



On Bayesian index policies for sequential resource allocation

Emilie Kaufmann

► To cite this version:

Emilie Kaufmann. On Bayesian index policies for sequential resource allocation. *Annals of Statistics*, 2018, 46 (2), pp.842-865. 10.1214/17-AOS1569 . hal-01251606v3

HAL Id: hal-01251606

<https://hal.science/hal-01251606v3>

Submitted on 6 Nov 2017

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

On Bayesian index policies for sequential resource allocation

Emilie Kaufmann

CNRS & Univ. Lille, UMR 9189 (CRIStAL), Inria SequeL
emilie.kaufmann@univ-lille1.fr

Abstract

This paper is about index policies for minimizing (frequentist) regret in a stochastic multi-armed bandit model, inspired by a Bayesian view on the problem. Our main contribution is to prove that the Bayes-UCB algorithm, which relies on quantiles of posterior distributions, is asymptotically optimal when the reward distributions belong to a one-dimensional exponential family, for a large class of prior distributions. We also show that the Bayesian literature gives new insight on what kind of exploration rates could be used in frequentist, UCB-type algorithms. Indeed, approximations of the Bayesian optimal solution or the Finite Horizon Gittins indices provide a justification for the kl-UCB⁺ and kl-UCB-H⁺ algorithms, whose asymptotic optimality is also established.

1 Introduction

This paper presents new analyses of Bayesian flavored strategies for sequential resource allocation in an unknown, stochastic environment modeled as a multi-armed bandit. A *stochastic multi-armed bandit model* is a set of K probability distributions, $\mathcal{V}_1, \dots, \mathcal{V}_K$, called arms, with which an agent interacts in a sequential way. At round t , the agent, who does not know the arms' distributions, chooses an arm A_t . The draw of this arm produces an independent sample X_t from the associated probability distribution $\mathcal{V}_{A_t}$, often interpreted as a reward. Indeed, the arms can be viewed as those of different slot machines, also called *one-armed bandits*, generating rewards according to some underlying probability distribution.

In several applications that range from the motivating example of clinical trials [38] to the more modern motivation of online advertisement (e.g., [16]), the goal of the agent is to adjust his strategy $\mathcal{A} = (A_t)_{t \in \mathbb{N}}$, also called a *bandit algorithm*, in order to maximize the rewards accumulated during his interaction with the bandit model. The adopted strategy has to be sequential, in the sense that the next arm to play is chosen based on past observations: letting $\mathcal{F}_t = \sigma(A_1, X_1, \dots, A_t, X_t)$ be the σ -field generated by the observations up to round t , A_t is $\sigma(\mathcal{F}_{t-1}, U_t)$ -measurable, where U_t is a uniform random variable independent from $\mathcal{F}_{t-1}$ (as algorithms may be randomized).

More precisely, the goal is to design a sequential strategy maximizing the expectation of the sum of rewards up to some horizon T . If $\mu_1, \dots, \mu_K$ denote the means of the arms, and $\mu^* = \max_a \mu_a$, this is equivalent to minimizing the *regret*, defined as the expected difference between the reward accumulated by an oracle strategy always playing the best arm, and the reward accumulated by a strategy $\mathcal{A}$:

$$R(T, \mathcal{A}) := \mathbb{E} \left[T\mu^* - \sum_{t=1}^T X_t \right] = \mathbb{E} \left[\sum_{t=1}^T (\mu^* - \mu_{A_t}) \right]. \quad (1)$$

The expectation is taken with respect to the randomness in the sequence of successive rewards from each arm a , denoted by $(Y_{a,s})_{s \in \mathbb{N}}$, and the possible randomization of the algorithm, $(U_t)_t$. We denote by $N_a(t) = \sum_{s=1}^t \mathbb{1}_{(A_s=a)}$ the number of draws from arm a at the end of round t , so that $X_t = Y_{A_t, N_{A_t}(t)}$.

This paper focuses on good strategies in *parametric* bandit models, in which the distribution of arm a depends on some parameter θ_a : we write $\mathcal{V}_a = \nu_{\theta_a}$. Like in every parametric model, two different views

can be adopted. In the frequentist view, $\boldsymbol{\theta} = (\theta_1, \dots, \theta_K)$ is an unknown parameter. In the Bayesian view, $\boldsymbol{\theta}$ is a random variable, drawn from a prior distribution Π . More precisely, we define $\mathbb{P}_{\boldsymbol{\theta}}$ (resp. $\mathbb{E}_{\boldsymbol{\theta}}$) the probability (resp. expectation) under the probabilistic model in which for all a , $(Y_{a,s})_{s \in \mathbb{N}}$ is i.i.d. distributed under ν_{θ_a} and $\mathbb{P}^{\Pi}$ (resp. $\mathbb{E}^{\Pi}$) the probability (resp. expectation) under the probabilistic model in which for all a $(Y_{a,s})_{s \in \mathbb{N}}$ is i.i.d. conditionally to θ_a with conditional distribution ν_{θ_a} , and $\boldsymbol{\theta} \sim \Pi$. The expectation in (1) can thus be taken under either of these two probabilistic models. In the first case this leads to the notion of frequentist regret, which depends on $\boldsymbol{\theta}$:

$$R_{\boldsymbol{\theta}}(T, \mathcal{A}) := \mathbb{E}_{\boldsymbol{\theta}} \left[\sum_{t=1}^T (\mu^* - \mu_{A_t}) \right] = \sum_{a=1}^K (\mu^* - \mu_a) \mathbb{E}_{\boldsymbol{\theta}}[N_a(T)]. \quad (2)$$

In the second case, this leads to the notion of Bayesian regret, sometimes called *Bayes risk* in the literature (see [27]), which depends on the prior distribution Π :

$$\mathcal{R}_{\Pi}(T, \mathcal{A}) := \mathbb{E}^{\Pi} \left[\sum_{t=1}^T (\mu^* - \mu_{A_t}) \right] = \int R_{\boldsymbol{\theta}}(T, \mathcal{A}) d\Pi(\boldsymbol{\theta}). \quad (3)$$

The first bandit strategy was introduced by Thompson in 1933 [38] in a Bayesian framework, and a large part of the early work on bandit models is adopting the same perspective [10, 7, 19, 8]. Indeed, as Bayes risk minimization has an *exact*—yet often intractable—solution, finding ways to efficiently compute this solution has been an important line of research. Since 1985 and the seminal work of Lai and Robbins [28], there is also a precise characterization of good bandit algorithms in a frequentist sense. They show that for any *uniformly efficient policy* $\mathcal{A}$ (i.e. such that for all $\boldsymbol{\theta}$, $R_{\boldsymbol{\theta}}(T, \mathcal{A}) = o(T^{\alpha})$ for all $\alpha \in]0, 1]$), the number of draws of any sub-optimal arm a ($\mu_a < \mu^*$) is asymptotically lower bounded as follows:

$$\liminf_{T \rightarrow \infty} \frac{\mathbb{E}_{\boldsymbol{\theta}}[N_a(T)]}{\log T} \geq \frac{1}{\text{KL}(\nu_{\theta_a}, \nu_{\theta^*})}, \quad (4)$$

where $\text{KL}(\nu, \nu')$ denotes the Kullback-Leibler divergence between the distributions ν and ν' . From (2), this yields a lower bound on the regret.

This result holds for simple parametric bandit models, including exponential family bandit models presented in Section 2, that will be our main focus in this paper. It paved the way to a new line of research, aimed at building *asymptotically optimal* strategies, that is, strategies matching the lower bound (4) for some classes of distributions. Most of the algorithms proposed since then belong to the family of *index policies*, that compute at each round one index per arm, depending on the history of rewards observed from this arm only, and select the arm with largest index. More precisely, they are UCB-type algorithms, building confidence intervals for the means of the arms and choosing as an index for each arm the associated Upper Confidence Bound (UCB). The design of the confidence intervals has been successively improved [27, 1, 6, 5, 4, 21, 14] so as to obtain simple index policies for which non-asymptotic upper bound on the regret can be given. Among them, the kl-UCB algorithm [14] matches the lower bound (4) for exponential family bandit models. As they use confidence intervals on unknown parameters, all these index policies are based on *frequentist tools*. Nevertheless, it is interesting to note that the first index policy was introduced by Gittins in 1979 [19] to solve a Bayesian multi-armed bandit problem and is based on *Bayesian tools*, i.e. on exploiting the posterior distribution on the parameter of each arm.

However, tools and objectives can be separated: one can compute the Bayes risk of an algorithm based on frequentist tools, or the (frequentist) regret of an algorithm based on Bayesian tools. In this paper, we focus on the latter and advocate the use of index policies inspired by Bayesian tools for minimizing regret, in particular the Bayes-UCB algorithm [24], which is based on quantiles of the posterior distributions on the means. Our main contribution is to prove that this algorithm is asymptotically optimal, i.e. that it matches the lower bound (4), for any exponential bandit model and for a large class of prior distributions. Our analysis relies on two new ingredients: tight bounds on the tail of

posterior distributions (Lemma 4), and a self-normalized deviation inequality featuring an exploration rate that decreases with the number of observations (Lemma 5). This last tool also allows us to prove the asymptotic optimality of two variants of kl-UCB, called kl-UCB⁺ and kl-UCB-H⁺, that display improved empirical performance. Interestingly, the alternative exploration rate used by these two algorithms is already suggested by asymptotic approximations of the Bayesian exact solution or the Finite-Horizon Gittins indices.

The paper is structured as follows. Section 2 introduces the class of exponential family bandit models that we consider in the rest of the paper, and the associated frequentist and Bayesian tools. In Section 3, we present the Bayes-UCB algorithm, and give a proof of its asymptotic optimality. We introduce kl-UCB⁺ and kl-UCB-H⁺ in Section 4, in which we prove their asymptotic optimality and also exhibit connections with existing Bayesian policies. In Section 5, we illustrate numerically the good performance of our three asymptotically optimal, Bayesian-flavored index policies in terms of regret. We also investigate their ability to attain an optimal rate in terms of Bayes risk. Some proofs are provided in the supplemental paper [23].

Notation Recall that $N_a(t) = \sum_{s=1}^t \mathbb{1}_{(A_s=a)}$ is the number of draws from arm a at the end of round t . Letting $\hat{\mu}_{a,s} = \frac{1}{s} \sum_{k=1}^s Y_{a,k}$ be the empirical mean of the first s rewards from a , the empirical mean of arm a after t rounds of the bandit algorithm, $\hat{\mu}_a(t)$, satisfies $\hat{\mu}_a(t) = 0$ if $N_a(t) = 0$, $\hat{\mu}_a(t) = \hat{\mu}_{a,N_a(t)}$ otherwise.

2 (Bayesian) exponential family bandit models

In the rest of the paper, we consider the important class of *exponential family bandit models*, in which the arms belong to a one-parameter canonical exponential family.

2.1 Exponential family bandit model

A one-parameter canonical exponential family is a set $\mathcal{P}$ of probability distributions, indexed by a real parameter θ called the natural parameter, that is defined by

$$\mathcal{P} = \{\nu_\theta, \theta \in \Theta : \nu_\theta \text{ has a density } f_\theta(x) = \exp(\theta x - b(\theta)) \text{ w.r.t } \xi\},$$

where $\Theta = (\theta^-, \theta^+) \subseteq \mathbb{R}$ is an open interval, b a twice-differentiable and convex function (called the log-partition function) and ξ a reference measure. Examples of such distributions include Bernoulli distributions, Gaussian distributions with known variance, Poisson distributions, or Gamma distributions with known shape parameter.

If $X \sim \nu_\theta$, it can be shown that $\mathbb{E}[X] = \dot{b}(\theta)$ and $\text{Var}[X] = \ddot{b}(\theta) > 0$, where $\dot{b}$ (resp. $\ddot{b}$) is the derivative (resp. second derivative) of b with respect to the natural parameter θ . Thus there is a one-to-one mapping between the natural parameter θ and the mean $\mu = \dot{b}(\theta)$, and distributions in an exponential family can be alternatively parametrized by their mean. Letting $\mathbf{J} := \dot{b}(\Theta)$, for $\mu \in \mathbf{J}$ we denote by ν^μ the distribution in $\mathcal{P}$ that has mean μ : $\nu^\mu = \nu_{\dot{b}^{-1}(\mu)}$. The variance $V(\mu)$ of the distribution ν^μ is related to its mean in the following way:

$$V(\mu) = \ddot{b}(\dot{b}^{-1}(\mu)). \quad (5)$$

In the sequel, we fix an exponential family $\mathcal{P}$ and consider a bandit model $\nu^\mu = (\nu^{\mu_1}, \dots, \nu^{\mu_K})$, where ν^{μ_a} belongs to $\mathcal{P}$ and has mean μ_a . When considering Bayesian bandit models, we restrict our attention to product prior distributions on $\boldsymbol{\mu} = (\mu_1, \dots, \mu_K)$, such that μ_a is drawn from a prior distribution on $\mathbf{J} = \dot{b}(\Theta)$ that has density f_a with respect to the Lebesgue measure. We let π_a^t be the posterior distribution on μ_a after the first t rounds of the bandit game. With a slight abuse of notation, we will identify π_a^t with its density, for which a more precise expression is provided in Section 2.3.

2.2 Kullback-Leibler divergence and confidence intervals

For distributions that belong to a one-parameter exponential family, the large deviation rate function has a simple and explicit form, featuring the Kullback-Leibler (KL) divergence, and one can build tight confidence intervals on their means. The KL-divergence between two distributions ν_θ and ν_λ in an exponential family has a closed form expression as a function of the natural parameters θ and λ , given by

$$K(\theta, \lambda) := \text{KL}(\nu_\theta, \nu_\lambda) = \dot{b}(\theta)(\theta - \lambda) - b(\theta) + b(\lambda). \quad (6)$$

We also introduce $d(\mu, \mu')$ as the KL-divergence between the distributions of means μ and μ' :

$$d(\mu, \mu') := \text{KL}(\nu^\mu, \nu^{\mu'}) = K(\dot{b}^{-1}(\mu), \dot{b}^{-1}(\mu')).$$

Applying the Cramér-Chernoff method (see e.g. [9]) in an exponential family yields an explicit deviation inequality featuring this divergence function: if $\hat{\mu}_s$ is the empirical mean of s samples from ν^μ and $x > \mu$, one has $\mathbb{P}(\hat{\mu}_s > x) \leq \exp(-sd(x, \mu))$. This inequality can be used to build a confidence interval for μ based on a *fixed number of observations* s . Inside a bandit algorithm, computing a confidence interval on the mean of an arm a requires to take into account the *random number of observations* $N_a(t)$ available at round t . Using a self-normalized deviation inequality (see [14] and references therein), one can show that, at any round t of a bandit game, the kl-UCB index, defined as

$$u_a(t) := \sup \{q \in \mathcal{J} : N_a(t)d(\hat{\mu}_a(t), q) \leq \log(t \log^c(t))\}, \quad (7)$$

where $c \geq 3$ is a real parameter, satisfies $\mathbb{P}(u_a(t) > \mu_a) \gtrsim 1 - 1/(t \log^{c-2} t)$ and is thus an upper confidence bound on μ_a . The *exploration rate*, which is here $\log(t \log^c(t))$, controls the coverage probability of the interval.

Closed-form expressions for the divergence function d in the most common examples of exponential families are available (see [14]). Using the fact that $y \mapsto d(x, y)$ is increasing when $y > x$, an approximation of $u_a(t)$ can then be obtained using, for example, binary search.

2.3 Posterior distributions in Bayesian exponential family bandits

It is well-known that the posterior distribution on the mean of a distribution that belongs to an exponential family depends on two sufficient statistics: the number of observations and the empirical means of these observations. With f_a the density of the prior distribution on μ_a , introducing

$$\pi_{a,n,x}(u) := \frac{\exp\left(n \left[\dot{b}^{-1}(u)x - b(\dot{b}^{-1}(u)) \right]\right) f_a(u)}{\int_{\mathcal{J}} \exp\left(n \left[\dot{b}^{-1}(u)x - b(\dot{b}^{-1}(u)) \right]\right) f_a(u) du} \quad \text{for } u \in \mathcal{J},$$

the density of the posterior distribution on μ_a after t rounds of the bandit game can be written

$$\pi_a^t = \pi_{a, N_a(t), \hat{\mu}_a(t)}.$$

While our analysis holds for any choice of prior distribution, in practice one may want to exploit the existence of families of conjugate priors (e.g. Beta distributions for Bernoulli rewards, Gaussian distributions for Gaussian rewards, Gamma distributions for Poisson rewards). With a prior distribution chosen in such a family, the associated posterior distribution is well-known and its quantiles are easy to compute, which is of particular interest for the Bayes-UCB algorithm, described in the next section.

Finally, we give below a rewriting of the posterior distribution that will be very useful in the sequel to obtain tight bounds on its tails.

Lemma 1.

$$\pi_{a,n,x}(u) = \frac{\exp(-nd(x, u)) f_a(u)}{\int_{\mathcal{J}} \exp(-nd(x, u)) f_a(u) du}, \quad \text{for all } u \in \mathcal{J}.$$

Proof Let $u \in \mathcal{J}$. One has

$$\begin{aligned}
\pi_{a,n,x}(u) &= \frac{\exp\left(n\left[\dot{b}^{-1}(u)x - b(\dot{b}^{-1}(u))\right]\right) f_a(u)}{\int_{\mathcal{J}} \exp\left(n\left[\dot{b}^{-1}(u)x - b(\dot{b}^{-1}(u))\right]\right) f_a(u) du} \times \frac{e^{-n[x\dot{b}^{-1}(x) - b(\dot{b}^{-1}(x))]}}{e^{-n[x\dot{b}^{-1}(x) - b(\dot{b}^{-1}(x))]}} \\
&= \frac{\exp\left(-n\left[x(\dot{b}^{-1}(x) - \dot{b}^{-1}(u)) - b(\dot{b}^{-1}(x)) + b(\dot{b}^{-1}(u))\right]\right) f_a(u)}{\int_{\mathcal{J}} \exp\left(-n\left[x(\dot{b}^{-1}(x) - \dot{b}^{-1}(u)) - b(\dot{b}^{-1}(x)) + b(\dot{b}^{-1}(u))\right]\right) f_a(u) du} \\
&= \frac{\exp(-nd(x, u)) f_a(u)}{\int_{\mathcal{J}} \exp(-nd(x, u)) f_a(u) du},
\end{aligned}$$

using the closed form expression (6) and the fact that $\theta = \dot{b}^{-1}(\mu)$.

3 Bayes-UCB: a simple and optimal Bayesian index policy

3.1 Algorithm and main result

The Bayes-UCB algorithm is an index policy that was introduced by [24] in the context of parametric bandit models. Given a prior distribution on the parameters of the arms, the index used for each arm is a well-chosen quantile of the (marginal) posterior distributions of its mean. For exponential family bandit models, given a product prior distribution on the means, the Bayes-UCB index is

$$q_a(t) := Q\left(1 - \frac{1}{t(\log t)^c}; \pi_a^t\right) = Q\left(1 - \frac{1}{t(\log t)^c}; \pi_{a, N_a(t), \hat{\mu}_a(t)}\right),$$

where $Q(\alpha; \pi)$ is the quantile of order α of the distribution π (that is, $\mathbb{P}_{X \sim \pi}(X \leq Q(\alpha; \pi)) = \alpha$) and c is a real parameter. In the particular case of bandit models with Gaussian arms, [33] have introduced a variant of Bayes-UCB with a slightly different tuning of the confidence level, under the name UCL (for Upper Credible Limit).

While the efficiency of Bayes-UCB has been demonstrated even beyond bandit models with independent arms, regret bounds are available only in very limited cases. For Bernoulli bandit models asymptotic optimality is established by [24] when a uniform prior distribution on the mean of each arm is used. For Gaussian bandit models [33] give a logarithmic regret bound when an uninformative prior is used. In this section, we provide new finite-time regret bounds that hold in general exponential family bandit models, showing that a slight variant of Bayes-UCB is asymptotically optimal for a large class of prior distributions.

We fix an exponential family, characterized by its log-partition function b and the interval $\Theta =]\theta^-, \theta^+[$ of possible natural parameters. We let $\mu^- = \dot{b}(\theta^-)$ and $\mu^+ = \dot{b}(\theta^+)$ (μ^- may be equal to $-\infty$ and μ^+ to $+\infty$). We analyze Bayes-UCB for exponential bandit models satisfying the following assumption.

Assumption 2. *There exists $\mu_0^- > \mu^-$ and $\mu_0^+ < \mu^+$ such that $\forall a \in \{1, \dots, K\}$, $\mu_0^- \leq \mu_a \leq \mu_0^+$.*

For Poisson or Exponential distributions, this assumption requires that the means of all arms are different from zero, while they should be included in $]0, 1[$ for Bernoulli distributions. We now introduce a regularized version of the Bayes-UCB index that relies on the knowledge of μ_0^- and μ_0^+ , as

$$\bar{q}_a(t) := Q\left(1 - \frac{1}{t(\log t)^c}; \pi_{a, N_a(t), \bar{\mu}_a(t)}\right), \quad (8)$$

where $\bar{\mu}_a(t) = \min(\max(\hat{\mu}_a(t), \mu_0^-), \mu_0^+)$. Note that μ_0^- and μ_0^+ can be chosen arbitrarily close to μ^- and μ^+ respectively, in which case $\bar{q}_a(t)$ often coincides with the original Bayes-UCB index $q_a(t)$.

Theorem 3. Let ν^μ be an exponential bandit model satisfying Assumption 2. Assume that for all a , π_a^0 has a density f_a with respect to the Lebesgue measure such that $f_a(u) > 0$ for all $u \in \mathcal{J} = \dot{\mathcal{b}}(\Theta)$. Let $c \geq 7$. The algorithm that draws each arm once and for $t \geq K$ selects at time $t + 1$

$$A_{t+1} = \operatorname{argmax}_a \bar{q}_a(t),$$

with $\bar{q}_a(t)$ defined in (8) satisfies, for all $\varepsilon > 0$,

$$\forall a \neq a^*, \quad \mathbb{E}[N_a(T)] \leq \frac{1 + \varepsilon}{d(\mu_a, \mu^*)} \log(T) + o_\varepsilon(\log(T)).$$

From Theorem 3, taking the lim sup and letting ε go to zero show that (this slight variant of) Bayes-UCB satisfies

$$\forall a \neq a^*, \quad \limsup_{T \rightarrow \infty} \frac{\mathbb{E}[N_a(T)]}{\log(T)} \leq \frac{1}{d(\mu_a, \mu^*)}.$$

Thus this index policy is asymptotically optimal, as it matches Lai and Robbins' lower bound (4). As we shall see in Section 5, from a practical point of view Bayes-UCB outperforms kl-UCB and performs similarly (sometimes slightly better, sometimes slightly worse) as Thompson Sampling, another popular Bayesian algorithm that we now discuss.

3.2 Posterior quantiles versus posterior samples

Over the past few years, another Bayesian algorithm, Thompson Sampling, has become increasingly popular for its good empirical performance, and we explain how Bayes-UCB is related to this alternative, randomized, Bayesian approach.

The Thompson Sampling algorithm, that draws each arm according to its posterior probability of being optimal, was introduced in 1933 as the very first bandit algorithm [38] and re-discovered recently for its good empirical performance [36, 16]. Thompson Sampling can be implemented in virtually any Bayesian bandit model in which one can sample the posterior distribution, by drawing *one* sample from the posterior on each arm and selecting the arm that yields the largest sample. In any such case, Bayes-UCB can be implemented as well and may appear as a more robust alternative as the quantiles can be estimated based on *several* samples in case there is no efficient algorithm to compute them.

Our experiments of Section 5 show that Bayes-UCB as well as the other Bayesian-flavored index policies presented in Section 4 are competitive with Thompson Sampling in general one-dimensional exponential families. Compared to Bayes-UCB, the theoretical understanding of Thompson Sampling is more limited: this algorithm is known to be asymptotically optimal in exponential family bandit models, yet only for specific choices of prior distributions [25, 3, 26].

In more complex bandit models, there are situations in which Bayes-UCB is indeed used over Thompson Sampling. When there is a potentially infinite number of arms and the mean reward function is assumed to be drawn from a Gaussian Process, the GP-UCB of [37], that coincides with Bayes-UCB, is very popular in the Bayesian optimization community [11].

3.3 Tail bounds for posterior distributions

Just like the analysis of [24], the analysis of Bayes-UCB that we give in the next section relies on tight bounds on the tails of posterior distributions that permit to control quantiles. These bounds are expressed with the Kullback-Leibler divergence function d . Therefore, an additional tool in the proof is the control of the deviations of the empirical mean rewards from the true mean reward, measured with this divergence function, which follows from the work of [14].

In the particular case of Bernoulli bandit models, Bayes-UCB uses quantiles of Beta posterior distributions. In that case a specific argument, namely the fact that $\text{Beta}(a, b)$ is the distribution of the a -th order statistic among $a + b - 1$ uniform random variables, relates a Beta distribution (and its tails) to a

Binomial distribution (and its tails). This ‘Beta-Binomial trick’ is also used extensively in the analysis of Thompson Sampling for Bernoulli bandits proposed by [2, 25, 3]. Note that this argument can only be used for Beta distributions with integer parameters, which rules out many possible prior distributions. The analysis of [33] in the Gaussian case also relies on specific tails bounds for the Gaussian posterior distributions. For exponential family bandit models, an upper bound on the tail of the posterior distribution was obtained by [26] using the Jeffrey’s prior.

Lemma 4 below present more general results that hold for any class of exponential family bandit models and any prior distribution with a density that is positive on $\mathcal{J} = \dot{b}(\Theta)$. For such (proper) prior distributions, we give deterministic upper and lower bounds on the corresponding posterior probabilities $\pi_{a,n,x}([v, \mu^+])$. Compared to the result of [26], which is not presented in this deterministic way, Lemma 4 is based on a different rewriting of the posterior distribution, given in Lemma 1.

Lemma 4. *Let μ_0^-, μ_0^+ be defined in Assumption 2.*

1. *There exist two positive constants A and B such that for all x, v that satisfy $\mu_0^- < x < v < \mu_0^+$, for all $n \geq 1$, for all $a \in \{1, \dots, K\}$,*

$$An^{-1}e^{-nd(x,v)} \leq \pi_{a,n,x}([v, \mu^+]) \leq B\sqrt{n}e^{-nd(x,v)}.$$

2. *There exists a constant C such that for all x, v that satisfy $\mu_0^- < v \leq x < \mu_0^+$, for all $n \geq 1$, for all $a \in \{1, \dots, K\}$,*

$$\pi_{a,n,x}([v, \mu^+]) \geq \frac{C}{\sqrt{n}}.$$

The constants A, B, C depend on μ_0^-, μ_0^+ , b and the prior densities.

This result permits in particular to show that the quantile $\bar{q}_a(t)$ defined in (8) satisfies $\underline{u}_a(t) \leq \bar{q}_a(t) \leq \bar{u}_a(t)$, with

$$\begin{aligned} \underline{u}_a(t) &= \sup \left\{ q < \mu_0^+ : N_a(t)d(\bar{\mu}_a(t), q) \leq \log((At \log^c(t))/N_a(t)) \right\}, \\ \bar{u}_a(t) &= \sup \left\{ q < \mu_0^+ : N_a(t)d(\bar{\mu}_a(t), q) \leq \log\left(Bt \log^c(t) \sqrt{N_a(t)}\right) \right\}. \end{aligned}$$

Hence, despite their Bayesian nature, the indices used in Bayes-UCB are strongly related to frequentist kl-UCB type indices. However, compared to the index $u_a(t)$ defined in (7), the exploration rate that appears in $\underline{u}_a(t)$ and $\bar{u}_a(t)$ also features the current number of draws $N_a(t)$. Lai gives in [27] an asymptotic analysis of any index strategy of the above form with an exploration function $g(T/N_a(t))$, where $g(t) \sim \log(t)$ when t goes to infinity. Yet neither $\underline{u}_a(t)$ nor $\bar{u}_a(t)$ are not exactly of that form, and we propose below a finite-time analysis that relies on new, non-asymptotic, tools.

3.4 Finite-time analysis

We give here the proof of Theorem 3. To ease the notation, assume that arm 1 is an optimal arm, and let a be a suboptimal arm.

$$\mathbb{E}[N_a(T)] = \mathbb{E} \left[\sum_{t=0}^{T-1} \mathbb{1}_{(A_{t+1}=a)} \right] = 1 + \mathbb{E} \left[\sum_{t=K}^{T-1} \mathbb{1}_{(A_{t+1}=a)} \right].$$

We introduce a truncated version of the KL-divergence, $d^+(x, y) := d(x, y)\mathbb{1}_{(x < y)}$ and let g_t be a decreasing sequence to be specified later.

Using that, by definition of the algorithm, if a is played at round $t + 1$, it holds in particular that $\bar{q}_a(t) \geq \bar{q}_1(t)$, one has

$$\begin{aligned} (A_{t+1} = a) &\subseteq (\mu_1 - g_t \geq \bar{q}_1(t)) \bigcup (\mu_1 - g_t \leq \bar{q}_1(t), A_{t+1} = a) \\ &\subseteq (\mu_1 - g_t \geq \bar{q}_1(t)) \bigcup (\mu_1 - g_t \leq \bar{q}_a(t), A_{t+1} = a). \end{aligned}$$

This yields

$$\mathbb{E}[N_a(T)] \leq 1 + \sum_{t=K}^{T-1} \mathbb{P}(\mu_1 - g_t \geq \bar{q}_1(t)) + \sum_{t=K}^{T-1} \mathbb{P}(\mu_1 - g_t \leq \bar{q}_a(t), A_{t+1} = a).$$

The posterior bounds established in Lemma 4 permit to further upper bound the two sums in the right-hand side of the above inequality. With C defined in Lemma 4, we introduce t_0 , defined by

$$t \geq t_0 \Rightarrow (\mu_1 - g_t \geq \mu_0^- \text{ and } C^2 t \log(t)^{2c} > 1).$$

On the one hand, for $t \geq t_0$,

$$\begin{aligned} (\mu_1 - g_t \geq \bar{q}_1(t)) &= \left(\pi_{1, N_1(t), \bar{\mu}_1(t)}([\mu_1 - g_t, \mu^+]) \leq \frac{1}{t \log^c t} \right) \\ &= \left(\pi_{1, N_1(t), \bar{\mu}_1(t)}([\mu_1 - g_t, \mu^+]) \leq \frac{1}{t \log^c t}, \bar{\mu}_1(t) \leq \mu_1 - g_t \right), \end{aligned}$$

since by the lower bound in the second statement of Lemma 4,

$$\begin{aligned} &\left(\pi_{1, N_1(t), \bar{\mu}_1(t)}([\mu_1 - g_t, \mu^+]) \leq \frac{1}{t \log^c t}, \bar{\mu}_1(t) \geq \mu_1 - g_t \right) \\ &\subset \left(\frac{C}{\sqrt{N_1(t)}} \leq \frac{1}{t \log^c t} \right) \subset (N_1(t) \geq C^2 t^2 \log^{2c} t) \subset (N_1(t) > t) = \emptyset. \end{aligned}$$

Now using the lower bound in the first statement of Lemma 4,

$$\begin{aligned} (\mu_1 - g_t \geq \bar{q}_1(t)) &\subset \left(\frac{A e^{-N_1(t) d(\bar{\mu}_1(t), \mu_1 - g_t)}}{N_1(t)} \leq \frac{1}{t \log^c t}, \bar{\mu}_1(t) \leq \mu_1 - g_t \right) \\ &\subset \left(N_1(t) d^+(\hat{\mu}_1(t), \mu_1 - g_t) \geq \log \left(\frac{A t \log^c t}{N_1(t)} \right) \right). \end{aligned}$$

On the other hand,

$$\begin{aligned} &\sum_{t=K}^{T-1} \mathbb{P}(\mu_1 - g_t \leq \bar{q}_a(t), A_{t+1} = a) \\ &= \sum_{t=K}^{T-1} \mathbb{P} \left(\pi_{a, N_a(t), \bar{\mu}_a(t)}([\mu_1 - g_t, \mu^+]) \geq \frac{1}{t \log^c t}, A_{t+1} = a \right) \\ &\leq \sum_{t=K}^{T-1} \mathbb{P} \left(\bar{\mu}_a(t) < \mu_1 - g_t, \pi_{a, N_a(t), \bar{\mu}_a(t)}([\mu_1 - g_t, \mu^+]) \geq \frac{1}{t \log^c t}, A_{t+1} = a \right) \\ &\quad + \sum_{t=K}^{T-1} \mathbb{P}(\bar{\mu}_a(t) \geq \mu_1 - g_t, A_{t+1} = a). \end{aligned} \tag{9}$$

Using Lemma 4, the first sum in (9) is upper bounded by

$$\begin{aligned}
& \sum_{t=K}^{T-1} \mathbb{P} \left(B \sqrt{N_a(t)} e^{-N_a(t) d^+(\bar{\mu}_a(t), \mu_1 - g_t)} \geq \frac{1}{t \log^c t}, A_{t+1} = a \right) \\
& \leq \sum_{t=K}^{T-1} \sum_{s=1}^t \mathbb{P} \left(B \sqrt{s} e^{-s d^+(\bar{\mu}_{a,s}, \mu_1 - g_t)} \geq \frac{1}{t \log^c t}, N_a(t) = s, A_{t+1} = a \right) \\
& \leq \sum_{t=K}^{T-1} \sum_{s=1}^t \mathbb{P} \left(s d^+(\bar{\mu}_{a,s}, \mu_1 - g_s) \leq \log(T \log^c T) + \log(B) + \frac{1}{2} \log s, \right. \\
& \qquad \qquad \qquad \left. N_a(t) = s, A_{t+1} = a \right) \\
& \leq \sum_{s=1}^T \mathbb{P} \left(s d^+(\bar{\mu}_{a,s}, \mu_1 - g_s) \leq \log T + c \log \log T + \log(B) + \frac{1}{2} \log s \right) \\
& \leq \sum_{s=1}^T \mathbb{P} \left(s d^+(\hat{\mu}_{a,s}, \mu_1 - g_s) \leq \log T + c \log \log T + \log(B) + \frac{1}{2} \log s \right) \\
& \quad + \sum_{s=1}^T \mathbb{P}(\hat{\mu}_{a,s} < \mu_0^-).
\end{aligned}$$

To third inequality follows from exchanging the sums over s and t and using that $\sum_{t=1}^N \mathbb{1}_{(N_a(t)=s) \cap (A_{t+1}=a)}$ is smaller than 1 for all s . The last inequality uses that if $\hat{\mu}_{a,s} \geq \mu_0$, $\bar{\mu}_{a,s} \leq \hat{\mu}_{a,s}$ and $d^+(\bar{\mu}_{a,s}, \mu_1 - g_s) \geq d^+(\hat{\mu}_{a,s}, \mu_1 - g_s)$. Then by Chernoff inequality,

$$\sum_{s=1}^T \mathbb{P}(\hat{\mu}_{a,s} < \mu_0^-) \leq \sum_{s=1}^{\infty} \exp(-s d(\mu_0^-, \mu_a)) = \frac{1}{1 - e^{-d(\mu_0^-, \mu_a)}}.$$

Still using Chernoff inequality, the second sum in (9) is upper bounded by

$$\begin{aligned}
& \sum_{t=K}^{T-1} \mathbb{P}(\hat{\mu}_a(t) \geq \mu_1 - g_t, A_{t+1} = a) \leq \sum_{t=K}^{T-1} \mathbb{P}(\hat{\mu}_a(t) \geq \mu_1 - g_{N_a(t)}, A_{t+1} = a) \\
& \leq \sum_{t=K}^{T-1} \sum_{s=1}^t \mathbb{P}(\hat{\mu}_{a,s} \geq \mu_1 - g_s, N_a(t) = s, A_{t+1} = a) \\
& \leq \sum_{s=1}^T \mathbb{P}(\hat{\mu}_{a,s} \geq \mu_1 - g_s) \leq \sum_{s=1}^{\infty} \exp(-s d(\mu_1 - g_s, \mu_a)) := N_0 < +\infty.
\end{aligned}$$

Putting things together, we showed that there exists some constant $N = \max(t_0, N_0 + (1 - e^{-d(\mu_0^-, \mu_a)})^{-1}) + 1$ such that

$$\begin{aligned}
\mathbb{E}[N_a(T)] & \leq N + \underbrace{\sum_{t=K}^{T-1} \mathbb{P} \left(N_1(t) d^+(\hat{\mu}_1(t), \mu_1 - g_t) \geq \log \left(\frac{A t \log^c t}{N_1(t)} \right) \right)}_{T_1} \\
& \quad + \underbrace{\sum_{s=1}^T \mathbb{P} \left(s d^+(\hat{\mu}_{a,s}, \mu_1 - g_s) \leq \log T + c \log \log T + \log(B) + \frac{1}{2} \log s \right)}_{T_2}
\end{aligned}$$

Term T_1 is shown below to be of order $o(\log(T))$, as $\hat{\mu}_1(t)$ cannot be too far from $\mu_1 - g_t$. Note however that the deviation is expressed with $\log(t/N_1(t))$ in place of the traditional $\log(t)$, which makes the proof of Lemma 5 more intricate. In particular, Lemma 5 applies to a specific sequence (g_t) defined therein, and a similar result could not be obtained for the choice $g_t = 0$, unlike Lemma 6 below.

Lemma 5. *Let g_t be such that $d(\mu_1 - g_t, \mu_1) = \frac{1}{\log(t)}$. If $c \geq 7$, for all A , if t is larger than $\exp(\max(\sqrt{3}, A^{-1/7}))$,*

$$\begin{aligned} \mathbb{P} \left(N_1(t) d^+(\hat{\mu}_1(t), \mu_1 - g_t) \geq \log \frac{At \log^c t}{N_1(t)} \right) \\ \leq e \left(\frac{1}{At \log t} + \frac{3 \log \log t + \log A}{At \log^2 t} + \frac{1}{At \log^3 t} \right) + \frac{1}{t^2}. \end{aligned}$$

From Lemma 5, one has

$$\begin{aligned} (T_1) &\leq e \sum_{t=K}^{T-1} \frac{\log^2 t + 3(\log t) \log \log(t) + \log A \log t + 1}{At(\log^3 t)} + \sum_{t=K}^{T-1} \frac{1}{t^2} \\ &\leq \frac{e}{A} \left(2 + \frac{3}{e} + \frac{\log A}{\log K} \right) \sum_{t=K}^{T-1} \frac{1}{t \log(t)} + \frac{\pi^2}{6} \\ &\leq \frac{e}{A} \left(2 + \frac{3}{e} + \frac{\log A}{\log K} \right) \log \log T + \frac{\pi^2}{6}. \end{aligned}$$

The following lemma permits to give an upper bound on Term T2.

Lemma 6. *Let f, g, h be three functions such that*

$$f(s) \xrightarrow{s \rightarrow \infty} \infty, \quad g(s) \xrightarrow{s \rightarrow \infty} 0 \quad \text{and} \quad \frac{h(s)}{s} \xrightarrow{s \rightarrow \infty} 0,$$

with g and $s \mapsto h(s)/s$ non-increasing for s large enough.

For all $\varepsilon > 0$ there exists a (problem-dependent) constant $N_a(\varepsilon)$ such that for all $T \geq N_a(\varepsilon)$,

$$\begin{aligned} \sum_{s=1}^T \mathbb{P} (sd^+(\hat{\mu}_{a,s}, \mu_1 - g(s)) \leq f(T) + h(s)) \\ \leq \frac{1 + \varepsilon}{d(\mu_a, \mu_1)} f(T) + \sqrt{f(T)} \sqrt{\frac{8V_a^2 \pi (1 + \varepsilon)^3 d'(\mu_a, \mu_1)^2}{d(\mu_a, \mu_1)^3}} \\ + 8(1 + \varepsilon)^2 V_a^2 \left(\frac{d'(\mu_a, \mu_1)}{d(\mu_a, \mu_1)} \right)^2 \frac{1}{1 - e^{-d(\mu_0^-, \mu_a)}} + 1, \end{aligned}$$

with $V_a = \sup_{\mu \in [\mu_a, \mu_1]} V(\mu)$, where the variance function is defined in (5).

Let $\varepsilon > 0$. Using Lemma 6, with $f(s) = \log(s) + c \log \log(s) + \log(B)$, $g(s) = g_s$ defined in Lemma 5 and $h(s) = \frac{1}{2} \log(s)$, there exists problem dependent constants C_0 and $D_0(\varepsilon)$ such that

$$(T_2) \leq \frac{1 + \varepsilon}{d(\mu_a, \mu_1)} (\log T + c \log \log T) + C_0 \sqrt{\log T + c \log \log T} + D_0(\varepsilon).$$

Putting together the upper bounds on (T1) and (T2) yields the conclusion: for all $\varepsilon > 0$,

$$\mathbb{E}[N_a(T)] \leq \frac{1 + \varepsilon}{d(\mu_a, \mu^*)} \log(T) + O_\varepsilon(\sqrt{\log(T)}).$$

4 A Bayesian insight on alternative exploration rates

The kl-UCB index of an arm, $u_a(t)$, introduced in (7), uses the exploration rate $\log(t \log^c(t))$, that does not depend on arm a . Some alternatives to this universal exploration rate have been suggested in the literature, and we formally introduce two variants of kl-UCB, called kl-UCB⁺ and kl-UCB-H⁺ using an exploration rate that decreases with the number of draws of arm a . The tools developed for the analysis of Bayes-UCB allow us to prove the asymptotic optimality of both algorithms. We then show that the Bayesian literature on the multi-armed bandit problem provides a natural justification for these algorithms, that are related to approximations of the Bayesian optimal solution or the Gittins indices.

4.1 The kl-UCB⁺ and kl-UCB-H⁺ algorithms

We introduce in Definition 7 two new index policies, and prove their asymptotic optimality. The indices $u_a^{H,+}(t)$ and $u_a^{+}(t)$ both rely on an exploration rate that decreases with the number of plays of arm a . kl-UCB-H⁺ additionally requires the knowledge of the horizon T . In practice, both algorithms outperform kl-UCB, as can be seen in Section 5.

Definition 7. Let $c \geq 0$. We define kl-UCB-H⁺ and kl-UCB⁺ with parameter $c \geq 0$ as the index policies respectively based on the indices

$$u_a^{H,+}(t) = \sup \left\{ q : N_a(t) d(\hat{\mu}_a(t), q) \leq \log \left(\frac{T \log^c T}{N_a(t)} \right) \right\}, \quad (10)$$

$$u_a^{+}(t) = \sup \left\{ q : N_a(t) d(\hat{\mu}_a(t), q) \leq \log \left(\frac{t \log^c t}{N_a(t)} \right) \right\}. \quad (11)$$

A key step in the analysis of Bayes-UCB is the control of the probability of the event

$$\left(N_1(t) d^+(\hat{\mu}_1(t), \mu_1 - g_t) \geq \log \left(\frac{At \log^c t}{N_1(t)} \right) \right),$$

in which an exploration rate of order $\log(t/N_1(t))$ appears. This control is obtained in Lemma 5 which can also be used to analyze the kl-UCB-H⁺ and kl-UCB⁺ algorithms, that are based on such alternative exploration rates. The following theorem proves the asymptotic optimality of these two index policies. The proof is provided in Appendix B.

Theorem 8. Let $c \geq 7$. Each of the index policy associated to the indices defined by (11) and (10) satisfies, for all $\varepsilon > 0$,

$$\mathbb{E}[N_a(T)] \leq \frac{1 + \varepsilon}{d(\mu_a, \mu^*)} \log(T) + O_\varepsilon(\sqrt{\log(T)}).$$

The use of alternative exploration rates in UCB-type algorithms has appeared before in the bandit literature. For example the MOSS algorithm [4], based on the index

$$\hat{\mu}_a(t) + \sqrt{\frac{\log(T/(KN_a(t)))}{N_a(t)}},$$

is designed to be optimal in a minimax sense for bandit models with sub-gaussian rewards: the algorithm achieves a $O(\sqrt{KT})$ distribution-independent upper bound on the regret. Besides, it was already noted by [17] that the use of the exploration rate $\log(t/N_a(t))$ in place of $\log(t)$ in the kl-UCB algorithm leads to better empirical performance. In this paper, additionally to proving the asymptotic optimality of these approaches, we now provide a new insight on the use of such alternative exploration rates by relating the kl-UCB-H⁺ algorithm to other Bayesian policies.

4.2 Bayesian optimal solution and Gittins indices

The alternative exploration rate discussed in Section 4.1 happens to be related to two other Bayesian strategies for the multi-armed bandit problem: the Bayesian optimal solution and the Finite-Horizon Gittins index policy, that we present here.

In a Bayesian framework, the interaction of an agent with a multi-armed bandit can be modeled by a Markov Decision Process (MDP) in which the state Π_t is the current posterior distribution over the parameter of the arms. In exponential bandit models, the posterior over μ is $\Pi_t = \bigotimes \pi_a^t$. There are K possible actions and when action A_t is chosen in state Π_t , the observed reward X_t is a sample from arm A_t , that satisfies, conditionally to the past, $X_t \sim \nu^\mu$ and $\mu \sim \Pi_t(A_t)$. The new state is $\Pi^{t+1} = \bigotimes \pi_a^{t+1}$ with $\pi_a^{t+1} = \pi_a^t$ for all $a \neq A_t$ and the density of $\pi_{A_t}^{t+1}$ gets updated according to

$$\pi_{A_t}^{t+1}(u) \propto \exp(-(\dot{b}^{-1}(u)X_t - b(\dot{b}^{-1}(u))))\pi_{A_t}^t(u).$$

Bayes risk minimization, or reward maximization under the Bayesian probabilistic model, is equivalent to solving this MDP for the finite-horizon criterion, which boils down to finding a strategy of the form $A_t = g(\Pi_t)$ for some deterministic function g , that maximizes

$$\mathbb{E}^\Pi \left[\sum_{t=1}^T X_t^g \right], \quad (12)$$

where $(X_t^g)_t$ is the sequence of rewards obtained under policy g . From the theory of MDPs (see e.g., [32]), the optimal policy is solution of dynamic programming equations and can be computed by induction. However, due to the very large, if not infinite, state space (the set of possible posterior distributions over μ), the computation is often intractable.

In a slightly different setting, Gittins proved in 1979 [19] that the apparently intractable optimal policy reduces to an index policy, with corresponding indices later called the *Gittins indices*. He considers the discounted Bayesian multi-armed bandit problem, in which the goal is to find a policy g that minimizes

$$\mathbb{E}^\Pi \left[\sum_{t=1}^{\infty} \alpha^{t-1} X_t^g \right],$$

for some discount parameter $\alpha \in]0, 1[$. Interestingly, it was proved in [8] that the discount is necessary for this reduction to hold: in particular, the policy maximizing (12) is *not* an index policy. However, the notion of Gittins indices is a powerful concept that can also be defined in a finite horizon multi-armed bandit. The Finite-Horizon Gittins index of an arm depends on the current posterior distribution on its mean ($\pi = \pi_a^t$) and on the remaining time to play ($r = T - t$). It can be interpreted as the price worth paying for playing an arm with posterior π at most r times. Indeed, for $\lambda > 0$ consider the following game, called $\mathcal{C}_\lambda$, in which a player can either pay λ and draw the arm to receive a sample Y_t , which results in a reward $Y_t - \lambda$, or stop playing, which yields no reward. As precisely defined below, the Gittins index is the critical value of λ for which the optimal policy in $\mathcal{C}_\lambda$ is to stop playing the arm from the beginning. This definition transposes to the non-discounted case one of the equivalent definitions of the discounted Gittins index that can be found in [20].

Definition 9. *The Finite-Horizon Gittins index for a current posterior π and remaining time r is $G(\pi, r) = \inf\{\lambda \in \mathbb{R} : V_\lambda^*(\pi, r) = 0\}$, with*

$$V_\lambda^*(\pi, r) = \sup_{0 \leq \tau \leq r} \mathbb{E}_{Y_t \stackrel{i.i.d.}{\sim} \nu^\mu, \mu \sim \pi} \left[\sum_{t=1}^{\tau} (Y_t - \lambda) \right],$$

where the supremum is taken over all stopping time τ smaller than r a.s., with the convention $\sum_{t=1}^0 \cdot = 0$.

Computing the FH-Gittins indices requires to compute $V_\lambda^*(\pi, r)$ for several values of λ in order to find the critical value (using, e.g., binary search). Each computation requires solving a MDP, but on a smaller state space: the possible posterior distributions on the mean of a single arm. Hence the FH-Gittins algorithm, that is the index policy based on the Finite-Horizon Gittins indices,

$$A_{t+1} = \operatorname{argmax}_{a=1,\dots,K} G(\pi_a^t, T-t),$$

is a more practical algorithm than the Bayesian optimal solution. Although FH-Gittins does not coincide with the Bayesian optimal solution, we believe it is a good approximation. This is supported by simulations performed in a two-armed Bernoulli bandit problem, for which we compute the Bayes risk of the optimal strategy and that of the FH-Gittins algorithm up to horizon $T = 70$, as presented in Figure 1. For small horizons, [18] propose a comparison of different algorithms with the Bayesian optimal solution and similarly notice that the Bayes risk of FH-Gittins (called Λ -strategy) is very close to the optimal value, for various choices of prior and horizons.

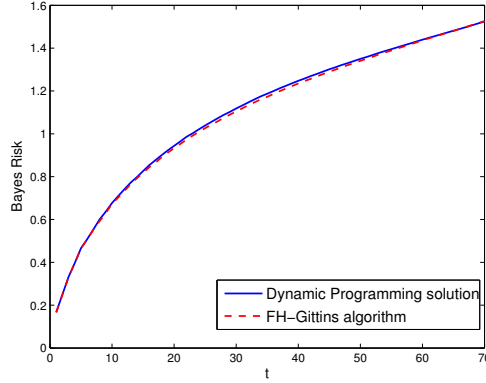


Figure 1: Bayes risk of the optimal strategy (blue) and FH-Gittins (dashed red) estimated using $N = 10^6$ replications of a bandit game, for which the means are drawn from $\mathcal{U}([0, 1])$

Compared to a simple index policy like Bayes-UCB, the computational cost of the FH-Gittins algorithm (not to mention that of the Bayesian optimal strategy) is still very high. In particular, the complexity of these two approaches grows dramatically when the horizon T increases, which motivates some approximations that have been proposed for large horizons, described in the next sections.

However, when the FH-Gittins algorithm is efficiently implementable (that is, for relatively small horizons), we would like to advocate its use for minimizing the frequentist regret. Indeed our experiments of Section 5 report good empirical performance in Bernoulli bandit models. In this particular case, using a uniform prior on the means, the set of (Beta) posterior is parametrized by two integers (the number of zeros and ones observed so far), and we could implement FH-Gittins up to horizon $T = 1000$. An efficient implementation of FH-Gittins for Gaussian bandits, up to horizon $T = 10000$, has been recently given by [29]. More generally, finding efficient methods to compute Finite-Horizon Gittins indices is still an area of investigation [31]. Interestingly, [29] provides the first theoretical elements supporting the use of FH-Gittins for regret minimization, by giving the first logarithmic upper bound on its regret in the particular case of Gaussian bandit models. However, the asymptotic optimality of this algorithm for Gaussian bandits and more general models remains a conjecture.

4.3 Approximation of the Bayesian optimal solution

In the paper [27], Lai shows that, in exponential family bandit models, the Bayes risk of *any* strategy is asymptotically lower bounded by $C_0(\pi) \log^2(T)$, when $C_0(\pi)$ is a prior-dependent constant. He also provides matching strategies, which implies in particular that the Bayes risk of the Bayesian optimal

solution is of order $\log^2(T)$. Any strategy matching this lower bound can be viewed as an asymptotic approximation of the Bayesian optimal solution.

In the particular case of product prior distributions, we provide in Theorem 10 a Bayes risk lower bound that is slightly more general than Lai's result in the sense that it does not require the prior distribution on the natural parameter of each arm to have a compact support. The proof of this result, provided in Appendix D, follows however closely that of [27]. The lower bound is expressed in terms of the prior distribution on the natural parameters $\theta = (\theta_1, \dots, \theta_K)$ of the arms, with the following notation. For $a = 1, \dots, K$, we let $\theta_{-a} = (\theta_1, \dots, \theta_{a-1}, \theta_{a+1}, \dots, \theta_K)$ be the vector of Θ^{K-1} that consists of all components of θ except component number a . We let $\theta_a^* = \max_{i \neq a} \theta_i$, so that θ_a^* only depends on θ_{-a} .

Theorem 10. *Let H be a prior distribution on Θ^K that has a product form, such that each marginal has a density h_a with respect to the Lebesgue measure λ that satisfies $h_a(\theta) > 0$ for all $\theta \in \Theta$. Letting H_{-a} be the marginal distribution of θ_{-a} , that has density $\prod_{i \neq a} h_i(\theta_i)$ with respect to $\lambda^{\otimes K-1}$, one assumes that*

$$\forall a = 1, \dots, K, \quad \int_{\Theta^{K-1}} h_a(\theta_a^*) dH_{-a}(\theta_{-a}) < \infty.$$

Under the prior distribution H , the Bayes risk of any strategy $\mathcal{A}$ satisfies

$$\liminf_{T \rightarrow \infty} \frac{\mathcal{R}^H(T, \mathcal{A})}{\log^2(T)} \geq \frac{1}{2} \sum_{a=1}^K \int_{\Theta^{K-1}} h_a(\theta_a^*) dH_{-a}(\theta_{-a}).$$

For exponential family bandit models with a product prior, Lai provides the first (asymptotic) prior-dependent Bayes risk upper bounds, when Θ is compact. Letting $[\mu_0^-, \mu_0^+] = \dot{b}(\Theta)$, he shows in particular that the index policy based on

$$I_a(t) = \sup \left\{ q \in [\mu_0^-, \mu_0^+] : N_a(t) \bar{d}(\hat{\mu}_a(t), q) \leq \log \left(\frac{T}{N_a(t)} \right) \right\}, \quad (13)$$

where $\bar{d}(x, y) = d(\max(\mu_0^-, \min(\mu_0^+, x)), y)$, has a Bayes risk that asymptotically matches the lower bound of Theorem 10. This index policy is very similar to kl-UCB- H^+ and differs only from the use of a regularized version of the divergence function d .

While a recent line of research on Bayesian randomized algorithms (e.g. Thompson Sampling) has provided Bayes risk upper bounds in quite general settings ([35, 34]), to the best of our knowledge, no upper bound scaling in $\log^2(T)$ has been obtained for exponential family bandit models since the work of Lai. [12, 30] give the first prior-dependent upper bounds on the Bayes risk of Thompson Sampling, in a particular case quite different from our setting: a two-armed bandit model in which the means of the arms are known up to a permutation. The joint prior distribution is thus supported on (μ_1, μ_2) and (μ_2, μ_1) . In Section 5.2, we investigate numerically the optimality of the Bayesian index policies discussed in the paper with respect to the lower bound of Theorem 10.

4.4 Approximation of the Finite-Horizon Gittins indices

As discussed Section 4.2, the FH-Gittins algorithm, that is the index policy associated to

$$J_a(t) = G(\pi_a^t, T - t),$$

is conjectured to be a good approximation of the Bayesian optimal policy, yet the above indices remain difficult to compute. Building on approximations of the Finite-Horizon Gittins indices that can be extracted from the literature permits to obtain a related *efficient* index policy.

Recall from Definition 9 that the Finite-Horizon Gittins index takes the form

$$G(\pi, r) = \inf \{ \lambda \in \mathbb{R} : V_\lambda^*(\pi, r) = 0 \},$$

where $V_\lambda^*(\pi, r)$ corresponds to the optimal value function associated to a calibration game $\mathcal{C}_\lambda$. In the paper [13], Burnetas and Katehakis propose tight bounds on the value function $V_\lambda^*(\pi_{a,n,x}, r)$ for exponential family bandits. These bounds permit to derive asymptotic approximations of the FH-Gittins indices, when r is large, and to show that, for large values of the remaining time $T - t$,

$$J_a(t) \simeq \sup \left\{ q \in [\mu^-, \mu^+] : N_a(t) \tilde{d}(\hat{\mu}_a(t), q) \leq \log \left(\frac{T-t}{N_a(t)} \right) \right\}. \quad (14)$$

This approximation is valid under the assumption that Θ is compact: $[\mu^-, \mu^+] = \dot{b}(\Theta)$ and $\tilde{d}$ is another regularization of the divergence function d , such that, for any y , $\tilde{d}(x, y) = d(x, y)$ for $x > \mu^-$ and for $x \leq \mu^-$,

$$\tilde{d}(x, y) = d(\mu^-, y) + (\dot{b}^{-1}(y) - \dot{b}^{-1}(\mu^-))(\mu^- - x).$$

In the particular case of Gaussian bandit models, the work of Chang and Lai [15] on the approximation of discounted Gittins indices can also be adapted to obtain approximations of the Finite-Horizon Gittins indices, showing the same tendency as in (14): compared to the corresponding kl-UCB index, here the $\log t$ is replaced by $\log((T-t)/N_a(t))$. This alternative exploration rate also appears in the non-asymptotic lower bound on the Gaussian Gittins index obtained by [29].

These approximations of the Finite-Horizon Gittins indices provide another justification for exploration rates of the form $\log(h(t, T)/N_a(t))$, with some function h , which are also used by the kl-UCB-H⁺ and kl-UCB⁺ algorithms. These two algorithms can thus be viewed as Bayesian (inspired) index policies.

5 Numerical experiments

5.1 Regret minimization

We first perform experiments with a moderate horizon $T = 1000$, which permits to include the Finite-Horizon Gittins algorithm discussed in Section 4.2. Figure 2 displays the regret of kl-UCB, Thompson Sampling and the four Bayesian (or Bayesian inspired) index policies discussed in this paper, in two instances of two-armed Bernoulli bandit problems. The Bayesian index policies display comparable, if not better, performance than kl-UCB and Thompson Sampling. In particular, FH-Gittins appears to be significantly better than the other algorithms on the instance with small rewards.

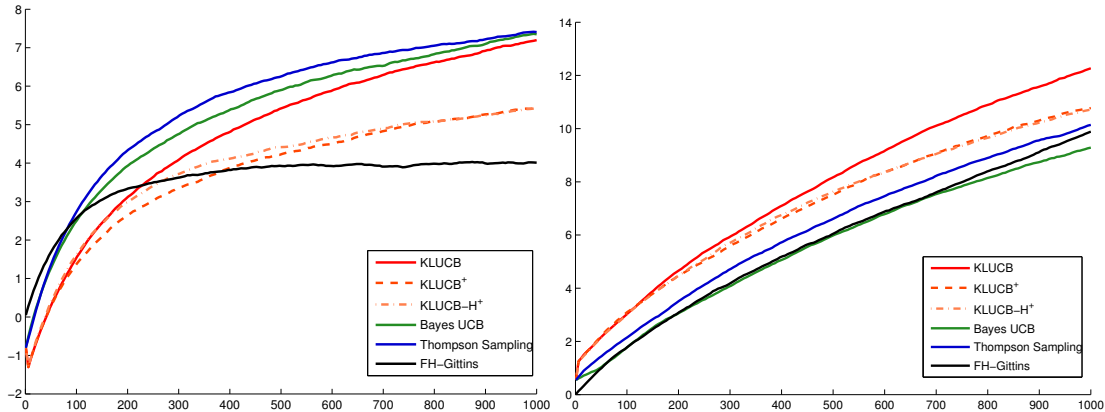


Figure 2: Regret on two-armed Bernoulli bandits ($\mu = [0.05 \ 0.15]$ (left) $\mu = [0.75 \ 0.8]$ (right)) up to horizon $T = 1000$, averaged over $N = 10000$ simulations

For a larger horizon $T = 20000$, we then run experiments on a bandit model in which rewards follow an exponential distribution (which is a particular Gamma distribution). Bayes-UCB and Thompson

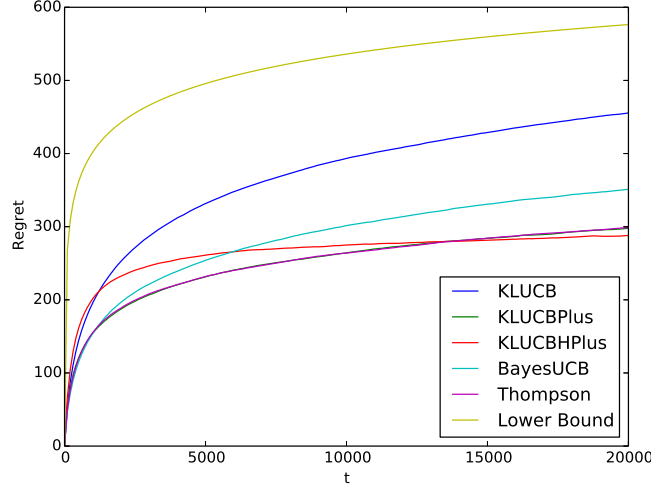


Figure 3: Regret on a five-armed bandit with Exponential distributions with means $\mu = [1 \ 1.5 \ 2 \ 2.5 \ 3]$ up to horizon $T = 20000$, averaged over $N = 50000$ simulations

Sampling are implemented using a conjugate $\text{InvGamma}(1, 1)$ prior on the means. Results are displayed in Figure 3. In this setting, Bayes-UCB, kl-UCB^+ and kl-UCB-H^+ improve over kl-UCB , and are also competitive with Thompson Sampling. As already noted in several works (e.g. [14]), the Lai and Robbins lower bound, that is asymptotic, is quite pessimistic for finite (even large) horizons.

5.2 Bayes risk minimization

In this paper, Bayes risk minimization and its exact solution is mostly presented as a justification for improved algorithms for regret minimization. However, it is also interesting to understand whether the proposed algorithm are good approximations of the Bayesian solution, i.e. whether they match the asymptotic lower bound of Theorem 10.

We report here results of experiments in Bernoulli bandit models with a uniform prior on the means. In this setting, some computations (that are detailed in Appendix D.4) show that the lower bound rewrites

$$\liminf_{T \rightarrow \infty} \frac{\mathcal{R}(T, \mathcal{A})}{\log^2(T)} \geq \frac{K-1}{K+1}.$$

In particular, we see that the asymptotic rate of the Bayesian regret is (almost) independent of the number of arms. For several values of K , we display on Figure 4 the Bayes risk $\mathcal{R}_T(\mathcal{A}_T)$ of several algorithms, together with the theoretical lower bound, as a function of $\log^2(T)$.

For each value of K , we observe that all the algorithms have a Bayes risk that seems to be affine in $\log^2(T)$. For Thompson Sampling, kl-UCB^+ and kl-UCB-H^+ the slope is close to $(K-1)/(K+1)$, whereas for kl-UCB and Bayes-UCB it is strictly larger. This leads to the conjecture that the first three algorithms are asymptotically optimal in a Bayesian sense. It is to be noted that, while the Bayes risk of these algorithms seems to be of order $(K-1)/(K+1)\log^2(T) + C(K)$ for large values of T , the second-order term $C(K)$ appears to be increasing significantly with the number of arms. Compared to Lai and Robbins' lower bound on the regret, this lower bound does not appear to be over-pessimistic in finite time.

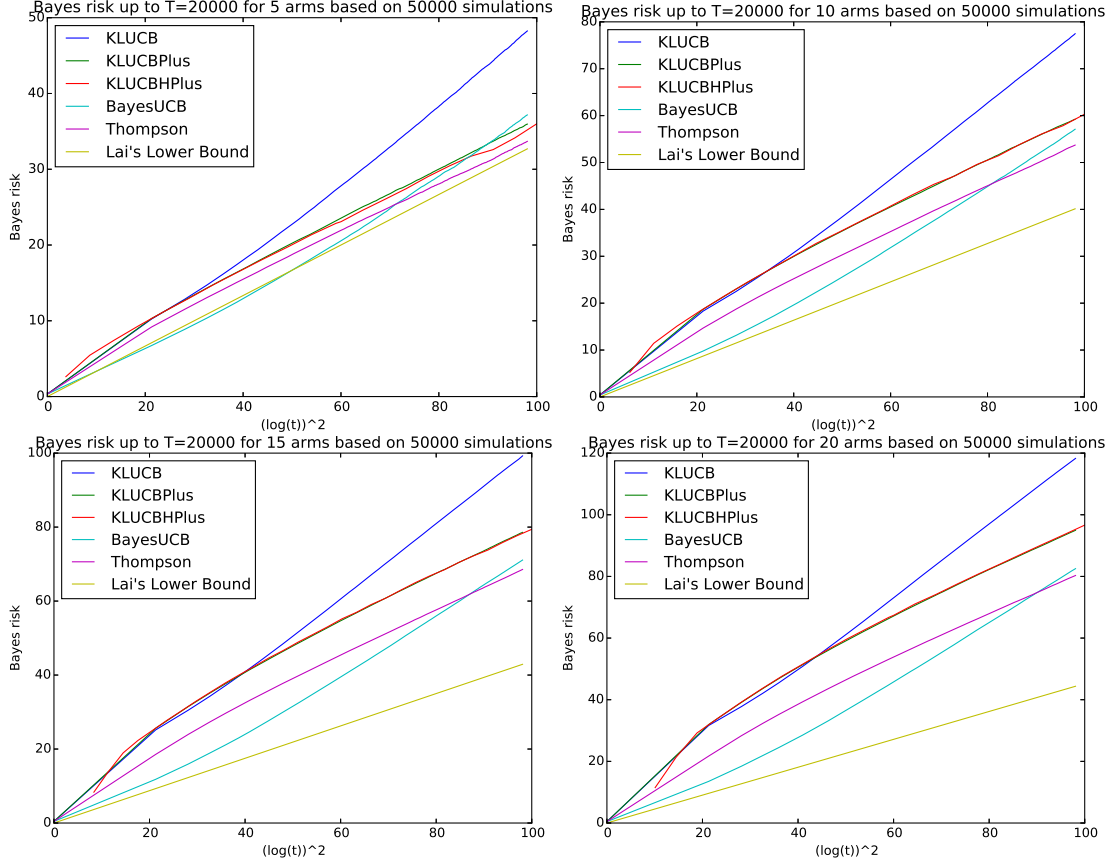


Figure 4: Bayes risk up to $T = 20000$ on a Bernoulli bandit model with a uniform prior on the K arms, for $K = 5, 10, 15, 20$, averaged over $N = 50000$ simulations.

6 Conclusion

This paper provides an analysis of the Bayes-UCB algorithm that does not rely on arguments specific to Bernoulli or Gaussian distributions, and is valid in any exponential family bandit model. It also brings theoretical justifications for the use of the kl-UCB-H^+ and kl-UCB^+ algorithms together with a new insight on the alternative exploration rate used by these algorithms. Finally, the proposed analysis holds for a wide class of prior distributions, namely all distributions that have positive density with respect to the Lebesgue measure. This shows that the choice of prior has no impact on the asymptotic optimality of Bayes-UCB, unlike what happens for Thompson Sampling in Gaussian bandit with unknown mean and variance [22]. Beyond asymptotic optimality, an interesting direction of future work would be to quantify the impact of the prior on second-order terms in the regret. Another important research direction is to better understand the Finite-Horizon Gittins strategy, which performs well in practice, but whose asymptotic optimality is still to be established.

Acknowledgement The author acknowledges the support of the French Agence Nationale de la Recherche (ANR) under grants ANR-13-BS01-0005 (SPADRO project) and ANR-11-JS02-005-01 (GAP project).

References

- [1] Agrawal, R. (1995). Sample mean based index policies with $O(\log n)$ regret for the multi-armed bandit problem. *Advances in Applied Probability*, 27(4):1054–1078.
- [2] Agrawal, S. and Goyal, N. (2012). Analysis of Thompson Sampling for the multi-armed bandit problem. In *Proceedings of the 25th Conference On Learning Theory*.
- [3] Agrawal, S. and Goyal, N. (2013). Further Optimal Regret Bounds for Thompson Sampling. In *Proceedings of the 16th Conference on Artificial Intelligence and Statistics*.
- [4] Audibert, J.-Y. and Bubeck, S. (2010). Regret Bounds and Minimax Policies under Partial Monitoring. *Journal of Machine Learning Research*.
- [5] Audibert, J.-Y., Munos, R., and Szepesvári, C. (2009). Exploration-exploitation trade-off using variance estimates in multi-armed bandits. *Theoretical Computer Science*, 410(19).
- [6] Auer, P., Cesa-Bianchi, N., and Fischer, P. (2002). Finite-time analysis of the multiarmed bandit problem. *Machine Learning*, 47(2):235–256.
- [7] Bellman, R. (1956). A problem in the sequential design of experiments. *The indian journal of statistics*, 16(3/4):221–229.
- [8] Berry, D. and Fristedt, B. (1985). *Bandit Problems. Sequential allocation of experiments*. Chapman and Hall.
- [9] Boucheron, S., Lugosi, G., and Massart, P. (2013). *Concentration inequalities. A non asymptotic theory of independence*. Oxford University Press.
- [10] Bradt, R., Johnson, S., and Karlin, S. (1956). On sequential designs for maximizing the sum of n observations. *Annals of Mathematical Statistics*, 27(4):1060–1074.
- [11] Brochu, E., Cora, V., and De Freitas, N. (2010). A Tutorial on Bayesian Optimization of Expensive Cost Functions, with Application to Active User Modeling and Hierarchical Reinforcement Learning. Technical report, University of British Columbia.
- [12] Bubeck, S. and Liu, C.-Y. (2013). Prior-free and prior-dependent regret bounds for Thompson Sampling. In *Advances in Neural Information Processing Systems*.
- [13] Burnetas, A. and Katehakis, M. (2003). Asymptotic Bayes Analysis for the finite horizon one armed bandit problem. *Probability in the Engineering and Informational Sciences*, 17:53–82.
- [14] Cappé, O., Garivier, A., Maillard, O.-A., Munos, R., and Stoltz, G. (2013). Kullback-Leibler upper confidence bounds for optimal sequential allocation. *Annals of Statistics*, 41(3):1516–1541.
- [15] Chang, F. and Lai, T. (1987). Optimal stopping and dynamic allocation. *Advances in Applied Probability*, 19:829–853.
- [16] Chapelle, O. and Li, L. (2011). An empirical evaluation of Thompson Sampling. In *Advances in Neural Information Processing Systems*.
- [17] Garivier, A. and Cappé, O. (2011). The KL-UCB algorithm for bounded stochastic bandits and beyond. In *Proceedings of the 24th Conference on Learning Theory*.
- [18] Ginebra, J. and Clayton, M. (1999). Small-sample performance of Bernoulli two-armed bandit Bayesian strategies. *Journal of Statistical Planning and Inference*, 79(1):107–122.

- [19] Gittins, J. (1979). Bandit processes and dynamic allocation indices. *Journal of the Royal Statistical Society, Series B*, 41(2):148–177.
- [20] Gittins, J., Glazebrook, K., and Weber, R. (2011). *Multi-armed bandit allocation indices (2nd Edition)*. Wiley.
- [21] Honda, J. and Takemura, A. (2010). An Asymptotically Optimal Bandit Algorithm for Bounded Support Models. In *Proceedings of the 23rd Conference on Learning Theory*.
- [22] Honda, J. and Takemura, A. (2014). Optimality of Thompson Sampling for Gaussian Bandits depends on priors. In *Proceedings of the 17th conference on Artificial Intelligence and Statistics*.
- [23] Kaufmann, E. (2016). Supplement to ”on bayesian index policies for sequential resource allocation”. *arXiv:1601.01190v2*.
- [24] Kaufmann, E., Cappé, O., and Garivier, A. (2012a). On Bayesian Upper-Confidence Bounds for Bandit Problems. In *Proceedings of the 15th conference on Artificial Intelligence and Statistics*.
- [25] Kaufmann, E., Korda, N., and Munos, R. (2012b). Thompson Sampling : an Asymptotically Optimal Finite-Time Analysis. In *Proceedings of the 23rd conference on Algorithmic Learning Theory*.
- [26] Korda, N., Kaufmann, E., and Munos, R. (2013). Thompson Sampling for 1-dimensional Exponential family bandits. In *Advances in Neural Information Processing Systems*.
- [27] Lai, T. (1987). Adaptive treatment allocation and the multi-armed bandit problem. *Annals of Statistics*, 15(3):1091–1114.
- [28] Lai, T. and Robbins, H. (1985). Asymptotically efficient adaptive allocation rules. *Advances in Applied Mathematics*, 6(1):4–22.
- [29] Lattimore, T. (2015). Regret analysis of the finite-horizon gittins index strategy for multi-armed bandits. *arXiv:1511.06014*.
- [30] Liu, C.-Y. and Li, L. (2015). On the prior sensitivity of thompson sampling. *arXiv:1506.03378*.
- [31] Nino-Mora, J. (2011). Computing a Classic Index for Finite-Horizon Bandits. *INFORMS Journal of Computing*, 23(2):254–267.
- [32] Puterman, M. (1994). *Markov Decision Processes. Discrete Stochastic. Dynamic Programming*. Wiley.
- [33] Reverdy, P., Srivastava, V., and Leonard, N. E. (2014). Modeling human decision making in generalized gaussian multiarmed bandits. volume 102, pages 544–571.
- [34] Russo, D. and Van Roy, B. (2014). Learning to optimize via information direct sampling. In *Advances in Neural Information Processing Systems (NIPS)*.
- [35] Russo, D. and Van Roy, B. (2014). Learning to optimize via posterior sampling. *Mathematics of Operations Research (to appear)*.
- [36] Scott, S. (2010). A modern Bayesian look at the multi-armed bandit. *Applied Stochastic Models in Business and Industry*, 26:639–658.
- [37] Srinivas, N., Krause, A., Kakade, S., and Seeger, M. (2010). Gaussian Process Optimization in the Bandit Setting : No Regret and Experimental Design. In *Proceedings of the International Conference on Machine Learning*.
- [38] Thompson, W. (1933). On the likelihood that one unknown probability exceeds another in view of the evidence of two samples. *Biometrika*, 25:285–294.

The appendix is structured as follows. Appendix A presents Pinsker-like inequalities, that is quadratic approximations of the Kullback-Leibler divergence functions, when the natural parameters of the distributions belong to some compact interval. These inequalities will be useful at several places of this supplementary material. Appendix B gathers the proof of Theorem 8 and the proofs of the lemmas introduced in the finite-time analysis of Bayes-UCB. Appendix C and Appendix D are respectively dedicated to establishing the posterior tail bounds of Lemma 4 and proving the asymptotic lower bound on the Bayes risk stated in Theorem 10.

A Pinsker-like inequalities

For on any compact $\mathcal{C} \subset \Theta$, one can obtain quadratic approximations of the KL-divergence as a function of either the natural parameters or the means. These useful inequalities are stated in Proposition 11

Proposition 11. *Let $\mathcal{C}$ be a compact subset of Θ . Introducing*

$$c_1 := \inf_{\theta \in \mathcal{C}} \ddot{b}(\theta) > 0 \quad \text{and} \quad c_2 := \sup_{\theta \in \mathcal{C}} \ddot{b}(\theta) < \infty, \quad (15)$$

one has

$$\forall (\theta, \theta') \in (\mathcal{C})^2, \quad \frac{c_1}{2}(\theta - \theta')^2 \leq K(\theta, \theta') \leq \frac{c_2}{2}(\theta - \theta')^2, \quad (16)$$

$$\forall (x, v) \in (\dot{b}(\mathcal{C}))^2, \quad \frac{1}{2c_2}(x - v)^2 \leq d(x, v) \leq \frac{1}{2c_1}(x - v)^2. \quad (17)$$

If $(x, v) \in (\dot{b}(\mathcal{C}))^2$ are such that $x < v$, one has

$$\dot{b}^{-1}(v) - \dot{b}^{-1}(x) \leq \frac{1}{c_1}(v - x). \quad (18)$$

Proof These three statements follow from Lagrange formulas. For example to derive (17), given that $d(x, y) = K(\dot{b}^{-1}(x), \dot{b}^{-1}(y))$, it can be shown, using the close form expression (6), that

$$\frac{d}{dx}d(x, v) = \dot{b}^{-1}(x) - \dot{b}^{-1}(v) \quad \text{and} \quad \frac{d^2}{dx^2}d(x, v) = \frac{1}{\ddot{b}(\dot{b}^{-1}(x))}.$$

From the second-order Lagrange formula applied to $x \mapsto d(x, v)$, there exists $c \in]x, v[$ (or $]v, x[$) such that

$$d(x, v) = \frac{1}{2} \frac{1}{\ddot{b}(\dot{b}^{-1}(c))} (x - v)^2 \leq \frac{1}{2c_1} (x - v)^2.$$

The other inequalities are obtained using similar arguments.

B Finite-time analysis

B.1 Proof of Theorem 8

We first give an analysis of the index policy associated to $u_a^+(t)$. Introducing g_t defined by $d(\mu_1 - g_t, \mu_1) = \frac{1}{\log(t)}$, one can write a decomposition similar to that used in the proof of Theorem 3:

$$\begin{aligned} \mathbb{E}[N_a(T)] &\leq 1 + \sum_{t=K}^{T-1} \mathbb{P}(\mu_1 - g_t \geq u_1^+(t)) \\ &\quad + \sum_{t=K}^{T-1} \mathbb{P}(\mu_1 - g_t \leq u_a^+(t), A_{t+1} = a) \\ &\leq 1 + \sum_{t=K}^{T-1} \mathbb{P}\left(N_1(t)d^+(\hat{\mu}_1(t), \mu_1 - g_t) \geq \log\left(\frac{t \log^c(t)}{N_a(t)}\right)\right) \end{aligned} \quad (19)$$

$$+ \sum_{t=K}^{T-1} \mathbb{P}(N_a(t)d^+(\hat{\mu}_a(t), \mu_1 - g_t) \leq \log(T \log^c(T)), A_{t+1} = a), \quad (20)$$

using the definition of $u_a^+(t)$ and the fact that $t \log^c t / N_a(t) \leq T \log^c T$. Lemma 5 can be applied (with $A = 1$) to show that the sum in (19) is of order $o(\log(T))$, while the sum in (20) can be rewritten and upper bounded using Lemma 6: for all $\varepsilon > 0$, the result follows from

$$\begin{aligned} &\mathbb{E} \sum_{t=K}^{T-1} \sum_{s=1}^t \mathbb{1}_{(sd^+(\hat{\mu}_{a,s}, \mu_1 - g_t) \leq \log(T \log^c T))} \mathbb{1}_{(A_{t+1}=a, N_a(t)=s)} \\ &\leq \sum_{s=1}^{T-1} \mathbb{P}(sd^+(\hat{\mu}_{a,s}, \mu_1 - g_t) \leq \log(T \log^c(T))) \\ &\leq \frac{1 + \varepsilon}{d(\mu_a, \mu_1)} \log(T \log^c(T)) + o_\varepsilon(\log(T)). \end{aligned}$$

For the index policy associated to $u_a^{H,+}(t)$, using a similar decomposition,

$$\begin{aligned} \mathbb{E}[N_a(T)] &\leq 1 + \sum_{t=K}^{T-1} \mathbb{P}(\mu_1 - g_t \geq u_1^{H,+}(t)) \\ &\quad + \sum_{t=K}^{T-1} \mathbb{P}(\mu_1 - g_t \leq u_a^{H,+}(t), A_{t+1} = a) \\ &\leq 1 + \sum_{t=K}^{T-1} \mathbb{P}\left(N_1(t)d^+(\hat{\mu}_1(t), \mu_1 - g_t) \geq \log\left(\frac{T \log^c(T)}{N_a(t)}\right)\right) \end{aligned} \quad (21)$$

$$+ \sum_{t=K}^{T-1} \mathbb{P}(N_a(t)d^+(\hat{\mu}_a(t), \mu_1 - g_t) \leq \log(T \log^c(T)), A_{t+1} = a). \quad (22)$$

The sums in (22) and (20) are the same, and lower bounding $T \log^c T$ by $t \log^c t$ in each term of the sum in (21) shows that it is upper bounded by (19). Thus, this index policy is also asymptotically optimal.

B.2 Proof of Lemma 5

To upper bound

$$(A) := \mathbb{P}\left(N_1(t)d^+(\hat{\mu}_1(t), \mu_1 - g_t) \geq \log \frac{At \log^c t}{N_1(t)}\right),$$

we consider two cases in which arm 1 has or not been drawn a lot.

$$(A) \leq \underbrace{\mathbb{P} \left(N_1(t) d^+ (\hat{\mu}_1(t), \mu_1 - g_t) \geq \log \left(\frac{At \log^c t}{N_1(t)} \right), N_1(t) \leq \log^4(t) \right)}_{A_1} \\ + \underbrace{\mathbb{P} \left(N_1(t) d^+ (\hat{\mu}_1(t), \mu_1 - g_t) \geq \log \left(\frac{At \log^c t}{N_1(t)} \right), N_1(t) > \log^4(t) \right)}_{A_2}$$

To upper bound term A_1 , we write

$$\begin{aligned} & \left(N_1(t) d^+ (\hat{\mu}_1(t), \mu_1 - g_t) \geq \log \left(\frac{At \log^c t}{N_1(t)} \right), N_1(t) \leq \log^4(t) \right) \\ & \subseteq \left(N_1(t) d^+ (\hat{\mu}_1(t), \mu_1) \geq \log(At) + c \log \log(t) - 4 \log \log(t), \right. \\ & \qquad \qquad \qquad \left. N_1(t) \leq \log^4(t) \right) \\ & \subseteq \left(N_1(t) d^+ (\hat{\mu}_1(t), \mu_1) \geq \log(At) + 3 \log \log t \right), \end{aligned}$$

using that $c \geq 7$. The self-normalized concentration inequality proved in [14] and stated in Lemma 12 permits to further upper bound A_1 :

$$(A_1) \leq e^{\frac{\log^2 t + 3(\log t) \log \log(t) + \log(A) \log t + 1}{At \log^3 t}}.$$

Lemma 12.

$$\mathbb{P} (\exists s \in \{1, \dots, t\} : s d^+ (\mu_{1,s}, \mu_1) \geq \delta) \leq (\delta \log(t) + 1) \exp(-\delta + 1).$$

To upper bound term A_2 , if t is such that $\log^7 t \geq A^{-1}$, we write

$$\begin{aligned} & \left(N_1(t) d^+ (\hat{\mu}_1(t), \mu_1 - g_t) \geq \log \left(\frac{At \log^c t}{N_1(t)} \right), N_1(t) \geq \log^4(t) \right) \\ & \subseteq \left(N_1(t) d^+ (\hat{\mu}_1(t), \mu_1 - g_t) \geq 0, N_1(t) \geq \log^4(t) \right) \\ & \subseteq \left(\hat{\mu}_1(t) \leq \mu_1 - g_t, N_1(t) \geq \log^4(t) \right), \end{aligned}$$

Thus, if t is such that $t \geq \exp(\sqrt{3})$ (which implies $\log^3 t \geq 3 \log t$),

$$\begin{aligned} (A_2) & \leq \mathbb{P} (\hat{\mu}_1(t) \leq \mu_1 - g_t, N_1(t) \geq \log^4(t)) \\ & \leq \mathbb{P} (\exists s \in [\log(t)^4; t] : \hat{\mu}_{1,s} \leq \mu_1 - g_t) \\ & \leq \sum_{s=\lceil \log(t)^4 \rceil}^t \mathbb{P} (\hat{\mu}_{1,s} \leq \mu_1 - g_t) \leq \sum_{s=\lceil \log(t)^4 \rceil}^t e^{-s d(\mu_1 - g_t, \mu_1)} \\ & \leq t e^{-(\log t)^4 d(\mu_1 - g_t, \mu_1)} = t e^{-(\log t)^3} \leq t e^{-3 \log t} = \frac{1}{t^2}. \end{aligned}$$

Combining this with the upper bound on A_1 yields the result.

B.3 Proof of Lemma 6

The quantity to be upper bounded is

$$(B) := \sum_{s=1}^T \mathbb{P} (s d^+ (\hat{\mu}_{a,s}, \mu_1 - g(s)) \leq f(T) + h(s)).$$

The function $w(q) = d^+(\hat{\mu}_{a,s}, q)$ is convex. Moreover it is differentiable on the interval $]\hat{\mu}_{a,s}, \mu^+[$ with $w'(q) = \frac{q - \hat{\mu}_{a,s}}{V(q)} \mathbb{1}_{(\hat{\mu}_{a,s} \leq q)}$. Thus, if $\hat{\mu}_{a,s} \geq \mu_0^-$,

$$d^+(\hat{\mu}_{a,s}, \mu_1 - g(s)) \geq d^+(\hat{\mu}_{a,s}, \mu_1) - g(s) \frac{\mu_1 - \hat{\mu}_{a,s}}{V(\mu_1)} \geq d^+(\hat{\mu}_{a,s}, \mu_1) - g(s) \frac{\mu_1 - \mu_0^-}{V(\mu_1)}.$$

Therefore

$$\begin{aligned} (B) &\leq \sum_{s=1}^T \mathbb{P} \left(d^+(\hat{\mu}_{a,s}, \mu_1) \leq \frac{f(T)}{s} + g(s) \frac{\mu_1 - \mu_0^-}{V(\mu_1)} + \frac{h(s)}{s} \right) + \sum_{s=1}^T \mathbb{P}(\hat{\mu}_{a,s} < \mu_0^-) \\ &\leq \sum_{s=1}^T \mathbb{P} \left(d^+(\hat{\mu}_{a,s}, \mu_1) \leq \frac{f(T)}{s} + r(s) \right) + \frac{1}{1 - e^{-d(\mu_0^-, \mu_a)}}, \end{aligned}$$

using Chernoff inequality and introducing

$$r(s) := g(s) \frac{\mu_1 - \mu_0^-}{V(\mu_1)} + \frac{h(s)}{s}.$$

Let $\varepsilon > 0$. One also introduce

$$K_T(\varepsilon) := \left\lceil \frac{(1 + \varepsilon)f(T)}{d(\mu_a, \mu_1)} \right\rceil.$$

From the assumptions on f, g and h , there exists s_0 such that r is non-increasing for $s \geq s_0$ and one has

$$K_T(\varepsilon) \xrightarrow{T \rightarrow \infty} \infty \quad \text{and} \quad r(s) \xrightarrow{s \rightarrow \infty} 0.$$

For T such that $K_T \geq s_0$,

$$(B) \leq K_T + \sum_{s=K_T+1}^T \mathbb{P} \left(d^+(\hat{\mu}_{a,s}, \mu_1) \leq \frac{f(T)}{s} + r(K_T) \right) + C_a,$$

with $C_a = 1 / (1 - e^{-d(\mu_0^-, \mu_a)})$. As $r(K_T) \rightarrow 0$, there exists $N_a(\varepsilon)$ such that

$$T \geq N_a(\varepsilon) \Rightarrow r(K_T) \leq d(\mu_a, \mu_1) \frac{\varepsilon}{1 + \varepsilon}.$$

Then, if $T \geq N_a(\varepsilon)$, one has, for all $s \geq K_T + 1$,

$$\frac{f(T)}{s} + r(K_T) \leq d(\mu_a, \mu_1)$$

and there exists $\mu^*(s) \in]\mu_a; \mu_1[$ such that $d(\mu^*(s), \mu_1) = \frac{f(T)}{s} + r(K_T)$. Then, using Chernoff inequality and the inequality

$$\forall \mu > \mu', \quad d(\mu, \mu') \geq \frac{1}{2 \sup_{\mu \in [\mu', \mu]} V(\mu)} (\mu - \mu')^2,$$

stated in [14] and that follows from Lagrange equality, one can write

$$\begin{aligned} (B) &\leq K_T + \sum_{s=K_T+1}^T \mathbb{P}(\hat{\mu}_{a,s} > \mu^*(s)) + C_a \leq K_T + \sum_{s=K_T+1}^T e^{-sd(\mu^*(s), \mu_a)} + C_a \\ &\leq K_T + \sum_{s=K_T+1}^T e^{-s \frac{(\mu^*(s) - \mu_a)^2}{2V_a^2}} + C_a \leq K_T + \int_{K_T}^{\infty} e^{-s \frac{(\mu^*(s) - \mu_a)^2}{2V_a^2}} ds + C_a, \end{aligned}$$

where $V_a = \sup_{\mu \in]\mu_a, \mu_1[} V(\mu)$. Using the convexity of $x \mapsto d(x, \mu_1)$, a lower bound on $\mu^*(s) - \mu_a$ can be obtained, as in Appendix 2 of [14]:

$$\mu^*(s) - \mu_a \geq \frac{d(\mu_a, \mu_1) - \left[\frac{f(T)}{s} + r(K_T) \right]}{-d'(\mu_a, \mu_1)}$$

[14] also provide tight upper bound on the resulting integrals, and following a similar approach allows us to conclude the proof:

$$\begin{aligned} & \int_{K_T}^{\infty} e^{-s \frac{(\mu^*(s) - \mu_a)^2}{2V_a^2}} ds \\ & \leq \int_{K_T}^{\infty} \exp \left(-\frac{s}{2V_a^2 d'(\mu_a, \mu_1)^2} \left(d(\mu_a, \mu_1) - \left(\frac{f(T)}{s} + r(K_T) \right) \right)^2 \right) ds \\ & \leq f(T) \int_{\frac{1+\varepsilon}{d(\mu_a, \mu_1)}}^{\infty} \exp \left(-\frac{uf(T)}{2V_a^2 d'(\mu_a, \mu_1)^2} \left(d(\mu_a, \mu_1) - \left(\frac{1}{u} + r(K_T) \right) \right)^2 \right) du \\ & \leq f(T) \int_{\frac{1+\varepsilon}{d(\mu_a, \mu_1)}}^{\frac{2(1+\varepsilon)}{d(\mu_a, \mu_1)}} \exp \left(-\frac{(1+\varepsilon) \left(d(\mu_a, \mu_1) - \left(\frac{1}{u} + r(K_T) \right) \right)^2}{2V_a^2 d(\mu_a, \mu_1) d'(\mu_a, \mu_1)^2} f(T) \right) du \\ & \quad + f(T) \int_{\frac{2(1+\varepsilon)}{d(\mu_a, \mu_1)}}^{\infty} \exp \left(-\frac{uf(T)}{2V_a^2 d'(\mu_a, \mu_1)^2} \frac{d(\mu_a, \mu_1)^2}{4(1+\varepsilon)^2} \right) du \\ & \leq f(T) \frac{4(1+\varepsilon)^2}{d(\mu_a, \mu_1)^2} \int_0^{\infty} \exp \left(-\frac{(1+\varepsilon)v^2 f(T)}{2V_a^2 d(\mu_a, \mu_1) d'(\mu_a, \mu_1)^2} \right) dv \\ & \quad + 8(1+\varepsilon)^2 V_a^2 \left(\frac{d'(\mu_a, \mu_1)}{d(\mu_a, \mu_1)} \right)^2 \\ & \leq \sqrt{f(T)} \sqrt{\frac{8V_a^2 \pi (1+\varepsilon)^3 d'(\mu_a, \mu_1)^2}{d(\mu_a, \mu_1)^3}} + 8(1+\varepsilon)^2 V_a^2 \left(\frac{d'(\mu_a, \mu_1)}{d(\mu_a, \mu_1)} \right)^2. \end{aligned}$$

C Posterior tail bounds

Let μ^-, μ^+ be such that $J := \dot{b}(\Theta) = (\mu^-, \mu^+)$. We give here the proof of Lemma 4, that follows directly from bounds on

$$\pi_{n,x}([v, \mu^+]) := \frac{\int_v^{\mu^+} \exp(-nd(x, u)) f_0(u) du}{\int_J \exp(-nd(x, u)) f_0(u) du}$$

for a density function f_0 satisfying $f_0(u) > 0$ for all $u \in J$. We fix μ_0^-, μ_0^+ : $\dot{b}(\theta^-) < \mu_0^- < \mu_0^+ < \dot{b}(\theta^+)$.

First, we fix a compact $\mathcal{C}$ included in Θ such that $[\mu_1^-, \mu_1^+] := \dot{b}(\mathcal{C})$ satisfy

$$\dot{b}(\theta^-) < \mu_1^- < \mu_0^- < \mu_0^+ < \mu_1^+ < \dot{b}(\theta^+).$$

We let $J_{\mathcal{C}} = [\mu_1^-, \mu_1^+]$ and c_1 and c_2 be the upper and lower bounds on $\ddot{b}$ on $\mathcal{C}$, defined as (15). We will often use the quadratic bounds on the Kullback-Leibler divergence on this compact, that are stated in Proposition 11. Also, we will use monotonicity properties of the divergence function: for all $y \in J$, $x \mapsto d(x, y)$ is decreasing on $]\mu^-, y]$ and increasing on $]y, \mu^+[$ whereas for all $x \in J$, $y \mapsto d(x, y)$ is decreasing on $]\mu^-, x]$ and increasing on $]x, \mu^+[$.

Let x, v such that $\mu_0^- < x < v < \mu_0^+$. One has

$$\pi_{n,x}([v, \mu^+]) = \frac{\int_v^{\mu^+} e^{-nd(x, u)} f_0(u) du}{\int_{\mu_1^-}^{\mu_1^+} e^{-nd(x, u)} f_0(u) du}. \quad (23)$$

For any $V_{n,x} \subset J$,

$$\pi_{n,x}([v, \mu^+]) \leq \frac{e^{-nd(x,v)} \int_v^{\mu^+} f_0(u) du}{\int_{V_{n,x}} e^{-nd(x,u)} f_0(u) du} \leq \frac{e^{-nd(x,v)}}{\int_{V_{n,x}} e^{-nd(x,u)} f_0(u) du}.$$

We now choose $V_{n,x} = \{u \in J_C : \frac{n(x-u)^2}{2c_1} \leq 1\}$. From (17), $nd(x, u) \leq 1$ on $V_{n,x}$. Hence

$$\int_{V_{n,x}} e^{-nd(x,u)} f_0(u) du \geq e^{-1} \inf_{u \in J_C} f_0(u) \int_{V_{n,x}} 1 du$$

and

$$\begin{aligned} \int_{V_{n,x}} 1 du &= \lambda \left([\mu_1^-, \mu_1^+] \cap \left[x - \sqrt{\frac{2c_1}{n}}, x + \sqrt{\frac{2c_1}{n}} \right] \right) \\ &\geq \min \left(\sqrt{\frac{2c_1}{n}}, \mu_1^+ - \mu_1^- \right). \end{aligned}$$

The following inequality yields the upper bound in statement 1:

$$\pi_{n,x}([v, \mu^+]) \leq \frac{e}{\sqrt{2c_1} \inf_{u \in J_C} f_0(u)} \max \left(\sqrt{n}, \frac{\sqrt{2c_1}}{(\mu_1^+ - \mu_1^-)} \right) e^{-nd(x,v)}.$$

As $e^{-nd(x,u)} \leq 1$, the denominator in (23) is upper bounded by 1, thus

$$\pi_{n,x}([v, \mu^+]) \geq \int_v^{\mu^+} e^{-nd(x,u)} f_0(u) du \geq \int_v^{\mu_1^+} e^{-nd(x,u)} f_0(u) du.$$

This last integral can be lower bounded in the following way:

$$\begin{aligned} \int_v^{\mu_1^+} e^{-nd(x,u)} f_0(u) du &= e^{-nd(x,v)} \int_v^{\mu_1^+} e^{-n[d(x,u)-d(x,v)]} f_0(u) du \\ &= e^{-nd(x,v)} \int_v^{\mu_1^+} e^{-n[d(v,u)+(b^{-1}(u)-b^{-1}(v))(v-x)]} f_0(u) du \\ &\geq e^{-nd(x,v)} \int_v^{\mu_1^+} e^{-n[\frac{1}{2c_1}(u-v)^2 + \frac{1}{c_1}(u-v)(v-x)]} f_0(u) du, \end{aligned} \quad (24)$$

where the last inequality follows from (17) and (18). We let

$$\phi(u) = \frac{1}{2c_1} [(u-v)^2 + 2(u-v)(v-x)].$$

One has $\phi'(u) = (u-x)/c_1$, thus ϕ is strictly increasing on $[v, \mu_1^+]$ and it can be checked that $\phi^{-1}(y) = x + \sqrt{(v-x)^2 + 2c_1 y}$. Thus letting $y = n\phi(u)$, one has

$$du = \frac{c_1}{n\sqrt{(v-x)^2 + 2c_1 y/n}} dy,$$

and

$$\begin{aligned} (24) &= \frac{c_1 e^{-nd(x,v)}}{n} \int_0^{\frac{n}{2c_1}[(\mu_1^+ - v)^2 + 2(\mu_1^+ - v)(v-x)]} \frac{e^{-y}}{\sqrt{(v-x)^2 + 2c_1 y/n}} f_0(y) dy \\ &\geq \frac{c_1 e^{-nd(x,v)} \min_{J_C} f_0}{n} \int_0^{\frac{n}{2c_1}[(\mu_1^+ - v)^2 + 2(\mu_1^+ - v)(v-x)]} \frac{e^{-y} dy}{\sqrt{(v-x)^2 + (\mu_1^+ - v)^2 + 2(\mu_1^+ - v)(v-x)}} \\ &\geq \frac{c_1 e^{-nd(x,v)} \min_{J_C} f_0}{n(\mu_1^+ - x)} \left(1 - e^{-\frac{n}{2c_1}(\mu_1^+ - v)^2} \right). \end{aligned}$$

Finally, using that $\mu_0^- < x$ and $v \leq \mu_0^+$, one obtains

$$\pi_{n,x}([v, \mu^+]) \geq \left(\frac{c_1(1 - e^{-\frac{(\mu_1^+ - \mu_0^+)^2}{2c_1}}) \min_{J_C} f_0}{(\mu_1^+ - \mu_0^-)} \right) \frac{1}{n} e^{-nd(x,v)},$$

which yields the lower bound in statement 1.

We now prove statement 2. Let x, v such that $\mu_0^- < v \leq x < \mu_0^+$. As $[x, \mu_1^+] \subset [v, \mu^+]$, one has

$$\pi_{n,x}([v, \mu^+]) \geq \int_x^{\mu_1^+} e^{-nd(x,u)} f_0(u) du$$

$F_x : u \mapsto \sqrt{d(x,u)}$ is a one-to-one mapping between $[x, \mu_1^+]$ and $[0, \sqrt{d(x, \mu_1^+)})$. Moreover, letting $d'(x, u) = \frac{d}{du} d(x, u) = \frac{u-x}{b(b^{-1}(u))} = \frac{u-x}{V(u)}$,

$$F'_x(u) = \frac{d'(x, u)}{2\sqrt{d(x, u)}} \underset{u \rightarrow x}{\sim} \frac{(u-x)/V(u)}{2\sqrt{\frac{1}{2}(x-u)^2/V(x)}} \underset{u \rightarrow x}{\rightarrow} \sqrt{\frac{V(x)}{2}}.$$

F'_x is continuous on $[x, \mu_1^+]$ and strictly positive, thus the inverse mapping $\phi_x : [0, \sqrt{d(x, \mu_1^+)}) \rightarrow [x, \mu_1^+]$ is well defined and differentiable. Letting $u = \phi_x(y)$, one has

$$\begin{aligned} \int_x^{\mu_1^+} e^{-nd(x,u)} f_0(u) du &= \int_0^{\sqrt{d(x, \mu_1^+)}} e^{-ny^2} f_0(\phi_x(y)) \frac{2\sqrt{d(x, \phi_x(y))}}{d'(x, \phi_x(y))} dy \\ &\geq \inf_{u \in J_C} f_0(u) \frac{2}{\sqrt{n}} \int_0^{\sqrt{nd(x, \mu_1^+)}} e^{-y^2} \frac{\sqrt{d(x, \phi_x(\frac{y}{\sqrt{n}}))}}{d'(x, \phi_x(\frac{y}{\sqrt{n}}))} dy. \end{aligned}$$

The mapping $(x, u) \mapsto \sqrt{d(x, u)}/d'(x, u)$ is continuous and strictly positive on the compact set $\mathcal{S} = \{(x, u) \in [\mu_0^-, \mu_0^+] \times [\mu_1^-, \mu_1^+] : x \leq u \leq \mu_1^+\}$ therefore, one can define

$$c = \inf_{(x, u) \in \mathcal{S}} \frac{\sqrt{d(x, u)}}{d'(x, u)} > 0.$$

For $n \geq 1$, one has

$$\int_x^{\mu_1^+} e^{-nd(x,u)} f_0(u) du \geq \frac{1}{\sqrt{n}} \left(2c \inf_{u \in J_C} f_0(u) \int_0^{\sqrt{d(\mu_0^+, \mu_1^+)}} e^{-y^2} dy \right).$$

Thus there exists a constant $C = C(\mu_1^-, \mu_1^+, f_0) > 0$ such that

$$\pi_{n,x}([v, \mu^+]) \geq \frac{C}{\sqrt{n}},$$

which concludes the proof.

D Lower bound on the Bayesian regret

D.1 Proof of Theorem 10

Let $\mathcal{A}$ be a bandit algorithm. Introducing

$$C_{\text{opt}} = \frac{1}{2} \sum_{a=1}^K \int_{\Theta^{K-1}} h_a(\theta_a^*) dH_{-a}(\boldsymbol{\theta}_{-a}),$$

we assume that $\mathcal{A}$ satisfies the following: there exists constants $C > C_{\text{opt}}$ and $T_0 > 0$ such that

$$\forall T \geq T_0, \mathcal{R}^H(T, \mathcal{A}) \leq C(\log T)^2. \quad (25)$$

Note that if $\mathcal{A}$ does not satisfy the above assumption, the desired conclusion follows directly:

$$\liminf_{T \rightarrow \infty} \frac{\mathcal{R}^H(T, \mathcal{A})}{\log^2(T)} \geq C_{\text{opt}}.$$

In the sequel denote by θ^- and θ^+ the lower and upper bounds of the interval $\Theta : \Theta =]\theta^-, \theta^+[$. The Bayes risk of $\mathcal{A}$ rewrites

$$\begin{aligned} \mathcal{R}^H(T, \mathcal{A}) &= \mathbb{E}[R_{\boldsymbol{\theta}}(T, \mathcal{A})] = \mathbb{E} \left[\sum_{a=1}^K (\dot{b}(\theta_a^*) - \dot{b}(\theta_a)) \mathbb{E}_{\boldsymbol{\theta}}[N_a(T)] \right] \\ &= \sum_{a=1}^K \int_{\{\boldsymbol{\theta} \in \Theta^K : \theta_a < \theta_a^*\}} (\dot{b}(\theta_a^*) - \dot{b}(\theta_a)) \mathbb{E}_{\boldsymbol{\theta}}[N_a(T)] dH(\boldsymbol{\theta}) \end{aligned}$$

Letting $\mathcal{T}_a$ be the a -th term in this last sum, one has

$$\begin{aligned} \mathcal{T}_a &= \int_{\Theta^{K-1}} \int_{\{\theta_a \in \Theta : \theta_a < \theta_a^*\}} (\dot{b}(\theta_a^*) - \dot{b}(\theta_a)) \mathbb{E}_{\boldsymbol{\theta}}[N_a(T)] h_a(\theta_a) d\theta_a dH_{-a}(\boldsymbol{\theta}_{-a}) \\ &= \int_{\Theta^{K-1}} \int_0^{\theta_a^* - \theta^-} (\dot{b}(\theta_a^*) - \dot{b}(\theta_a^* - t)) \mathbb{E}_{\boldsymbol{\theta}_{a,t}}[N_a(T)] h_a(\theta_a^* - u) du dH_{-a}(\boldsymbol{\theta}_{-a}), \end{aligned}$$

where $\boldsymbol{\theta}_{a,u} := (\theta_1, \dots, \theta_{a-1}, \theta_a^* - u, \theta_{a+1}, \dots, \theta_K)$.

Let $\gamma \in]0, 1[$ and let $B = [b^-, b^+]$ be a compact subset of Θ . For T large enough, such that

$$1/(b^- - \theta^-) < \log T < T^{\frac{1-\gamma}{2}},$$

reducing the integration domain by first letting $\boldsymbol{\theta}_{-a} \in B^{K-1}$ and then $u \in [T^{-(1-\gamma)/2}, (\log T)^{-1}]$, one has

$$\begin{aligned} \mathcal{T}_a &\geq \int_{B^{K-1}} \int_{T^{-(1-\gamma)/2}}^{(\log T)^{-1}} (\dot{b}(\theta_a^*) - \dot{b}(\theta_a^* - u)) \mathbb{E}_{\boldsymbol{\theta}_{a,u}}[N_a(T)] h_a(\theta_a^* - u) du dH_{-a}(\boldsymbol{\theta}_{-a}) \\ &\geq (1 - \gamma) \int_{B^{K-1}} \int_{T^{-(1-\gamma)/2}}^{(\log T)^{-1}} \frac{2h_a(\theta_a^*) K(\theta_a^* - u, \theta_a^* + \zeta u) \mathbb{E}_{\boldsymbol{\theta}_{a,t}}[N_a(T)]}{u} du dH_{-a}(\boldsymbol{\theta}_{-a}). \end{aligned}$$

The last inequality follows from the technical lemma stated below, in which the constant ζ is defined.

Lemma 13. *Let $\gamma > 0$. There exists $\zeta \in]0, 1[$ and $u_0 > 0$ such that for all $\boldsymbol{\theta}_{-a} \in B^{K-1}$ and $0 \leq u \leq u_0$,*

$$\forall \theta \in B, \quad \frac{(\dot{b}(\theta) - \dot{b}(\theta - u)) h_a(\theta - u)}{K(\theta - u, \theta + \zeta u)} \geq (1 - \gamma) \frac{2h_a(\theta)}{u}.$$

Now we need to give a lower bound on $\mathbb{E}_{\boldsymbol{\theta}_{a,u}}[N_a(T)]$, that will subsequently be integrated over $B^{K-1} \times [T^{-(1-\gamma)/2}, (\log T)^{-1}]$. Lai and Robbins provide such a lower bound in [28], but under the assumption (not satisfied here) that, for all $\alpha \in]0, 1[$, $\mathcal{A}$ has a $o(T^\alpha)$ regret on every bandit model. Moreover their lower bound is asymptotic, which makes it more complicated to integrate. Lemma 14 below provides a non-asymptotic lower bound on $\mathbb{E}_{\boldsymbol{\theta}_{a,u}}[N_a(T)]$, that also follows from a change of distribution argument.

Lemma 14. *Let $\zeta \in]0, 1[$ and $B = [b^-, b^+] \subset \Theta$. Introducing*

$$e_{T,u}(\boldsymbol{\theta}_{-a}) = \inf \{ \mathbb{E}_{\boldsymbol{\theta}}[T - N_a(T)] : \theta_a \in \Theta, \theta_a^* + \zeta u/2 \leq \theta_a \leq \theta_a^* + \zeta u \},$$

for every $\gamma \in]0, 1[$ there exists positive constants C_1, u_1 and T_1 (that depend on B, γ and ζ) such that if $u \leq u_1$ and $Tu^2 > T_1$, $\forall \boldsymbol{\theta}_{-a} \in B^{K-1}$,

$$\mathbb{E}_{\boldsymbol{\theta}_{a,u}}[N_a(T)] \geq \frac{(1-\gamma)\log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)} \left(1 - e^{-C_1 \log(Tu^2)} - \frac{2u^2 e_{T,u}(\boldsymbol{\theta}_{-a})}{(Tu^2)^{\frac{\gamma}{2}}} \right).$$

Using Lemma 14, if T satisfies moreover $\log(T) \geq 1/\min(u_0, u_1, \gamma/\log(T_1))$,

$$\mathcal{T}_a \geq 2(1-\gamma)^2(\mathcal{I}_1(T) - \mathcal{I}_2(T) - 2\mathcal{I}_3(T)),$$

where

$$\begin{aligned} \mathcal{I}_1(T) &:= \int_{B^{K-1}} h_a(\theta_a^*) \int_{T^{-(1-\gamma)/2}}^{(\log T)^{-1}} \frac{\log(Tu^2)}{u} du dH_{-a}(\boldsymbol{\theta}_{-a}), \\ \mathcal{I}_2(T) &:= \int_{B^{K-1}} h_a(\theta_a^*) \int_{T^{-(1-\gamma)/2}}^{(\log T)^{-1}} \frac{\log(Tu^2)}{u} \frac{1}{(Tu^2)^{C_1}} du dH_{-a}(\boldsymbol{\theta}_{-a}), \\ \mathcal{I}_3(T) &:= \int_{B^{K-1}} h_a(\theta_a^*) \int_{T^{-(1-\gamma)/2}}^{(\log T)^{-1}} \frac{\log(Tu^2)}{u} t^2 (Tu^2)^{-\frac{\gamma}{2}} e_{T,u}(\boldsymbol{\theta}_{-a}) du dH_{-a}(\boldsymbol{\theta}_{-a}). \end{aligned}$$

First, an explicit calculation yields

$$\mathcal{I}_1(T) = \frac{1}{4} \left(\left(1 - \frac{2\log \log(T)}{\log(T)} \right)^2 - \gamma^2 \right) \log^2(T) \int_{B^{K-1}} h_a(\theta_a^*) dH_{-a}(\boldsymbol{\theta}_{-a}),$$

which shows that

$$\mathcal{I}_1(T) \underset{T \rightarrow \infty}{\sim} \frac{1}{4} (1 - \gamma^2) \left(\int_{B^{K-1}} h_a(\theta_a^*) dH_{-a}(\boldsymbol{\theta}_{-a}) \right) \log^2(T).$$

Then, for every $\varepsilon > 0$, there exists $T_2(\varepsilon)$ such that for all $T \geq T_2(\varepsilon)$, for all $u \geq T^{-(1-\gamma)/2}$, $1/(Tu^2)^{C_1} \leq \varepsilon$. Hence, for $T \geq T_2(\varepsilon)$,

$$\mathcal{I}_2(T) \leq \varepsilon \mathcal{I}_1(T).$$

This proves that $\mathcal{I}_2(T) = \underset{T \rightarrow \infty}{o}(\log^2(T))$.

Finally, to prove that $\mathcal{I}_3(T) = \underset{T \rightarrow \infty}{o}(\log^2(T))$, we start by writing

$$\mathcal{I}_3(T) = \int_{T^{-(1-\gamma)/2}}^{(\log T)^{-1}} \frac{\log(Tu^2)}{u} (Tu^2)^{-\frac{\gamma}{2}} u^2 \left(\int_{B^{K-1}} e_{T,u}(\boldsymbol{\theta}_{-a}) h_a(\theta_a^*) dH_{-a}(\boldsymbol{\theta}_{-a}) \right) du.$$

and we provide an upper bound on the inner integral. First note that if $\boldsymbol{\theta}$ is such that $\theta_a > \theta_a^*$, one has

$$R_{\boldsymbol{\theta}}(T, \mathcal{A}) \geq (\dot{b}(\theta_a) - \dot{b}(\theta_a^*)) \mathbb{E}_{\boldsymbol{\theta}}[T - N_a(T)].$$

Using (25) together with this last inequality, one obtains, for every u ,

$$\begin{aligned} C \log^2(T) &\geq \int_{\{\boldsymbol{\theta} \in B^K : \theta_a^* + \zeta u/2 < \theta_a < \theta_a^* + \zeta u\}} R_T(\mathcal{A}, \boldsymbol{\theta}) dH(\boldsymbol{\theta}) \\ &\geq \int_{B^{K-1}} \int_{\theta_a^* + \zeta u/2}^{\theta_a^* + \zeta u} (\dot{b}(\theta_a) - \dot{b}(\theta_a^*)) \mathbb{E}_{\boldsymbol{\theta}}[T - N_a(T)] h_a(\theta_a) d\theta_a dH_{-a}(\boldsymbol{\theta}_{-a}) \\ &\geq \int_{B^{K-1}} e_{T,u}(\boldsymbol{\theta}_{-a}) \int_{\theta_a^* + \zeta u/2}^{\theta_a^* + \zeta u} (\dot{b}(\theta_a) - \dot{b}(\theta_a^*)) h_a(\theta_a) d\theta_a dH_{-a}(\boldsymbol{\theta}_{-a}). \end{aligned}$$

With $B = [b^-, b^+]$, let u_2 be such that the compact $B' = [b^- + \zeta u_2/2, b^+ + \zeta u_2]$ is included in Θ . As h_a is uniformly continuous and bounded on B' , there exists u_2 such that and for all $\theta_a^* \in B$, for all $u \leq u_2$,

$$\inf_{[\theta_a^* + \zeta u/2, \theta_a^* + \zeta u]} h_a(\theta) \geq \frac{2}{3} h_a(\theta_a^*).$$

Let $u \leq u_2$. Introducing $c_1 = \inf_{\theta \in B'} \ddot{b}(\theta) > 0$, using the Lagrange formula,

$$\begin{aligned} C \log^2(T) &\geq \frac{2c_1}{3} \int_{B^{K-1}} e_{T,u}(\boldsymbol{\theta}_{-a}) \int_{\theta_a^* + \zeta u/2}^{\theta_a^* + \zeta u} (\theta_a - \theta_a^*) h_a(\theta_a^*) d\theta_a dH_{-a}(\boldsymbol{\theta}_{-a}) \\ &= \frac{c_1}{4} \zeta^2 u^2 \int_{B^{K-1}} e_{T,u}(\boldsymbol{\theta}_{-a}) h_a(\theta_a^*) dH_{-a}(\boldsymbol{\theta}_{-a}). \end{aligned}$$

Finally, if $T^{-\frac{1-\gamma}{2}} \leq u \leq u_2$,

$$\int_{B^{K-1}} e_{T,u}(\boldsymbol{\theta}_{-a}) h_a(\theta_a^*) dH_{-a}(\boldsymbol{\theta}_{-a}) \leq \frac{4C}{c_1 \zeta^2} \frac{\log^2(T)}{u^2} \leq \frac{4C}{c_1 \zeta^2 \gamma^2} \frac{(\log(Tu^2))^2}{u^2}.$$

For T satisfying $\log(T)^{-1} \leq u_2$, $[T^{-(1-\gamma)/2}, (\log T)^{-1}] \subseteq [T^{-\frac{1-\gamma}{2}}, u_2]$ and

$$\begin{aligned} \mathcal{I}_3(T) &\leq \int_{T^{-(1-\gamma)/2}}^{(\log T)^{-1}} \frac{\log(Tu^2)}{u} (Tu^2)^{-\frac{\gamma}{2}} u^2 \left(\frac{4C}{c_1 \zeta^2 \gamma^2} \frac{(\log(Tu^2))^2}{u^2} \right) du \\ &= \frac{4C}{c_1 \zeta^2 \gamma^2} \int_{T^{-(1-\gamma)/2}}^{(\log T)^{-1}} \frac{\log(Tu^2)}{u} \left(\frac{(\log(Tu^2))^2}{(Tu^2)^{\frac{\gamma}{2}}} \right) du. \end{aligned}$$

Let $\varepsilon > 0$. As $x \mapsto \log^2(x)/(x^{\gamma/2})$ tends to zero when x tends to infinity, and $Tu^2 \geq T^\gamma$ for $u \geq T^{-(1-\gamma)/2}$, there exists a constant $T_3(\varepsilon)$ such that

$$\text{for } T \geq T_3(\varepsilon), \text{ for } t \geq T^{-(1-\gamma)/2}, \quad \frac{(\log(Tu^2))^2}{(Tu^2)^{\frac{\gamma}{2}}} \leq \varepsilon.$$

Hence, for $T \geq T_3(\varepsilon)$,

$$\begin{aligned} \mathcal{I}_3(T) &\leq \varepsilon \frac{4C}{c_1 \zeta^2 \gamma^2} \int_{T^{-(1-\gamma)/2}}^{(\log T)^{-1}} \frac{\log(Tu^2)}{u} du \\ &= \varepsilon \frac{C}{c_1 \zeta^2 \gamma^2} \left(\left(1 - \frac{2 \log \log(T)}{\log(T)} \right)^2 - \gamma^2 \right) \log^2(T), \end{aligned}$$

which proves that $\mathcal{I}_3(T) = o(\log^2(T))$.

Putting everything together, we proved that, for every algorithm $\mathcal{A}$, for every $\gamma > 0$, for every compact $B \subset \Theta$,

$$\liminf_{T \rightarrow \infty} \frac{\mathcal{R}_T(\mathcal{A}, H)}{\log^2(T)} \geq (1-\gamma)^2 (1-\gamma^2) \frac{1}{2} \sum_{a=1}^K \int_{B^{K-1}} h_a(\theta_a^*) dH_{-a}(\boldsymbol{\theta}_{-a}).$$

Taking the supremum over all compact set B yields, for every $\gamma > 0$,

$$\liminf_{T \rightarrow \infty} \frac{\mathcal{R}_T(\mathcal{A}, H)}{\log^2(T)} \geq (1-\gamma)^2 (1-\gamma^2) \frac{1}{2} \sum_{a=1}^K \int_{\Theta^{K-1}} h_a(\theta_a^*) dH_{-a}(\boldsymbol{\theta}_{-a}),$$

provided the integral in the right-hand side is finite. Letting γ go to zero concludes the proof.

D.2 Proof of Lemma 13

Let $\zeta \in]0, 1[$ be fixed. As $B = [b^-, b^+]$ is strictly included in Θ , there exists u_1 such that $\mathcal{C} := [b^- - u_1, b^+ + \zeta u_1]$ is included in Θ . For $(\theta, u) \in B \times [0, u_1]$ we define

$$f(\theta, u) = \frac{(1 + \zeta)^2 u}{2} \frac{(\dot{b}(\theta) - \dot{b}(\theta - u))h_a(\theta - u)}{K(\theta - u, \theta + \zeta u)h_a(\theta)}.$$

f is continuous on $B \times [0, t_1]$ and it can be checked that

$$\lim_{(\theta, u) \rightarrow (\theta_0, 0)} f(\theta, u) = 1.$$

As f is uniformly continuous, there exists $u_0 \leq u_1$, such that for all $u \leq u_0$, for all $\theta \in B$,

$$|f(\theta, u) - 1| \leq \frac{\gamma}{2},$$

which rewrites

$$\left| \frac{(\dot{b}(\theta) - \dot{b}(\theta - u))h_a(\theta - u)}{K(\theta - u, \theta + \zeta u)} - \frac{2h_a(\theta)}{(1 + \zeta)^2 u} \right| \leq \frac{\gamma}{2} \frac{2h_a(\theta)}{(1 + \zeta)^2 u}$$

hence, for $u \leq u_0$, one has

$$\frac{(\dot{b}(\theta) - \dot{b}(\theta - u))h_a(\theta - u)}{K(\theta - u, \theta + \zeta u)} \geq \frac{1 - \frac{\gamma}{2}}{(1 + \zeta)^2} \frac{2h_a(\theta)}{u}.$$

Applying this to ζ such that $1 + \zeta = \sqrt{\frac{1 - \frac{\gamma}{2}}{1 - \gamma}}$ concludes the proof.

D.3 Proof of Lemma 14

Let $\zeta \in]0, 1[$ be fixed and define u_1 and $\mathcal{C} = [b^- - u_1, b^+ + \zeta u_1] \subset \Theta$ as in the proof of Lemma 13. Let $u \leq u_1$ and fix $\theta_{-a} \in B^{K-1}$. First, using Markov inequality,

$$\mathbb{E}_{\theta_{a,u}}[N_a(T)] \geq \frac{(1 - \gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)} \mathbb{P}_{\theta_{a,u}} \left(N_a(T) \geq \frac{(1 - \gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)} \right).$$

Thus it is sufficient to prove that there exists $C_1 > 0$ such that

$$\mathbb{P}_{\theta_{a,u}} \left(N_a(T) \leq \frac{(1 - \gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)} \right) \leq e^{-C_1 \log(Tu^2)} + \frac{2u^2 e_{T,u}(\theta_{-a})}{(Tu^2)^{\frac{\gamma}{2}}}. \quad (26)$$

As $u \leq u_1$, the set $\{\theta_a : \theta_a^* + \zeta u/2 \leq \theta_a \leq \theta_a^* + \zeta u\}$ is a compact set included in $\mathcal{C}$, therefore there exists $\lambda \in B^{K-1} \times \mathcal{C}$ that attains the infimum in the definition of $e_{T,u}(\theta_{-a})$:

$$e_{T,u}(\theta_{-a}) = \mathbb{E}_\lambda[T - N_a(T)],$$

with $\lambda_{-a} = \theta_{-a}$ and $\lambda_a = \theta_a^* + \varepsilon u$, for some $\varepsilon \in [\zeta/2, \zeta]$. Using Markov inequality,

$$\mathbb{P}_\lambda \left(N_a(T) \leq \frac{(1 - \gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)} \right) \leq \frac{\mathbb{E}_\lambda[T - N_a(T)]}{T - \frac{(1 - \gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)}} = \frac{e_{T,u}(\theta_{-a})}{T \left(1 - \frac{(1 - \gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)T} \right)}.$$

Introducing $c_1 = \inf_{\theta \in \mathcal{C}} \ddot{b}(\theta)$, using (16) in Proposition 11, for $u \leq u_1$,

$$\frac{(1 - \gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)T} \leq \frac{2(1 - \gamma) \log(Tu^2)}{c_1(1 + \zeta)^2(Tu^2)} \leq \frac{1}{2},$$

where the last inequality holds to Tu^2 large enough. Thus there exists $T_1 > 0$ such that for $u \leq u_1$ and $Tu^2 \geq T_1$,

$$\mathbb{P}_\lambda \left(N_a(T) \leq \frac{(1-\gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)} \right) \leq \frac{2e_{T,u}(\boldsymbol{\theta}_{-a})}{T}. \quad (27)$$

Introducing the log likelihood ratio $L_n = \sum_{s=1}^n \log \frac{f_{\theta_a^* - t}(Y_{a,s})}{f_{\lambda_a}(Y_{a,s})}$, where $Y_{a,s}$ are i.i.d. samples of the distribution of arm a , one can write

$$\begin{aligned} & \mathbb{P}_{\boldsymbol{\theta}_{a,u}} \left(N_a(T) \leq \frac{(1-\gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)} \right) \\ & \leq \mathbb{P}_{\boldsymbol{\theta}_{a,u}} \left(N_a(T) \leq \frac{(1-\gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)}, L_{N_a(T)} \leq \left(1 - \frac{\gamma}{2}\right) \log(Tu^2) \right) \end{aligned} \quad (28)$$

$$+ \mathbb{P}_{\boldsymbol{\theta}_{a,u}} \left(\max_{n \leq \frac{(1-\gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)}} L_n \geq \left(1 - \frac{\gamma}{2}\right) \log(Tu^2) \right) \quad (29)$$

An upper bound on Term (28) follows from a change of distribution argument. Let $\mathcal{E}$ be the event

$$\mathcal{E} := \left\{ N_a(T) \leq \frac{(1-\gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)}, L_{N_a(T)} \leq \left(1 - \frac{\gamma}{2}\right) \log(Tu^2) \right\}$$

As $\mathcal{E} \in \mathcal{F}_{N_a(T)}$, one has

$$\mathbb{P}_\lambda(\mathcal{E}) = \mathbb{E}_{\boldsymbol{\theta}_{a,u}} [\mathbb{1}_{\mathcal{E}} \exp(-L_{N_a(T)})] \geq \exp\left(-\left(1 - \frac{\gamma}{2}\right) \log(Tu^2)\right) \mathbb{P}_{\boldsymbol{\theta}_{a,u}}(\mathcal{E}).$$

Thus, using moreover (27),

$$\begin{aligned} (28) & \leq (Tu^2)^{1-\frac{\gamma}{2}} \mathbb{P}_\lambda(\mathcal{E}) \leq (Tu^2)^{1-\frac{\gamma}{2}} \mathbb{P}_\lambda \left(N_a(T) \leq \frac{(1-\gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)} \right) \\ & \leq 2u^2 (Tu^2)^{-\frac{\gamma}{2}} e_{T,u}(\boldsymbol{\theta}_{-a}). \end{aligned}$$

An upper bound of Term (29) follows from a concentration inequality specific to exponential families, stated as Lemma 15, whose proof is provided below for the sake of completeness.

Lemma 15. *Let the (Y_i) be i.i.d with distribution ν_θ and mean $\mu = \dot{b}(\theta)$.*

$$\mathbb{P} \left(\max_{n \leq N} \sum_{s=1}^n (\mu - Y_i) \geq x \right) \leq \exp \left(-Nd \left(\mu - \frac{x}{N}, \mu \right) \right)$$

Introducing the notation $\bar{\theta}_a = \theta_a^* - u$ and

$$K_T = \frac{(1-\gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)} = \frac{(1-\gamma) \log(Tu^2)}{K(\bar{\theta}_a, \theta_a^* + \zeta u)},$$

the log likelihood ratio can be made explicit, and satisfies, for $n \leq K_T$,

$$\begin{aligned} L_n &= \sum_{s=1}^n (\bar{\theta}_a - \lambda_a) Y_{a,s} - b(\bar{\theta}_a) + b(\lambda_a) \\ &= (\bar{\theta}_a - \lambda_a) \sum_{s=1}^n (Y_{a,s} - \dot{b}(\bar{\theta}_a)) + nK(\bar{\theta}_a, \lambda_a). \\ &\leq (\lambda_a - \bar{\theta}_a) \sum_{s=1}^n (\dot{b}(\bar{\theta}_a) - Y_{a,s}) + (1-\gamma) \log(Tu^2). \end{aligned}$$

Term (29) is upper bounded by

$$\begin{aligned} & \mathbb{P}_{\theta_{a,u}} \left(\max_{n \leq K_T} \left[(\lambda_a - \bar{\theta}_a) \sum_{s=1}^n (\dot{b}(\bar{\theta}_a) - Y_{a,s}) + (1 - \gamma) \log(Tu^2) \right] \geq \left(1 - \frac{\gamma}{2}\right) \log(Tu^2) \right) \\ & \leq \mathbb{P}_{\theta_{a,u}} \left(\max_{n \leq K_T} \sum_{s=1}^n (\dot{b}(\bar{\theta}_a) - Y_{a,s}) \geq \frac{\gamma}{2} \frac{\log(Tu^2)}{\lambda_a - \bar{\theta}_a} \right). \end{aligned}$$

Under $\theta_{a,u}$, the sequence $Y_{a,s}$ is i.i.d with distribution ν_{θ_a} . Therefore, using Lemma 15 one obtains, with the notation $\bar{\mu}_a = \dot{b}(\bar{\theta}_a)$,

$$(29) \leq \exp \left(-K_T d \left(\bar{\mu}_a - \frac{\gamma K(\theta_a^* - u, \theta_a^* + \zeta u)}{2(1 - \gamma)(\varepsilon + 1)u}, \bar{\mu}_a \right) \right).$$

Letting $c_1 = \inf_{\theta \in \mathcal{C}} \ddot{b}(\theta)$ and $c_2 = \inf_{\theta \in \mathcal{C}} \ddot{b}(\theta)$, from (16) in Proposition 11,

$$\frac{\gamma K(\theta_a^* - u, \theta_a^* + \zeta u)}{2(1 - \gamma)(\varepsilon + 1)u} \in \left[\frac{\gamma}{2(1 - \gamma)} \frac{c_1}{2} (\zeta + 1)^2 u^2; \frac{\gamma}{2(1 - \gamma)} \frac{c_2}{2} (\zeta + 1)^2 u^2 \right].$$

Thus, for u small enough, $\bar{\mu}_a$ and $\bar{\mu}_a - \frac{\gamma K(\theta_a^* - u, \theta_a^* + \zeta u)}{2(1 - \gamma)(\varepsilon + 1)u}$ belong to a compact $\mathcal{C}'$ satisfying $\mathcal{C} \subseteq \mathcal{C}' \subseteq \Theta$. Letting $c'_2 = \sup_{\theta \in \mathcal{C}'} \ddot{b}(\theta)$, using (17),

$$\begin{aligned} (29) & \leq \exp \left(-\frac{K_T}{2c'_2} \left(\frac{\gamma K(\theta_a^* - u, \theta_a^* + \zeta u)}{2(1 - \gamma)(\varepsilon + 1)u} \right)^2 \right) \\ & = \exp \left(-\log(Tu^2) \frac{\gamma^2}{8(1 - \gamma)c'_2} \frac{K(\theta_a^* - u, \theta_a^* + \zeta u)}{(1 + \varepsilon)^2 u^2} \right) \\ & \leq \exp \left(-\log(Tu^2) \frac{\gamma^2 c_1}{8(1 - \gamma)c'_2} \frac{c_1(1 + \zeta)^2}{(1 + \varepsilon)^2} \right). \end{aligned}$$

Letting $C_1 = \frac{\gamma^2 c_1}{8(1 - \gamma)c'_2} \frac{c_1(1 + \zeta)^2}{(1 + \varepsilon)^2}$, from the upper bounds obtained on (28) and (29), it follows that

$$\mathbb{P}_{\theta_{a,u}} \left(N_a(T) \leq \frac{(1 - \gamma) \log(Tu^2)}{K(\theta_a^* - u, \theta_a^* + \zeta u)} \right) \leq 2u^2 (Tu^2)^{-\frac{\gamma}{2}} e_{T,u}(\theta_{-a}) + e^{-C_1 \log(Tu^2)},$$

provided that $u \leq u_1$ and $Tu^2 \geq T_1$, which concludes the proof. $\square$

Proof of Lemma 15 The proof follows from the Chernoff technique and a maximal inequality.

Let $S_n = \sum_{s=1}^n (\mu - Y_i)$. For every $\lambda > 0$,

$$\mathbb{P} \left(\max_{n \leq N} S_n \geq x \right) = \mathbb{P} \left(\max_{n \leq N} e^{\lambda S_n} \geq e^{\lambda x} \right) \leq e^{-\lambda x} \mathbb{E} [e^{\lambda S_N}], \quad (30)$$

where the last inequality is a consequence of Doob's maximal inequality applied to $M_n = e^{\lambda S_n}$, which is a sub-martingale with respect to the filtration generated by the (Y_i) . Indeed, using the convexity of the mapping $x \mapsto e^{\lambda x}$,

$$\begin{aligned} \mathbb{E} [M_n - M_{n-1} | \mathcal{F}_{n-1}] & = e^{\lambda S_n} \mathbb{E} [e^{\lambda(S_n - S_{n-1})} - 1 | \mathcal{F}_{n-1}] \\ & \geq e^{\lambda S_n} \lambda \mathbb{E} [S_n - S_{n-1} | \mathcal{F}_{n-1}] = 0. \end{aligned}$$

Using the independence of the Y_i and $\mathbb{E}[e^{\lambda Y_i}] = \exp(b(\theta + \lambda) - b(\theta))$ for any $\lambda \in \mathbb{R}$, it can be show that

$$e^{-\lambda x} \mathbb{E} [e^{\lambda S_N}] = \exp \left(-N \left[\lambda \left(\frac{x}{N} - \dot{b}(\theta) \right) + b(\theta) - b(\theta - \lambda) \right] \right).$$

The exponent is minimized of λ^* satisfying $\dot{b}(\theta - \lambda^*) = \dot{b}(\theta) - x/N$ and

$$\begin{aligned} e^{-\lambda^* x} \mathbb{E} \left[e^{\lambda^* S_N} \right] &= \exp \left(-N \left[\dot{b}(\theta - \lambda^*)(-\lambda^*) - \dot{b}(\theta - \lambda^*) + b(\theta) \right] \right) \\ &= \exp(-NK(\theta - \lambda^*, \theta)) = \exp \left(-Nd \left(\mu - \frac{x}{N}, \mu \right) \right). \end{aligned}$$

The conclusion follows by plugging λ^* in (30).

D.4 The lower bound for Bernoulli bandits

As pointed out by [27], in the particular case in which $h_a(\theta) = q(\theta)$ for all $a = 1, \dots, K$, using the fact that the distribution of $\max_{a \in \mathcal{S}} \theta_a$ has density $kq(\theta)Q^{k-1}(\theta)$ where Q is the c.d.f. of the distribution with density q and $k = |\mathcal{S}|$, the constant in the lower bound can be expressed

$$\frac{1}{2} \sum_{a=1}^K \int_{\Theta^{K-1}} h_a(\theta_a^*) dH_{-a}(\theta_{-a}) = \frac{K(K-1)}{2} \int_{\Theta} q^2(\theta) Q^{K-2}(\theta) d\theta. \quad (31)$$

Now consider a Bernoulli bandit model with K arms, with a uniform prior distribution on the mean of each arm. The set of Bernoulli distribution of means $\mu \in [0, 1]$ form an exponential family when each distribution is parametrized by the natural parameter $\theta = \log(\mu/(1-\mu))$. This exponential family is characterized by

$$\Theta = \mathbb{R}, \quad b(\theta) = \log(1 + e^\theta),$$

and the reference measure is the Lebesgue measure. As each mean μ_a is drawn from a uniform distribution on $[0, 1]$, the associated natural parameter θ_a is drawn from a distribution on $\mathbb{R}$ having respectively density and c.d.f.

$$q(\theta) = \frac{e^\theta}{(1 + e^\theta)^2} \quad \text{and} \quad Q(\theta) = \frac{e^\theta}{1 + e^\theta}.$$

Using the formula (31), the constant of the lower bound is

$$\begin{aligned} \frac{K(K-1)}{2} \int_{-\infty}^{+\infty} \frac{e^{K\theta}}{(1 + e^\theta)^{K+2}} d\theta &= \frac{K(K-1)}{2} \int_0^1 \frac{x^{K-1}}{(1+x)^{K+2}} dx \\ &= \frac{K(K-1)}{2} \frac{1}{K(K+1)}, \end{aligned}$$

where the integral is computed using by inducting, using a by part integration. Finally, the asymptotic rate of the Bayes risk for a Bernoulli bandit model with K arms and a uniform prior on their means is

$$\frac{1}{2} \frac{K-1}{K+1} \log^2(T).$$