

Unbiased Smoothing using Particle Independent Metropolis–Hastings

Lawrence Middleton
University of Oxford

George Deligiannidis
University of Oxford

Arnaud Doucet
University of Oxford

Pierre E. Jacob
Harvard University

Abstract

We consider the approximation of expectations with respect to the distribution of a latent Markov process given noisy measurements. This is known as the smoothing problem and is often approached with particle and Markov chain Monte Carlo (MCMC) methods. These methods provide consistent but biased estimators when run for a finite time. We propose a simple way of coupling two MCMC chains built using Particle Independent Metropolis–Hastings (PIMH) to produce unbiased smoothing estimators. Unbiased estimators are appealing in the context of parallel computing, and facilitate the construction of confidence intervals. The proposed scheme only requires access to off-the-shelf Particle Filters (PF) and is thus easier to implement than recently proposed unbiased smoothers. The approach is demonstrated on a Lévy-driven stochastic volatility model and a stochastic kinetic model.

1 Introduction

1.1 State-space models, problem statement and review

Let t denote a discrete-time index. State-space models are defined by a latent Markov process $(X_t)_{t \geq 1}$ and observation process $(Y_t)_{t \geq 1}$, (X_t, Y_t) taking values in a measurable space $(\mathbf{X} \times \mathbf{Y}, \mathcal{B}(\mathbf{X}) \otimes \mathcal{B}(\mathbf{Y}))$ and satisfying

$$X_{t+1} | \{X_t = x\} \sim f(\cdot | x), \quad Y_t | \{X_t = x\} \sim g(\cdot | x),$$

for $t \geq 1$ with $X_1 \sim \mu(\cdot)$. In the following we assume that $\mathbf{X} \subseteq \mathbb{R}^{d_x}$ and $\mathbf{Y} \subseteq \mathbb{R}^{d_y}$, and use $f(\cdot | x)$, $g(\cdot | x)$ and $\mu(\cdot)$ to denote densities with respect to the corresponding Lebesgue measure. State inference given

a realization of the observations $Y_{1:T} = y_{1:T}$ for some fixed $T \in \mathbb{N}$ requires the posterior density

$$\pi(x_{1:T}) := p(x_{1:T} | y_{1:T}) \propto \mu(x_1) g(y_1 | x_1) \prod_{t=2}^T f(x_t | x_{t-1}) g(y_t | x_t),$$

and expectations w.r.t. to this density. This is known as smoothing in the literature. For non-linear non-Gaussian state-space models, this problem is complex as this posterior and its normalizing constant $p(y_{1:T})$, often called somewhat abusively likelihood, are intractable. We provide here a means to obtain unbiased estimators of $\pi(h) := \int h(x_{1:T}) \pi(x_{1:T}) dx_{1:T}$ for some function $h : \mathbf{X}^T \rightarrow \mathbb{R}$.

Particle methods return asymptotically consistent estimators which are however biased for a finite number of particles. Similarly MCMC kernels, such as the iterated Conditional Particle Filter (i-CPF) and PIMH [1], can be used for approximating smoothing expectations consistently but are also biased for a finite number of iterations. Additionally, although theoretical bounds on the bias are available for particle [12], PIMH [1] and i-CPF [7, 2, 27] estimators, these bounds are usually not sharp and/or rely on strong mixing assumptions which are not satisfied by most realistic models. Unbiased estimators of $\pi(h)$ computed on parallel machines can be combined into asymptotically valid confidence intervals as either the number of machines, the time budget, or both go to infinity [17].

Recently, it has been shown in [22, 27] that it is possible to obtain such unbiased estimators by combining the i-CPF algorithm with a debiasing scheme for MCMC algorithms proposed initially in [18] and further developed in [24]. These unbiased smoothing schemes couple two i-CPF kernels using common random numbers and a coupled resampling scheme. After a brief review of particle methods and of PIMH, we propose in Section 2 an alternative methodology relying on coupling two PIMH kernels. The method is easily implementable as it does not require any modification of the PF algorithm. It can also be used in scenarios where simulation from the Markov transition

kernel of the latent process involves a random number of random variables, whereas the methods proposed in [22, 27] would not be directly applicable to these settings. Additionally it does not require being able to evaluate pointwise the transition density contrary to the coupled conditional backward sampling PF scheme of [27]. Section 3 presents an analysis of the methodology when T is large. In Section 4, the method is demonstrated on a Lévy-driven stochastic volatility model and a stochastic kinetic model.¹

1.2 Particle methods

Particle methods are often used to approximate smoothing expectations [13, 25]. Such methods rely on sampling, weighting and resampling a set of N weighted particles (X_t^i, W_t^i) , where $X_t^i \in \mathbf{X}$ denotes the value of the i^{th} particle at iteration t and W_t^i its corresponding normalized weight, i.e. $\sum_{i=1}^N W_t^i = 1$. Letting $q_1(x_1)$, $q_t(x_t|x_{t-1})$, denote the proposal density at time $t = 1$ and at time $t \geq 2$ respectively, weighting occurs according to the following ‘incremental weights’:

$$w_1(x_1) := \frac{g(y_1|x_1)\mu(x_1)}{q_1(x_1)} \quad \text{for } t = 1,$$

$$w_t(x_{t-1}, x_t) := \frac{g(y_t|x_t)f(x_t|x_{t-1})}{q_t(x_t|x_{t-1})} \quad \text{for } t \geq 2.$$

We assume that $w_1(x_1) > 0$ and $w_t(x_{t-1}, x_t) > 0$ for $t = 2, \dots, T$ and all $x_{1:T}$. Pseudo-code for a standard PF is presented in Algorithm 1 where we let $r(\cdot|\mathbf{W}_t)$, with $\mathbf{W}_t := (W_t^1, \dots, W_t^N)$, denote the resampling distribution, a probability distribution on $[N]^N$ where $[N] := \{1, \dots, N\}$. We say that a resampling scheme is unbiased if $\sum_{i=1}^N r(A_t^i = k|\mathbf{W}_t) = NW_t^k$. All standard resampling schemes -multinomial, residual and systematic- are unbiased [13]. This PF procedure outputs an approximation $p_N(y_{1:T})$ of $p(y_{1:T})$ and an approximation $\pi_N(dx_{1:T})$ of the smoothing distribution $\pi(dx_{1:T})$. Under weak assumptions, it can be shown that $p_N(y_{1:T})$, resp. $\pi_N(h) := \sum_{i=1}^N W_T^i h(X_{1:T}^i)$, is an asymptotically consistent (in N) estimator of $p(y_{1:T})$, resp. of $\pi(h)$. However, whereas $p_N(y_{1:T})$ is unbiased ([9], Section 7.4.1), $\pi_N(h)$ admits an asymptotic bias of order C/N for a constant C which is typically impossible to evaluate and for which only loose bounds are available under realistic assumptions [12]. In the following by a call to the PF, $(X_{1:T}, p_N) \sim \text{PF}$, we mean a procedure which runs Algorithm 1 and returns $p_N := p_N(y_{1:T})$ (dependence on observations is notationally omitted) and a sample from the approximate smoothing distribution $X_{1:T} \sim \pi_N(\cdot)$, i.e. output $X_{1:T}^i$ with probability W_T^i .

¹Code to reproduce figures is provided at https://github.com/loimid/coupled_pimh

Algorithm 1 Particle Filter

For $i \in [N]$, sample $X_1^i \stackrel{i.i.d.}{\sim} q_1(\cdot)$, compute weights $W_1^i \propto w_1(X_1^i)$ and set $p_N(y_1) = \frac{1}{N} \sum_{i=1}^N w_1(X_1^i)$. For $2 \leq t \leq T$:

1. Sample particle ancestors $A_{t-1}^{1:N} \sim r(\cdot|\mathbf{W}_{t-1})$.
2. For $i \in [N]$, sample $X_t^i \sim q_t(\cdot|X_{t-1}^{A_{t-1}^i})$ and set $X_{1:t}^i = \{X_{1:t-1}^{A_{t-1}^i}, X_t^i\}$.
3. Compute weights $W_t^i \propto w_t(X_{t-1}^{A_{t-1}^i}, X_t^i)$ and set $p_N(y_{1:t}) = p_N(y_{1:t-1}) \cdot \frac{1}{N} \sum_{i=1}^N w_t(X_{t-1}^{A_{t-1}^i}, X_t^i)$.

Output $\pi_N(dx_{1:T}) := \sum_{i=1}^N W_T^i \delta_{X_{1:T}^i}(dx_{1:T})$ and $p_N(y_{1:T})$.

Algorithm 2 PIMH kernel $P((X_{1:T}^{(n)}, p_N^{(n)}), (\cdot, \cdot))$

1. Sample $(X_{1:T}^*, p_N^*) \sim \text{PF}$.
 2. Set $(X_{1:T}^{(n+1)}, p_N^{(n+1)}) = (X_{1:T}^*, p_N^*)$ with probability $\alpha(p_N^{(n)}, p_N^*) := 1 \wedge \frac{p_N^*}{p_N^{(n)}}$.
 3. Otherwise set $(X_{1:T}^{(n+1)}, p_N^{(n+1)}) = (X_{1:T}^{(n)}, p_N^{(n)})$.
-

1.3 PIMH method

An alternative way to estimate $\pi(h)$ consists of using an MCMC scheme targeting π . The PIMH algorithm achieves this by building a Markov chain on an extended space admitting a stationary distribution $\tilde{\pi}$ with marginal π using the PF described in Algorithm 1 as proposal distribution [1]. Algorithm 2 provides pseudo-code for sampling the PIMH kernel, where $(X_{1:T}^{(n)}, p_N^{(n)})$ denotes the current state of this Markov chain and $a \wedge b$ means $\min(a, b)$.

Validity and convergence properties of PIMH rely on viewing it as an Independent Metropolis–Hastings (IMH) sampler on an extended space. For ease of presentation, we detail this construction for an unbiased resampling scheme satisfying additionally $r(A_t^i = k|\mathbf{W}_t) = W_t^k$ for all $i, k \in [N]$. The PF of Algorithm 1 implicitly defines a distribution ψ over $N \times T$ particle coordinates and $N \times (T-1)$ ancestors. We use ζ to denote a sample from ψ , where $\zeta \in \mathcal{X} := \mathbf{X}^{NT} \times \{1, \dots, N\}^{N(T-1)}$, and the density of ψ is given by

$$\psi(\zeta) = \left(\prod_{i=1}^N q_1(x_1^i) \right) \prod_{t=2}^T \left(r(a_{t-1}^{1:N}|\mathbf{w}_{t-1}) \prod_{i=1}^N q(x_t^i|x_{t-1}^{a_{t-1}^i}) \right).$$

We let b_t^j denote the index of the ancestor particle of $x_{1:T}^j$ at generation t , which may be obtained deterministically from the ancestry, using $b_t^j = b_{t+1}^{b_t^j}$ with $b_T^j = j$. From [1], for $(j, \zeta) \in \{1, \dots, N\} \times \mathcal{X}$, we express

the target $\bar{\pi}(j, \zeta)$ of the resulting IMH sampler as

$$\bar{\pi}(j, \zeta) = \frac{\pi(x_{1:T}^j)}{N^T} \frac{\psi(\zeta)}{q_1(x_1^{b_1^j}) \prod_{t=2}^T r(b_{t-1}^j | \mathbf{w}_{t-1}) q_t(x_t^{b_t^j} | x_{t-1}^{b_{t-1}^j})}$$

Running a PF and sampling from π_N corresponds to sampling from the proposal $\bar{q}(j, \zeta) = \psi(\zeta) W_T^j$ and [1, Theorem 2] shows that $\bar{\pi}(j, \zeta) / \bar{q}(j, \zeta) = p_N(y_{1:T}) / p(y_{1:T})$ where for a given (j, ζ) we understand $p_N(y_{1:T})$ as a deterministic map from $\{1, \dots, N\} \times \mathcal{X}$ to $\mathbb{R}^+$. As a result, samples from $\bar{\pi}$ can be obtained asymptotically through accepting proposals with probability $\alpha(p_N^{(n)}, p_N^*)$. We can thus estimate $\pi(h)$ by averaging over iterations $h(X_{1:T}^n)$. As shown in [1, Theorem 6], we can also estimate $\pi(h)$ by averaging over iterations $\pi_N^{(n)}(h) = \sum_{i=1}^N W_T^{i,n} h(X_{1:T}^{i,n})$, hence reusing all the particle system used to generate the *accepted* proposal at iteration n as $\bar{\pi}(\pi_N(h)) = \pi(h)$. As such, although Algorithms 2 and 3 are stated in terms of proposing and accepting $(X_{1:T}, p_N)$, it is possible to consider them instead as proposing and accepting (J, ζ) . In [26], it is shown that $r(A_t^i = k | \mathbf{W}_t) = W_t^k$ is unnecessary. We only need to use an unbiased resampling scheme to obtain a valid PIMH scheme. This is achieved by defining an alternative target $\bar{\pi}$ on $\{1, \dots, N\} \times \mathcal{X}$ such that $X_{1:T}^J \sim \pi$ under $\bar{\pi}$ and $\bar{\pi}(j, \zeta) / \bar{q}(j, \zeta) = p_N(y_{1:T}) / p(y_{1:T})$ also hold. It is also established in [26] that one can even use adaptive resampling procedures [13, 11].

If we denote by Z the error of the log-likelihood estimator, i.e. $Z := \log\{p_N(y_{1:T}) / p(y_{1:T})\}$, the PIMH algorithm induces a Markov chain with transition kernel $Q(z, dz')$ given by

$$\{1 \wedge \exp(z' - z)\} g(dz') + \{1 - \alpha(z)\} \delta_z(dz'), \quad (1)$$

where $g(dz)$ is the distribution of Z under the law of the particle filter and $\alpha(z) := \int \{1 \wedge \exp(z' - z)\} g(dz')$ is the average acceptance probability from state z . Although not emphasized notationally, both $g(z)$ and $\alpha(z)$ are functions of N . Through an abuse of notation, we denote the invariant density of the above chain with $\pi(z) = g(z) \exp(z)$. Such a reparameterization has also been used in previous work to analyze the PIMH and the related particle marginal MH algorithm [31, 14].

2 Methodology

2.1 Unbiased MCMC via couplings

‘Exact estimation’ methods provide unbiased estimators of expectations with respect to the stationary distribution of a Markov chain using coupling techniques [18, 24]. In the following we use these tools to couple

PIMH kernels and estimate smoothing expectations, after briefly reviewing the general approach.

Consider two U -valued Markov chains, $U := (U^{(n)})_{n \geq 0}$ and $\tilde{U} := (\tilde{U}^{(n)})_{n \geq 0}$, each evolving marginally according to a kernel $K(u, du')$ with stationary distribution λ and initialized from η so that $U^{(n)} \stackrel{d}{=} \tilde{U}^{(n)}$ for all $n \geq 1$, where $\stackrel{d}{=}$ denotes equality in distribution. We couple these two chains so that $U^{(n)} = \tilde{U}^{(n-1)}$ for $n \geq \tau$, where τ is an almost surely finite meeting time. In this case, we see that for non-negative integers $k < t$ we have the following telescoping-sum decomposition

$$\begin{aligned} \mathbb{E}[h(U^{(t)})] &= \mathbb{E}[h(U^{(k)}) + \sum_{l=k+1}^t (h(U^{(l)}) - h(U^{(l-1)}))] \\ &= \mathbb{E}[h(U^{(k)}) + \sum_{l=k+1}^{t \wedge (\tau-1)} (h(U^{(l)}) - h(\tilde{U}^{(l-1)}))]. \end{aligned}$$

As a result, under integrability conditions which will be made precise in the sequel, taking the limit as $t \rightarrow \infty$ suggests an estimator $H_k(U, \tilde{U}) := h(U^{(k)}) + \sum_{l=k+1}^{\tau-1} (h(U^{(l)}) - h(\tilde{U}^{(l-1)}))$ with expectation $\lambda(h)$. A way to construct such chains considered in [18, 24, 21, 29] relies on sampling independently $(U^{(0)}, U^{(1)}) \sim \eta(du_0)K(u_0, du_1)$ and $\tilde{U}^{(0)} \sim \eta(du_0)$, and then successively sampling $(U^{(n+1)}, \tilde{U}^{(n)})$ from $\bar{K}((U^{(n)}, \tilde{U}^{(n-1)}), (\cdot, \cdot))$ such that $U^{(n+1)} = \tilde{U}^{(n)}$ with positive probability, ensuring that both chains evolve marginally according to K . Furthermore, averaging $H_l(U, \tilde{U})$ over a range of values $l \in \{k, k+1, \dots, m\}$, for some $m \geq k$, preserves unbiasedness, suggesting the following ‘time-averaged’ unbiased estimator of $\mu(h)$

$$\begin{aligned} H_{k:m} &= \frac{1}{m-k+1} \sum_{l=k}^m h(U^{(l)}) + \\ &\sum_{l=k+1}^{\tau-1} \min \left(1, \frac{l-k}{m-k+1} \right) (h(U^{(l)}) - h(\tilde{U}^{(l-1)})) \\ &:= \text{MCMC}_{k:m}(h) + \text{BC}_{k:m}(h). \quad (2) \end{aligned}$$

We view $\text{MCMC}_{k:m}(h)$ as the standard MCMC sample average up to time m , discarding the first k iterates as ‘burn-in’. The second term, $\text{BC}_{k:m}(h)$, is the ‘bias correction’ term with $\text{BC}_{k:m}(h) := 0$ if $\tau-1 < k+1$. The estimator requires only that two chains be simulated until meeting at time τ , after which, if $\tau < m$, only one chain must be simulated up to m . By [24, Proposition 3.1], the validity of the resulting estimators is guaranteed under the following assumptions. We discuss these assumptions for PIMH in Section 2.3.

Assumption 1. *Each chain is initialized marginally from a distribution η , evolves according to a kernel K , and is such that $\mathbb{E}[h(U^{(n)})] \rightarrow \lambda(h)$ as $n \rightarrow \infty$.*

Furthermore, there exist constants $\eta > 0$ and $D < \infty$ such that $\mathbb{E}[|h(U^{(n)})|^{2+\eta}] < D$ for all $n \geq 0$.

Assumption 2. *The two chains are such that the meeting time $\tau = \inf\{n \geq 1 : U^{(n)} = \tilde{U}^{(n-1)}\}$ satisfies $\mathbb{P}[\tau > n] \leq D\delta^n$, for some constants $0 < D < \infty$ and $\delta \in (0, 1)$. The chains stay together after meeting, i.e. $U^{(n)} = \tilde{U}^{(n-1)}$ for all $n \geq \tau$.*

Proposition 3. [24] *Under Assumptions 1 and 2, the estimator $H_{k:m}$ obtained by these coupled chains is unbiased, has finite variance and finite expected cost.*

2.2 Coupled PIMH

To obtain unbiased estimators of $\pi(h)$, we use the framework detailed in Section 2.1 for $K = P$ the transition kernel of PIMH and $\lambda = \bar{\pi}$ the corresponding invariant distribution of PIMH defined on $\mathcal{U} := \{1, \dots, N\} \times \mathcal{X}$. This requires introducing a coupling of PIMH kernels. For IMH chains, a natural choice of initialization of the chain is from the proposal distribution. We adopt this in the following, sampling both $(X_{1:T}^{(0)}, p_N^{(0)}) \sim \text{PF}$ and $(\tilde{X}_{1:T}^{(0)}, \tilde{p}_N^{(0)}) \sim \text{PF}$. However, in contrast to previous constructions [24, 21, 29], our method allows the chains to couple at time $\tau = 1$ through re-using the initial value $(\tilde{X}_{1:T}^{(0)}, \tilde{p}_N^{(0)})$ as a proposal for the first chain.

We summarize the resulting procedure in Algorithm 3. To simplify presentation of the algorithm, we use the convention $\alpha(\tilde{p}_N^{(-1)}, p_N^*) := 1$. Finally, although the coupling scheme is framed in terms of PIMH, it is clear that this algorithm and estimators apply equally to any IMH algorithm.

Algorithm 3 Coupled PIMH

1. Sample $(X_{1:T}^{(0)}, p_N^{(0)}) \sim \text{PF}$, set $n = 1$ and $\tau = \infty$.
2. While $n < \max(m, \tau)$
 - (a) Sample $(X_{1:T}^*, p_N^*) \sim \text{PF}$.
 - (b) Sample $u \sim \mathcal{U}[0, 1]$.
 - (c) If $u \leq \alpha(p_N^{(n-1)}, p_N^*)$ set

$$(X_{1:T}^{(n)}, p_N^{(n)}) = (X_{1:T}^*, p_N^*)$$

$$\text{else } (X_{1:T}^{(n)}, p_N^{(n)}) = (X_{1:T}^{(n-1)}, p_N^{(n-1)}).$$

- (d) If $u \leq \alpha(\tilde{p}_N^{(n-2)}, p_N^*)$ set

$$(\tilde{X}_{1:T}^{(n-1)}, \tilde{p}_N^{(n-1)}) = (X_{1:T}^*, p_N^*)$$

$$\text{else } (\tilde{X}_{1:T}^{(n-1)}, \tilde{p}_N^{(n-1)}) = (\tilde{X}_{1:T}^{(n-2)}, \tilde{p}_N^{(n-2)}).$$

- (e) If $u \leq \alpha(p_N^{(n-1)}, p_N^*) \wedge \alpha(\tilde{p}_N^{(n-2)}, p_N^*)$ set $\tau = n$.

- (f) Set $n \leftarrow n + 1$.

3. Return $H_{k:m}$ as in (2).
-

At each iteration n of Algorithm 3 the proposal $(X_{1:T}^*, p_N^*)$ can be accepted by one, both or neither chains with positive probability. The meeting time τ corresponds to the first time the proposal $(X_{1:T}^*, p_N^*)$ is accepted by both chains. We can also return another unbiased estimator $\bar{H}_{k:m}$ of the form (2) with $h(X_{1:T})$ replaced by $\pi_N(h)$ as we have previously seen that $\bar{\pi}(\pi_N(h)) = \pi(h)$. The Rao-Blackwellised estimator $\bar{H}_{k:m}$ will typically outperform significantly $H_{k:m}$ when we are interested in smoothing expectations of functions of states close to T . For functions of states close to the origin, e.g. $h(x_{1:T}) = x_1$, we have $\bar{H}_{k:m} = H_{k:m}$ with high probability if N is moderate because of the particle degeneracy problem [13, 23, 25]. This is illustrated in Appendix A.3.

2.3 Validity and meeting times

The validity of our unbiased estimators is ensured if the following weak assumptions are satisfied.

Assumption 4. *There exist constants $\eta > 0$ and $D < \infty$ such that $\mathbb{E}[|h(X_{1:T}^{(n)})|^{2+\eta}] < D$ for all $n \geq 0$.*

Assumption 5. *There exist constants $\eta > 0$ and $D < \infty$ such that $\mathbb{E}[|\sum_{i=1}^N W_T^{i,(n)} h(X_{1:T}^{i,(n)})|^{2+\eta}] < D$ for all $n \geq 0$.*

Assumption 6. *The resampling scheme is unbiased and there exist finite constants $(\bar{w}_t)_{t=1}^T$ such that $\sup_{x \in \mathcal{X}} w_1(x) \leq \bar{w}_1$ and $\sup_{(x,x') \in \mathcal{X} \times \mathcal{X}} w_t(x, x') \leq \bar{w}_t$ for $t \in \{2, \dots, T\}$.*

Proposition 7. *Under Assumptions 4 and 6, resp. Assumptions 5 and 6, the estimator $H_{k:m}$, resp. $\bar{H}_{k:m}$, of $\pi(h)$ obtained from Algorithm 3 is unbiased and has finite variance and finite expected cost.*

To establish the result for $H_{k:m}$, note that Assumption 4 and our construction implies Assumption 1 for $\lambda = \bar{\pi}$. Assumption 6 provides verifiable and sufficient conditions to ensure uniform ergodicity of PIMH [1, Theorem 3]. Hence it follows from [24, Proposition 3.4] that the geometric bound on the tails of τ is satisfied. Additionally, Algorithm 3 ensures that the chains stay together for all $n \geq \tau$ so Assumption 2 is satisfied. A similar reasoning provides the result for $\bar{H}_{k:m}$.

We have the following precise description of the distribution of the meeting time τ for Algorithm 3. Let $\text{Geo}(\gamma)$ denote the geometric distribution on the strictly positive integers with success probability γ .

Proposition 8. *The meeting time satisfies $\tau|Z^{(0)} \sim \text{Geo}(\alpha(Z^{(0)}))$ and $\mathbb{P}[\tau = 1] \geq \frac{1}{2}$. Additionally, we have $\lim_{N \rightarrow \infty} \mathbb{P}[\tau = 1] = 1$ under Assumption 6.*

We recall here that $\alpha(z)$ is the average acceptance probability from state z . The geometric distribution

result follows from Proposition 10 in Appendix. The rest of the proposition follows from

$$\begin{aligned}\mathbb{P}[\tau = 1] &= \int \alpha(z)g(dz) \\ &= \iint \{1 \wedge \exp(z' - z)\} g(dz)g(dz').\end{aligned}$$

It entails trivially that $\mathbb{P}[\tau = 1] \geq \frac{1}{2}$. Noting that $\lim_{N \rightarrow \infty} Z = 0$ a.s. under g (which is dependent on N) under Assumption 6, see e.g. [9], we obtain $\lim_{N \rightarrow \infty} \mathbb{P}[\tau = 1] = 1$ by dominated convergence, thus $\lim_{N \rightarrow \infty} \text{BC}_{k:m}(h) = 0$.

2.4 Unbiased filtering

A PF generates estimates $p_N(y_{1:t})$ of $p(y_{1:t})$ at all times $t = 1, \dots, T$. These estimates can be used to perform one step of coupled PIMH for T pairs of Markov chains, each pair corresponding to one of the smoothing distributions $p(x_{1:t}|y_{1:t})$. This requires some additional bookkeeping to keep track of the meeting times and unbiased estimators associated with each $p(x_{1:t}|y_{1:t})$ but can be used to unbiasedly estimate expectations with respect to all filtering distributions $p(x_t|y_{1:t})$. This was not directly feasible with coupled i-CPF in [22]. In particular, this allows us to estimate unbiasedly the predictive likelihood terms $p(y_t|y_{1:t-1})$ which can be used for a goodness-of-fit test [15].

3 Analysis

3.1 Meeting time: large sample approximation

We investigate here the distribution of the meeting time in the large sample regime, i.e. in the interesting scenarios where T is large. Under strong mixing assumptions, [5] showed that letting $\frac{T}{N} = \gamma$ for some $\gamma > 0$ then the following Central Limit Theorem (CLT) holds: $Z \xrightarrow{d} \mathcal{N}(-\frac{1}{2}\sigma^2, \sigma^2)$ as $T \rightarrow \infty$ where $\sigma^2 = \gamma\bar{\sigma}^2$ for some $\bar{\sigma}^2 > 0$. Empirically, this CLT appears to hold for many realistic models not satisfying these strong mixing assumptions [31, 14]. Under this CLT, using similar regularity conditions as in [32], it can be shown that the Markov kernel Q defined in (1) converges in some suitable sense towards the Markov kernel Q_σ given by

$$\{1 \wedge \exp(z' - z)\} g_\sigma(z')dz' + \{1 - \alpha_\sigma(z)\} \delta_z(dz'),$$

where $g_\sigma(z) = \mathcal{N}(z; -\frac{1}{2}\gamma\bar{\sigma}^2, \gamma\bar{\sigma}^2)$ and $\alpha_\sigma(z) := \int \{1 \wedge \exp(z' - z)\} g_\sigma(z')dz'$ denotes the average acceptance probability in z . For this limiting kernel², the following result holds.

²Note that Q_σ is not uniformly ergodic whereas Q is under Assumption 6.

Proposition 9. [14, Corollary 3] *The invariant density of Q_σ is $\pi_\sigma(z) = \mathcal{N}(z; \frac{1}{2}\sigma^2, \sigma^2)$ and its average acceptance probability is given by*

$$\alpha_\sigma(z) := 1 - \Phi\left(\frac{z + \frac{\sigma^2}{2}}{\sigma}\right) + e^{-z}\Phi\left(\frac{z - \frac{\sigma^2}{2}}{\sigma}\right),$$

where Φ denotes the cumulative distribution function of the standard Normal distribution.

Under this large sample approximation, the probability $\mathbb{P}[\tau = n]$ can be written as

$$\mathbb{E}[\mathbb{P}[\tau = n|Z^{(0)}]] = \int \alpha_\sigma(z)(1 - \alpha_\sigma(z))^{n-1}g_\sigma(z)dz, \quad (3)$$

thus $\mathbb{P}[\tau = 1] = \frac{1}{2} \{1 + \exp(\sigma^2)\text{Erfc}(\sigma)\}$, Erfc denoting the complementary error function. Similarly we see that the expected meeting time is given by $\mathbb{E}[\tau] = \mathbb{E}_{g_\sigma}[\alpha_\sigma(Z)^{-1}]$. This allows us to approximate numerically expectations, quantiles and probabilities of τ as a function of σ , the standard deviation of $\log p_N(y_{1:T})$ under the law of the particle filter. Figure 1 displays $\mathbb{P}[\tau = 1]$ and $\mathbb{E}[\tau]$ as a function of σ . We see that $\mathbb{P}[\tau = 1] = 0.71$ for $\sigma = 1$ rising to $\mathbb{P}[\tau = 1] = 0.95$ for $\sigma = 0.1$. The expected meeting time $\mathbb{E}[\tau]$ is also the expected number of iterations to return an unbiased estimator for $m = 0$ and grows relatively benignly with σ .

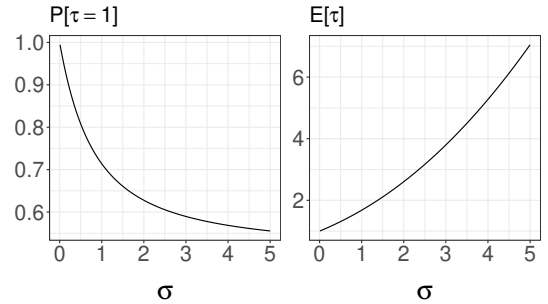


Figure 1: $\mathbb{P}[\tau = 1]$ (left) and $\mathbb{E}[\tau]$ (right) as a function of the standard deviation σ of $\log p_N(y_{1:T})$.

We compare the distribution of the meeting times with its large sample approximation (3) on a stationary auto-regressive (AR) model $X_t \sim \mathcal{N}(aX_{t-1}, 1)$ and $Y_t \sim \mathcal{N}(X_t, \sigma_y^2)$ with $a = 0.5$ and $\sigma_y^2 = 10$ on a simulated dataset of $T = 100$. Estimates of $\mathbb{P}[\tau \geq n]$ were obtained empirically for a range of values of N and compared with the predicted values based on the estimated variance σ^2 of $\log p_N(y_{1:T})$ using 10^5 runs of coupled PIMH. Varying N between 10 and 110, σ^2 was between 0.2 and 3.0 in this range. The tail probabilities of the meeting time over this range are shown in Figure 2. Confidence intervals for the estimates of $\mathbb{P}[\tau \geq n]$ are shown in Figure 2 with error bars indicating ± 2 standard deviations. We see that there

is a satisfactory agreement, with predicted tail probabilities closely matching confidence intervals for each value, and larger values of N leading as expected to shorter meeting times on average.

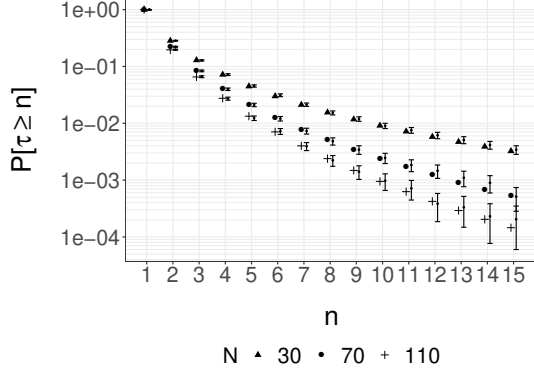


Figure 2: Empirical estimates of $\mathbb{P}[\tau \geq n]$ (with ± 2 standard errors) and comparison with the large sample approximation for toy example.

3.2 On the selection of N for large m

We provide here a heuristic for selecting N to minimize the variance of $H_{k:m}$ at fixed computational budget when both m and T are large. We will minimize the computational inefficiency defined by

$$C[H] := \mathbb{V}[H_{k:m}] \times N$$

as the running time is proportional to N . For m sufficiently large, we expect that $\tau < m$ with very high probability and so the time to obtain a single unbiased estimator is m . We note that increasing N typically leads to a decreasing asymptotic variance of the ergodic averages associated with the PIMH chain and from Proposition 8 a reduction in the bias correction, however at the cost of more computation. For large m , we also expect that the dominant term of $\mathbb{V}[H_{k:m}]$ will arise from $\text{MCMC}_{k:m}(h)$ and will be essentially the asymptotic variance of h given by $\mathbb{V}_\pi[h] \text{IF}(h)$ divided by $(m - k + 1)$, where $\text{IF}(h)$ is the Integrated Auto-correlation Time (IACT) of h for the PIMH kernel. For N sufficiently large, we expect that $X_{1:T}$ and p_N are approximately independent under the PF proposal. By a reasoning similar to the proof of [31, Lemma 4], $\text{IF}(h)$ will then be approximately proportional to $\text{IF}(\sigma)$ defined in Eq. (11) in [31]. This is illustrated in Appendix A.2 where we plot $\text{IF}(h)$ for a variety of test functions for a range of N against $\text{IF}(\sigma)$ over the corresponding range of σ and show that they are indeed approximately proportional. Minimizing $\text{IF}[H]$ w.r.t. N is then approximately equivalent to minimizing $\text{IF}(\sigma)/\sigma^2$ as σ^2 is typically inversely proportional to N . This minimization has already been carried out

in [31] where it was found that the minimizing argument is $\sigma = 0.92$. Practically, this means that one should select N to ensure that the standard deviation of $\log p_N(y_{1:T})$ is equal approximately to this value. The resulting value of N is expected to be close to the value of N minimizing $C[H]$, which is approximately proportional to $\text{IF}(h)/\sigma^2$. This is verified in Figure 3 on the AR example of Section 3.1. Note that these guidelines do not apply to $\bar{H}_{k:m}$ for h a function of states close to T as it is not true that $\pi_N(h)$ and p_N are approximately independent.

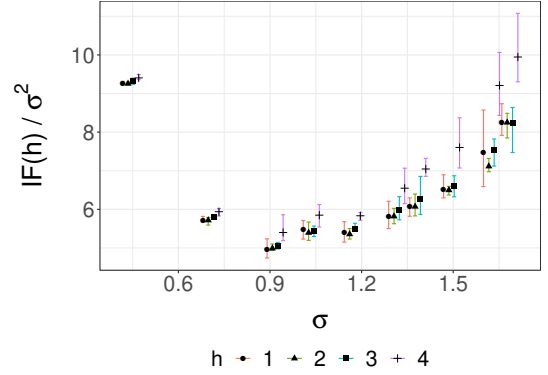


Figure 3: $\text{IF}(h_i)/\sigma^2$ for $h_1 = x_1$, $h_2 = x_T$, $h_3 = \sum_t x_t$ and $h_4 = \sum_t x_t^2$ as a function of σ^2 .

4 Numerical experiments

We apply the methodology to two models where the transition density of the latent process is analytically intractable but simulation from it is possible. However, this involves sampling a random number of random variables. The unbiased smoother proposed in [22] relies on common random numbers and it is unclear how one could implement it in this context. The coupled conditional backward sampling PF scheme proposed in [27] does not apply as we cannot evaluate the transition density pointwise. In both scenarios, we use the bootstrap PF with multinomial resampling, that is $q_1(x_1) = \mu(x_1)$ and $q_t(x_t|x_{t-1}) = f(x_t|x_{t-1})$, and Assumption 6 is satisfied.

4.1 Stochastic kinetic model

We consider a stochastic kinetic model represented by a jump Markov process introduced in [20]. Such models describe a system of chemical reactions in continuous time, with a reaction occurring under the collision of two species at random times. The discrete number of each species describes the state with jumps representing the change in a particular species. The discrete valued state $(X_{t,q})_{t \geq 0, 1 \leq q \leq Q}$ comprises a Q -vector of species at each time, where one of R reactions may

occur at any random time, given by hazard functions f_r for $r \in \{1, \dots, R\}$. The effect of such a reaction is described by a stoichiometry matrix S , where the instantaneous change in the number of species q for a certain reaction r out of a possible R different reactions is encoded in element $S_{q,r}$. For the prokaryotic autoregulation model and parameterisation considered in [19], the state is a four dimensional vector evolving according to 8 possible reactions, for which the stoichiometry matrix is given by

$$S = \begin{pmatrix} 0 & 0 & 1 & 0 & 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & 1 & -2 & 2 & 0 & -1 \\ -1 & 1 & 0 & 0 & 1 & -1 & 0 & 0 \\ -1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix},$$

$$f(X, c) = (c_1 X_4 X_3, c_2 (k - X_4), c_3 X_4, c_4 X_1, c_5 X_2 (X_2 - 1)/2, c_6 X_3, c_7 X_1, c_8 X_2)',$$

Coefficients $c_{1:8}$ and k are parameters of the model, given by $c = (0.1, 0.7, 0.35, 0.2, 0.1, 0.9, 0.3, 0.1)$ and $k = 10$. We collect $T = 100$ noisy observations of the latent process at regular intervals of length $\Delta = 0.1$, i.e.

$$Y_t = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 2 & 0 \end{pmatrix} X_{\Delta t} + \epsilon_t, \quad \epsilon_t \stackrel{i.i.d.}{\sim} \mathcal{N}(0, I_2).$$

We fix the initial condition $X_0 = (8, 8, 8, 5)$.

To simulate synthetic data and to run the bootstrap PF, we sample the latent process X_t using Gillespie's direct method [16], whereby the time to the next event is exponential with rate $\sum_{r=1}^R f_r(X, c)$ and reaction r occurs with probability $f_r(X, c) / \sum_{r=1}^R f_r(X, c)$. The estimated survival probabilities of the meeting time τ , with ± 2 standard errors, computed using 500 independent runs of coupled PIMH are plotted in Figure 4, along with the probabilities obtained from the large sample approximation in Section 3.1, showing good agreement between the two. In Figure 5 we display the unbiased smoothing estimators obtained by averaging the unbiased estimators obtained over 500 independent runs for $N = 1,000, k = m = 0$ and the corresponding confidence intervals. Alternative choices of k, m can lead to improved performance at fixed computational budget [24, 29].

4.2 Lévy-driven stochastic volatility

Introduced in [3], Lévy-driven stochastic volatility models provide a flexible model for the log-returns of a financial asset. Letting (Y_t) denote the log-return process, we have

$$Y_t = \mu + \beta V_t + V_t^{1/2} \epsilon_t, \quad \epsilon_t \stackrel{i.i.d.}{\sim} \mathcal{N}(0, 1),$$

where V_t (termed the 'actual volatility') is treated as a stationary stochastic process. The latent state comprises the pair of actual and spot volatility $X_t =$

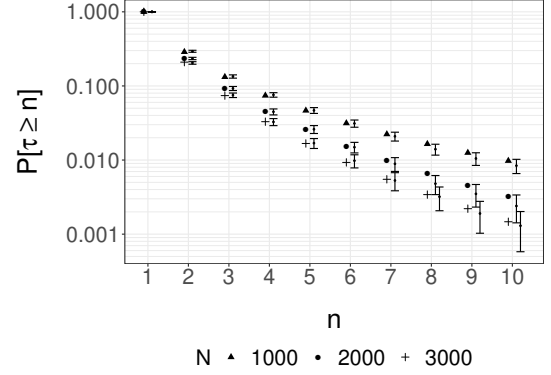


Figure 4: Empirical estimates of $\mathbb{P}[\tau \geq n]$ (with ± 2 standard errors) and comparison with those implied by the large sample approximation for stochastic kinetic model.

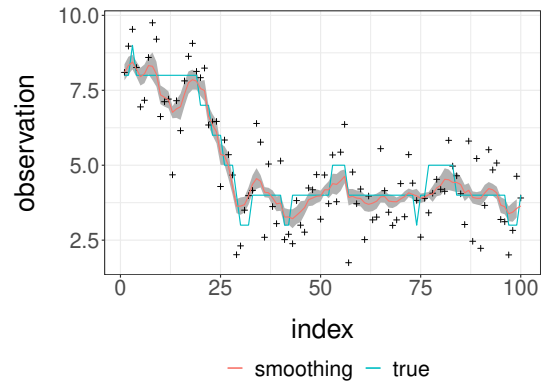


Figure 5: Unbiased estimates of $\mathbb{E}(X_{1,t} | y_{1:T})$ (red) (with ± 3 pointwise standard errors) and $X_{1,t}$ (blue) for Markov jump process.

(V_t, W_t) . In the terminology of [4], the integrated volatility is the integral of the spot volatility and the actual volatility is an increment of the integrated volatility over some unit time. Initializing $Z_0 \sim \Gamma(\xi^2/\omega^2, \xi^2/\omega^2)$, the process evolves through sampling the following random variables and recursing the state $X_t = (V_t, W_t)$ according to

$$\begin{aligned} K &\sim \text{Poisson}(\lambda \xi^2 / \omega^2), \quad C_{1:K} \stackrel{i.i.d.}{\sim} \mathcal{U}[t-1, t], \\ E_{1:K} &\stackrel{i.i.d.}{\sim} \text{Exp}(\xi / \omega^2), \quad W_t = e^{-\lambda} W_{t-1} + \sum_{j=1}^K e^{-\lambda(t-C_j)} E_j, \\ V_t &= \frac{1}{\lambda} (W_{t-1} - W_t + \sum_{j=1}^K E_j). \end{aligned}$$

In particular, conditionally on X_{t-1} , simulation of X_t requires a random number of random numbers sampled at each iteration. The parameters (ξ, ω^2) denote the stationary mean and variance of the spot volatility respectively, λ describes the exponential decay of

autocorrelations, β denotes the risk premium for excess volatility and μ the drift of the log-return. In the following we perform unbiased smoothing of X_t using $T = 500$ data from the S&P 500 index used in [6]. A summary of the data and parameter inference is included in Appendix A.4.

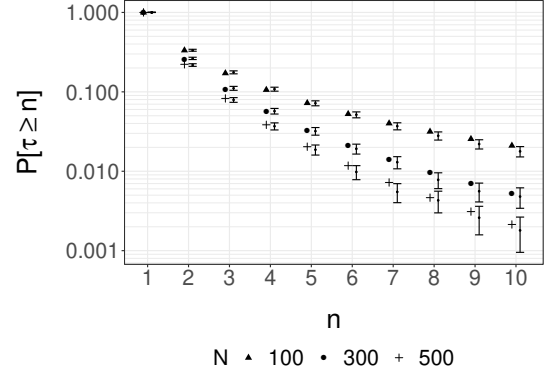
The empirical distributions of the meeting time for $N \in \{100, \dots, 500\}$ obtained using 1,000 runs are shown in Figure 6, again showing good agreement with the large-sample approximation. We then obtain unbiased estimators $H_{k:m}$ for $h(x_{1:T}) = \sum_{t=1}^T v_t$ using $k = 20, m = 512$ over this grid of N . Figure 6b plots $(m - k + 1)\mathbb{V}[H_{k:m}]/(\mathbb{V}_\pi[h]\sigma^2)$. We expect this function to be close to $\text{IF}(h)/\sigma^2$ as $m \rightarrow \infty$ and $\text{IF}(h)/\sigma^2$ to be minimized around 0.92. The experiments are consistent with this result. Figure 6c presents the unbiased estimates of the spot volatility W_t obtained by averaging the unbiased estimates obtained over 1,000 runs and the corresponding confidence intervals for $k = m = 0$ and $N = 100$. Different choices of k, m could lead to improved performance at fixed computational budget.

5 Discussion

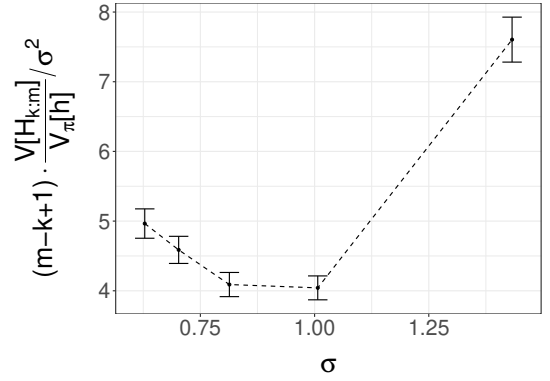
We have introduced a simple approach to perform unbiased smoothing in state-space models and we have provided guidance on the choice of tuning parameters through appealing to a large sample approximation.

We have established the validity of the estimators when the incremental weights are bounded (Assumption 6) which ensures uniform ergodicity of the PIMH. Rejection sampling is possible under a similar assumption and would provide exact samples from the smoothing distribution. However the expected number of trials before acceptance of such a rejection scheme increases typically exponentially fast with T . If we are only interested in obtaining unbiased smoothing estimators and if the CLT discussed in Section 3.1 holds, we expect our coupling scheme to only require increasing N linearly with T to control σ and thus the corresponding expectation of the meeting time.

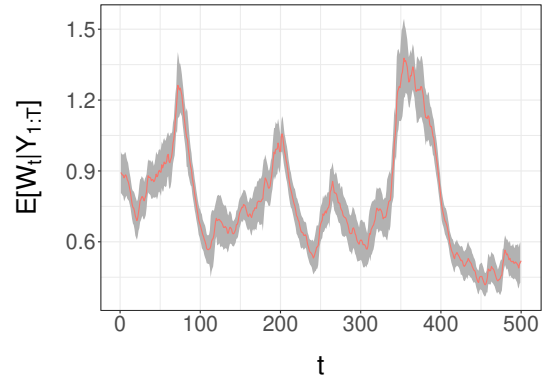
Finally, the scheme proposed here can be extended to obtain unbiased estimators of expectations with respect to any posterior distribution by replacing the particle filter proposal within the IMH by Annealed Importance Sampling [30] or a sequential Monte Carlo sampler [10]. This is illustrated in Appendix B.



(a) Empirical (± 2 standard errors) and predicted tail probabilities implied by large-sample approximation $\mathbb{P}[\tau \geq n]$.



(b) Estimates of $(m - k + 1)\mathbb{V}[H_{k:m}]/(\mathbb{V}_\pi[h]\sigma^2)$ (± 3 standard errors) computed using 10,000 runs.



(c) Unbiased estimates of $\mathbb{E}(W_t|y_{1:T})$ of the S&P 500 over 2005–2007 and confidence intervals at ± 3 standard errors.

Figure 6: Distribution of meeting times (top), computational inefficiency (bottom) and smoothing state estimators for Lévy-driven stochastic volatility applied to real data (bottom).

References

- [1] Christophe Andrieu, Arnaud Doucet, and Roman Holenstein. Particle Markov chain Monte Carlo methods. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 72(3):269–342, 2010.
- [2] Christophe Andrieu, Anthony Lee, and Matti Viola. Uniform ergodicity of the iterated conditional SMC and geometric ergodicity of particle Gibbs samplers. *Bernoulli*, 24(2):842–872, 2018.
- [3] Ole E Barndorff-Nielsen and Neil Shephard. Non-Gaussian Ornstein–Uhlenbeck-based models and some of their uses in financial economics. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 63(2):167–241, 2001.
- [4] Ole E Barndorff-Nielsen and Neil Shephard. Econometric analysis of realized volatility and its use in estimating stochastic volatility models. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 64(2):253–280, 2002.
- [5] Jean Bérard, Pierre Del Moral, and Arnaud Doucet. A lognormal central limit theorem for particle approximations of normalizing constants. *Electronic Journal of Probability*, 19, 2014.
- [6] Nicolas Chopin, Pierre E Jacob, and Omiros Papaspiliopoulos. SMC2: an efficient algorithm for sequential analysis of state space models. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 75(3):397–426, 2013.
- [7] Nicolas Chopin and Sumeetpal S Singh. On particle Gibbs sampling. *Bernoulli*, 21(3):1855–1883, 2015.
- [8] Jem N Corcoran and Richard L Tweedie. Perfect sampling from independent Metropolis-Hastings chains. *Journal of Statistical Planning and Inference*, 104(2):297–314, 2002.
- [9] Pierre Del Moral. *Feynman-Kac Formulae: Genealogical and Interacting Particle Systems with Applications*. Springer-Verlag, New York, 2004.
- [10] Pierre Del Moral, Arnaud Doucet, and Ajay Jasra. Sequential Monte Carlo samplers. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 68(3):411–436, 2006.
- [11] Pierre Del Moral, Arnaud Doucet, and Ajay Jasra. On adaptive resampling strategies for sequential Monte Carlo methods. *Bernoulli*, 18(1):252–278, 2012.
- [12] Pierre Del Moral, Arnaud Doucet, and Gareth W Peters. Sharp propagation of chaos estimates for Feynman–Kac particle models. *Theory of Probability & Its Applications*, 51(3):459–485, 2007.
- [13] Arnaud Doucet and Adam M. Johansen. A tutorial on particle filtering and smoothing: Fifteen years later. In D. Crisan and B. Rozovsky, editors, *The Oxford Handbook of Nonlinear Filtering*, pages 656–704. Oxford University Press, 2011.
- [14] Arnaud Doucet, Michael K Pitt, George Deligiannidis, and Robert Kohn. Efficient implementation of Markov chain Monte Carlo when using an unbiased likelihood estimator. *Biometrika*, 102(2):295–313, 2015.
- [15] Richard Gerlach, Chris Carter, and Robert Kohn. Diagnostics for time series analysis. *Journal of Time Series Analysis*, 20(3):309–330, 1999.
- [16] Daniel T Gillespie. Exact stochastic simulation of coupled chemical reactions. *The Journal of Physical Chemistry*, 81(25):2340–2361, 1977.
- [17] Peter W. Glynn and Philip Heidelberger. Analysis of parallel replicated simulations under a completion time constraint. *ACM Transactions on Modeling and Computer Simulations*, 1(1):3–23, 1991.
- [18] Peter W Glynn and Chang-Han Rhee. Exact estimation for Markov chain equilibrium expectations. *Journal of Applied Probability*, 51(A):377–389, 2014.
- [19] Andrew Golightly and Theodore Kypraios. Efficient SMC2 schemes for stochastic kinetic models. *Statistics and Computing*, 28(6):1215–1230, 2018.
- [20] Andrew Golightly and Darren J Wilkinson. Bayesian inference for stochastic kinetic models using a diffusion approximation. *Biometrics*, 61(3):781–788, 2005.
- [21] Jeremy Heng and Pierre E Jacob. Unbiased Hamiltonian Monte Carlo with couplings. *Biometrika*, (forthcoming), 2018.
- [22] Pierre E Jacob, Fredrik Lindsten, and Thomas B Schön. Smoothing with couplings of conditional particle filters. *Journal of the American Statistical Association*, (forthcoming), 2018.
- [23] Pierre E Jacob, Lawrence M Murray, and Sylvain Rubenthaler. Path storage in the particle filter. *Statistics and Computing*, 25(2):487–496, 2015.

- [24] Pierre E Jacob, John O’Leary, and Yves F Atchadé. Unbiased Markov chain Monte Carlo with couplings. *arXiv preprint arXiv:1708.03625*, 2017.
- [25] Nikolas Kantas, Arnaud Doucet, Sumeetpal S Singh, Jan Maciejowski, and Nicolas Chopin. On particle methods for parameter estimation in state-space models. *Statistical Science*, 30(3):328–351, 2015.
- [26] Anthony Lee, Lawrence Murray, and Adam M Johansen. Resampling in conditional SMC algorithms. *Technical Report*, 2019.
- [27] Anthony Lee, Sumeetpal S Singh, and Matti Vi-hola. Coupled conditional backward sampling particle filter. *arXiv preprint arXiv:1806.05852*, 2018.
- [28] Anthony Lee, Christopher Yau, Michael B Giles, Arnaud Doucet, and Christopher C Holmes. On the utility of graphics cards to perform massively parallel simulation of advanced Monte Carlo methods. *Journal of Computational and Graphical Statistics*, 19(4):769–789, 2010.
- [29] Lawrence Middleton, George Deligiannidis, Arnaud Doucet, and Pierre E Jacob. Unbiased Markov chain Monte Carlo for intractable target distributions. *arXiv preprint arXiv:1807.08691*, 2018.
- [30] Radford M Neal. Annealed importance sampling. *Statistics and Computing*, 11(2):125–139, 2001.
- [31] Michael K Pitt, Ralph dos Santos Silva, Paolo Giordani, and Robert Kohn. On some properties of Markov chain Monte Carlo simulation methods based on the particle filter. *Journal of Econometrics*, 171(2):134–151, 2012.
- [32] Sebastian M Schmon, George Deligiannidis, Arnaud Doucet, and Michael K Pitt. Large sample asymptotics of the pseudo-marginal method. *arXiv preprint arXiv:1806.10060*, 2018.
- [33] Yan Zhou, Adam M Johansen, and John AD Aston. Toward automatic model comparison: an adaptive sequential Monte Carlo approach. *Journal of Computational and Graphical Statistics*, 25(3):701–726, 2016.

A Appendix

A.1 Proof of Proposition 8

We establish a monotonicity property of the proposed coupling for PIMH. Monotonicity properties of IMH samplers were exploited in an exact simulation context in [8] without this explicit construction.

Proposition 10. *Under the proposed coupling scheme the sequence of likelihood estimates $(p_N^{(n+1)})_{n \geq 0}$ stochastically dominates $(\tilde{p}_N^{(n)})_{n \geq 0}$ in the sense that for any $n \geq 1$, $s \geq 0$*

$$p_N^{(n)} \geq \tilde{p}_N^{(n-1)} \Rightarrow p_N^{(n+s)} \geq \tilde{p}_N^{(n+s-1)} \quad a.s.$$

Proof. The coupling procedure in Algorithm 3 uses a single proposal p_N^* and samples $u \sim \mathcal{U}[0, 1]$, with proposals being accepted according to:

if $u \leq 1 \wedge \frac{p_N^*}{p_N^{(n)}}$ then $p_N^{(n+1)} = p_N^*$, else $p_N^{(n+1)} = p_N^{(n)}$,

if $u \leq 1 \wedge \frac{p_N^*}{\tilde{p}_N^{(n-1)}}$ then $\tilde{p}_N^{(n)} = p_N^*$, else $\tilde{p}_N^{(n)} = \tilde{p}_N^{(n-1)}$.

We see that if $p_N^{(n)} \geq \tilde{p}_N^{(n-1)}$ then either

1. $p_N^{(n+1)} = p_N^*$ in which case

$$u \leq 1 \wedge \frac{p_N^*}{p_N^{(n)}} \implies u \leq 1 \wedge \frac{p_N^*}{\tilde{p}_N^{(n-1)}}$$

as $p_N^{(n)} \geq \tilde{p}_N^{(n-1)}$ so that $p_N^{(n+1)} = \tilde{p}_N^{(n)} = p_N^*$ (i.e. the chains meet).

2. $p_N^{(n+1)} = p_N^{(n)}$ and so $p_N^* \leq p_N^{(n)}$. In this case, either $\tilde{p}_N^{(n)} = p_N^*$, and so $p_N^{(n+1)} \geq \tilde{p}_N^{(n)}$, or $\tilde{p}_N^{(n)} = \tilde{p}_N^{(n-1)}$ in which case both chains have rejected p_N^* and the ordering is preserved.

Finally, from the initialization of the procedure, we have $p_N^{(1)} \geq \tilde{p}_N^{(0)}$ because the initial state of the second chain is used as a proposal in the first iteration of the first chain. $\square$

From the above reasoning we see that the chains meet when the first chain accepts its proposal for the first time, as the second chain then necessarily accepts the same proposal.

From the initial state with likelihood estimate $p_N^{(0)}$, the acceptance probability of the first chain is $\int 1 \wedge (p_N/p_N^{(0)})\bar{g}(p_N)dp_N$, with $\bar{g}$ denoting the density of the PF likelihood estimator p_N . Thus, the time to the first acceptance follows a Geometric distribution with success probability $\int 1 \wedge (p_N/p_N^{(0)})\bar{g}(p_N)dp_N$. The result stated in Proposition 8 follows when rewriting the problem using the error of the log-likelihood estimator $\log\{p_N(y_{1:T})/p(y_{1:T})\}$.

A.2 Integrated autocorrelation time for various test functions

We show here experimentally that $\text{IF}(h)$ is approximately proportional to $\text{IF}(\sigma)$ for various test functions: $h_1 : x_{1:T} \mapsto x_1$, $h_2 : x_{1:T} \mapsto x_T$, $h_3 : x_{1:T} \mapsto \sum_t x_t$ and $h_4 : x_{1:T} \mapsto \sum_t x_t^2$. This is illustrated in Figure 7 where $\text{IF}(h)$ is displayed for a range of N against $\text{IF}(\sigma)$ over the corresponding range of σ .

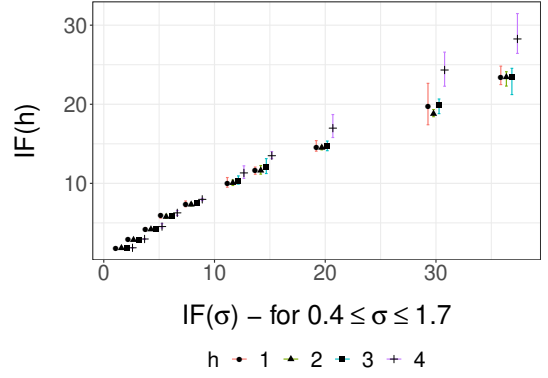


Figure 7: Inefficiency $\text{IF}[h_i]$ versus $\text{IF}[\sigma]$, with markers indicating the test functions h_1, h_2, h_3, h_4 . The vertical axis scale is relative, depending on the test function.

A.3 Rao-Blackwellisation for stochastic kinetic model

We demonstrate here the gains arising from the use of a Rao-Blackwellized estimator detailed in Section 2.2. We display in Figure 8 the variance of the two unbiased estimators of $\mathbb{E}(X_{1,\Delta t}|y_{1:T})$ for $t = \Delta, \dots, T\Delta$ and $T = 100$ for the latent Markov jump process and two different values of N , 200 and 3,000. In both cases we set $k = m = 0$. As expected $\bar{H}_{0:0}$ outperforms

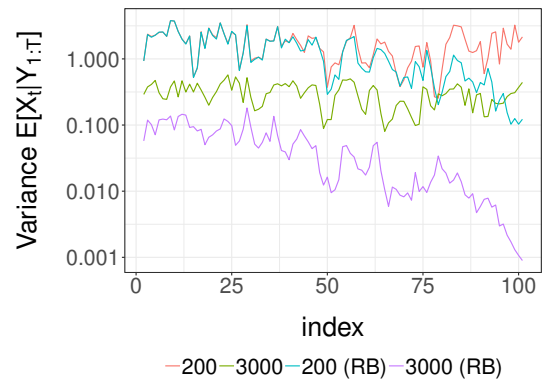
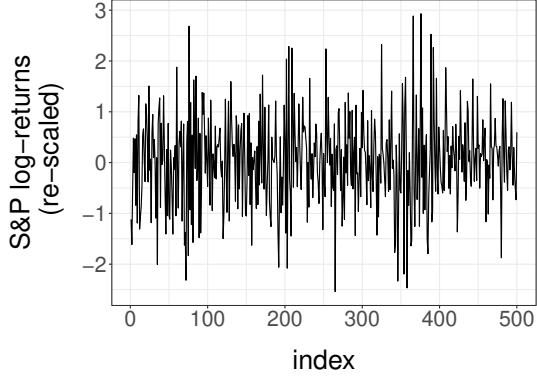
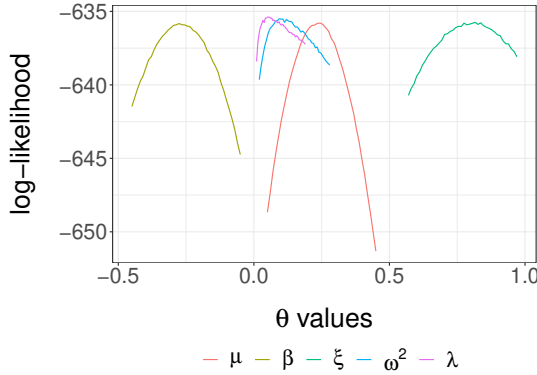


Figure 8: Empirical variance of unbiased estimators of $\mathbb{E}(X_{1,\Delta t}|y_{1:T})$: $H_{0:0}$ and Rao-Blackwellised (RB) estimator $\bar{H}_{0:0}$ for stochastic kinetic model.



(a) Daily returns of S&P 500 data



(b) Particle estimates of the log-likelihood function

Figure 9: Data and parameter estimation for Lévy-driven stochastic volatility model.

$H_{0:0}$ but the benefits are much higher for t close to T than when t is close to 1. For example, we see that for $N = 200$ the estimators $H_{0:0}$ and $\bar{H}_{0:0}$ coincide for $t \leq 35$. This is an expected consequence of the particle path degeneracy problem [13, 23, 25], with many particles $(X_{1:T}^i)_{i \in [N]}$ obtained by the PF at time T sharing common ancestors for t close to 1 when N is too small; see [23] for results on the corresponding coalescent time.

A.4 Data and parameter estimation

The data used comprised of $T = 500$ log-returns of the S&P 500, starting from the 3rd January 2005, with the scaled data taken from the stochastic volatility example used to demonstrate SMC² in [6]. We plot the raw data in Figure 9a. Parameters were estimated using a two-stage procedure, with SMC² used to find a region of high marginal likelihood under the model. Further refinement was performed to compute the maximum likelihood estimator (MLE) $\hat{\theta}$ of $\theta = (\mu, \beta, \xi, \omega^2, \lambda)$ using a grid search around values close to the optimum using $N = 10,000$ particles. We obtained $\hat{\theta} = (0.24, -0.28, 0.82, 0.09, 0.05)$. Likelihood

curves around the optimal values are shown in Figure 9b, where for each parameter component the log-likelihood was varied while keeping the other parameters fixed at $\hat{\theta}$.

B Application of unbiased estimation to SMC samplers

SMC samplers are a class of SMC algorithms that can be used in Bayesian inference to approximate expectations w.r.t. complex posteriors for static models [10]. We show here how we can directly use the methodology proposed in this paper to obtain unbiased estimators of these expectations.

B.1 Bayesian computation using SMC samplers

Assume one is interested in sampling from the posterior density $\pi(x) \propto \nu(x)L(x)$ where $\nu(x)$ the prior density w.r.t. a suitable dominating measure and $L(x)$ is the likelihood. We also assume that one can sample from ν . To approximate π , a specific version of SMC samplers introduces a sequence of $T - 1$ intermediate densities π_t for $t = 2, \dots, T$ bridging ν to π using

$$\gamma_t(x) = \nu(x)L^{\beta_t}(x), \quad \pi_t(x) = \frac{\gamma_t(x)}{\mathcal{Z}_t},$$

where $\beta_1 = 0 < \beta_2 < \dots < \beta_T = 1$. The choice of the sequence $\{\beta_t : t = 2, \dots, T - 1\}$ can be guided using a preliminary adaptive SMC scheme, as in [33], which should subsequently be fixed to preserve unbiasedness of the normalizing constant estimate and validity of the resulting PIMH. In SMC samplers, particles are initialized at time $t = 1$ by sampling from the prior ensuring $w_1(x_1) = 1$. At time $t \geq 2$, particles are sampled according to an MCMC kernel leaving π_{t-1} invariant and are then weighted according to

$$w_t(x_{t-1}, x_t) = \frac{\gamma_t(x_{t-1})}{\gamma_{t-1}(x_{t-1})} = L^{\beta_t - \beta_{t-1}}(x_{t-1}).$$

Particles are resampled according to these weights and we set

$$\mathcal{Z}_{t,N} = \mathcal{Z}_{t-1,N} \cdot \frac{1}{N} \sum_{i=1}^N w_t(X_{t-1}^{A_{t-1}^i}, X_t^i),$$

with $\mathcal{Z}_{1,N} = 1$. At time T , $\pi_N(dx_T) := \sum_{i=1}^N W_T^i \delta_{X_T^i}(dx_T)$ provides a Monte Carlo approximation of the distribution π and $\mathcal{Z}_{T,N}$ plays the role of $p_N(y_{1:T})$, approximating the normalizing constant $\mathcal{Z}_T$ of $\pi = \pi_T$. If no resampling is used, this specific version of SMC samplers coincides with AIS [30] in which case $\mathcal{Z}_{T,N}$ is given by the average of the product of the incremental weights from time $t = 1$ to

$t = T$ instead of the product of the averaged incremental weights. We can use this SMC sampler algorithm or AIS directly within the coupled PIMH scheme, replacing $p_N(y_{1:T})$ by $\mathcal{Z}_{T,N}$ in the acceptance probabilities. We see that $\sup_{(x,x') \in \mathcal{X}^2} w_t(x,x') < \infty$ provided that $\sup_{x \in \mathcal{X}} L(x) < \infty$. Under this condition, if Assumption 5 is satisfied then the estimator $\bar{H}_{k:m}$ of $\pi(h)$ is unbiased and has finite variance and finite expected cost.

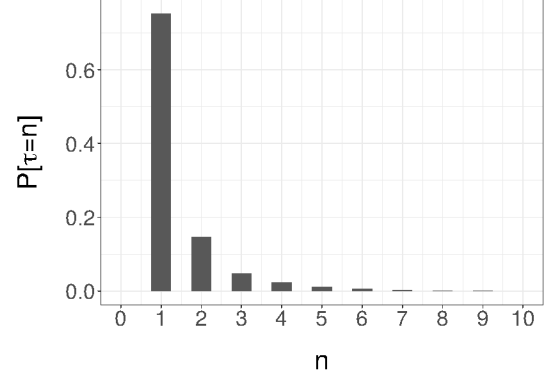
B.2 Numerical example

We use here coupled PIMH to debias expectations w.r.t. the posterior distribution for a Bayesian mixture model discussed in [28]. We have

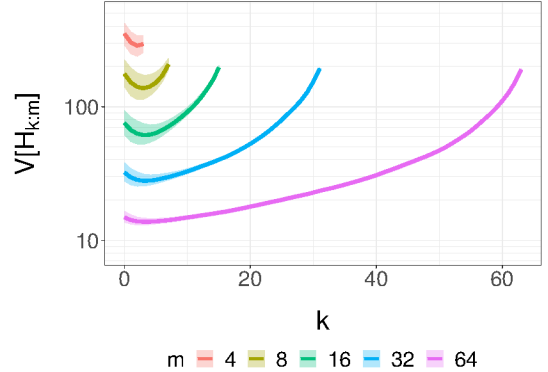
$$L(x) = \prod_{n=1}^M \left(\frac{1}{D} \sum_{i=1}^D \mathcal{N}(y_n; x^i, \sigma^2) \right)$$

with $x := (x^1, \dots, x^D) \in \mathbb{R}^D$ constituting the unknown mean components. We consider here $D = 4$ mixture components and $M = 100$ observations. A uniform prior distribution is placed on x over the hypercube $[-10, 10]^D$. The resulting posterior distribution is multimodal. We set $\sigma = 1$ and simulate observations from the model with true values $x^* = (-3, 0, 3, 6)$. We adopt a symmetric random walk for the Metropolis–Hastings proposals with identity covariance and pick $\beta_t = \left(\frac{t-1}{T-1} \right)^2$ for $T = 200$.

We simulate 10,000 estimators with $m = 64$ and $N = 100$, after which we are able to estimate variance of test functions for a range of values of k and m noting that there will be some correlation introduced between estimators. The results are shown in Figure 10 where we plot the meeting times of the unbiased estimators in Figure 10a and the variance of the estimators for a range of values of m in Figure 10b using $h : x \mapsto x^1 + x^2 + (x^1)^2 + (x^2)^2$. Uncertainty in the estimated values of $\mathbb{V}[\bar{H}_{k:m}(h)]$ was obtained using 1,000 bootstrap samples, resampling 10,000 of the 10,000 unbiased estimators with replacement. The figure shows the variance of these estimators for a range of values of $m \in \{4, 8, \dots, 64\}$ while varying $k \in \{0, \dots, m-1\}$. We see, firstly, as expected that as m increases the variance of the estimators reduces. Secondly, for each value of m we see that there exists an optimal value of k , however, as m increases the optimum becomes less pronounced, suggesting that as m increases there is a degree of insensitivity to the choice of k . Finally, for the range of m considered, the optimal values of k appear in a comparatively small interval close to the origin, suggesting that it is not necessary to use large values of k to reduce the variance contribution arising from the bias correction.



(a) Empirical distribution of meeting times for SMC sampler with $N = 100$ over 10,000 independent runs. The estimated 95th and 99th percentiles were 6 and 13 respectively.



(b) $\mathbb{V}[\bar{H}_{k:m}]$ of Rao-Blackwellised unbiased estimators of $\pi(h)$ for SMC sampler as a function of k for a range of values of m . The shaded regions correspond to the 1st and 99th percentiles of the variance estimator.

Figure 10: SMC sampler unbiased estimators