0(\pi)}{\int_{\Theta \setminus \mathcal{M}_\varepsilon(f_t)} \exp(-\mathcal{H}(f_t, \pi)t) d\mu_0(\pi)}. \end{aligned}$$

Next, we provide a lower bound on this likelihood ratio:

$$\begin{aligned} &\frac{\int_{\mathcal{M}_\varepsilon(f_t)} \exp(-\mathcal{H}(f_t, \pi)t) d\mu_0(\pi)}{\int_{\Theta \setminus \mathcal{M}_\varepsilon(f_t)} \exp(-\mathcal{H}(f_t, \pi)t) d\mu_0(\pi)} \\ &\geq \frac{\int_{\mathcal{M}_{((1-\alpha)\varepsilon)}(f_t)} \exp(-\mathcal{H}(f_t, \pi)t) d\mu_0(\pi)}{\exp\left(-\left[\min_{\pi \in \Theta} \mathcal{H}(f_t, \pi) + \varepsilon\right]t\right)} \\ &\geq \frac{\int_{\mathcal{M}_{((1-\alpha)\varepsilon)}(f_t)} \exp\left(-\left[\min_{\pi \in \Theta} \mathcal{H}(f_t, \pi) + (1-\alpha)\varepsilon\right]t\right) d\mu_0(\pi)}{\exp\left(-\left[\min_{\pi \in \Theta} \mathcal{H}(f_t, \pi) + \varepsilon\right]t\right)} \\ &\geq \frac{\int_{B_{((1-\alpha)\varepsilon/2R)}(\mathcal{M}(f_t))} \exp\left(-\left[\min_{\pi \in \Theta} \mathcal{H}(f_t, \pi) + (1-\alpha)\varepsilon\right]t\right) d\mu_0(\pi)}{\exp\left(-\left[\min_{\pi \in \Theta} \mathcal{H}(f_t, \pi) + \varepsilon\right]t\right)} \\ &\geq \phi((1-\alpha)\varepsilon/2R) \frac{\exp\left(-\left(\min_{\pi \in \Theta} \mathcal{H}(f_t, \pi) + (1-\alpha)\varepsilon\right)t\right)}{\exp\left(-\left[\min_{\pi \in \Theta} \mathcal{H}(f_t, \pi) + \varepsilon\right]t\right)} \\ &= \phi((1-\alpha)\varepsilon/2R) \exp(\alpha\varepsilon t). \end{aligned}$$

Here, the first and second inequalities follow from the definitions of $\mathcal{M}_\varepsilon$, the third inequality from Claim 4, and the fourth inequality from ϕ positivity on Θ . The result follows by setting $A(\varepsilon) = \phi((1 - \alpha)\varepsilon/2R)$.

A.7 Counterexamples to Proposition 1

EXAMPLE 4. (*Connected Θ is not sufficient for Proposition 1*) ◇

Let $Y = \{A, B, C\}$, $\Theta = \{p : p(A)p(B)p(C) = 0\} \setminus \{p : p(A), p(B) \subseteq (1/4, 3/4)\}$, and $\hat{\mu}_0$ be the uniform measure on Θ . Let $\mu'_0 = \delta_{(1/4, 3/4, 0)}/2 + \delta_{(3/4, 1/4, 0)}/2$ and $\mu_0 = \hat{\mu}_0/2 + \mu'_0/2$. Suppose $y_t = B$ if t is odd and $y_t = A$ if t is even. At every odd period $2t + 1$, $t \geq 1$, $M(f_{2t+1}) = \{(1/4, 3/4, 0)\}$, and for $\varepsilon < 1/12$,

$$\begin{aligned} \lim_{t \rightarrow \infty} \frac{\mu_{2t+1}(B_\varepsilon(M(f_{2t+1})))}{1 - \mu_{2t+1}(B_\varepsilon(M(f_{2t+1})))} &\leq \lim_{t \rightarrow \infty} \frac{\mu_{2t+1}(B_\varepsilon(1/4, 3/4, 0))}{\mu_{2t+1}(B_\varepsilon(3/4, 1/4, 0))} \\ &\leq \lim_{t \rightarrow \infty} \frac{\mu_0(B_\varepsilon(1/4, 3/4, 0))}{\mu_0(\{3/4, 1/4, 0\})} \frac{(1/4)^t (3/4)^{t+1}}{(1/4)^{t+1} (3/4)^t} \\ &= 3 \frac{\mu_0(B_\varepsilon(1/4, 3/4, 0))}{\mu_0(\{3/4, 1/4, 0\})} \end{aligned}$$

so beliefs do not concentrate on the KL minimizer.

EXAMPLE 5. (*Convergence is not uniform over all paths*) ◇

This example shows that even if Θ is convex, beliefs need not to converge to the KL minimizers along paths where the empirical distribution converges to the boundary.

Let $Y = \{A, B, C\}$, $\Theta = \{p : p(A) = 1/3\}$, and μ_0 be the uniform measure on Θ . When $f_{2n} = (1 - 1/n, 1/2n, 1/2n)$, $M(f_{2n}) = \{(1/3, 1/3, 1/3)\}$ for all $n \in \mathbb{N}$. However, fix an $\varepsilon \in (0, 1/12)$. Then

$$\begin{aligned} &\frac{\mu_{2n}(B_\varepsilon(M(f_{2n})))}{1 - \mu_{2n}(B_\varepsilon(M(f_{2n})))} \\ &\leq \frac{\int_{B_\varepsilon(M(f_{2n}))} \exp(-H(f_{2n}, p)2n) d\mu_0(p)}{\int_{P \setminus B_\varepsilon(M(f_{2n}))} \exp(-H(f_{2n}, p)2n) d\mu_0(p)} \\ &\leq \frac{\mu_0(B_\varepsilon(M(f_{2n}))) \exp(-H(f_{2n}, (1/3, 1/3, 1/3))2n)}{\mu_0(B_{2\varepsilon}(M(f_{2n})) \setminus B_\varepsilon(M(f_{2n}))) \exp(-H(f_{2n}, (1/3, 1/3 + 2\varepsilon, 1/3 - 2\varepsilon))2n)} \\ &= \frac{\mu_0(B_\varepsilon(M(f_{2n})))}{\mu_0(B_{2\varepsilon}(M(f_{2n})) \setminus B_\varepsilon(M(f_{2n})))} \frac{(1/3)(1/3)}{(1/3 + 2\varepsilon)(1/3 - 2\varepsilon)} \\ &\rightarrow_n \frac{\mu_0(B_\varepsilon(\{(1/3, 1/3, 1/3)\}))}{\mu_0(B_{2\varepsilon}(\{(1/3, 1/3, 1/3)\}) \setminus B_\varepsilon(\{(1/3, 1/3, 1/3)\}))} \frac{(1/3)(1/3)}{(1/3 + 2\varepsilon)(1/3 - 2\varepsilon)}, \end{aligned}$$

so beliefs do not concentrate.

A.8 Proof of Lemma 2

LEMMA 2. For every $\varepsilon, \varepsilon', \kappa \in \mathbb{R}_+, t \in \mathbb{N}, f_t \in \Delta^\Theta(Y), \bar{q} \in M_{\varepsilon'}(f_t)$, with $\varepsilon' + \kappa \leq \varepsilon$,

$$\frac{\mu_t(M_\varepsilon(f_t))}{1 - \mu_t(M_\varepsilon(f_t))} \geq \mu_0(B_{\kappa/R(f_t, \kappa, \bar{q})}(\bar{q})) \exp((\varepsilon - \kappa - \varepsilon')t),$$

where

$$R(f_t, \kappa, \bar{q}) = \max \left\{ \max_{\substack{q \in \Pi \cap B_\kappa(\bar{q}) \\ z \in \text{supp } f_t}} \frac{1}{q(z)}, 1 \right\}.$$

PROOF. The proof follows from the following claims.

CLAIM 5. For all $p, \hat{p} \in \Pi, q \in P_s$,

$$|L(q||p) - L(q||\hat{p})| \leq \max_{z \in \text{supp } q} \max \left\{ \frac{1}{p(z)}, \frac{1}{\hat{p}(z)} \right\} \|p - \hat{p}\|_\infty.$$

PROOF.

$$\begin{aligned} & |L(q||p) - L(q||\hat{p})| \\ &= \left| \sum_{z \in \text{supp } q} (\log(p(z)) - \log(\hat{p}(z))) q(z) \right| \\ &= \left| \sum_{z \in \text{supp } q} \int_{\hat{p}(z)}^{p(z)} \frac{1}{r} dr q(z) \right| \leq \sum_{z \in \text{supp } q} \max \left\{ \frac{1}{\hat{p}(z)}, \frac{1}{p(z)} \right\} |p(z) - \hat{p}(z)| q(z) \\ &\leq \max_{z \in \text{supp } q} \max \left\{ \frac{1}{p(z)}, \frac{1}{\hat{p}(z)} \right\} \sum_{z \in \text{supp } q} |p(z) - \hat{p}(z)| q(z) \\ &\leq \max_{z \in \text{supp } q} \max \left\{ \frac{1}{p(z)}, \frac{1}{\hat{p}(z)} \right\} \|p - \hat{p}\|_\infty, \end{aligned}$$

where the last equality follows from the definition of the supremum distance. $\square$

CLAIM 6. For every $p' \in \Pi, f \in \Delta^\Theta(Y), \varepsilon', \kappa \in \mathbb{R}_+$, and $\bar{q} \in M_{\varepsilon'}(f)$,

$$p' \in B_{\kappa/R(f, \kappa, \bar{q})}(\bar{q}) \implies L(f||p') + \varepsilon' + \kappa \geq \max_{\theta \in \Theta} L(f||p_\theta).$$

PROOF. Since R is bounded below by 1, $p' \in B_{\kappa/R(f, \kappa, \bar{q})}(\bar{q})$ implies $p' \in B_\kappa(\bar{q})$. Therefore, both $p', \bar{q}$ are in $\Pi \cap B_\kappa(\bar{q})$, so from the definition of R , $\max_{z \in \text{supp } f} \max \left\{ \frac{1}{p'(z)}, \frac{1}{\bar{q}(z)} \right\} \leq R(f, \kappa, \bar{q})$. Moreover, $p' \in B_{\kappa/R(f, \kappa, \bar{q})}(\bar{q})$ implies that $\bar{q} \in B_{\kappa/R(f, \kappa, \bar{q})}(p') \cap M_{\varepsilon'}(f)$, so by Claim 5,

$$L(f||\bar{q}) - L(f||p') \leq \frac{\kappa}{R(f, \kappa, \bar{q})} \max_{z \in \text{supp } f} \max \left\{ \frac{1}{p'(z)}, \frac{1}{\bar{q}(z)} \right\} \leq \frac{\kappa}{R(f, \kappa, \bar{q})} R(f, \kappa, \bar{q}) = \kappa,$$

and hence

$$L(f||p') \geq L(f||\bar{q}) - \kappa \geq \max_{\theta \in \Theta} L(f||p_\theta) - \varepsilon' - \kappa.$$

□

To prove the lemma, observe that

$$\begin{aligned} \frac{\mu_t(M_\varepsilon(f_t))}{1 - \mu_t(M_\varepsilon(f_t))} &\geq \frac{\int_{M_{\kappa+\varepsilon'}(f_t)} \exp(L(f_t||p)t) d\mu_0(dp)}{\int_{\Pi \setminus M_\varepsilon(f_t)} \exp(L(f_t||p)t) d\mu_0(dp)} \\ &\geq \frac{\exp\left(\left(\max_{\theta \in \Theta} L(f_t||p_\theta) - \kappa - \varepsilon'\right)t\right)}{\exp\left(\left(\max_{\theta \in \Theta} L(f_t||p_\theta) - \varepsilon\right)t\right)} \frac{\mu_0(M_{\kappa+\varepsilon'}(f_t))}{\mu_0(\Pi \setminus M_\varepsilon(f_t))} \\ &= \exp((\varepsilon - \kappa - \varepsilon')t) \frac{\mu_0(M_{\kappa+\varepsilon'}(f_t))}{\mu_0(\Pi \setminus M_\varepsilon(f_t))} \\ &\geq \exp((\varepsilon - \kappa - \varepsilon')t) \mu_0(B_{\kappa/R(f_t, \kappa, \bar{q})}(\bar{q})). \end{aligned}$$

The first inequality follows from $\varepsilon' + \kappa \leq \varepsilon$, the second from pointwise bounding the integrands and the definition of M_ε , and the third from Claim 6. □

A.9 Proof of Proposition 2

PROPOSITION 2. *If L is uniformly strongly m -concave on $\mathcal{F}$, then for every $\phi : \mathbb{R}_{++} \rightarrow \mathbb{R}_{++}$,*

$$\frac{\mu_t(B_\varepsilon(\theta^*(f_t)))}{1 - \mu_t(B_\varepsilon(\theta^*(f_t)))} \geq \phi\left(\frac{\varepsilon}{4}\right) \exp\left\{t \frac{\varepsilon^2 m}{2}\right\}$$

for all μ_0 that are ϕ positive on $\Theta \subseteq \mathbb{R}$, $\varepsilon \in (0, 1)$, $t \in \mathbb{N}$, and $f_t \in \Delta^\Theta(Y) \cap \mathcal{F}$.

PROOF. We claim first that for every $\theta \in \Theta$ and $f \in \Delta^\Theta(Y) \cap \mathcal{F}$, $\nabla_\theta L(f||p_{\theta^*(f)})^T(\theta - \theta^*(f)) \leq 0$. If not,

$$0 < \nabla_\theta L(f||p_{\theta^*(f)})^T(\theta - \theta^*(f)) = \lim_{k \rightarrow 0} \frac{L(f||p_{\theta^*(f)+k(\theta-\theta^*(f))}) - L(f||p_{\theta^*(f)})}{k}.$$

But this means that there is $\hat{k} \in (0, 1)$ such that $L(f||p_{\theta^*(f)+\hat{k}(\theta-\theta^*(f))}) - L(f||p_{\theta^*(f)}) > 0$. As Θ is convex, $(1 - \hat{k})\theta^*(f) + \hat{k}\theta$ belongs to Θ , but then $\theta^*(f)$ would not be a likelihood maximizer.

Next, as L is uniformly strongly m -concave,

$$L(f||p_\theta) - L(f||p_{\theta^*(f)}) \leq \nabla_\theta L(f||p_{\theta^*(f)})(\theta - \theta^*(f)) - \frac{m}{2}|\theta - \theta^*(f)| \leq -\frac{m}{2}|\theta - \theta^*(f)|.$$

The statement is trivially true if $\Theta \subseteq B_\varepsilon(\theta^*(f))$. If not, since $L(f||p_{(\cdot)})$ is concave and Θ is convex, at least one of

$$\theta + \varepsilon \in \operatorname{argmax}_{\theta: |\theta - \theta^*(f)| \geq \varepsilon} L(f||p_\theta),$$

and

$$\theta - \varepsilon \in \operatorname{argmax}_{\theta: |\theta - \theta^*(f)| \geq \varepsilon} L(f||p_\theta)$$

holds. We prove the result in the first case, the proof for the other case is symmetric. Let $\bar{\theta}, \theta' \in \Theta$ be such that

$$\bar{\theta} \geq \theta^*(f) + \varepsilon \geq \theta^*(f) + \frac{\varepsilon}{2} \geq \theta' \geq \theta^*(f).$$

We have

$$\begin{aligned} L(f||p_{\bar{\theta}}) - L(f||p_{\theta'}) &\leq L_\theta(f||p_{\theta'}) (\bar{\theta} - \theta') - \frac{m}{2} (\bar{\theta} - \theta')^2 \\ &\leq L_\theta(f||p_{\theta^*(f)}) (\bar{\theta} - \theta') - \frac{m}{2} (\bar{\theta} - \theta')^2 \leq -\frac{m}{2} (\bar{\theta} - \theta')^2 \leq -\frac{\varepsilon^2 m}{2}, \end{aligned}$$

where the first inequality follows from the strong m -concavity of L , and the second by the concavity of L in θ . Therefore,

$$\begin{aligned} \frac{\mu_t(B_\varepsilon(\theta^*(f)))}{1 - \mu_t(B_\varepsilon(\theta^*(f)))} &\geq \frac{\mu_t\left(\left[\theta^*(f), \theta^*(f) + \frac{\varepsilon}{2}\right]\right)}{1 - \mu_t(B_\varepsilon(\theta^*(f)))} \\ &\geq \frac{\mu_0\left(\left[\theta^*(f), \theta^*(f) + \frac{\varepsilon}{2}\right]\right)}{1 - \mu_0(B_\varepsilon(\theta^*(f)))} \exp\left(t \frac{\varepsilon^2 m}{2}\right) \\ &\geq \mu_0\left(\left[\theta^*(f), \theta^*(f) + \frac{\varepsilon}{2}\right]\right) \exp\left(t \frac{\varepsilon^2 m}{2}\right) \\ &= \phi\left(\frac{\varepsilon}{4}\right) \exp\left(t \frac{\varepsilon^2 m}{2}\right). \end{aligned} \quad \square$$

A.10 Proof of Proposition 3

PROPOSITION 3. *Suppose that $\Theta \subseteq \mathbb{R}$ is convex and that μ_0 is ϕ positive on Θ . If L is uniformly strongly m -concave on $\mathcal{F}$, then for all $\varepsilon > 0$ there is a $T \in \mathbb{N}$ such that $d(A^*(\mu_t), A^*(M(f_t))) \leq \varepsilon$ for every $t > T$ and every $f_t \in \mathcal{F}$.*

PROOF. Since A^* is uniformly continuous, there exists $\varepsilon' \in \mathbb{R}_{++}$ such that for all $\nu \in \Delta(\Theta)$ and $\mu_t \in B_{\varepsilon'}(\nu)$, $d(A^*(\mu_t), A^*(\nu)) \leq \varepsilon$. Since $\Theta \subseteq \mathbb{R}^k$, the topology of weak convergence

on $\Delta(\Theta)$ is metrized by the Lévy–Prokhorov metric. By the definition of this metric, $\|\nu - \delta_p\|_{LP} \leq \varepsilon'$ whenever

$$\frac{\nu(B_{\varepsilon'/2}(p))}{1 - \nu(B_{\varepsilon'/2}(p))} \geq \frac{1 - \varepsilon'/2}{\varepsilon'/2}.$$

The statement follows from applying Proposition 2 and choosing

$$T \geq \frac{2 \log\left(\frac{1 - \varepsilon'/2}{\phi(\varepsilon'/8)\varepsilon'/2}\right)}{(\varepsilon')^2 m/8}. \quad \square$$

A.11 Example 6

EXAMPLE 6. (*Unlimited Bernoulli trials*) Suppose the outcome y corresponds to the number of Bernoulli trials needed to get one success. If the agent believes the trials are i.i.d. with parameter p_θ , their subjective distribution for outcome y is $p_\theta(y) = \theta(1 - \theta)^{y-1}$. Suppose that all success probabilities are considered possible, so that $\Pi = \{p_\theta : \theta \in [0, 1]\}$. Then

$$L(f||p_\theta) = \sum_{z=1}^{\infty} f(z) [(z-1) \log(1 - \theta) + \log(\theta)] = \log(1 - \theta) \bar{z} + \log(\theta) - \log(1 - \theta).$$

That L is uniformly strongly 1-concave immediately follows from taking derivatives:

$$\begin{aligned} \frac{\partial L(f||p_\theta)}{\partial \theta} &= -\frac{\bar{z}}{(1 - \theta)} + \frac{1}{\theta} + \frac{1}{1 - \theta}, \\ \frac{\partial L(f||p_\theta)^2}{\partial^2 \theta} &= -\frac{\bar{z}}{(1 - \theta)^2} - \frac{1}{\theta^2} + \frac{1}{(1 - \theta)^2}. \end{aligned}$$

The unique log-likelihood maximizing parameter is $\theta^*(f) = 1/\bar{z}$. Suppose that the prior belief μ is such that

$$\lim_{\theta \rightarrow 0} \frac{\mu\left(\left[0, \theta + \sqrt{\frac{2\theta}{1 - \theta}}\right]\right)}{\exp[L(\delta_{1/\theta}||p_{1/2}) - L(\delta_{1/\theta}||p_\theta)]} = 0. \quad (11)$$

An example that satisfies the restriction is given by the CDF

$$F(\theta) = \begin{cases} \left(\exp\left[-\log(1 - \theta)\frac{1}{\theta} - \log\left(\frac{\theta}{1 - \theta}\right) + \frac{1}{\theta} \log \frac{1}{2}\right]\right)^2 & \text{for } \theta \leq 1/10, \\ F\left(\frac{1}{10}\right) + \left(\frac{10\theta}{9} - \frac{1}{9}\right)\left(1 - F\left(\frac{1}{10}\right)\right) & \text{for } \theta > 1/10. \end{cases}$$

The restriction implies there cannot be A and g such that

$$\frac{\mu_1(M_{1/2}(\delta_c))}{1 - \mu_1(M_{1/2}(\delta_c))} \geq A \exp(g) \quad \forall c \in \mathbb{R}_{++}, \quad (12)$$

so that the conclusion of Theorem 1 does not hold for $t = 1$. To see why equation (12) cannot be satisfied, observe that for every $K > 0$, there is $c > 0$ such that

$$\begin{aligned} \frac{\mu_1(M_{1/2}(\delta_c))}{1 - \mu_1(M_{1/2}(\delta_c))} &\leq \frac{\mu_1\left(\left[0, \frac{1}{c} + \sqrt{\frac{2/c}{1-1/c}}\right]\right)}{\mu_1\left(\left[\frac{1}{4}, \frac{1}{2}\right]\right)} \leq \frac{\mu\left(\left[0, \frac{1}{c} + \sqrt{\frac{2/c}{1-1/c}}\right]\right)}{\mu\left(\left[\frac{1}{4}, \frac{1}{2}\right]\right)} \frac{\exp(L(\delta_c||p_{1/c}))}{\exp(L(\delta_c||p_{\frac{1}{2}}))} \\ &= \frac{\mu\left(\left[0, \frac{1}{c} + \sqrt{\frac{2/c}{1-1/c}}\right]\right)}{\mu\left(\left[\frac{1}{4}, \frac{1}{2}\right]\right)} \frac{\exp\left(c \log(1-1/c) - \log\left(\frac{1/c}{1-1/c}\right)\right)}{\exp\left(c \log \frac{1}{2}\right)} \leq K, \end{aligned}$$

where the last inequality follows because equation (11) implies the left-hand side is arbitrarily close to 0 for sufficiently high c . $\diamond$

REFERENCES

- Acemoglu, Daron, Victor Chernozhukov, and Muhamet Yildiz (2016), “Fragility of asymptotic agreement under Bayesian learning.” *Theoretical Economics*, 11, 187–225. [1587]
- Al-Najjar, Nabil and Eran Shmaya (2019), “Recursive utility and parameter uncertainty.” *Journal of Economic Theory*, 181, 274–288. [1587, 1589]
- Arrow, Kenneth J. and Jerry R. Green (1973), “Notes on expectations equilibria in Bayesian settings.” Working Paper No. 33, Stanford University. [1586]
- Bachelier, Louis (1900), “Théorie de la spéculation.” *Annales scientifiques de l’École normale supérieure.*, 17, 21–86. [1589]
- Berk, Robert H. (1966), “Limiting behavior of posterior distributions when the model is incorrect.” *The Annals of Mathematical Statistics*, 37, 51–58. [1585, 1586, 1594]
- Bohren, J. Aislinn and Daniel Hauser (2021), “Learning with model misspecification: Characterization and robustness.” *Econometrica*, 89, 3025–3077. [1586]
- Bowen, Renee, Danil Dmitriev, and Simone Galperti (2023), “Learning from shared news: When abundant information leads to belief polarization.” *Quarterly Journal of Economics*, 138, 955–1000. [1603]
- Clark, Daniel and Drew Fudenberg (2021), “Justified communication equilibrium.” *American Economic Review*, 111, 3004–3034. [1587]
- Clark, Daniel, Drew Fudenberg, and Kevin He (2022), “Observability, dominance, and induction in learning models.” *Journal of Economic Theory*, 206, 105569. [1587]
- Cogley, Timothy, Riccardo Colacito, Lars Peter Hansen, and Thomas J. Sargent (2008), “Robustness and US monetary policy experimentation.” *Journal of Money, Credit and Banking*, 40, 1599–1623. [1587, 1601]

- Cogley, Timothy, Riccardo Colacito, and Thomas J. Sargent (2007), “Benefits from US monetary policy experimentation in the days of Samuelson and Solow and Lucas.” *Journal of Money, Credit and Banking*, 39, 67–99. [1587, 1601]
- Cogley, Timothy and Thomas J. Sargent (2008), “Anticipated utility and rational expectations as approximations of Bayesian decision making.” *International Economic Review*, 49, 185–221. [1587, 1601]
- Diaconis, Persi and David Freedman (1990), “On the uniform consistency of Bayes estimates for multinomial probabilities.” *The Annals of Statistics*, 18, 1317–1327. [1585, 1586, 1587, 1588, 1589, 1590, 1591, 1595, 1596, 1603]
- Dupuis, Paul and Richard S. Ellis (2011), *A Weak Convergence Approach to the Theory of Large Deviations*, volume 902. John Wiley & Sons. [1610]
- Enke, Benjamin and Florian Zimmermann (2019), “Correlation neglect in belief formation.” *The Review of Economic Studies*, 86, 313–332. [1588]
- Esponda, Ignacio (2008), “Information feedback in first price auctions.” *The RAND Journal of Economics*, 39, 491–508. [1589]
- Esponda, Ignacio and Demian Pouzo (2016), “Berk–Nash equilibrium: A framework for modeling agents with misspecified models.” *Econometrica*, 84, 1093–1130. [1586]
- Esponda, Ignacio and Demian Pouzo (2021), “Equilibrium in misspecified Markov decision processes.” *Theoretical Economics*, 16, 717–757. [1586]
- Esponda, Ignacio, Demian Pouzo, and Yuichi Yamamoto (2021), “Asymptotic behavior of Bayesian learners with misspecified models.” *Journal of Economic Theory*, 195, 105260. [1586, 1602]
- Eusepi, Stefano and Bruce Preston (2018), “The science of monetary policy: An imperfect knowledge perspective.” *Journal of Economic Literature*, 56, 3–59. [1587]
- Fama, Eugene F. (1965), “The behavior of stock-market prices.” *The journal of Business*, 38, 34–105. [1589]
- Frankel, David M., Stephen Morris, and Ady Pauzner (2003), “Equilibrium selection in global games with strategic complementarities.” *Journal of Economic Theory*, 108, 1–44.
- Frick, Mira, Ryota Iijima, and Yuhta Ishii (2021), “A note on welfare comparisons for biased endogenous learning.” [1586]
- Frick, Mira, Ryota Iijima, and Yuhta Ishii (2022), “Welfare comparisons for biased learning.” [1586]
- Frick, Mira, Ryota Iijima, and Yuhta Ishii (2023), “Belief convergence under misspecified learning: A martingale approach.” *Review of the Economic Studies*, 90, 781–814. [1586]
- Fudenberg, Drew, Kevin He, and Lorens A. Imhof (2017), “Bayesian posteriors for arbitrarily rare events.” *Proceedings of the National Academy of Sciences*, 114, 4925–4929. [1601]

Fudenberg, Drew and Kevin He (2018), “Learning and type compatibility in signaling games.” *Econometrica*, 86, 1215–1255. [1587]

Fudenberg, Drew, Giacomo Lanzani, and Philipp Strack (2021a), “Limit points of endogenous misspecified learning.” *Econometrica*, 89, 1065–1098. [1586, 1588, 1601]

Fudenberg, Drew, Giacomo Lanzani, and Philipp Strack (2021b), “Selective memory equilibrium.” [1586]

Fudenberg, Drew, Giacomo Lanzani, and Philipp Strack (2022), “Pathwise concentration bounds for misspecified Bayesian beliefs.” Available at SSRN 3805083. [1588]

Fudenberg, Drew and Giacomo Lanzani (2023), “Which misperceptions persist?” *Theoretical Economics*, 18, 1271–1315. [1586]

Fudenberg, Drew and David K. Levine (1993), “Steady state learning and Nash equilibrium.” *Econometrica*, 547–573. [1587]

Fudenberg, Drew and David K. Levine (2006), “Superstition and rational learning.” *American Economic Review*, 96, 630–651. [1587]

Fudenberg, Drew, Gleb Romanyuk, and Philipp Strack (2017), “Active learning with a misspecified prior.” *Theoretical Economics*, 12, 1155–1189. [1586, 1602]

Gonçalves, Duarte (2020), “Sequential sampling and equilibrium.” <https://bit.ly/3oPuSFL>. [1587, 1588, 1589]

He, Kevin (2022), “Mislearning from censored data: The gambler’s fallacy in optimal-stopping problems.” *Theoretical Economics*, 17, 1269–1312. [1586]

He, Kevin and Jonathan Libgober (2021), “Evolutionarily stable (mis)specifications: Theory and application.” arXiv:2012.15007. [1586]

Heidhues, Paul, Botond Köszegi, and Philipp Strack (2018), “Unrealistic expectations and misguided learning.” *Econometrica*, 86, 1159–1214. [1603]

Heidhues, Paul, Botond Köszegi, and Philipp Strack (2021), “Convergence in models of misspecified learning.” *Theoretical Economics*, 16, 73–99. [1586]

Kleijn, Bas J. K. and Aad W. Van der Vaart (2012), “The Bernstein–von-Mises theorem under misspecification.” *Electronic Journal of Statistics*, 6, 354–381. [1587]

Kreps, David M. (1998), “Anticipated utility and dynamic choice.” *Econometric Society Monographs*, 29, 242–274. [1587, 1601]

Lanzani, Giacomo (2022), “Dynamic concern for misspecification.” [1602]

Levy, Gilat, Inés Moreno de Barreda, and Ronny Razin (2021), “Polarized extremes and the confused centre: Campaign targeting of voters with correlation neglect.” *Quarterly Journal of Political Science*, 16, 139–155. [1586]

Levy, Gilat, Ronny Razin, and Alwyn Young (2020), “Misspecified politics and the recurrence of populism.” [1602]

Mazumdar, Eric, Aldo Pacchiano, Yi-an Ma, Peter L. Bartlett, and Michael I. Jordan (2020), “On Thompson sampling with Langevin algorithms.” ArXiv preprint [arXiv:2002.10002](https://arxiv.org/abs/2002.10002). [1602]

Molavi, Pooya (2019), *Macroeconomics With Learning and Misspecification: A General Theory and Applications*. <https://pooyamolavi.com/CREE.pdf>. [1586]

Montiel Olea, José Luis, Pietro Ortoleva, Mallesh Pai, and Andrea Prat (2021), “Competing models.” ArXiv preprint [arXiv:1907.03809](https://arxiv.org/abs/1907.03809). [1603]

Nyarko, Yaw (1991), “Learning in mis-specified models and the possibility of cycles.” *Journal of Economic Theory*, 55, 416–427. [1586, 1602]

Preston, Bruce (2005), “Learning about monetary policy rules when long-horizon expectations matter.” *International Journal of Central Banking*, 1. [1587]

Sanov, Ivan Nicolaevich (1961), “On the probability of large deviations of random variables.” *Selected Translations in Mathematical Statistics and Probability*, 1, 213–244. [1610]

Schwartzstein, Joshua (2014), “Selective attention and learning.” *Journal of the European Economic Association*, 12, 1423–1452. [1587]

Schwartzstein, Joshua and Adi Sunderam (2021), “Using models to persuade.” *American Economic Review*, 111, 276–323. [1587]

Shen, Xiaotong and Larry Wasserman (2001), “Rates of convergence of posterior distributions.” *The Annals of Statistics*, 29, 687–714. [1587]

Spiegler, Ran (2020), “Behavioral implications of causal misperceptions.” *Annual Review of Economics*, 12, 81–106. [1588]

Co-editor Todd D. Sarver handled this manuscript.

Manuscript received 18 February, 2022; final version accepted 15 November, 2022; available online 6 December, 2022.