



# Computational aspects of Bayesian spectral density estimation

Nicolas Chopin, Judith Rousseau, Brunero Liseo

## ► To cite this version:

Nicolas Chopin, Judith Rousseau, Brunero Liseo. Computational aspects of Bayesian spectral density estimation. 2011. hal-00767466

**HAL Id: hal-00767466**

**<https://hal.science/hal-00767466>**

Preprint submitted on 19 Dec 2012

**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.

# COMPUTATIONAL ASPECTS OF BAYESIAN SPECTRAL DENSITY ESTIMATION

N. CHOPIN (CREST-ENSAE), J. ROUSSEAU (CREST-ENSAE AND UNIVERSITÉ PARIS  
DAUPHINE), AND B. LISEO (UNIVERSITA DI ROMA)

ABSTRACT. Gaussian time-series models are often specified through their spectral density. Such models pose several computational challenges, in particular because of the non-sparse nature of the covariance matrix. We derive a fast approximation of the likelihood for such models. We use importance sampling to correct for the approximation error. We show that the variance of the importance sampling weights vanishes as the sample size goes to infinity. We show that the posterior is typically multi-modal, and derive a Sequential Monte Carlo sampler based on an annealing sequence in order to sample from the approximate posterior. Performance of the overall approach is evaluated on simulated and real datasets.

## 1. INTRODUCTION

Several models in the time series literature are defined through their spectral density. Assuming Gaussianity, one observes a vector  $\mathbf{x}$  of length  $n$ , from the Gaussian distribution

$$\mathbf{x}|\mu, f \sim N(\mu\mathbf{1}, \mathbf{T}(f))$$

where  $\mathbf{1} = (1, \dots, 1)'$ , and  $\mathbf{T}(f)$  is the  $n \times n$  Toeplitz matrix associated to spectral density  $f$ , with entries  $\mathbf{T}(f)(l, m) = \gamma_f(l - m)$ ,

$$(1.1) \quad \gamma_f(l) = \int_{-\pi}^{\pi} f(\lambda) e^{il\lambda} d\lambda, \quad l = -(n-1), \dots, (n-1).$$

The models vary with respect to the specification of  $f$ . For instance, the FEXP parametrisation (Hurvich et al., 2002; Moulines and Soulier, 2003) assumes that

$$(1.2) \quad f(\lambda) = \frac{1}{2\pi} |1 - e^{-i\lambda}|^{-2d} g(\lambda), \quad g(\lambda) = \exp \left\{ \sum_{j=0}^k \xi_j \cos(j\lambda) \right\}.$$

This specification conveniently separates the long-range behaviour, as determined by parameter  $d \in [0, 1/2)$ , and the short-memory part  $g$ . By taking  $k$  large enough,  $g$  may be arbitrarily close to any function in a certain regularity class. Rousseau et al. (2012) show that, for a well chosen prior with respect to  $(k, d, \xi_1, \dots)$ , the corresponding posterior is consistent for both  $d$  and  $f$ , under semi-parametric settings; that is, assuming that the true spectral density belongs to a certain infinite-dimensional class of functions.

This FEXP parametrisation will be our running example. We note however that many other parametrisations of  $f$  are possible, which can also be tackled by the

methodology developed in this paper. One may take  $d = 0$ , and obtain a non-parametric procedure for estimating the spectral density, under a short-memory assumption. One may replace  $g$  by another type of expansion, a spline regression, and so on. Finally, one may also consider a purely parametric model, such as an ARFIMA( $p, q$ ) model, given by:

$$f(\lambda) = \frac{1}{2\pi} |1 - e^{-i\lambda}|^{-2d} \left| \frac{1 + \sum_{j=1}^q \xi_j e^{-ij\lambda}}{1 - \sum_{j=1}^p \phi_j e^{-ij\lambda}} \right|^2.$$

ARMA models may also be defined through their spectral density, but well-known specialised methods exist for such models, hence they fall outside the scope of this paper.

Whatever the specification of  $f$ , parametric or semiparametric, several computational difficulties arise regarding Bayesian inference for such models. First, the likelihood of the data

$$(1.3) \quad p(\mathbf{x}|\mu, f) = (2\pi)^{-n/2} |\mathbf{T}(f)|^{-1/2} \exp \left\{ -\frac{1}{2} (\mathbf{x} - \mu\mathbf{1})^T \mathbf{T}(f)^{-1} (\mathbf{x} - \mu\mathbf{1}) \right\}$$

involves a determinant and a quadratic form which are expensive to compute, i.e. the cost is  $O(n^3)$  if one uses off-the-shelf methods. Second, the entries of matrix  $\mathbf{T}(f)$  itself, that is, the Fourier integrals (1.1) cannot be computed reliably using standard quadrature methods, because of the many oscillations of the integrand for large values of  $k$ . Third, Gibbs sampling is generally not feasible for posterior distributions associated to this likelihood. One usually resorts to the Metropolis-Hastings sampler (see e.g. Robert and Casella, 2004, Chap. 7), which is difficult to tune in order to obtain reasonable performance. This is quite problematic in this context: since the likelihood is expensive, performing several pilot runs in order to progressively tune the sampler may be a long and tedious process for the user. This problem is compounded if  $f$  is specified through a trans-dimensional prior; for instance, in the FEXP model above, if  $k$  is random. Then one needs to implement an algorithm for trans-dimensional sampling spaces, such as Green's algorithm (Green, 1995; Richardson and Green, 1997), also known as the reversible-jump sampler, which is harder to tune yet.

This paper proposes to address these problems in an unified manner, and is organised as follows. Section 2 discusses the exact computation the likelihood. It is seen that the cost of this operation is  $O(n^3)$ , but the constant in  $O(n^3)$  is typically small, hence on modern hardware it remains possible to perform a reasonable number of such evaluations provided  $n$  is not too large.

Section 3 proposes a fast approximation of the likelihood, the cost of which is essentially  $O(n)$ . This approximation scheme motivates the following approach. In a first step, we perform Monte Carlo simulation of the approximate posterior, that is, the prior times the approximate likelihood. This step should cost  $O(n)$ , but typically with a large constant in front of  $n$ , due to the intensive nature of Monte Carlo algorithms. In a second step, but only if  $n$  is not too large, we correct for the approximation by doing importance sampling on a reasonable number of simulated samples; the cost of this second step is  $O(n^3)$  but this time with a small constant in front of  $n^3$ . Thus, what we observe in practice is that the cost of this second step is in fact negligible with respect to the first step, at least for say  $n \ll 10^4$ . On the other hand, if  $n$  is large, we show that the importance sampling step becomes

superfluous, because the variance of the importance sampling weights converge to 0 as  $n \rightarrow +\infty$  (under appropriate assumptions).

Section 4 discusses a Sequential Monte Carlo algorithm (SMC) for sampling from the approximate posterior; see Del Moral et al. (2006) for a general introduction to SMC. We shall see that the following advantages of SMC (relative to for instance MCMC) are particularly useful in this specific context. First, as explained above, one would like to apply the importance sampling step to a sample as small as possible, because the exact likelihood is expensive to compute. We shall see that our SMC sampler typically generates particles that are close to IID samples from its target distribution. Second, as also mentioned above, one has little prior information on the structure of the (approximate or exact) posterior. It is easy however to make SMC adaptive, by learning iteratively features of the target from the sample of simulated particles. Third, we observe in certain settings that semi-parametric models such as the FEXP model, may generate multimodal posteriors. The SMC sampler we propose is based on tempering ideas, and thus more able to escape from minor local modes.

Section 4 illustrates the proposed approach on simulated and real data.

We shall use the following notations: vectors and matrices are always in bold face, e.g.  $\mathbf{x}$  and  $\mathbf{\Sigma}$ , the determinant, transpose, and trace of  $\mathbf{\Sigma}$  are denoted respectively  $|\mathbf{\Sigma}|$ ,  $\mathbf{\Sigma}^T$  and  $\text{tr}(\mathbf{\Sigma})$ .

## 2. EXACT COMPUTATION OF THE LIKELIHOOD

**2.1. Marginalisation.** As a preliminary, we note that  $f$  is often parametrised in such a way that  $f = \sigma^2 \bar{f}_{\boldsymbol{\theta}}$ , where  $\sigma$  is a scale parameter, and  $\boldsymbol{\theta}$  is the vector of all remaining parameters (except  $\mu$ ). For instance, in the FEXP case, we may set  $\sigma^2 = \exp(\xi_0)$ ,  $\boldsymbol{\theta} = (k, \boldsymbol{\theta}_k)$ ,  $\boldsymbol{\theta}_k = (\text{logit}(2d), \xi_1, \dots, \xi_k)$ , and

$$(2.1) \quad \bar{f}_{\boldsymbol{\theta}}(\lambda) = \frac{1}{2\pi} |1 - e^{-i\lambda}|^{-2d} \bar{g}_{\boldsymbol{\theta}}(\lambda), \quad \bar{g}_{\boldsymbol{\theta}}(\lambda) = \exp \left\{ \sum_{j=1}^k \xi_j \cos(j\lambda) \right\}.$$

(The function  $\text{logit}$  is defined as  $\text{logit}(x) = \log(x) - \log(1 - x)$ , and is used here to facilitate the construction of proposal distributions, see Section 4). It is then possible to marginalise out both  $\mu$  and  $\sigma^2$  from the likelihood, provided these two parameters are assigned the following standard  $g$ -prior distribution, independently from  $\boldsymbol{\theta}$ :  $1/\sigma^2 \sim \text{Gamma}(a, b)$ ,  $\mu|\sigma^2 \sim N(m_{\mu}, \sigma^2/g_{\mu})$ . The marginal likelihood reads:

$$(2.2) \quad \begin{aligned} p(\mathbf{x}|\boldsymbol{\theta}) &= \int p(\mathbf{x}|\mu, \sigma^2, \boldsymbol{\theta}) p(\mu, \sigma^2) d\mu d\sigma^2 \\ &\propto \left| \mathbf{T}(\bar{f}_{\boldsymbol{\theta}}) + \frac{1}{g_{\mu}} \mathbf{E} \right|^{-1/2} \times \\ &\quad \left\{ b + \frac{1}{2} (\mathbf{x} - m_{\mu} \mathbf{1})^T \left( \mathbf{T}(\bar{f}_{\boldsymbol{\theta}}) + \frac{1}{g_{\mu}} \mathbf{E} \right)^{-1} (\mathbf{x} - m_{\mu} \mathbf{1}) \right\}^{-a-n/2}, \end{aligned}$$

where  $\mathbf{E}$  is the  $n \times n$  matrix filled with ones.

Marginalising out parameters usually improves the performance of the sampling algorithm. We note however that the approach developed in this paper would work with little modification for the unmarginalised likelihood  $p(\mathbf{x}|\mu, \sigma^2, \boldsymbol{\theta})$ . These

two likelihood functions suffer from the same computational difficulties, which are described in the next section.

**2.2. Computational difficulties associated to likelihood evaluation.** We review in this section the specific difficulties that arise when evaluating either the standard likelihood function (1.3), or the marginal version (2.2). First, both likelihood functions include some quadratic form  $\mathbf{y}\Sigma^{-1}\mathbf{y}$  involving the inverse of a  $n \times n$  symmetric matrix  $\Sigma$ ; in (1.3),  $\Sigma = \mathbf{T}(f)$ , and in (2.2),  $\Sigma = \mathbf{T}(\bar{f}_\theta) + \frac{1}{g_\mu}\mathbf{E}$ . Second, both functions involve the determinant of the same matrix  $\Sigma$ . Third, in both cases, evaluating the entries of  $\Sigma$  requires computing simultaneously  $n$  Fourier integrals, see (1.1).

Our solution to the third difficulty is described in the two next Sections. Regarding the two first points, the most direct solution is to compute the Cholesky lower triangle of  $\Sigma$ ,  $\Sigma = \mathbf{C}\mathbf{C}^T$ . Then, one obtains the determinant by taking the square of the product of the diagonal elements of  $\mathbf{C}$ , and one computes  $\mathbf{y}^T\Sigma^{-1}\mathbf{y} = (\mathbf{C}^{-1}\mathbf{y})^T\mathbf{C}^{-1}\mathbf{y}$  as the norm of the solution (in  $\mathbf{z}$ ) of the linear system  $\mathbf{C}\mathbf{z} = \mathbf{y}$ , which is quickly obtained by back-substitution. The Cholesky decomposition is a  $O(n^3)$  operation.

For the sake of completeness, we mention briefly faster, but specialised, algorithms for solving directly the system  $\Sigma\mathbf{z} = \mathbf{y}$ , in order to compute  $\mathbf{y}^T\Sigma^{-1}\mathbf{y}$ . Given that  $\Sigma$  is Toeplitz, one may use Levinson's algorithm (Levinson, 1949; Press et al., 2007, p. 96), which is  $O(n^2)$ . Alternatively, Chen et al. (2006) have developed a variant of the conjugate gradient method, based on a particular preconditioned matrix. Their algorithm requires  $O(\log^{3/2}n)$  iterations, each involving a FFT transform over  $2n$  points. The overall cost is therefore  $O(n\log^{5/2}n)$ .

Unfortunately, these alternative approaches do not provide an evaluation of the determinant as a by-product. Chen et al. (2006), Holan et al. (2009) approximate the determinant by using a particular asymptotic approximation, which we describe later, but obviously this approach is not entirely satisfactory when used within a non-asymptotic, and in particular, a Bayesian, approach.

Perhaps more importantly, we note that the constant in front of the  $O(n^3)$  cost of the Cholesky decomposition is typically very small, due to very efficient implementation in most scientific software. To give an order of magnitude, for  $n = 10^3$ , 1000 of such operations takes about one minute on the first author's computer. This observation underpins the strategy laid out in the introduction: to run some Monte Carlo algorithm so as to sample from an approximation of the posterior, then, provided  $n$  is not too large, to correct for the approximation using importance sampling on a moderate (possibly sub-sampled from the first step) Monte Carlo sample.

**2.3. Computing the Fourier coefficients.** In this section and the following, we consider the problem of evaluating simultaneously the  $n$  Fourier integrals defined in (1.1). As noted in the introduction, using standard quadrature would work very poorly, because the integrand in (1.1) strongly oscillates when  $n$  gets large. The solution we describe here seems well known in the numerical mathematics literature (e.g. Press et al., 2007, Chap. 13, Sect. 9), yet, to the best of our knowledge, it has not been used before in the time series literature. Instead, previous approaches (e.g. Chen et al., 2006; Holan et al., 2009) rely on more specific algorithms such as the splitting algorithm of Bertelli and Caporin (2002). The approach described

here has the same computational cost as such alternative approaches, that is, that of a FFT (Fast Fourier Transform), i.e.  $O(n \log n)$ . However, we find our approach slightly more convenient, for the following reasons: (a) this is a generic approach, which requires only pointwise evaluation of the spectral density, whereas the splitting algorithm requires exact expressions for the moving average coefficients (of the short memory part) which are specific to the considered class of spectral densities; for instance, it is unclear how one could use these methods if  $f$  would be specified through splines; (b) it is characterised by only one level of approximation (see below), whereas the splitting algorithm expresses first the moving averages coefficients as an infinite sum, which must be truncated, then plug these coefficients into another infinite sum, which must be truncated again, so assessing the numerical error is slightly more delicate; and (c) in long memory settings, the terms of these infinite sums are supposed to decay slowly, hence one may need to truncate to a large number of terms.

Let  $M$  a power of two, such that  $M \geq 2n$ . We explain first how to compute efficiently and simultaneously the Fourier integrals

$$\gamma_g(l) = \int_{-\pi}^{\pi} g(\lambda) e^{il\lambda} d\lambda$$

for a given bounded function  $g$ , and  $0 \leq l \leq M/2$ . (If  $n < M/2$ , simply discard the extra values.) If the spectral density  $f$  itself is bounded, that is, if it corresponds to a short memory process, then the following method may be used directly by taking  $g = f$ . If  $f$  corresponds to a long-memory process, then  $f$  diverges at 0, and a minor modification is required to use the following method, see next section.

The idea is to replace  $g$  by a linear interpolation  $\tilde{g}$ :

$$\tilde{g}(\lambda) = \sum_{j=0}^{M-1} g_j \psi\left(\frac{\lambda - \lambda_j}{\Delta}\right) + g_0 \varphi_0\left(\frac{\lambda - \lambda_0}{\Delta}\right) + g_M \varphi_M\left(\frac{\lambda - \lambda_M}{\Delta}\right)$$

where  $\Delta = 2\pi/M$ ,  $\lambda_j = -\pi + j\Delta$ ,  $g_j = \tilde{g}(\lambda_j)$   $j = 0, \dots, M$ , and  $\psi$  is the linear interpolation kernel, i.e.  $\psi(\lambda) = (1 - |\lambda|)^+$ ;  $\varphi_0$  and  $\varphi_M$  are boundary corrections, the expression of which may be skipped for the rest of the discussion. Note  $\tilde{g}$  is defined on the entire real-line, and is zero outside  $[-\pi, \pi]$ . Applying the operator  $\int (\cdot) e^{il\lambda} d\lambda$  yields:

$$\begin{aligned} \int_{-\pi}^{\pi} \tilde{g}(\lambda) e^{il\lambda} d\lambda &= \int_{-\infty}^{+\infty} \tilde{g}(\lambda) e^{il\lambda} d\lambda \\ &= \sum_{j=0}^{M-1} g_j \int_{-\infty}^{+\infty} \psi\left(\frac{\lambda - \lambda_j}{\Delta}\right) e^{il\lambda} d\lambda \\ &\quad + g_0 \int_{-\infty}^{+\infty} \varphi_0\left(\frac{\lambda - \lambda_0}{\Delta}\right) e^{il\lambda} d\lambda + g_M \int_{-\infty}^{+\infty} \varphi_M\left(\frac{\lambda - \lambda_M}{\Delta}\right) e^{il\lambda} d\lambda \\ &= \Delta(-1)^l \left\{ W(l\Delta) \sum_{j=0}^{M-1} g_j e^{ijl\frac{2\pi}{M}} + g_0 \alpha_0(l\Delta) + g_M \alpha_M(l\Delta) \right\} \end{aligned}$$

where  $W(\lambda) = \int_{-\infty}^{\infty} \psi(u) e^{i\lambda u} du$ ,  $\alpha_0(\lambda) = \int_{-\infty}^{\infty} \varphi_0(u) e^{i\lambda u} du$ ,  $\alpha_M(\lambda) = \int_{-\infty}^{\infty} \varphi_M(u - M) e^{i\lambda u} du$ . The first sum may be computed using a FFT, in  $O(M \log(M))$  time.

All the other functions admit close-form expressions, given by Press et al. (2007, Chap. 13, Sect. 9):

$$W(\lambda) = \frac{2(1 - \cos \lambda)}{\lambda^2}, \quad \alpha_0 = \alpha_M = -\frac{W}{2}.$$

The scheme above may be adapted so as to rely on a cubic (rather than linear) interpolation. This point is interesting in settings where the spectral density is parametrised in terms of cubic splines. Then, the method described here becomes exact. Otherwise, the accuracy of the method is determined by the size of the grid,  $M + 1$ . We follow Press et al. (2007) and takes  $M$  to be the smallest power of two such that  $M \geq 2n$ , and we observe in practice that it gives very accurate results (in the sense that larger value of  $M$  give sensibly the same values).

**2.4. Computing the Fourier coefficients when  $f$  diverges at 0.** In this Section, we explain how to adapt the method above for computing Fourier integrals, in the situations where  $f$  is a long-range dependent spectral density,

$$f(\lambda) = \frac{1}{2\pi} |1 - e^{-i\lambda}|^{-2d} g(\lambda)$$

where  $g$  is a bounded function, and  $0 < d < 1/2$ . In this case,  $f$  diverges at 0, hence it may not be well approximated by a piecewise linear function. To address this problem, we simply decompose  $f$  in two terms:

$$2\pi f(\lambda) = |1 - e^{-i\lambda}|^{-2d} g(0) + |1 - e^{-i\lambda}|^{-2d} \{g(\lambda) - g(0)\},$$

and compute each Fourier integral as a sum of two Fourier integrals corresponding to each terms. The Fourier integrals of the first term correspond to the autocovariance function of a fractionally integrated noise, which admits a close-form expression (Brockwell and Davis, 2009, Chap. 13):

$$\frac{1}{2\pi} \int_{-\pi}^{\pi} |1 - e^{-i\lambda}|^{-2d} e^{il\lambda} d\lambda = \begin{cases} \frac{\Gamma(l+d)\Gamma(1-d)}{\Gamma(l-d+1)\Gamma(d)} & \text{if } l \geq 1, \\ \frac{\Gamma(1-2d)}{\Gamma(1-d)^2} & \text{if } l = 0. \end{cases}$$

The Fourier integrals of the second term may be computed directly using the method of the previous section: assuming  $g$  is differentiable at 0, the second term vanishes at 0, since  $|1 - e^{-i\lambda}|^{-2d} \{g(\lambda) - g(0)\} \sim g'(0)\lambda^{1-2d}$ , with  $1 - 2d \geq 0$ .

### 3. APPROXIMATED LIKELIHOOD

**3.1. Principle.** For convenience, we consider only the standard likelihood function (1.3) in this section, and furthermore we assume that  $\mu = 0$ , i.e. a model with zero mean. The main idea of our approximation scheme is to replace  $\mathbf{T}(f)^{-1}$  by  $\mathbf{T}(1/4\pi^2 f)$  in the quadratic form, yielding

$$\tilde{p}(\mathbf{x}|f) = (2\pi)^{-n/2} |\mathbf{T}(f)|^{-1/2} \exp \left\{ -\frac{1}{2} \mathbf{x}^T \mathbf{T} \left( \frac{1}{4\pi^2 f} \right) \mathbf{x} \right\}.$$

It is easy to see that the quadratic form within the exponential may now be computed in  $O(n \log n)$  time. We return to this point and other implementation aspects in the next section.

The idea of approximating  $\mathbf{T}(f)^{-1}$  by  $\mathbf{T}(1/4\pi^2 f)$  is related to the asymptotic theory of Toeplitz matrices. It is in fact a common technical tool in the asymptotic theory of long-memory processes (Dahlhaus, 1989; Rousseau et al., 2012), but to the best of our knowledge has not been used for computational reasons before.

Now assume that  $f$  is parametrised in some way,  $f = f_{\theta}$ , and denote  $p(\mathbf{x}|f) = p(\mathbf{x}|\theta)$ ,  $\tilde{p}(\mathbf{x}|f) = \tilde{p}(\mathbf{x}|\theta)$ . As explained in the introduction, our strategy boils down to sample from the approximate posterior  $\pi_n(\theta|\mathbf{x}) \propto p(\theta)\tilde{p}(\mathbf{x}|\theta)$ , then to perform importance sampling from the approximate posterior to the true posterior, that is, to assign some weight to any simulation from the approximated posterior, with a weight function defined as:

$$w_{\text{Corr}}(\theta) \triangleq \frac{p(\mathbf{x}|\theta)}{\tilde{p}(\mathbf{x}|\theta)} = \exp \left\{ -\frac{1}{2} \mathbf{x}^T \left[ \mathbf{T}(f_{\theta})^{-1} - \mathbf{T}\left(\frac{1}{4\pi^2 f}\right) \right] \mathbf{x} \right\}.$$

The following theorem justifies this strategy.

**Theorem 1.** *Consider the FEXP model, as defined by (1.2), and let  $\pi_n$  denote the approximated posterior distribution defined as  $\pi_n(\theta) \propto p(\theta)\tilde{p}(\mathbf{x}|\theta)$ , where  $p(\theta)$  is some prior density with respect to parameter  $\theta$ , then, under certain conditions (given in the Appendix) on the prior distribution  $p(\theta)$ , and the true distribution of  $\mathbf{x}$ , with associated spectral density  $f_o$ , one has*

$$\mathbb{E}^{\pi_n} [w_{\text{Corr}}(\theta)] = w_{\text{Corr}}^0 \{1 + o_P(1)\}, \quad \mathbb{E}^{\pi_n} [w_{\text{Corr}}(\theta)^2] = (w_{\text{Corr}}^0)^2 \{1 + o_P(1)\}$$

where  $w_{\text{Corr}}^0 = p(\mathbf{x}|f_o)/\tilde{p}(\mathbf{x}|f_o)$ .

The proof and the technical conditions on the prior and the true spectral density are given in the Appendix. Because this theorem relies heavily on technical results of Rousseau et al. (2012), it is restricted to the semi-parametric FEXP model presented in the introduction, but with zero mean. We believe it could be extended to other classes of models with some extra effort.

In practical terms, this theorem says that the variance of the weights goes to zero as  $n$  goes to infinity, or, in other words, that the importance weights become nearly constant as  $n$  goes to infinity.

**3.2. Practical implementation.** In this section, we work out a practical implementation of the approximation scheme proposed in the previous section. We now turn our attention to the marginalised likelihood defined in (2.2).

First, we simplify the quadratic form by ignoring the uncertainty with respect to  $\mu$ , i.e. by taking  $m_{\mu} = \bar{\mathbf{x}}$ ,  $g_{\mu} = +\infty$ , so that the likelihood simplifies to

$$\tilde{p}(\mathbf{x}|\theta) \propto |\mathbf{T}(\bar{f}_{\theta})|^{-1/2} \left\{ b + \frac{1}{2} \tilde{\mathbf{x}}^T \mathbf{T}(\bar{f}_{\theta})^{-1} \tilde{\mathbf{x}} \right\}^{-a+n/2}, \quad \tilde{\mathbf{x}} = \mathbf{x} - \bar{\mathbf{x}}\mathbf{1}.$$

We observe that this particular approximation is quite accurate in practice, which relates to the fact, in long-memory scenarios,  $\bar{\mathbf{x}}$  is the standard estimator of  $\mu$ , and typically converges faster than other features (e.g.  $d$ ) of the model.

Second, as explained in the previous section, we replace the inverse of  $\mathbf{T}(\bar{f}_{\theta})$  by  $\mathbf{T}(4\pi^2/\bar{f}_{\theta})$  which leads to the following approximation of the quadratic form:

$$(3.1) \quad \tilde{\mathbf{x}}^T \mathbf{T}(\bar{f}_{\theta})^{-1} \tilde{\mathbf{x}} \approx \tilde{\mathbf{x}}^T \left( \frac{4\pi^2}{\bar{f}_{\theta}} \right) \tilde{\mathbf{x}} = \sum_{j=0}^{n-1} c_j(\tilde{\mathbf{x}}) \gamma_h(j)$$

where the coefficients  $c_k(\tilde{\mathbf{x}})$  may be pre-computed, once and for all, from the data (denoting  $\tilde{x}_i = x_i - \bar{x}$  the components of  $\tilde{\mathbf{x}}$ ):

$$c_j(\tilde{\mathbf{x}}) = \begin{cases} \sum_{i=1}^n \tilde{x}_i^2 & \text{if } j = 0 \\ 2 \sum_{i=1}^{n-k} \tilde{x}_i \tilde{x}_{i+j} & \text{if } j = 1, \dots, n-1 \end{cases}$$

and the  $\gamma_h(j)$ 's are the coefficients of the Toeplitz matrix  $\mathbf{T}(4\pi^2/\bar{f}_\theta)$ , that is, the Fourier integrals corresponding to function  $h = 1/(4\pi^2\bar{f})$ . Using the method described in Section 2.3, one obtains a  $O(n \log n)$  overall cost for evaluating the quadratic form.

Third, we approximate the determinant using a  $O(1)$  asymptotic approximation, as explained in the next section.

**3.3. Determinant approximation.** Approximations of determinants of the form  $|\mathbf{T}(\bar{f}_\theta)|$  typically rely on asymptotic expansions. For instance, Whittle's approximation consists in replacing  $\log |\mathbf{T}(\bar{f}_\theta)|/n$  by its limit,

$$\frac{1}{n} \log |\mathbf{T}(\bar{f}_\theta)| \rightarrow \int_{-\pi}^{\pi} \log \{2\pi \bar{f}_\theta(\lambda)\} d\lambda.$$

Chen et al. (2006) use more refined asymptotic results on Toeplitz matrices (e.g. Böttcher and Silbermann, 1999, p. 177) to obtain a more accurate approximation. Using their approach, one obtains in the FEXP case

$$\log |\mathbf{T}(\bar{f})| \approx D_n(\bar{f}_\theta) \triangleq d^2 \log n + \frac{1}{4} \sum_{j=1}^k j \xi_j^2 + d \sum_{j=1}^k j \xi_j + \log \frac{G(1-d)^2}{G(1-2d)},$$

where  $G$  is Barnes' function (Adamchik, 2001). This approach is easily adapted to other time series models, such as ARFIMA; we refer to Chen et al. (2006) for more details and possible extensions to other classes of models.

**3.4. Further approximation.** The cost of the approximation described in the previous section is that of a FFT operation, that is  $O(n \log n)$ . This cost may be further reduced by remarking that the approximation of the quadratic form may be rewritten as (see e.g. Palma, 2007, Chap. 4):

$$\tilde{\mathbf{x}}^T \mathbf{T}(\bar{f}_\theta)^{-1} \tilde{\mathbf{x}} \approx \tilde{\mathbf{x}}^T \mathbf{T}\left(\frac{1}{4\pi^2 \bar{f}_\theta}\right) \tilde{\mathbf{x}} = \frac{n}{2\pi} \int_{-\pi}^{\pi} \frac{I(\lambda)}{\bar{f}_\theta(\lambda)} d\lambda, \quad I(\lambda) = \left| \sum_{j=1}^n \tilde{x}_j e^{ij\lambda} \right|^2,$$

that is,  $I(\lambda)$  is the periodogram of the (centred) dataset  $\tilde{\mathbf{x}}$ .

It is relatively easier to evaluate  $I(\lambda)$  at the Fourier frequencies  $\lambda_j = 2\pi j/n$ , which suggests a further approximation, where this integral is replaced by a Riemann sum computed over the  $\lambda_j$ :

$$\frac{n}{2\pi} \int_{-\pi}^{\pi} \frac{I(\lambda)}{\bar{f}_\theta(\lambda)} d\lambda \approx \sum_{j=1}^n \frac{I(\lambda_j)}{\bar{f}_\theta(\lambda_j)}.$$

Computing simultaneously the  $I(\lambda_j)$ 's requires performing a FFT transform. The cost is  $O(n \log n)$ , but this needs to be done only once, for a given dataset. Then the approximate likelihood may be evaluated for many different values of  $\theta$ , at a  $O(n)$  cost. Since typically about  $10^4 - 10^6$  such evaluations are performed when Monte Carlo sampling from the approximate posterior, one may for practical purposes ignore the pre-computation time, and consider this further approximation as a  $O(n)$  operation.

To conclude, our final approximation takes the following form:

$$(3.2) \quad \tilde{p}(\mathbf{x}|\boldsymbol{\theta}) \propto D_n(\boldsymbol{\theta})^{-1/2} \left\{ b + \frac{1}{2} \sum_{j=1}^n \frac{I(\lambda_j)}{\bar{f}_{\boldsymbol{\theta}}(\lambda_j)} \right\}^{-a-n/2}.$$

This is very close in spirit to Whittle's approximation of the likelihood, which is based on the idea that the  $I(\lambda_j)$ 's are nearly independent, with variance  $\bar{f}_{\boldsymbol{\theta}}(\lambda_j)$ ; see again e.g. Palma (2007, Chap. 4). Rigorously speaking, we have not been able to establish that this further approximation is valid in the sense defined in Section 3.1, that is, that the variance of an importance sampling step from the approximated to the true likelihood converges to zero. We merely observe empirically that this approximate likelihood is nearly indistinguishable numerically to the more principled approximation developed in the previous Sections, that is:

$$\tilde{p}(\mathbf{x}|\boldsymbol{\theta}) \propto D_n(\boldsymbol{\theta})^{-1/2} \left\{ b + \frac{1}{2} \tilde{\mathbf{x}}' \mathbf{T} \left( \frac{1}{4\pi^2 \bar{f}_{\boldsymbol{\theta}}} \right) \tilde{\mathbf{x}} \right\}^{-a+n/2}.$$

On the other hand, we also observe that the speed improvement brought by this further approximation is rather modest, which is in line with the respective theoretical costs of  $O(n)$  and  $O(n \log n)$ . From now on, we do not distinguish between these two approximations. We only note that an additional advantage of (3.2) is that its gradient is easy to compute, which would make it possible to use Langevin-type MCMC moves.

#### 4. MONTE CARLO SAMPLING

**4.1. Background.** This section discusses Monte Carlo sampling from the approximate posterior, that is,

$$\tilde{\pi}_n(\boldsymbol{\theta}) \propto p(\boldsymbol{\theta}) \tilde{p}(\mathbf{x}|\boldsymbol{\theta})$$

where  $p(\boldsymbol{\theta})$  is some prior density defined with respect to parameter  $\boldsymbol{\theta}$ , and  $\tilde{p}(\mathbf{x}|\boldsymbol{\theta})$  is the approximate likelihood defined in Section 3. Although this discussion is, as before, not specific to the FEXP model, we shall assume for notational simplicity that the considered model is parametrised as follows:

$$\boldsymbol{\theta} = (k, \boldsymbol{\theta}_k) \in \cup_{i \in \mathbb{N}} \{i\} \times \mathbb{R}^{i+1}.$$

For instance, in the FEXP case, we have seen that  $\boldsymbol{\theta}_k = (\text{logit}(2d), \xi_1, \dots, \xi_k)$ . We shall also assume a nested structure for the  $\boldsymbol{\theta}'_k$ s, that is,  $\boldsymbol{\theta}_0 = (\theta_0)$ ,  $\boldsymbol{\theta}_{k+1}^T = (\boldsymbol{\theta}_k^T, \theta_{k+1})^T$ . Again, in the FEXP case,  $\theta_0 = \text{logit}(2d)$ ,  $\theta_j = \xi_j$  for  $j \geq 1$ . Finally, we denote  $p_k(\boldsymbol{\theta}_k) = p(\boldsymbol{\theta}_k|k)$ , the prior of  $\boldsymbol{\theta}_k$ , conditional of  $k$ ,  $p_{k+1}(\theta_{k+1}|\boldsymbol{\theta}_k)$ , the prior of component  $\theta_{k+1}$ , conditional on the first  $k$  components of  $\boldsymbol{\theta}_{k+1}$  being equal to those of vector  $\boldsymbol{\theta}_k$ , and  $\tilde{p}(\mathbf{x}|k, \boldsymbol{\theta}_k)$  the approximate likelihood  $\tilde{p}(\mathbf{x}|\boldsymbol{\theta})$  for  $\boldsymbol{\theta} = (k, \boldsymbol{\theta}_k)$ .

A common approach to this problem would be to alternate the steps described below as Algorithms 1 and 2, that is a random walk Metropolis step with respect to  $\boldsymbol{\theta}_k$ , conditional on  $k$ , and a birth-and-death step, that is, a particular instance of Green (1995)'s reversible jump sampler, which attempts at either incrementing (birth) or decrementing (death)  $k$ . In case a birth is proposed, the current vector  $\boldsymbol{\theta}_k$  is completed with a draw from the conditional prior  $p_{k+1}(\theta_{k+1}|\boldsymbol{\theta}_k)$ . In case a death step is proposed, component  $\theta_k$  is deleted from the current vector  $\boldsymbol{\theta}_k$ .

Of course, many variants of this basic MCMC strategy could be considered. However, such variants are unlikely to directly address the following two limitations

**Algorithm 1** Gaussian random walk Metropolis step (conditional on  $k$ )Input:  $\theta = (k, \theta_k)$ Output:  $\theta' = (k, \theta'_k)$ 

1. Sample  $\theta_k^* \sim N(\theta_k, \Sigma_k)$ .
2. With probability  $1 \wedge r$ , take  $\theta'_k = \theta_k^*$ , otherwise  $\theta'_k = \theta_k$ , with

$$r = \frac{p_k(\theta_k^*)\tilde{p}(\mathbf{x}|k, \theta_k^*)}{p_k(\theta_k)\tilde{p}(\mathbf{x}|k, \theta_k)}.$$

**Algorithm 2** Birth and death stepInput:  $\theta = (k, \theta_k)$ Output:  $\theta' = (k', \theta'_{k'})$ Constants:  $\rho_{k \rightarrow k+1} = \rho_{k \rightarrow k-1} = 1/2$  for  $k \geq 1$ ,  $\rho_{0 \rightarrow 1} = 1$ ,  $\rho_{0 \rightarrow -1} = 0$ .

1. Let  $\theta^* = (k^*, \theta_{k^*}^*)$ , where, with probability  $\rho_{k \rightarrow k+1}$ ,  $k^* = k + 1$ ,  $\theta_{k^*}^* = (\theta_k, \theta_{k+1})$ , and  $\theta_{k+1} \sim p_{k+1}(\xi_{k+1}|\theta_k)$  (birth step); otherwise, with probability  $\rho_{k \rightarrow k-1} = 1 - \rho_{k \rightarrow k+1}$ ,  $k^* = k - 1$ ,  $\theta_{k^*}^* = (\theta_0, \dots, \theta_{k-1})$  (death step).
2. Set  $\theta' = \theta^*$  with probability  $1 \wedge r$ , otherwise  $\theta' = \theta$ , with

$$r = \frac{\rho_{k^* \rightarrow k} p(k^*) \tilde{p}(\mathbf{x}|\theta^*)}{\rho_{k \rightarrow k^*} p(k) \tilde{p}(\mathbf{x}|\theta)}.$$

of standard MCMC: (a) calibration; and (b) local behaviour. Calibration refers to the difficulty of choosing certain tuning parameters, such as that the scales  $\Sigma_k$  in the random walk step. It is well known that the choice of such tuning parameters may have a critical impact on the performance of the algorithm. We shall see in our numerical study, Section (5), that tuning manually the  $\Sigma_k$ 's is rather tedious. A more principled approach would be to use some form of adaptive MCMC sampling, see Andrieu and Thoms (2008) for a review, where one uses past samples to iteratively adapt the tuning parameters to the target distribution. However, adaptive MCMC is less straightforward to use in trans-dimensional settings (as an infinite number of tuning parameters  $\Sigma_k$  must be learnt), and also does not address the second difficulty, which we now comment upon.

The phrase “local behaviour” refers to the fact that MCMC samplers have difficulties exploring multi-modal posteriors, as they tend to get trapped in the attraction of meaningless modes. One may wonder if multi-modality is an issue for the classes of models considered in this paper.

In our simulations, we observed that the FEXP model may indeed generate multimodal posteriors, and we propose the following explanation. Figure 4.1 plots a spectral density of FEXP type, see (2.1), with  $d = 0.4$ ,  $k = 3$ , and  $(\xi_1, \xi_2, \xi_3) = (1, -1, 1)$ . Overlaid are two other FEXP spectral densities, but with  $k = 10$  in both cases,  $d = 0.2$  (dotted line),  $d = 0$  (dashed line), and coefficients  $\xi_j$  fit by least-squares on the Fourier frequencies  $\pi j/100$ ,  $j = 1, \dots, 100$ . These three densities are hard to distinguish except maybe in the close vicinity of 0. This indicates that, unless the sample size is very large, the posterior distribution may be multi-modal, with modes that may correspond to values of  $d$  which are very far from the true value, while  $k$  is set to a larger value so as to introduce additional components to compensate the bias in  $d$ .

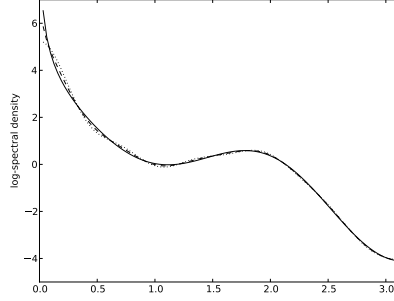


FIGURE 4.1. Graphical illustration of a source of multimodality for the FEXP model: FEXP log spectral densities, with  $k = 3$ ,  $d = 0.4$ ,  $(\xi_1, \xi_2, \xi_3) = (1, -1, 1)$  (solid line), compared to  $k = 10$ ,  $d = 0$  (dashed line), and  $d = 0.2$  (dotted line), and coefficients  $\xi_j$  adjusted by least-squares estimation.

---

**Algorithm 3** Generic SMC sampler

---

All operations are for  $j = 1, \dots, N$ .

- (1) Sample  $\theta^j \sim \eta_0(\theta)$ .
- (2) For  $t = 1, \dots, T$ , do:
  - (a) Reweighting step: Compute and normalise weights

$$w_t(\theta^j) \propto \frac{\eta_t(\theta^j)}{\eta_{t-1}(\theta^j)}, \quad W_t^j = \frac{w_t(\theta^j)}{\sum_{l=1}^N w_t(\theta^l)}.$$

- (b) Resampling step: sample  $\hat{\theta}^j$  from the multinomial distribution that assigns probability  $W_t^l$  to value  $\theta_t^l$ ,  $l = 1, \dots, N$ .
  - (c) Move step: regenerate  $\theta^j$  as

$$\theta^j \sim K_t(\hat{\theta}^j, d\theta)$$

where  $K_t$  is a Markovian kernel that leaves  $\eta_t$  invariant.

---

**4.2. Sequential Monte Carlo.** Algorithm 3 describes a generic SMC sampler. Compared to the more general framework of Del Moral et al. (2006), this algorithm is specialised to the case where the sequence of distributions  $(\eta_t)$  is defined on a common sampling space,  $\Theta$ , and also to a specific choice of the backward kernel, see Del Moral et al. (2006) for more details. The algorithm depends on the specification of a sequence of distributions  $(\eta_t)$  and a sequence of Markovian kernels  $(K_t)$ . The former must be such that it is easy to sample from  $\eta_0$ , and that  $\eta_T = \tilde{\pi}_n$ , the target distribution (the approximate posterior, as defined in the previous section). The latter must be such that  $K_t$  leaves invariant  $\eta_t$ , and is typically a MCMC kernel.

For  $(\eta_t)$ , we take a particular annealing sequence (Neal, 2001), that is, a geometric bridge between the prior and the approximate posterior

$$(4.1) \quad \eta_t(\theta) \propto p(\theta) \{\tilde{p}(\mathbf{x}|\theta)\}^{\gamma_t}$$

with  $\gamma_0 = 0 < \dots < \gamma_T = 1$ . The weight function is then:

$$w_t(\theta) = \tilde{p}(\mathbf{x}|\theta)^{\alpha_t}, \quad \alpha_t = (\gamma_t - \gamma_{t-1}).$$

We follow Jasra et al. (2011), Schäfer and Chopin (2011), and adjust dynamically the annealing coefficients  $\gamma_t$  by solving at iteration  $t$ , with respect to variable  $\alpha_t$ , the equation

$$\frac{\left(\sum_{j=1}^N w_t(\boldsymbol{\theta}^j)\right)^2}{\sum_{j=1}^N w_t(\boldsymbol{\theta}^j)^2} = cN$$

for some fixed  $c \in (0, 1)$ ; in our simulations, we took the default value  $c = 1/2$ , and we used Brent's root-find algorithm (Press et al., 2007, Section 9.3) to solve numerically this equation. The left hand side is the effective sample size of Kong et al. (1994); it takes values in  $[1, N]$ , and is a convenient measure of the weight degeneracy.

For  $(K_t)$ , we use  $M$  steps of the MCMC sampler described in the previous section, that is, we repeat  $M$  times the following sequence: Algorithm 1 (random walk Metropolis), then Algorithm 2 (birth-and-death). A big advantage of the SMC framework is that we can use the current particle system to calibrate these MCMC steps. Specifically, before each move step we set  $\boldsymbol{\Sigma}_l = \tau_l \mathbf{S}_l$ ,  $\tau_l = 2.38^2/(l+1)$ , where  $\mathbf{S}_l$  is the covariance matrix of those resampled particles  $\hat{\boldsymbol{\theta}}^{(j)} = (\hat{k}^{(j)}, \boldsymbol{\theta}_{\hat{k}^{(j)}}^{(j)})$  such that  $\hat{k}^{(j)} = l$ ; in case this set is empty, we take instead  $\mathbf{S}_l = \mathbf{I}_l$ . This particular choice is motivated by the theory on optimal scaling of Hastings-Metropolis kernel, as developed in Roberts and Rosenthal (2001).

The only tuning parameters of this algorithm are the number of particles  $N$ , and the number of MCMC steps performed at each move step,  $M$ . In practice, taking  $M$  in the range 5 – 20 seems to lead to reasonable performance, but of course, it may be adjusted more precisely by for instance looking at the acceptance rate of the MCMC steps in a preliminary run; see Section 5.

Again, one may propose many variants of this SMC sampler. For instance, one could also use adaptive reversible jump steps, where the jump probabilities  $\rho_{k \rightarrow k+1}$ ,  $\rho_{k \rightarrow k-1}$  (currently set to 1/2) could be optimised using the particle sample, or even design more elaborate proposals based on a family of linear transforms, as in Green (2003), where the corresponding matrices may also be learnt from the particle sample. (We did some experiments with this strategy, but did not obtain significantly better results in the examples we looked at.) However, our main focus here is to develop a black box algorithm, which requires as little input from the user as possible, while giving reasonable performance, even in the presence of multimodality. Our numerical experiments seem to indicate this is the case, see Section 5.

## 5. NUMERICAL EXPERIMENTS

**5.1. Settings.** The first part of this numerical study focus on the following simulated example. We sampled a long time series from an ARFIMA(1,  $d$ , 1) model, with  $d = 0.45$ , and AR (resp. MA) coefficient  $-0.9$  (resp.  $-0.2$ ), and applied the FEXP model described in the Introduction. We find this example to be challenging, because the spectral density of the simulated process has a particular shape which is not easily approximated by a FEXP spectral density, at least for small values of  $k$ ; see e.g. Fig. 5.4 and additional comments around this Figure. (From a modelling perspective, the presence of an autoregressive root close to one makes

it difficult to determine whether the persistence in the data is really characteristic of long range behaviour.) Thus, and as described more properly in Rousseau et al. (2012), the FEXP model described in the introduction must be understood as a semi-parametric procedure, which makes it possible to consistently estimate the true spectral density (and related quantities, e.g. the long-range coefficient  $d$ ), provided this spectral density belongs to some infinite-dimensional class of functions (of a certain regularity). In particular, one expects the posterior of  $k$  to shift towards infinity as  $n$  goes to infinity, which adds to the computational difficulty. The second part of the study considers a real dataset, see Section 5.6.

Except in the consistency study (Section 5.4), the dataset consists of the  $n = 10^4$  first values of the simulated time-series. In the consistency study, different sample sizes are compared, by again taking the  $n$  first points for different values of  $n$ . Except in the prior sensitivity section (Section 5.5), we always take the following prior: independently,  $d \sim \text{Uniform}[0, 1/2]$ ,  $k \sim \text{Geometric}(1/5)$  (with support  $\{0, 1, 2, \dots\}$ ),  $\xi_j \sim N(0, 100j^{-2\beta})$ ,  $\mu|\sigma \sim N(0, \sigma^2/g_\mu)$ ,  $1/\sigma^2 \sim \text{Gamma}(a, b)$ , with hyper-parameters  $\beta = 1$ ,  $a = b = 0.5$ ,  $g_\mu = 0.1$ . Note that this prior falls slightly outside the class of prior distributions that ensures consistency, according to Rousseau et al. (2012). In fact, we found it interesting to see whether the conditions to ensure consistency of Rousseau et al. (2012) may be too strong; our numerical study seems to indicate it may indeed be the case.

All the simulations presented in this Section were performed on a standard 3GHz desktop computer, without resort to any form of parallelisation, and implemented in the Python language (using the NumPy and SciPy libraries); the programs are available on the first author's web-page.

**5.2. Performance of SMC.** This section is concerned with the performance of the SMC sampler, regarding sampling the approximate posterior. Note that the final correction step (i.e. the importance sampling step from the approximate to the true posterior) is therefore not performed; see Section 5.4 for an evaluation of the performance of the correction step.

We run the SMC sampler 10 times, for  $N = 1000$  particles, and  $M = 5$  MCMC steps; CPU time is 20 minutes per run. Variability over these ten runs is small for posterior expectations of the  $\xi_j$ 's (not shown) and  $d$  (see top left plot of Figure 5.1), but is a bit larger for the posterior distribution of  $k$  (see same plot). The right plot of Figure 5.2 reveals that the acceptance rate of the birth and death step is often below 10%, which suggests  $M = 5$  steps may not be enough to successfully move the particles with respect to component  $k$ . In that plot, the  $x$ -axis gives the annealing coefficient; that is  $\gamma_t$  at iteration  $t$ , which goes from  $\gamma_0 = 0$  to  $\gamma_T = 1$ , see (4.1). Note that this axis is a non-linear function of elapsed CPU time, as the sequence  $\gamma_t$  grows slowly at the beginning, and progressively accelerates, whereas the cost per iteration  $t$  should be roughly constant.

We run again the SMC sampler with  $N = 1000$  particles, but this time with  $M = 20$  steps; CPU time is then 1 hour 20 minutes. We obtain very satisfactory results; see the box-plots in Figure 5.1. In particular, regarding the posterior expectation  $\hat{d}$  (with respect to the approximated posterior) of component  $d$ , we observe that the empirical variance of the estimates of  $\hat{d}$  obtained from the ten runs is very close to  $\text{Var}^{\pi_n}(d)$ , that is, the variance would be obtained from  $N = 1000$  independent and identically distributed simulations from the posterior. Figure 5.1 also reports superimposed weighted histograms obtained from the second set of runs, and a

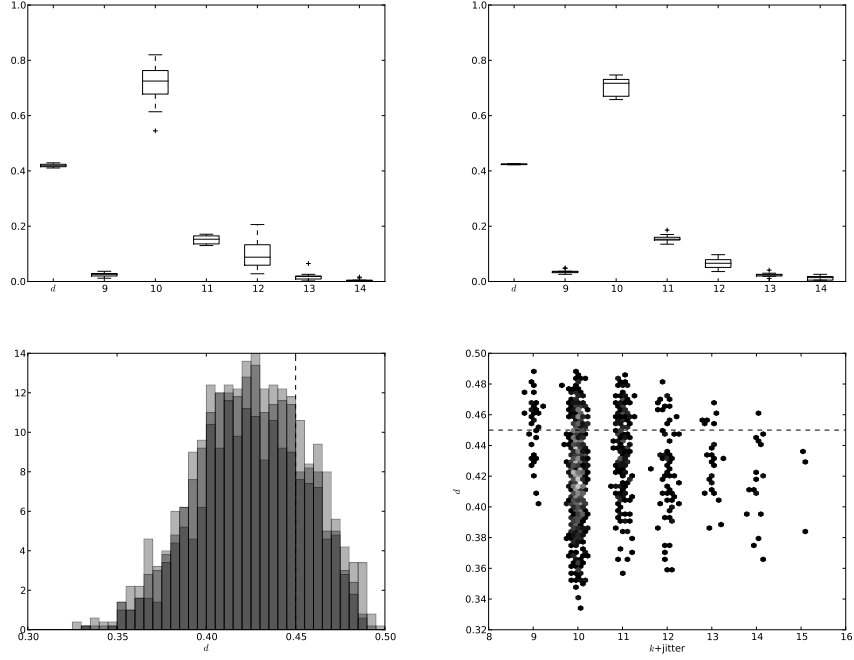


FIGURE 5.1. Top row: box-plots, over 10 repeated runs, of the SMC estimates of the posterior expectation of  $d$  and the posterior probabilities of  $k = 9, \dots, 14$  (Left:  $(N, M) = (10^3, 5)$ ; Right:  $(N, M) = (10^3, 20)$ ); Bottom Left: marginal posterior of  $d$ , as estimated by superimposed (with transparency effects) weighted histograms obtained by 3 first SMC runs; Bottom Right: hexagonal binning plot of  $k$  versus  $d$ , from the first SMC run; a  $N(0, 0.1^2)$  jitter is added to  $k$ , and colour is proportional to the sum of the weights of the particles falling in each hexagon.

particle approximation of the bivariate posterior distribution of  $(k, d)$ . (Some noise was added to  $k$  for the sake of visualisation.)

**5.3. Performance of MCMC.** This section compares the performance of our SMC sampler with standard MCMC. As the previous section, this comparison is in terms of sampling from the approximate posterior, and no correction step is performed at the end of the algorithm. The MCMC sampler we use is the one described in Section 4.1.

We simulate several MCMC chains with different tuning parameters, but always of size  $N = 1.6 \times 10^5$ , which gives a running time per chain of about 50 minutes (as compared to 20 minutes for the first sequence of SMC runs).

We start by 10 runs, with  $\Sigma_k$ , the scale of the random walk in the  $\theta_k$ -space, set to  $\tau \mathbf{I}_{k+1}$  ( $\tau$  times the identity matrix of rank  $k+1$ ); after some pilot runs, we take  $\tau = 0.015$  so as to obtain an acceptance rate for the random walk step close to 25%.

One chain seems to alternate between two local modes, one around  $(k, d) = (20, 10^{-3})$ , the other around  $(k, d) = (30, 10^{-3})$ ; see first row of Figure 5.3. The

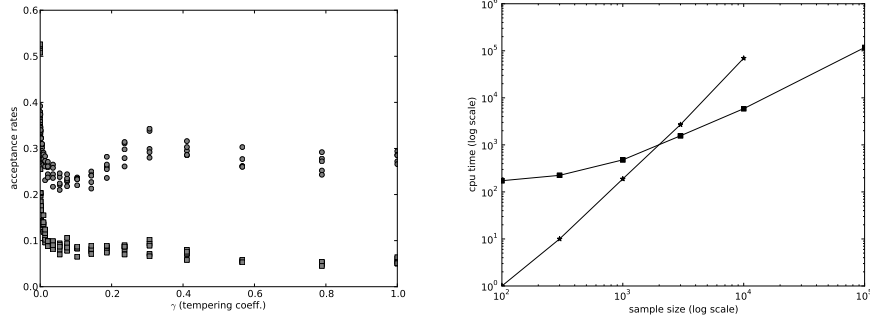


FIGURE 5.2. Left: acceptance rates of the random walk step (circles) and the birth and death step (squares) versus the tempering coefficient  $\gamma_t$ , which progresses from 0 to 1 during the algorithm ( $N = 1000$ ,  $M = 5$ ). Right: CPU times in seconds for the SMC sampler targeting the approximate posterior (squares) and the correction step (stars) versus the sample size of the data.

nine other chains do not get trapped in these local modes; however their mixing properties seem quite poor; see the MCMC traces for  $d$  the middle row plot. Note how the darkest trace never seems to visit the region  $d > 0.45$ , which accounts for 20% of the posterior mass, according to our SMC runs.

These results are hardly satisfactory. We also experimented with  $\Sigma_k$  set to  $\tau$  times the prior covariance matrix of  $\theta_k$ , but this did not give better results (not shown). Finally, we implemented a (crudely) adaptive version of our random walk Metropolis sampler, for the posterior conditional on  $k = 10$  (no reversible jump steps), where, during a burn-in period of  $5 \times 10^4$  iterations, the covariance matrix of the random walk proposal is re-calibrated every  $10^3$  iterations, using past samples; specifically  $\Sigma_k$  is set to  $(2.38^2/(k+1))\mathbf{S}_t$  (see Roberts and Rosenthal, 2001), where  $\mathbf{S}_t$  is the empirical covariance matrix computed on the  $t$  first iterations. One sees that learning the correlation structure of the posterior (conditional on  $k$ ) significantly improve the mixing of the chain, although the auto-correlation function (with respect to  $d$ , and computed post burn-in) decays slowly, see bottom row of 5.3. However, learning the posterior covariance matrix, conditional on  $k$ , for each  $k$  in a given range, would of course be a bit wasteful.

**5.4. Consistency study.** In this section, we show how both the approximate and the true posterior evolve as the sample size grows. The goal is to assess both the asymptotic properties of our FEXP semi-parametric procedure, and the effect of the correction step on inference. Results are obtained from a single SMC run.

Figures 5.4 and 5.5 summarise posterior inference for the sample sizes, from top to bottom of Figure 5.4,  $n = 100$ , 300, 1000, 3000, and, from top to bottom of Figure 5.5,  $n = 10^4$  and  $n = 10^5$ . On the left (resp. right) side, the marginal posterior distributions of  $d$  (resp. the 80% confidence bands for the spectral density) are compared; light grey is the approximate posterior, and dark grey is the true posterior. (Transparency effects are used.) One sees that the effect of the correction step becomes unnoticeable very quickly as far as the estimation of the spectral density (away from 0) is concerned. The effect on the marginal posterior of  $d$  takes

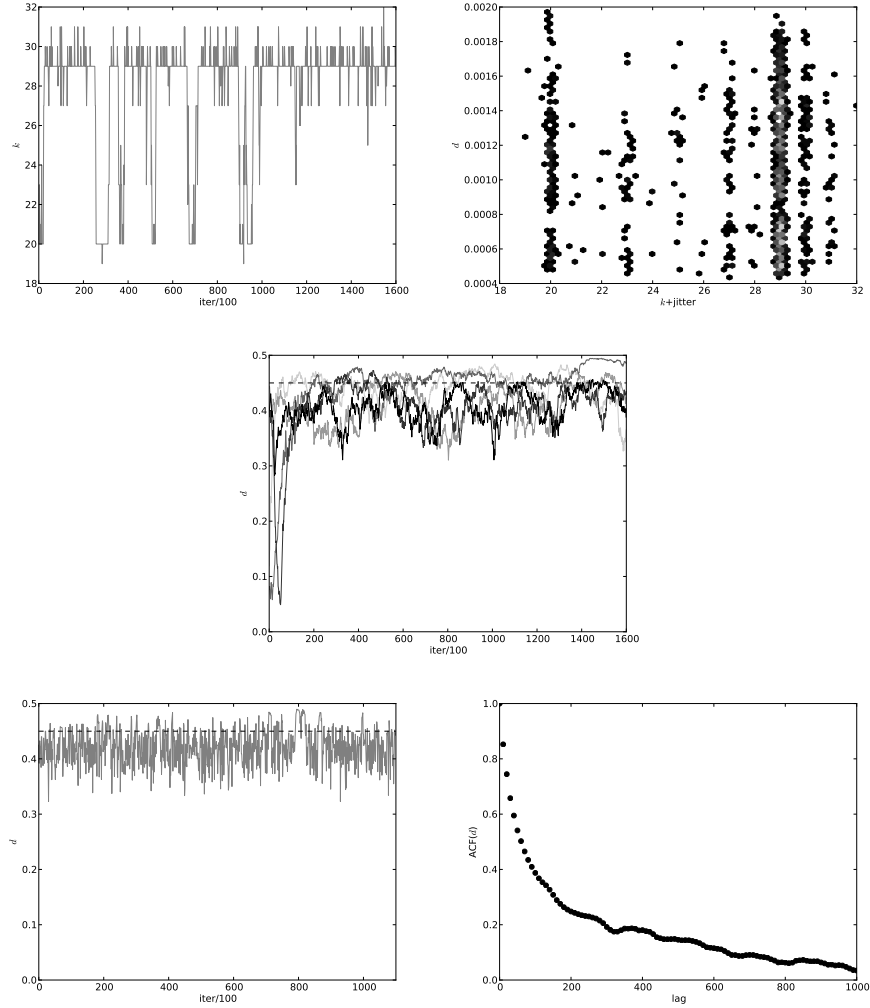


FIGURE 5.3. Top row: one MCMC run alternating between two minor modes (Left: trace of  $k$ , Right: scatter-plot of  $(k, d)$ , with a jitter term added to  $k$  as in Fig. ). Middle row: traces of  $d$  for five out of the nine remaining MCMC chains. Bottom row: adaptive MCMC; trace of  $d$  (left) and ACF plot (right), computed post burn-in.

a bit more time to vanish, but is already negligible for  $n$  around 3000; in fact, the effective sample size of the correction step is already above 900 (out of 1000 particles) for  $n = 3000$ .

The right panel of Figure 5.2 compares the CPU cost of the SMC sampler and the correction step. One sees that, unless  $n$  is of order of  $10^4$  or more, the overall cost of the complete procedure is dominated by that of the SMC sampler. For  $n \gg 10^4$ , the correction step becomes quickly too expensive, because of the  $O(n^3)$  cost, and was not performed for  $n = 10^5$ . (Which is why the dark grey plots are

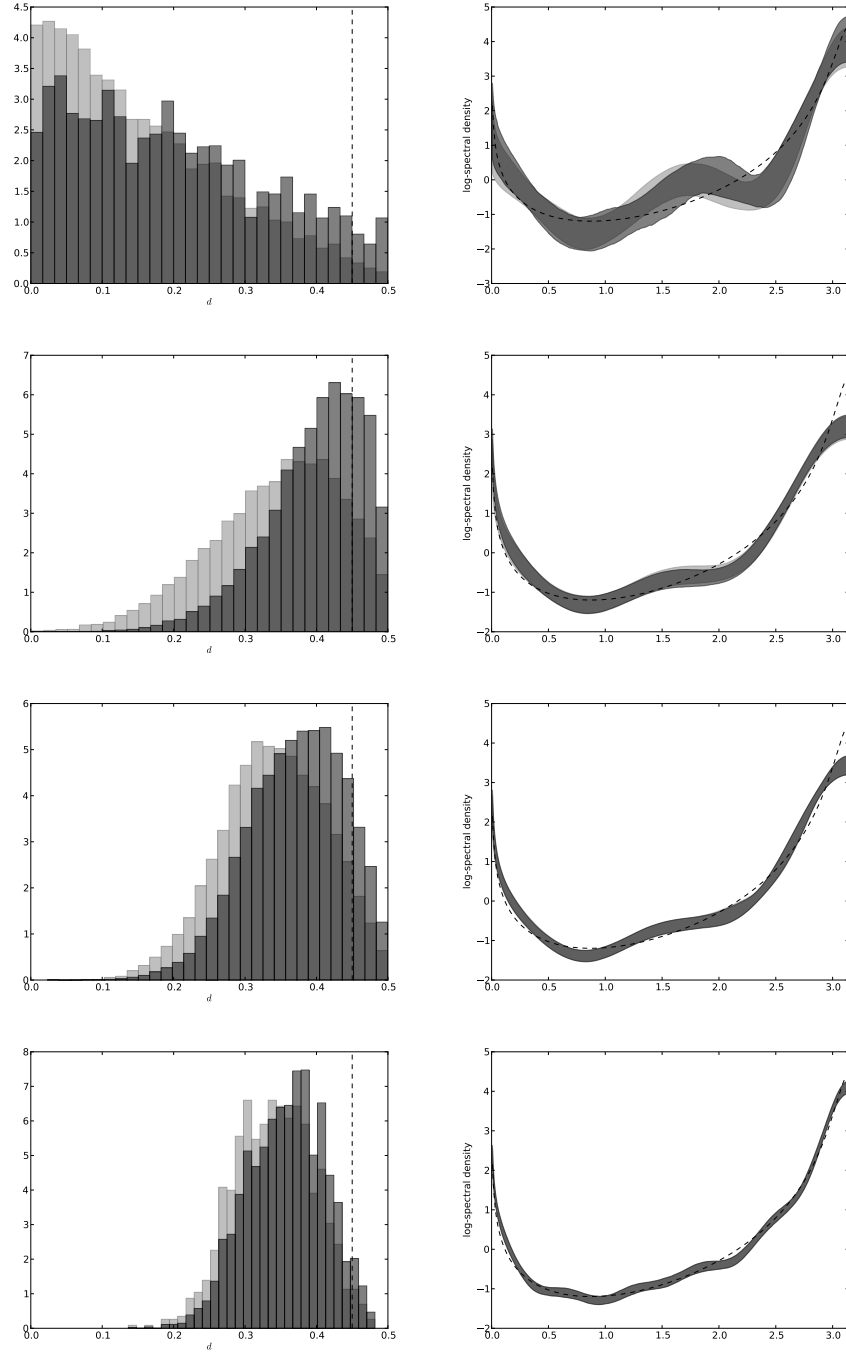


FIGURE 5.4. Consistency study for the ARFIMA data, Left: marginal posterior of  $d$ , approximate (light grey histogram), or exact (dark grey histogram); Right: 80% confidence bands for the log-spectral density (same color code), true spectral density (dashed line). From top to bottom, sample size is 100, 300, 1000, 3000.

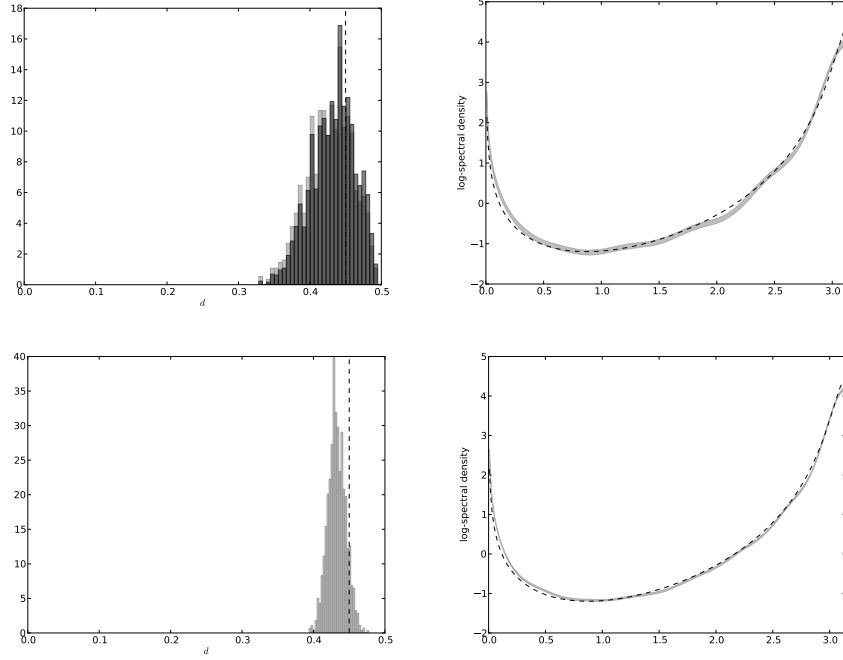


FIGURE 5.5. Consistency study for the ARFIMA data, same caption as Fig. 5.4, sample size is  $10^4$  (top) and  $10^5$  (bottom).

missing in the bottom row of Figure 5.5.) Fortunately, one sees that for  $n \gg 10^4$  the correction step seems to be superfluous.

**5.5. Prior sensitivity analysis.** We have seen in Section 5.1 that the prior for the FEXP coefficients  $\xi_j$  was set to  $N(0, 100j^{-2\beta})$ , with  $\beta = 1$ ; this hyper-parameter is related to assumptions on the smoothness of the true spectral density (the larger  $\beta$ , the smoother the true spectral density); see Rousseau et al. (2012). In this section, we consider briefly the effect of  $\beta$  on posterior inference.

Figure 5.6 gives the same types of posterior plots as in the previous sections, i.e. 80% confidence bands for the spectral density, and marginal posterior for the vector  $(k, d)$ , for  $\beta = 0$  (Left side) and  $\beta = 2$  (right side). This must be compared with the top right plot of Figure 5.5 and the bottom right plot of 5.1, for which  $\beta = 1$ . One sees that the choice of  $\beta$  has a strong impact on the posterior marginal distribution of  $k$ , but a rather moderate impact on either  $d$  or the spectral density itself, except maybe on the right edge of the spectral density plots (for  $\lambda$  close to  $\pi$ ). Since  $k$  is essentially a nuisance parameter, one sees that the choice of  $\beta$  does not seem to be too critical for inferential purposes. On the other hand, it is interesting to note the impact of  $\beta$  on the computational difficulty to explore the posterior. We observe that, for  $\beta = 2$ , it is even more difficult for a MCMC sampler to escape local modes such as those shown in the top row of Figure 5.3 (corresponding MCMC traces not shown).

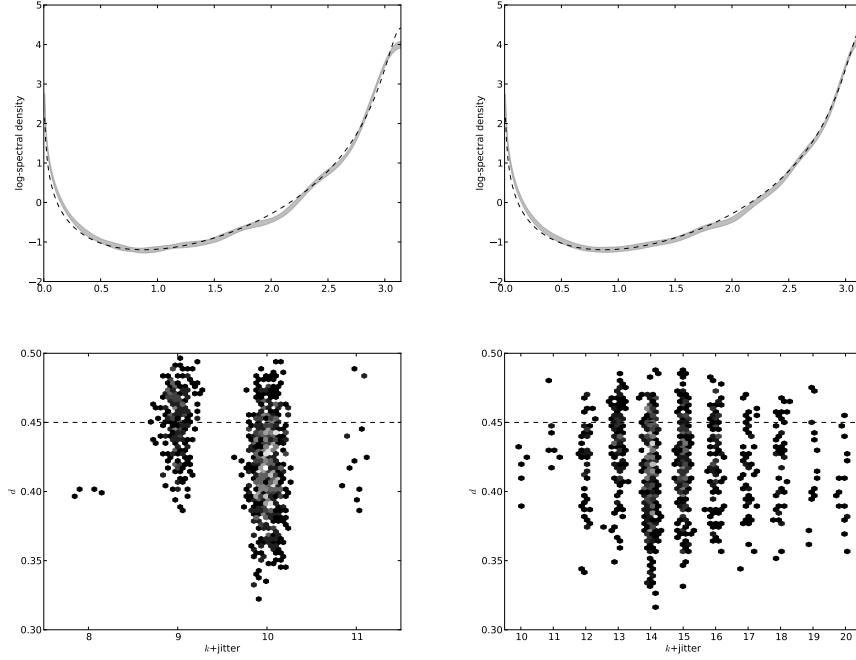


FIGURE 5.6. Prior sensitivity analysis: 80% posterior confidence bands for log-spectral density (top), hexagonal bin plot of  $(k + \text{jitter}, d)$ , where jitter is  $N(0, 0.1^2)$  noise (bottom), for two different priors:  $\beta = 0$ , i.e.  $\xi_j \sim N(0, 100)$  (left); and  $\beta = 2$ , i.e.  $\xi_j \sim N(0, 100j^{-4})$  (right).

**5.6. Real data study: Ethernet Traffic.** In a preliminary study, we applied our methodology to the popular Nile data set, but found the example not to be very challenging, either from a computational point of view ( $n = 663$ ), or from an inferential point of view (an ARFIMA(0,  $d$ , 0) model, or in other words, a fractional Gaussian noise model, seems to fit the data well). More details may be obtained from the first author.

In this paper, we consider instead the Ethernet Traffic dataset of Leland et al. (1994), which can be found in the `longmemo` R package. This is a time-series of length  $n = 4000$ , which records the number of packets passing through a particular network per time unit. (For convenience, we divided the data values by 1000, in order to use the same prior as for the simulated dataset.) The right side of Fig. 5.7 plots the empirical spectrum of this time series. Although the empirical spectrum is an asymptotically biased estimator of the true spectral density under long-range dependence, the bias is typically small for frequencies sufficiently far away from 0 (e.g. Moulines and Soulier, 2003). Thus comparing the empirical spectrum with a given estimator of the spectral density provides at least some guidance on the performance of the said estimator.

Interestingly, we find these data to be quite challenging for frequentist parametric procedures. The `FEXPest` command of the `longmemo` R package, which computes

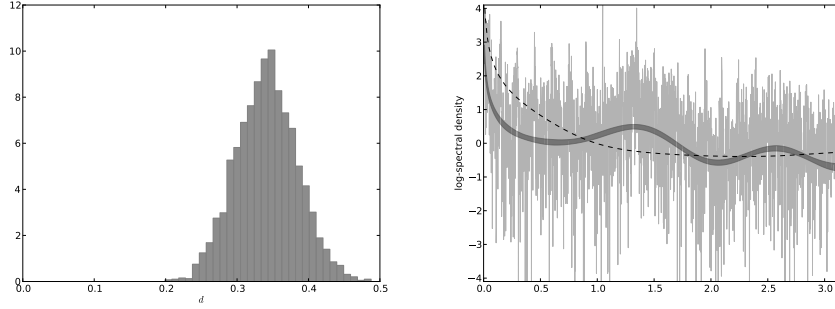


FIGURE 5.7. Ethernet data, Left: marginal posterior distribution of  $d$  from the FEXP semi-parametric model (as represented by a weighted histogram obtained from the SMC sampler); Right: Bayesian 80% confidence band for the spectral density (dark grey), spectral density corresponding to ML estimation of an ARFIMA model, with orders  $p = q = 3$  (dashed line).

Beran (1993)'s estimator for a parametric FEXP model, returns the estimate  $\hat{d} = 0.52$ , although  $d \in (0, 1/2)$ . The `fracdiff` command from the `fracdiff` package, which does maximum likelihood estimation for ARFIMA( $p, q$ ) models, returns  $\hat{d} = 0.3$  for  $p = q = 3$  but the estimated spectral density, plotted as a dashed line in the right panel of Figure 5.7 does not seem to fit the spectrum of the data even for higher frequencies. (According to the same procedure, the orders  $p = q = 3$  give the smallest AIC among ARFIMA( $p, p$ ) models, with  $p \leq 10$ ).

In contrast, our Bayesian semi-parametric procedure seems to fit the data rather well; see again the right side of Figure 5.7. We find strong evidence in favour of long-range dependence, as evidenced by the marginal posterior distribution of  $d$  in the left panel of Figure 5.7.

From a computational point of view, we mention briefly that we observe the same things as in the previous section; that is, that the correction step has a negligible impact on the results, and that the output SMC sampler shows little variability when run several times.

#### ACKNOWLEDGEMENTS

N. Chopin and J. Rousseau are supported by the ANR grant “Bandhit” of the French Ministry of research.

#### APPENDIX: PROOF OF THEOREM 1

The proof borrows several intermediate results from Rousseau et al. (2012). Recall that the model is  $\mathbf{x} \sim N(0, \mathbf{T}(f_\theta))$ , with  $f_\theta$  specified as in (1.2). We assume that the prior  $p(\theta)$  decomposes as  $p(d)p(k)p_k(\xi_k)$ , where the support of  $p(d)$  is  $[0, 1/2 - t]$  for some  $t > 0$ ,  $p(k)$  is a geometric distribution, and, given  $k$ ,  $\xi_j(j+1)^{\alpha+1/2} \sim \mathcal{N}(0, 1)$  independently except from the fact that it is truncated on

$$\Xi_k(\alpha, L) = \{\xi_k \in \mathbb{R}^{k+1}; \sum_{j=0}^k (j+1)^{2\alpha+1} \xi_j^2 \leq L\}$$

where  $L$  is a large positive constant.

Let  $P_o^n$  be the true distribution of  $\mathbf{x}$ , that is, a Gaussian distribution  $N(0 \times \mathbf{1}, \mathbf{T}(f_o))$ , where the true spectral density  $f_o$  admits a FEXP representation, as in (1.2), with parameters  $d_o > 0$ ,  $k_o = +\infty$ , and  $\xi_{0,o}, \xi_{1,o}, \dots$  belonging to the Sobolev ball of radius  $L$ ,  $\Xi(\beta, L) = \{(\xi_0, \xi_1, \dots); \sum_{j=0}^{\infty} (j+1)^{2\beta+1} \xi_j^2 \leq L\}$ . Throughout the proof we denote by  $|A|_{Fr} = \text{tr}[A^T A]^{1/2}$  the Froebinius norm of a matrix  $A$  and by  $\|A\|^2 = \sup\{\mathbf{x}^T A^T A \mathbf{x}; \|\mathbf{x}\| = 1\}$  its euclidian norm.

We now prove that, provided  $0 \leq d_o \leq 1/2 - t$ ,  $\beta \geq \alpha > 3/2$ , and for  $L_o$  small enough (compared to  $L$ ), one has:

$$\mathbb{E}^{\pi_n} [w^2(\boldsymbol{\theta})] = w_o^2 \{1 + o_p(1)\}, \quad w_o^2 = o_p(e^{u_n(\log n)^{3/2}})$$

for any sequence  $(u_n)$  such that  $u_n \rightarrow +\infty$ , and where  $w$  (resp.  $w_o$ ) is used as a short-hand for  $w_{\text{Corr}}$  (resp.  $w_{\text{Corr}}^o$ ) in the rest of the proof. One may obtain the same type of results for  $\mathbb{E}^{\pi_n} [w(\boldsymbol{\theta})]$  using the same calculations.

One has:

$$\begin{aligned} \mathbb{E}^{\pi_n} [w^2(\boldsymbol{\theta})] &= \frac{\int w^2(\boldsymbol{\theta}) \exp \left\{ -\frac{1}{2} \mathbf{x}^T \mathbf{T}(1/4\pi^2 f_{\boldsymbol{\theta}}) \mathbf{x} - l_{n,o} \right\} |\mathbf{T}(f_{\boldsymbol{\theta}})|^{-1/2} p(\boldsymbol{\theta}) d\boldsymbol{\theta}}{\int \exp \left\{ -\frac{1}{2} \mathbf{x}^T \mathbf{T}(1/4\pi^2 f_{\boldsymbol{\theta}}) \mathbf{x} - l_{n,o} \right\} |\mathbf{T}(f_{\boldsymbol{\theta}})|^{-1/2} p(\boldsymbol{\theta}) d\boldsymbol{\theta}} \\ &\triangleq \frac{N_n}{D_n} \end{aligned}$$

where  $l_{n,o} = -\{\mathbf{x}^T \mathbf{T}(f_o)^{-1} \mathbf{x} + \log \det \mathbf{T}(f_o)\} / 2$ .

Let  $\epsilon_n^2 = (n/\log n)^{-\beta/(2\beta+1)}$ . The idea of the proof is to first show that

$$(5.1) \quad \mathbb{E}^{\pi_n} \left[ \frac{w^2(\boldsymbol{\theta})}{w_o^2} \right] = \mathbb{E}^{\pi_n} \left[ \frac{w^2(\boldsymbol{\theta})}{w_o^2} \mathbb{1}_{S_{n,1}}(\boldsymbol{\theta}) \mathbb{1}_{k \leq k_{n,1}} \right] + o_p(1)$$

where  $S_{n,1} = \{\boldsymbol{\theta}; l(f_o, f_{\boldsymbol{\theta}}) \leq \epsilon_n\}$ ,  $k_{n,1} = k_o \epsilon_n^{-1/\alpha}$  for some  $k_o > 0$ , and  $l(f, f')$  is the  $L^2$  distance between spectral log-densities,  $l(f, f') = \int_{-\pi}^{\pi} \{\log(f/f')\}^2$ . Let  $S_{n,1,k} = \{\boldsymbol{\theta}_k; \boldsymbol{\theta} = (k, \boldsymbol{\theta}_k) \in S_{n,1}\}$ . In a second step, we show that

$$(5.2) \quad \sup_{\boldsymbol{\theta} \in S_{n,1,k}} |\log w(\boldsymbol{\theta}) - \log w_o| = o_p(1)$$

uniformly on  $k \leq k_{n,1}$ , which would then conclude the proof of Theorem 1.

We now prove (5.1). As proven in Rousseau et al. (2012) and Kruijer and Rousseau (2011) (for the true posterior, but the proof for the approximate posterior  $\pi_n$  follows the same lines), there exists  $c > 0$  such that

$$(5.3) \quad P_o^n \left[ D_n \geq e^{-c n \epsilon_n^2} \right] = o(1).$$

Indeed, let  $k_n = \lfloor k_0 (n/\log n)^{1/(2\beta+1)} \rfloor$  with  $k_0 > 0$  and

$$\mathcal{B}_n = \{\boldsymbol{\theta}; k = k_n, d_o \leq d \leq d_o + n^{-a}, |\xi_j - \xi_{j,o}| \leq n^{-a}, j = 1, \dots, k_n\},$$

then it is easy to see that, for some  $c_1 > 0$ ,

$$\mathbb{P}^p(\boldsymbol{\theta}) [\mathcal{B}_n] \geq e^{-c_1 k_n \log n},$$

where  $\mathbb{P}^p(\boldsymbol{\theta})$  denotes the prior probability. As in Rousseau et al. (2012) and Kruijer and Rousseau (2011), let  $\Omega_n(\boldsymbol{\theta}) = \{\mathbf{x}; \tilde{l}_n(\boldsymbol{\theta}) - l_{n,o} \geq -n\epsilon_n^2\}$ , with

$$\tilde{l}_n(\boldsymbol{\theta}) = -\frac{1}{2} \mathbf{x}^T \mathbf{T}(1/4\pi^2 f_{\boldsymbol{\theta}}) \mathbf{x} - \frac{1}{2} \log \det [\mathbf{T}(f_{\boldsymbol{\theta}})]$$

then we have that for all  $\boldsymbol{\theta} \in \mathcal{B}_n$ ,

$$(5.4) \quad P_o^n [\Omega_n(\boldsymbol{\theta})^c] = o(1)$$

if  $a$  is chosen large enough. Indeed let  $A(\boldsymbol{\theta}) = \mathbf{T}^{1/2}(f_o)\mathbf{T}(1/4\pi^2 f_{\boldsymbol{\theta}})\mathbf{T}^{1/2}(f_o) - \mathbf{I}_n$  and  $B(\boldsymbol{\theta}) = \mathbf{T}^{1/2}(f_o)\mathbf{T}(f_{\boldsymbol{\theta}})^{-1}\mathbf{T}^{1/2}(f_o) - \mathbf{I}_n$ . Using Lemma 2.4 in the supplement of Kruijer and Rousseau (2011), for all  $\boldsymbol{\theta} \in \mathcal{B}_n$ , we have that  $\text{tr} [\mathbf{T}^{-1/2}(f_o)\mathbf{T}(f_{\boldsymbol{\theta}})\mathbf{T}^{-1/2}(f_o)] \leq 2n$  for  $n$  large enough so that all eigenvalues of  $\mathbf{T}^{1/2}(f_o)\mathbf{T}(f_{\boldsymbol{\theta}})^{-1}\mathbf{T}^{1/2}(f_o)$  are bounded from below by  $1/2$  and

$$\log |\mathbf{I}_n + B(\boldsymbol{\theta})| = \text{tr} [B(\boldsymbol{\theta})] - \frac{1}{2} \text{tr} \left[ ((\mathbf{I}_n + \tau B(\boldsymbol{\theta}))^{-1} B(\boldsymbol{\theta}))^2 \right] \geq \text{tr} [B(\boldsymbol{\theta})] - 2 \text{tr} [B(\boldsymbol{\theta})^2].$$

Moreover using Lemma 2.4 in the supplement of Kruijer and Rousseau (2011), we have that  $\text{tr} [B(\boldsymbol{\theta})] = \text{tr} [A(\boldsymbol{\theta})] + O(n^\epsilon)$  for all  $\epsilon > 0$  so that, since for  $n$  large enough and  $\boldsymbol{\theta} \in \mathcal{B}_n$ ,  $2 \text{tr} [B(\boldsymbol{\theta})^2] + n^\epsilon \leq n\epsilon_n^2$ ,

$$P_o^n [\Omega_n(\boldsymbol{\theta})^c] \leq \mathbb{P}_{\mathbf{z} \sim \mathcal{N}(0, \mathbf{I}_n)} [\mathbf{z}^T A(\boldsymbol{\theta}) \mathbf{z} - \text{tr} [A(\boldsymbol{\theta})] \geq n\epsilon_n^2].$$

One also has:

$$|A(\boldsymbol{\theta})|_{Fr}^2 = \frac{n}{2\pi} \int_{-\pi}^{\pi} \left( \frac{f_o(\lambda)}{f_{\boldsymbol{\theta}}(\lambda)} - 1 \right) d\lambda + \text{Error},$$

where the Error is controlled by Lemma 2.6 in the supplement of Kruijer and Rousseau (2011), since we have

$$|A(\boldsymbol{\theta})|_{Fr}^2 = \text{tr} \left[ (\mathbf{T}(f_{\boldsymbol{\theta}}^{-1}) \mathbf{T}(f_o - f_{\boldsymbol{\theta}}))^2 \right]$$

and  $f_o = f_{\boldsymbol{\theta}} b$  where (taking  $\xi_j = 0$  for  $j > k$ )

$$b(\lambda) = (2 - 2 \cos \lambda)^{d-d_o} \exp \left\{ \sum_{j=0}^{\infty} (\xi_{o,j} - \xi_j) \cos j\lambda \right\}$$

and we note that, for  $\boldsymbol{\theta} \in \mathcal{B}_n$ , the function  $\exp \left\{ \sum_{j=0}^k \xi_j \cos j\lambda \right\}$  is Lipschitz with constant  $O(k_n^{(3/2-\beta)+}) = O(1)$ , for  $a$  is large enough (see Lemma 3.1 in the supplement of Kruijer and Rousseau (2011)). Therefore we obtain

$$\text{Error} = O(\log n \|b - 1\|_{\infty}) = O(\log n),$$

and  $|A(\boldsymbol{\theta})|_{Fr}^2 = O(n^{1-a} + \log n) = o(n^\tau n\epsilon_n^2)$ , for some  $\tau > 0$ , which combined with Lemma 1.3 in the supplement of Kruijer and Rousseau (2011), leads to, for some  $c > 0$ ,

$$P_o^n [\Omega_n(\boldsymbol{\theta})^c] \leq e^{-cn^\tau}.$$

We now turn to the second part of the proof. Let  $S_{n,j,k} = \{\boldsymbol{\theta}_k; l(f_o, f_{(k, \boldsymbol{\theta}_k)}) \in ((j-1)\epsilon_n, j\epsilon_n)\}$ , for  $j = 1, \dots, J_n$  where  $J_n = O(\epsilon_n^{-1})$ , and let

$$\begin{aligned} N_{n,j,k} &= w_o \int_{S_{n,j,k}} w(\boldsymbol{\theta})^2 \exp \left\{ \tilde{l}_n(\boldsymbol{\theta}) - l_{o,n} \right\} p(\boldsymbol{\theta}) d\boldsymbol{\theta} \\ &= w_o \int_{S_{n,j,k}} \frac{\exp \{l_n(\boldsymbol{\theta}) - l_{o,n}\}}{w(\boldsymbol{\theta})} p(\boldsymbol{\theta}) d\boldsymbol{\theta}. \end{aligned}$$

We prove that uniformly in  $J_0 \leq j \leq J_n$  and  $k \in \mathbb{N}$ ,  $N_{n,j,k} = o_p(e^{-cn\epsilon_n^2})$ , for some  $J_0 > 0$ . To do so we bound

$$P_o^n \left[ \sup_{\boldsymbol{\theta} \in S_{n,j,k}} |\log(w_o/w(\boldsymbol{\theta}))| > v_{n,j,k} \right]$$

for  $j \geq 2$  and for a properly chosen  $v_{n,j,k}$ . Let  $k_{n,j} = k_0(j\epsilon_n)^{-1/\alpha}$ ,  $J_{n,1} = J_1 n^{-\alpha\epsilon} (\log n)^\alpha \epsilon_n^{-1}$ , with  $J_1 > 0$ . Define for  $C > 0$  (large enough) and  $\tau > 0$  (small enough)

(5.5)

$$\begin{aligned} v_{n,j,k} &= k^{3/2} n^\epsilon j \epsilon_n, \quad \text{if } k \leq k_{n,j}, \quad j \geq 1 \\ v_{n,j,k} &= (C \log n (k \wedge n^{(1+\epsilon)/(\alpha-1/2)}) k^{-1/(\alpha-1/2)} + k^\tau) k_{n,j}^{-\alpha+1/2}, \quad \text{if } k > k_{n,j}. \end{aligned}$$

Note that if  $\epsilon$  is small,  $Ck \log n = o(n^\epsilon)$  for some  $k \geq k_{n,j}$  only if  $(j\epsilon_n)^{-1/\alpha} \lesssim n^\epsilon (\log n)^{-1}$ , i.e. if  $j \gtrsim n^{-\alpha\epsilon} (\log n)^\alpha \epsilon_n^{-1} := J_{n,1}$ . We write  $-\log w(\boldsymbol{\theta}) + \log(\xi_o) = \mathbf{z}' A(\boldsymbol{\theta}) \mathbf{z} / 2$  with  $\mathbf{z} = \mathbf{T}^{-1/2}(f_o) \mathbf{x} \sim \mathcal{N}(0, I_n)$ , and, under  $P_o$ , and we bound successively

$$\sup_{\boldsymbol{\theta} \in S_{n,j,k}} (\mathbf{z}' [A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)] \mathbf{z} + \text{tr}[A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)])$$

and

$$\sup_d |\mathbf{z}' [A_o - A(\boldsymbol{\theta}_d)] \mathbf{z} + \text{tr}[A_o - A(\bar{\boldsymbol{\theta}}_d)]|$$

with  $\bar{\boldsymbol{\theta}}_d = (d, k, \bar{\boldsymbol{\xi}}_{d,k})$  and  $\bar{\boldsymbol{\xi}}_{d,k} = \arg\min_{\boldsymbol{\xi} \in \mathbb{R}^{k+1}} l(f_o, f_{d,k}, \boldsymbol{\xi})$ , see Kruijer and Rousseau (2011). We now study

$$\sup_{\boldsymbol{\theta} \in S_{n,j,k}} (\mathbf{z}' [A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)] \mathbf{z} + \text{tr}[A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)]).$$

Note that  $n^{-1} |A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)|_{F_r}^2 = n^{-1} \text{tr} [(A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d))^2]$  converges towards 0, so we only need control the approximation error, which we split into the approximation error of  $\text{tr} [(\mathbf{T}(f_o) \mathbf{T}(f_{\boldsymbol{\theta}}^{-1} - f_{\bar{\boldsymbol{\theta}}_d}^{-1}))^2]$  and of  $\text{tr} [(\mathbf{T}(f_o) (\mathbf{T}(f_{\boldsymbol{\theta}})^{-1} - \mathbf{T}(f_{\bar{\boldsymbol{\theta}}_d})^{-1}))^2]$ . We now consider the first term. Note that

$$f_{\boldsymbol{\theta}}^{-1} - f_{\bar{\boldsymbol{\theta}}_d}^{-1} = f_{\bar{\boldsymbol{\theta}}_d}^{-1} (\exp[\sum_{j=0}^k (\xi_j - (\bar{\xi}_{d,k})_j) \cos(j\lambda)] - 1) := f_{\bar{\boldsymbol{\theta}}_d}^{-1} b_{\boldsymbol{\theta}}(\lambda)$$

where

$$\sup_{\lambda \in [-\pi, \pi]} |b_{\boldsymbol{\theta}}(\lambda)| \leq \sum_{j=0}^k |\xi_j - (\bar{\xi}_{d,k})_j| \leq \sqrt{k} \|\boldsymbol{\xi} - \bar{\boldsymbol{\xi}}_{d,k}\| \lesssim \sqrt{k} j \epsilon_n.$$

Note that if  $k \geq k_{n,j}$  then  $k^{1/2-\alpha} \leq \sqrt{k} j \epsilon_n$  and we bound instead

$$\sup_x |b_{\boldsymbol{\theta}}(x)| \leq \sum_{j=0}^k |\xi_j - (\bar{\xi}_{d,k})_j| \leq \sqrt{k_{n,j}} \|\boldsymbol{\xi} - \bar{\boldsymbol{\xi}}_{d,k}\| + k_{n,j}^{-\alpha+1/2}.$$

Lemma 2.1 of the supplement of Kruijer and Rousseau (2011) implies that the approximation error of the first term is bounded by  $O((\sum_{j=1}^k |\xi_j - (\bar{\xi}_{d,k})_j| n^\epsilon)^2) = O((n^{\epsilon'} \sqrt{k} j \epsilon_n)^2)$  for all  $\epsilon' > 0$  and all  $k \leq k_{n,j}$  and is bounded by  $O((k_{n,j}^{-\alpha+1/2} n^{\epsilon'})^2)$  for all  $k \geq k_{n,j}$  and all  $\epsilon' > 0$ . From Lemma 2.4 of the supplement of Kruijer and

Rousseau (2011), the same bounds apply to the second term. Finally we obtain that

$$\begin{aligned} |A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)|_{Fr} &= O(n^\epsilon \sqrt{k} j \epsilon_n) \quad \text{if } k \leq k_{n,j} \\ |A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)|_{Fr} &= O(n^\epsilon k_{n,j}^{-\alpha+1/2}) \quad \text{if } k > k_{n,j}. \end{aligned}$$

This implies that for all  $\xi \in S_{n,j,k}$ :

- If  $k \leq k_{n,j}$ , since  $v_{n,j,k}(\sqrt{k} j \epsilon_n)^{-1} = k n^\epsilon$ ,

$$\mathbb{P}(\mathbf{z}^T [A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)] \mathbf{z} + \text{tr}[A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)] > v_{n,j,k}) \leq e^{-ckn^\epsilon},$$

- If  $k > k_{n,j}$ , since  $v_{n,j,k} k_{n,j}^{\alpha-1/2} \geq n^\epsilon$ ,

$$Pr(\mathbf{z}^T [A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)] \mathbf{z} + \text{tr}[A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)] > v_{n,j,k}) \leq e^{-cv_{n,j,k} k_{n,j}^{\alpha-1/2}}.$$

Moreover a Taylor expansion of  $A(\xi')$  around  $A(\boldsymbol{\theta})$  implies that  $\forall \delta > 0$  and for all  $\xi, \xi'$

$$|\mathbf{z}^T [A(\boldsymbol{\theta}) - A(\boldsymbol{\theta}')] \mathbf{z} + \text{tr}[A(\boldsymbol{\theta}) - A(\boldsymbol{\theta}')]| \lesssim (\mathbf{z}^T \mathbf{z} + n) n^{2(d-d_o)+\delta} \left[ \sum_{j=0}^k |\xi_j - \xi'_j| + |d - d'| \right]. \quad (5.6)$$

Without loss of generality we can restrict ourselves on  $\Omega_n = \{\mathbf{z}^T \mathbf{z} \leq 2n\}$ , since  $\mathbb{P}[\Omega_n^c] = o(e^{-cn})$  for some positive  $c$ .

- If  $k \leq k_{n,j}$  and  $j \leq J_{n,1}$ , if  $\|\boldsymbol{\xi} - \boldsymbol{\xi}'\| \leq n^{-1-\epsilon} k^{-1/2}$  and  $|d - d'| \leq n^{-1-\epsilon}$ ,

$$|\mathbf{z}' [A(\boldsymbol{\theta}) - A(\boldsymbol{\theta}')] \mathbf{z} + \text{tr}[A(\boldsymbol{\theta}) - A(\boldsymbol{\theta}')]| = o(1), \quad \text{uniformly.}$$

If  $k \leq k_{n,j}$  and  $j \geq J_{n,1}$ , then

$$|\mathbf{z}' [A(\boldsymbol{\theta}) - A(\boldsymbol{\theta}')] \mathbf{z} + \text{tr}[A(\boldsymbol{\theta}) - A(\boldsymbol{\theta}')]| = o(n \epsilon_n^2 j^2), \quad \text{uniformly.}$$

Let  $E_{n,j,k}$  be the covering number of  $S_{n,j,k}$  by balls satisfying the above constraint then

$$E_{n,j,k} \leq e^{C' k \log n}$$

- If  $k > k_{n,j}$ . First if  $k > k_2 n \epsilon_n^2$ , for some  $k_2 > 0$  possibly large, then uniformly over  $\sum_{j=1}^k |\xi_j - \xi'_j| \leq n^{-1-\epsilon} k$  and  $|d - d'| \leq n^{-1-\epsilon}$ ,

$$|\mathbf{z}^T [A(\boldsymbol{\theta}) - A(\boldsymbol{\theta}')] \mathbf{z} + \text{tr}[A(\boldsymbol{\theta}) - A(\boldsymbol{\theta}')]| = o(k)$$

The covering number of  $S_{n,j,k}$  by the above constraints depends on  $k$ . Define  $K_n(k)$  such that  $K_n(k)^{-(\alpha-1/2)} = n^{-1-\epsilon} k$ , i.e.  $K_n(k) = n^{(1+\epsilon)/(\alpha-1/2)} k^{-1/(\alpha-1/2)}$ . Then

$$E_{n,j,k} \leq \exp(C \log n (k \wedge K_n(k)))$$

Now if  $k < k_2 n \epsilon_n^2$ , then uniformly over  $\sum_{j=1}^k |\xi_j - \xi'_j| \leq n^{-1-\epsilon} (n j^2 \epsilon_n^2)$  and  $|d - d'| \leq n^{-1-\epsilon}$ ,

$$|\mathbf{z}^T [A(\boldsymbol{\theta}) - A(\boldsymbol{\theta}')] \mathbf{z} + \text{tr}[A(\boldsymbol{\theta}) - A(\boldsymbol{\theta}')]| = o(n j^2 \epsilon_n^2)$$

and

$$E_{n,j,k} \leq \exp(C k \log n).$$

A simple chaining argument in  $S_{n,j,k}$  implies that for all  $j \leq J_{n,1}$  and all  $k \leq k_{n,j}$ , there exists  $M > 0$  such that

(5.7)

$$\mathbb{P} \left[ \Omega_n \cap \left\{ \sup_{\boldsymbol{\theta} \in S_{n,j,k}} (\mathbf{z}^T [A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)] \mathbf{z} + \text{tr}[A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)]) > v_{n,j,k} + M \right\} \right] \leq e^{-ckn^\epsilon}$$

and if  $j \geq J_{n,1}$ , for any  $\delta > 0$  and  $n$  large enough

$$(5.8) \quad \mathbb{P} \left[ \Omega_n \cap \left\{ \sup_{\boldsymbol{\theta} \in S_{n,j,k}} (\mathbf{z}^T [A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)] \mathbf{z} + \text{tr}[A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)]) > v_{n,j,k} + \delta n j^2 \epsilon_n^2 \right\} \right] \leq e^{-ckn^\epsilon}.$$

Note also that for all  $k \leq k_{n,j}$ ,  $v_{n,j,k} = k^{3/2} n^\epsilon j \epsilon_n \leq n^\epsilon (j \epsilon_n)^{1-3/(2\alpha)}$  so that whenever  $\alpha > 3/2$  and  $j \leq J_{n,1}$ ,  $v_{n,j,k} = o(1)$  and if  $j \geq J_{n,1}$ ,  $v_{n,j,k} = o(n j^2 \epsilon_n^2)$ .

If  $k > k_{n,j}$  and  $k > k_2 n \epsilon_n^2$ , for all  $j$ , and all  $\delta > 0$

$$(5.9) \quad \mathbb{P} \left[ \sup_{\boldsymbol{\theta} \in S_{n,j,k}} (\mathbf{z}^T [A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)] \mathbf{z} + \text{tr}[A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)]) > v_{n,j,k} + \delta k \right] \leq e^{-ck^\tau},$$

If  $k < k_2 n \epsilon_n^2$ , for all  $\delta > 0$  and  $n$  large enough,

$$(5.10) \quad \mathbb{P} \left[ \Omega_n \cap \left\{ \sup_{\boldsymbol{\theta} \in S_{n,j,k}} (\mathbf{z}^T [A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)] \mathbf{z} + \text{tr}[A(\boldsymbol{\theta}) - A(\bar{\boldsymbol{\theta}}_d)]) > v_{n,j,k} + \delta n j^2 \epsilon_n^2 \right\} \right] \leq e^{-c(k \log n + n^\epsilon)},$$

and for all  $k > k_{n,j}$ ,  $v_{n,j,k} = o(k + n j^2 \epsilon_n^2)$ . We now study

$$\sup_d |\mathbf{z}^T [A(\boldsymbol{\theta}_o) - A(\bar{\boldsymbol{\theta}}_d)] \mathbf{z} + \text{tr}[A(\boldsymbol{\theta}_o) - A(\bar{\boldsymbol{\theta}}_d)]|$$

We have

$$|A(\boldsymbol{\theta}_o) - A(\bar{\boldsymbol{\theta}}_d)|_{Fr}^2 = 0 + \text{Error},$$

where Error is the approximation error of the trace by its limiting integral. We bound separately the approximation error of  $\text{tr}[(\mathbf{T}(f_o)\mathbf{T}(f_o^{-1} - f_{\bar{\boldsymbol{\theta}}_d}^{-1}))^2]$  and of  $\text{tr}[(\mathbf{T}(f_o)(\mathbf{T}(f_o)^{-1} - \mathbf{T}(f_{\bar{\boldsymbol{\theta}}_d})^{-1})^2)]$ . We have, see Kruijer and Rousseau (2011),

$$f_{\bar{\boldsymbol{\theta}}_d} = f_o e^{(d-d_o)H_k - \Delta_{d_o,k}}, \quad H_k(x) = \sum_{j>k} \eta_j \cos(jx), \quad \eta_j = 2/j \quad \Delta_{d_o,k} = \sum_{j>k} \theta_{o,j} \cos(jx)$$

so that

$$f_{\bar{\boldsymbol{\theta}}_d} - f_o = f_o((d-d_o)H_k - \Delta_{d_o,k}) + O((|(d-d_o)H_k - \Delta_{d_o,k}|^2 x^{-2|d-d_o|})$$

Therefore, from Lemma 2.6 in Kruijer and Rousseau (2011) in the supplement,

$$(5.11) \quad \begin{aligned} & \text{error} \left( \text{tr} \left[ (\mathbf{T}(f_o)\mathbf{T}(f_o^{-1} - f_{\bar{\boldsymbol{\theta}}_d}^{-1}))^2 \right] \right) \\ &= O((|d-d_o| + \|\Delta_{d_o,k}\|_\infty)^2 \log n + (|d-d_o| + \|\Delta_{d_o,k}\|_\infty) l(f_o, f_{\bar{\boldsymbol{\theta}}_d})^{1/2} (\log n)^3). \end{aligned}$$

Similarly using Lemma 2.4 in the supplement of Kruijer and Rousseau (2011)

$$\text{error}(\text{tr}[(\mathbf{T}(f_o)(\mathbf{T}(f_o)^{-1} - \mathbf{T}(f_{\bar{\boldsymbol{\theta}}_d})^{-1})^2]) = O((|d-d_o| + \|\Delta_{d_o,k}\|_\infty)^2 n^\epsilon)$$

for all  $\epsilon > 0$ . Let  $(k, d)$  be such that  $\bar{\boldsymbol{\theta}}_d \in S_{n,j,k}$ , then

$$(5.12) \quad |A(\xi_o) - A(\bar{\boldsymbol{\theta}}_d)|_{Fr} = O((|d-d_o| + \|\Delta_{d_o,k}\|_\infty) n^\epsilon) = O(n^\epsilon (j \epsilon_n)^{(2\alpha-1)/(2\alpha)}), \quad \forall \epsilon > 0$$

see Lemma 3.1 in Kruijer and Rousseau (2011). We also have that for all  $|d-d'| \leq n^{-2}$ ,

$$|\mathbf{z}^T [A(\bar{\boldsymbol{\theta}}'_d) - A(\bar{\boldsymbol{\theta}}_d)] \mathbf{z} - \text{tr}[A(\bar{\boldsymbol{\theta}}'_d) - A(\bar{\boldsymbol{\theta}}_d)]| \leq (\mathbf{z}^T \mathbf{z} + n) n^{-1} = O_p(1)$$

uniformly, so that a simple chaining argument combined with (5.12) and Lemma 1.3 in the supplement of Kruijer and Rousseau (2011) implies that, for all  $\epsilon > 0$  (5.13)

$$\mathbb{P} \left( \sup_d |\mathbf{z}^T [A(\bar{\boldsymbol{\theta}}_d) - A(\boldsymbol{\theta}_o)] \mathbf{z} - \text{tr}[A(\bar{\boldsymbol{\theta}}_d) - A(\boldsymbol{\theta}_o)]| > n^\epsilon (j\epsilon_n)^{(2\alpha-1)/(2\alpha)} \right) \leq e^{-n^{\epsilon/2}}$$

Hence, combining (5.7), (5.9), (5.10) and (5.13) together with the fact that there exists  $J_0, C_0 > 0$  such that for all  $j \geq J_0$ , setting  $S_{n,j} = \{\boldsymbol{\theta}; (j-1)\epsilon_n \leq l(\boldsymbol{\theta}_o, \boldsymbol{\theta}) \leq j\epsilon_n\}$ ,

$$\int_{S_{n,j}} \exp \{l_n(\boldsymbol{\theta}) - l_n(\boldsymbol{\theta}_o)\} d\pi(\boldsymbol{\theta}) \leq e^{-C_0 j n \epsilon_n^2},$$

on a set having probability going to 1. Thus for all  $\delta > 0$ , there exists a set with  $P_o$  probability going to 1 such that

$$\begin{aligned} & \sum_{j=J_0}^{J_n} \sum_{k \in \mathbb{N}} p(k) N_{n,j,k} \\ & \leq \sum_{j=1}^{J_n} \sum_{k=1}^{k_{n,1}} \sup_{(d, \bar{\boldsymbol{\xi}}_{d,k}) \in S_{n,j,k}} \frac{w(\boldsymbol{\theta}_o)}{w(\bar{\boldsymbol{\theta}}_d)} \sup_{(d, \boldsymbol{\theta}) \in S_{n,j,k}} \frac{w(\bar{\boldsymbol{\theta}}_d)}{w(\boldsymbol{\theta})} p(k) \int_{S_{n,j,k}} e^{l_n(\boldsymbol{\theta}) - l_n(\boldsymbol{\theta}_o)} dp(\boldsymbol{\theta}) \\ & \leq \sum_{j=J_0}^{J_n} e^{n^\epsilon (j\epsilon_n)^{(2\alpha-1)/(2\alpha)}} \sum_{k=1}^{k_{n,1}} e^{\delta(nj^2\epsilon_n^2 + k)} p(k) \int_{S_{n,j,k}} e^{l_n(\boldsymbol{\theta}) - l_n(\boldsymbol{\theta}_o)} dp(\boldsymbol{\theta}) \\ & \leq e^{-2cn\epsilon_n^2} \end{aligned}$$

Finally note that if there exists  $\xi = (d, k, \boldsymbol{\theta}) \in \cup_{j < J_0} S_{n,j,k}$ , then  $\bar{\boldsymbol{\theta}}_{d,k} \in \cup_{j < J_0} S_{n,j,k}$  by definition of  $\bar{\boldsymbol{\theta}}_{d,k}$ . Replacing  $\epsilon_n$  by  $J_0\epsilon_n$ , with an abuse of notations, we write  $S_{n,1,k} := \cup_{j \leq J_0} S_{n,j,k}$  for all  $k$  and we split the set  $k$  into  $k \leq k_{n,1}$  and  $k > k_{n,1}$ .

$$\begin{aligned} & \mathbb{E}^{\pi_n} \left[ w^2(\boldsymbol{\theta}) \sum_{k=1}^{k_{n,1}} \mathbb{1}_{S_{n,1,k}} \right] \\ & \leq w_o^2 \sum_{k=1}^{k_{n,1}} \sup_{(d, \bar{\boldsymbol{\xi}}_{d,k}) \in S_{n,1,k}} \frac{w^2(\bar{\boldsymbol{\theta}}_d)}{w_o^2} \sup_{(d, \boldsymbol{\xi}) \in S_{n,1,k}} \frac{w^2(\boldsymbol{\theta})}{w^2(\bar{\boldsymbol{\theta}}_d)} p(k|x) \pi_n(S_{n,1,k}|k), \end{aligned}$$

and

$$\begin{aligned} & \mathbb{E}^{\pi_n} \left[ w^2(\boldsymbol{\theta}) \sum_{k=1}^{k_{n,1}} \mathbb{1}_{S_{n,1,k}} \right] \\ & \geq w_o^2 \sum_{k=1}^{k_{n,1}} \inf_{(d, \bar{\boldsymbol{\xi}}_{d,k}) \in S_{n,1,k}} \frac{w^2(\bar{\boldsymbol{\theta}}_d)}{w_o^2} \inf_{(d, \boldsymbol{\theta}) \in S_{n,1,k}} \frac{w^2(\boldsymbol{\theta})}{w^2(\bar{\boldsymbol{\theta}}_d)} p(k|x) \pi_n(S_{n,1,k}|k), \end{aligned}$$

we have using (5.13) associated with  $j = 1$ , with probability going to 1 uniformly in  $(d, \bar{\boldsymbol{\xi}}_{d,k}) \in \cup_{k < k_{n,1}} S_{n,1,k}$   $w^2(\bar{\boldsymbol{\theta}}_d)/w_o^2 = 1 + o_p(1)$ , and using (5.7) associated with  $j = 1$ ,  $w^2(\boldsymbol{\theta})/w^2(\bar{\boldsymbol{\theta}}_d) = 1 + o_p(1)$  uniformly in  $(d, \boldsymbol{\theta}) \in \cup_{k < k_{n,1}} S_{n,1,k}$ . Let

$k > k_{n,1} = k_0 \epsilon_n^{-1/\alpha} \gtrsim (n/\log n)^{\frac{\beta/\alpha}{2\beta+1}} \gg n \epsilon_n^2$ , then

$$\begin{aligned} & \mathbb{E}^{\pi_n} \left[ w^2(\boldsymbol{\theta}) \sum_{k=k_{n,1}+1}^{\infty} \mathbb{1}_{S_{n,1,k}} \right] \\ & \leq \frac{1}{w_o D_n} \sum_{k=k_{n,1}+1}^{\infty} p(k) \int_{S_{n,1,k}} \frac{w(\boldsymbol{\theta})^{-1}}{w(\bar{\boldsymbol{\theta}}_d)^{-1}} \frac{w(\bar{\boldsymbol{\theta}}_d)^{-1}}{w(\boldsymbol{\theta}_o)^{-1}} \exp \{l_n(\boldsymbol{\theta}) - l_n(\boldsymbol{\theta}_o)\} p(\boldsymbol{\theta}) d\boldsymbol{\theta} \end{aligned}$$

and on the set or  $x$  such that  $D_n \geq e^{-cn\epsilon_n^2}$ ,

$$\begin{aligned} & \mathbb{E}^{\pi_n} \left[ w^2(\boldsymbol{\theta}) \sum_{k=k_{n,1}+1}^{\infty} \mathbb{1}_{S_{n,1,k}} \right] \\ (5.14) \quad & \leq \frac{e^{cn\epsilon_n^2}}{w_o} \sum_{k=k_{n,1}+1}^{\infty} p(k) \sup_{(d,\theta) \in S_{n,1,k}} \frac{w(\boldsymbol{\theta})^{-1}}{w(\bar{\boldsymbol{\theta}}_d)^{-1}} \sup_{d:\bar{\boldsymbol{\theta}}_d \in S_{n,1,k}} \frac{w(\bar{\boldsymbol{\theta}}_d)^{-1}}{w(\boldsymbol{\theta}_o)^{-1}} \\ & \quad \times \int_{S_{n,1,k}} \exp \{l_n(\boldsymbol{\theta}) - l_n(\boldsymbol{\theta}_o)\} d\pi(d, \theta) \\ & \leq \frac{e^{cn\epsilon_n^2}}{w_o} \sum_{k=k_{n,1}+1}^{\infty} p(k) e^{\epsilon k} \int_{S_{n,1,k}} \exp \{l_n(\boldsymbol{\theta}) - l_n(\boldsymbol{\theta}_o)\} p(\boldsymbol{\theta}) d\boldsymbol{\theta}, \end{aligned}$$

for all  $\epsilon > 0$ , on a set of probability going to 1 using (5.9) and (5.10). A Markov inequality implies that

$$P_o^n \left[ \sum_{k=k_{n,1}+1}^{\infty} p(k) e^{\epsilon k} \int_{S_{n,1,k}} \exp \{l_n(\boldsymbol{\theta}) - l_n(\boldsymbol{\theta}_o)\} d\pi(d, \theta) > u_n \sum_{k=k_{n,1}+1}^{\infty} p(k) e^{\epsilon k} \right]$$

is  $O(1/u_n)$  for all  $u_n$  going to infinity, so that with probability going to 1,

$$\mathbb{E}^{\pi_n} \left[ w^2(\boldsymbol{\theta}) \sum_{k=k_{n,1}+1}^{\infty} \mathbb{1}_{S_{n,1,k}} \right] \leq e^{-c_1 k_{n,1} + cn\epsilon_n^2} = o(1)$$

for some  $c_1, c > 0$ . We finally obtain that

$$\mathbb{E}^{\pi_n} [w^2(\boldsymbol{\theta})] = w^2(\boldsymbol{\theta}_o)(1 + o_p(1)) + w^{-1}(\boldsymbol{\theta}_o) e^{-Cn\epsilon_n^2} o_p(1),$$

for some  $C > 0$ . Simple computations show that for all  $\epsilon > 0$ ,

$$P_o^n [|\log w(\boldsymbol{\theta}_o)| > n^\epsilon] = o(1),$$

we can also make precise the upper bound on  $|\log w(\boldsymbol{\theta}_o)|$  by controlling better  $|A(\boldsymbol{\theta}_o)|$ , leading to a term of order  $(\log n)^3$ , which terminates the proof.

## REFERENCES

- Adamchik, V. (2001). On the Barnes function. *Proceedings of the 2001 International Symposium on Symbolic and Algebraic Computation*, 101(474):15–20.
- Andrieu, C. and Thoms, J. (2008). A tutorial on adaptive MCMC. *Statist. Comput.*, 18(4):343–373.
- Beran, J. (1993). Fitting long-memory models by generalized linear regression. *Biometrika*, 80(4):817.

- Bertelli, S. and Caporin, M. (2002). A note on calculating autocovariances of long-memory processes. *J. Time Ser. Anal.*, 23(5):503–508.
- Böttcher, A. and Silbermann, B. (1999). *Introduction to large truncated Toeplitz matrices*. Universitext. Springer-Verlag, New York.
- Brockwell, P. and Davis, R. (2009). *Time series: theory and methods*. Springer Verlag, second ed. edition.
- Chen, W., Hurvich, C., and Lu, Y. (2006). On the correlation matrix of the discrete Fourier transform and the fast solution of large Toeplitz systems for long-memory time series. *J. Am. Statist. Assoc.*, 101:812–821.
- Dahlhaus, R. (1989). Efficient parameter estimation for self-similar processes. *Ann. Stat.*, 17(4):1749–1766.
- Del Moral, P., Doucet, A., and Jasra, A. (2006). Sequential Monte Carlo samplers. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 68(3):411–436.
- Green, P. J. (1995). Reversible jump Markov chain Monte Carlo computation and Bayesian model determination. *Biometrika*, 82:711–732.
- Green, P. J. (2003). Trans-dimensional markov chain monte carlo. In Green, P., Lid Hjort, N., and Richardson, S., editors, *Markov chain Monte Carlo in practice*, pages 179–198. Oxford University Press.
- Holan, S., McElroy, T., and Chakraborty, S. (2009). A Bayesian approach to estimating the long memory parameter. *Bayesian Analysis*, 4(1):159–190.
- Hurvich, C., Moulines, E., and Soulier, P. (2002). The FEXP estimator for potentially non-stationary linear time series. *Stoch. Processes and their applications*, 97(2):307–340.
- Jasra, A., Stephens, D., A. Doucet, A., and Tsagaris, T. (2011). Inference for Lévy driven stochastic volatility models via Sequential Monte Carlo. *Scand. J. of Statist.*, 38(1).
- Kong, A., Liu, J. S., and Wong, W. H. (1994). Sequential imputation and bayesian missing data problems. *J. Am. Statist. Assoc.*, 89:278–288.
- Kruijer, W. and Rousseau, J. (2011). Bayesian semi-parametric estimation of the long-memory parameter under FEXP-priors. *Arxiv 1202.4863*.
- Leland, W., Taqqu, M., Willinger, W., and Wilson, D. (1994). On the self-similar nature of ethernet traffic. *IEEE/ACM Transactions on Networking*, 2(1):1–15.
- Levinson, N. (1949). The Wiener RMS (root mean square) error criterion in filter design and prediction, appendix b of n. wiener, extrapolation, interpolation and smoothing of stationary time series.
- Moulines, E. and Soulier, P. (2003). Semiparametric spectral estimation for fractional processes. In Doukhan, P., Oppenheim, G., and Taqqu, M., editors, *Theory and applications of long-range dependence*, pages 251–301. Birkhäuser Boston.
- Neal, R. M. (2001). Annealed importance sampling. *Statist. Comput.*, 11:125–139.
- Palma, W. (2007). *Long-memory time series: theory and methods*, volume 662. Wiley-Blackwell.
- Press, W. H., Teukolsky, S. A., Vetterling, W. T., and Flannery, B. P. (2007). *Numerical Recipes: The Art of Scientific Computing*. Cambridge University Press.
- Richardson, S. and Green, P. J. (1997). On Bayesian analysis of mixtures with an unknown number of components. *J. R. Statist. Soc. B*, 59(4):731–792.
- Robert, C. P. and Casella, G. (2004). *Monte Carlo Statistical Methods*, 2nd ed. Springer-Verlag, New York.

- Roberts, G. O. and Rosenthal, J. S. (2001). Optimal Scaling for Various Metropolis-Hastings Algorithms. *Statist. Science*, 16(4):351–367.
- Rousseau, J., Chopin, N., and Liseo, B. (2012). Bayesian nonparametric estimation of the spectral density of a long or intermediate memory gaussian process. *Ann. Stat.*, (in press).
- Schäfer, C. and Chopin, N. (2011). Sequential Monte Carlo on large binary sampling spaces. *Statist. Comput.*, (in press). 10.1007/s11222-011-9299-z.