

An explicit variance reduction expression for the Rao-Blackwellised particle filter

Fredrik Lindsten* Thomas B. Schön* Jimmy Olsson**

* *Division of Automatic Control, Linköping University, Linköping, Sweden (e-mail: {lindsten, schon}@isy.liu.se)*

** *Center of Mathematical Sciences, Lund University, Lund, Sweden (e-mail: jimmy@maths.lth.se)*

Abstract: Particle filters (PFs) have shown to be very potent tools for state estimation in nonlinear and/or non-Gaussian state-space models. For certain models, containing a conditionally tractable substructure (typically conditionally linear Gaussian or with finite support), it is possible to exploit this structure in order to obtain more accurate estimates. This has become known as Rao-Blackwellised particle filtering (RBPF). However, since the RBPF is typically more computationally demanding than the standard PF per particle, it is not always beneficial to resort to Rao-Blackwellisation. For the same computational effort, a standard PF with an increased number of particles, which would also increase the accuracy, could be used instead. In this paper, we have analysed the asymptotic variance of the RBPF and provide an explicit expression for the obtained variance reduction. This expression could be used to make an efficient discrimination of when to apply Rao-Blackwellisation, and when not to.

Keywords: Particle filtering, Monte-Carlo methods, Rao-Blackwellised particle filter, Marginalised particle filter, Rao-Blackwellisation, Variance reduction.

1. INTRODUCTION AND RELATED WORK

Many important problems in various fields of science are related to state estimation in general state-space models, based on noisy observations. If a prior distribution is assumed for the initial state, the optimal filter is given by the Bayesian filtering recursions. In a few special cases, basically for linear Gaussian state-space (LGSS) models and finite state-space (FSS) models, the optimal filtering problem is analytically tractable. However, many interesting problems do not exhibit such nice properties, but are both nonlinear and/or non-Gaussian. In these cases, the optimal filter needs to be approximated in some way. Sequential Monte Carlo methods, or particle filters (PFs), have shown to be very powerful tools when addressing such intractable models. Since the introduction of the PF by Gordon et al. (1993), we have experienced a vast amount of research in the area. For instance, many improvements and extensions have been introduced to increase the accuracy of the filter, see e.g. Doucet and Johansen (2011) for an overview of recent developments.

One natural idea is to exploit any tractable substructure in the model, see e.g. Doucet et al. (2000b); Schön et al. (2005). More precisely, if the model, conditioned on one partition of the state, behaves like e.g. an LGSS or an FSS it is sufficient to employ particles for the intractable part and make use of the analytic tractability for the remaining part. Inspired by the Rao-Blackwell theorem,

this has become known as the Rao-Blackwellised particle filter (RBPF).

The motivation for the RBPF is to improve the accuracy of the filter, i.e. any estimator derived from the RBPF will intuitively have lower variance than the corresponding estimator derived from the standard PF. Informally, the reason for this is that in the RBPF, the particles are spread in a lower dimensional space, yielding a denser particle representation of the underlying distribution. The improved accuracy is also something that is experienced by practitioners. However, it can be argued that it is still not beneficial to resort to Rao-Blackwellisation in all cases. The reason is that the RBPF is typically more computationally expensive per iteration, compared to the standard PF (e.g. for an RBPF targeting a conditional LGSS model, each particle is equipped with a Kalman filter, all which needs to be updated at each iteration). Hence, for a fixed computational effort, we can choose to either use Rao-Blackwellisation or to run a standard PF, but instead increase the number of particles. Both these alternatives will reduce the variance of the estimators. Hence, it is important to understand and to be able to quantify how large variance reduction we can expect from the RBPF, in order to make suitable design choices for any given problem.

In this paper we shall study the asymptotic (in the number of particles) variances for the RBPF and the standard PF (we shall throughout the paper use the abbreviation SPF when referring to the non-Rao-Blackwellised PF). We provide an explicit expression for the difference between the variance of an estimator derived from the SPF and the variance of the corresponding estimator derived from

* This work was supported by: the project Calibrating Nonlinear Dynamical Models (Contract number: 621-2010-5876) funded by the Swedish Research Council and CADICS, a Linneaus Center also funded by the Swedish Research Council.

the RBPF. Doucet et al. (2000b) motivates the RBPF by concluding that the weight variance will be lower than for the SPF, but they do not consider the variances of any estimators. This is done by Chopin (2004), who, under certain assumptions, concludes that the variance of an estimator based on the SPF always is at least as high as for the RBPF. However, no explicit expression for the difference is given, and the test functions considered are restricted to one partition of the state-space. Doucet et al. (2000a) also analyse the RBPF and the reduction of asymptotic variance. However, they only consider an importance sampling setting and neglect the important resampling step. Karlsson et al. (2005) studies the problem empirically, by running simulations on a specific example. Here, they have also analysed the number of computations per iteration in the RBPF and SPF, respectively.

The rest of this paper is organised as follows. In Section 2 we present the PF framework and define the SPF and the RBPF. We thereafter introduce the natural estimators that can be derived from the two filters, and discuss their asymptotic properties in Section 3. The main result is given in Section 4, where we provide an explicit expression for the difference in asymptotic variance for the two estimators. In Section 5 this expression is studied for a special case and in Section 6 we discuss how it can be used in the design of a PF. Finally, in Section 7 we draw conclusions.

2. BACKGROUND

2.1 Notation

All random variables are defined on a common probability space $(\Omega, \mathcal{F}, P)$. For a measurable space $(X, \mathcal{X})$, we denote by $\mathbb{F}(X)$ the set of all $\mathcal{X}$ -measurable functions from X to $\mathbb{R}$. For a measure μ on $\mathcal{X}$ and $f \in \mathbb{F}(X)$, satisfying $\int |f| d\mu < \infty$, we denote by $\mu(f)$ the integral $\int f d\mu$. For $p \geq 1$, $L^p(X, \mu)$ denotes the set of functions $f \in \mathbb{F}(X)$ such that $\int |f|^p d\mu < \infty$.

Let $(X, \mathcal{X})$ and $(Y, \mathcal{Y})$ be two measurable spaces. A kernel V from X to Y is a map $V : (X, \mathcal{X}) \rightarrow \mathbb{R}_+$, such that (i) for each $x \in X$, the map $A \mapsto V(x, A)$ is a measure on $\mathcal{Y}$ (ii) for each $A \in \mathcal{Y}$, the map $x \mapsto V(x, A)$ is $\mathcal{X}$ -measurable. A kernel is called a *transition kernel* if $V(x, Y) = 1$ for any $x \in X$. We shall sometimes write $V(A | x)$ instead of $V(x, A)$. With $f : X \times Y \rightarrow \mathbb{R}$, $f(x, \cdot) \in L^1(Y, V(x, \cdot))$ we let $V(f)$ denote the function $V(f)[x] = \int f(x, y) V(x, dy)$.

For two measures μ and ν , we say that ν is absolutely continuous with respect to μ (written $\nu \ll \mu$) if $\mu(A) = 0 \Rightarrow \nu(A) = 0$. By $\mathcal{N}(m, \sigma^2)$ we denote the Gaussian distribution with mean m and variance σ^2 . Finally, convergence in distribution and convergence in probability are denoted by $\xrightarrow{D}$ and $\xrightarrow{P}$, respectively.

2.2 Particle filtering

Let $\{X_t, t \in \mathbb{N}\}$ be a discrete-time Markov process on the state-space $(X, \mathcal{X})$ (typically some power of the real line with the corresponding Borel σ -algebra). The process is hidden, but observed through the measurement sequence $\{Y_t, t \in \mathbb{N}\}$ defined on $(Y, \mathcal{Y})$. For a given sequence of

measurements $Y_{1:t} \triangleq \{Y_1, \dots, Y_t\}$ (a similar notation shall be used for other sequences as well), the joint smoothing distribution is defined as

$$\phi_t(A) \triangleq P(X_{1:t} \in A | Y_{1:t}), \quad A \in \mathcal{X}^t. \quad (1)$$

In particle filtering, we seek to approximate this distribution by a weighted particle system $\{x_{1:t}^{N,i}, w_t^{N,i}\}_{i=1}^N$. Observe that this is a triangular array of random variables and, though cumbersome, it is important to keep the dependence on N in the notation. There are many different ways to generate such a particle system (see e.g. Doucet and Johansen (2011) for an overview), but common to most are the concepts of sequential importance sampling and resampling. Importance sampling is used to propagate a weighted sample, targeting the smoothing distribution at time $t-1$, into a weighted sample targeting the smoothing distribution at time t . This is done by sampling new particles from a proposal kernel (a transition kernel from X^{t-1} to X^t),

$$R_{t-1}(\tilde{x}_{1:t-1}, dx_{1:t}). \quad (2)$$

The proposal kernel is chosen such that $\phi_t \ll \pi_t$, where we have defined the probability measure on $\mathcal{X}^t$,

$$\pi_t(dx_{1:t}) \triangleq \int_{X^{t-1}} \phi_{t-1}(d\tilde{x}_{1:t-1}) R_{t-1}(\tilde{x}_{1:t-1}, dx_{1:t}). \quad (3)$$

Besides from the absolute continuity condition, we shall not assume any restrictions on the choice of proposal kernel. For instance, it can depend on the measurement sequence, but we shall not make such dependence explicit. The distribution π_t can be seen as the “proposed smoothing distribution” at time t under the proposal kernel R_{t-1} .

To compensate for the fact that we sample from the wrong distribution, the samples are weighted. For this purpose we introduce the weight function

$$W_{t-1}(x_{1:t}) \triangleq \frac{d\phi_t}{d\pi_t}(x_{1:t}). \quad (4)$$

A well known problem in particle filtering is weight depletion, see e.g. Doucet et al. (2000b); Doucet and Johansen (2011). To remedy this the particle system can be resampled, i.e. particles with high weights are duplicated and particles with low weights are discarded. As previously mentioned, there are many options available to the user when designing a particle filter, e.g. in the choice of resampling scheme, if the number of particles shall be fixed or varying and when the resampling shall be performed. In this paper, the results will be given for the particle filter presented in Algorithm 1 below. Briefly, we consider arbitrary proposal kernels (2) (under the absolute continuity assumption) and multinomial resampling which is performed at each iteration of the algorithm. However, results similar to those presented in Section 4 could be obtained for other types of PFs as well, such as the auxiliary PF (Pitt and Shephard, 1999) and PFs with more sophisticated resampling schemes.

As previously mentioned, the PF presented in this section, i.e. which targets the “full” joint smoothing distribution (1), will throughout the remainder of this paper be denoted the standard particle filter (SPF).

Algorithm 1 Particle filter

Input: A weighted sample $\{x_{1:t-1}^{N,i}, w_{t-1}^{N,i}\}_{i=1}^N$ targeting ϕ_{t-1} .

Resampling: Sample N indices from a discrete distribution, i.e. for $i = 1, \dots, N$,

$$P(I^{N,i} = j \mid \{x_{1:t-1}^{N,k}, w_{t-1}^{N,k}\}_{k=1}^N) = w_{t-1}^{N,j} / \sum_{l=1}^N w_{t-1}^{N,l}.$$

Set $\tilde{x}_{1:t-1}^{N,i} = x_{1:t-1}^{N,I^{N,i}}$, $i = 1, \dots, N$. The equally weighted sample $\{\tilde{x}_{1:t-1}^{N,i}, 1\}_{i=1}^N$ targets ϕ_{t-1} .

Importance sampling: Choose a proposal kernel according to (2). Sample new particles according to

$$x_{1:t}^{N,i} \sim R_{t-1}(\tilde{x}_{1:t-1}^{N,i}, dx_{1:t}), \quad i = 1, \dots, N.$$

Compute the weights, using the weight function (4),

$$w_t^{N,i} = W_{t-1}(x_{1:t}^{N,i}).$$

Output: A weighted sample $\{x_{1:t}^{N,i}, w_t^{N,i}\}_{i=1}^N$ targeting ϕ_t .

2.3 Rao-Blackwellised particle filter

The idea underlying the Rao-Blackwellised particle filter is to exploit any tractable substructure in the targeted Markov process, if such a structure is present. In doing so, one can hope to obtain better particle approximations for a fixed number of particles. Hence, assume that each X_t can be partitioned according to $X_t = \{\Xi_t, Z_t\}$ and $\mathbf{X} = \mathbf{X}_\xi \times \mathbf{X}_z$. Furthermore, assume that ϕ_t factorises according to

$$\phi_t(dx_{1:t}) = \phi_t^m(d\xi_{1:t})\phi_t^c(dz_{1:t} \mid \xi_{1:t}), \quad (5)$$

where $\{\xi_t, z_t\}$ identifies to x_t . Here, ϕ_t^m is the marginal smoothing distribution of $\Xi_{1:t}$ and ϕ_t^c is the conditional smoothing distribution of $Z_{1:t}$ given $\Xi_{1:t} = \xi_{1:t}$. The conditional distribution is assumed to be analytically tractable, typically Gaussian or with finite support.

Remark 1. More precisely, ϕ_t^c is a kernel from $\mathbf{X}_\xi^t$ to $\mathbf{X}_z^t$. For each fixed $\xi_{1:t}$, $\phi_t^c(\cdot \mid \xi_{1:t})$ is a measure on $\mathbf{X}_z^t$, and it can hence be viewed as a conditional distribution. In the notation introduced in (5), the meaning is that ϕ_t is the product of the measure ϕ_t^m and the kernel ϕ_t^c . In the remainder of this paper we shall make frequent use of a Fubini like theorem for such products, see e.g. Uglov (1991).

Instead of running the SPF targeting the “full” joint smoothing distribution (1), we have the option to target the marginal distribution ϕ_t^m with a PF and then make use of an analytical expression for ϕ_t^c . Hence, we choose a proposal kernel $R_{t-1}^m(\xi_{1:t-1}, d\xi_{1:t})$ (from $\mathbf{X}_\xi^{t-1}$ to $\mathbf{X}_\xi^t$) such that $\phi_t^m \ll \pi_t^m$ and define a weight function $W_{t-1}^m(\xi_{1:t})$ analogously to (4). The measure π_t^m is defined analogously to (3).

A weighted particle system $\{\xi_{1:t}^{N,i}, \omega_t^{N,i}\}_{i=1}^N$, targeting ϕ_t^m , can then be generated in the same manner as in Algorithm 1. We simply replace $\{x_{1:t-1}^{N,i}, w_{t-1}^{N,i}\}_{i=1}^N$ with $\{\xi_{1:t-1}^{N,i}, \omega_{t-1}^{N,i}\}_{i=1}^N$ and ϕ, R, W with ϕ^m, R^m, W^m , respectively (again, superscript m for marginal). This will be referred to as the Rao-Blackwellised particle filter (RBPF).

Remark 2. The most common way to present the RBPF is for conditional LGSS models. In this case, the conditional distribution ϕ_t^c is Gaussian, which means that it can be computed using the Kalman filter recursions. Con-

sequently, the Kalman filter updates are often shown as intrinsic steps in the presentation of the RBPF algorithm, see e.g. Schön et al. (2005). In this paper, we adopt a more general view and simply see the RBPF as a regular PF, targeting the marginal distribution ϕ_t^m . We then assume that the conditional distribution ϕ_t^c is available by some means (for the conditional LGSS case, this would of course be by the Kalman filter), but it is not important for our results what those means are.

3. PROBLEM FORMULATION

The SPF and the RBPF can both be used to estimate expectations under the joint smoothing distribution. Assume that we, for some function $f \in L^1(\mathbf{X}^t, \phi_t)$, seek the expectation $\phi_t(f)$. For the SPF we use the natural estimator,

$$\hat{f}_{\text{SPF}}^N \triangleq \sum_{i=1}^N \frac{w_t^{N,i}}{\sum_{j=1}^N w_t^{N,j}} f(x_{1:t}^{N,i}). \quad (6)$$

For the RBPF we use the fact that $\phi_t(f) = \phi_t^m(\phi_t^c(f))$, and define the estimator,

$$\hat{f}_{\text{RBPF}}^N \triangleq \sum_{i=1}^N \frac{\omega_t^{N,i}}{\sum_{j=1}^N \omega_t^{N,j}} \phi_t^c \left(f(\{\xi_{1:t}^{N,i}, \cdot\}) \mid \xi_{1:t}^{N,i} \right). \quad (7)$$

The question then arise, how much better is (7) compared to (6)?

One analysis of this question, sometimes seen in the literature, is to simply consider a decomposition of variance,

$$\underbrace{\text{Var}(f)}_{\text{SPF}} = \underbrace{\text{Var}(\mathbb{E}[f \mid \Xi_{1:t}])}_{\text{RBPF}} + \underbrace{\mathbb{E}[\text{Var}(f \mid \Xi_{1:t})]}_{\geq 0}. \quad (8)$$

Here, the last term is claimed to be the variance reduction obtained in the RBPF. The decomposition is of course valid, the problem is that it does not answer our question. What we have in (8) is simply an expression for the variance of the test function f , it does not apply to the estimators (6) and (7).

Remark 3. It is not hard to see why the “simplified” analysis (8) has been considered. If the PF would produce independent and identically distributed (i.i.d.) samples from the target distribution (which it does not), then the analysis would be correct. More precisely, for i.i.d. samples, the central limit theorem states that the asymptotic variance of an estimator of a test function f , coincides with the variance of the test function itself (up to a factor $1/N$). However, as we have already pointed out, the PF does not produce i.i.d. samples. This is due to the resampling step, in which a dependence between the particles is introduced. At the end of Section 6, one of the inadequacies of (8) will be pointed out.

Hence, we are interested in the asymptotic variance of (6) and (7), respectively. To analyse this we shall borrow the concept of asymptotic normality from Douc and Moulines (2008).

Definition 1. (Asymptotic normality) Let $(\mathbf{X}, \mathcal{X})$ be a measurable space, $\mathbf{A}$ and $\mathbf{W}$ subsets of $\mathbb{F}(\mathbf{X})$, μ a probability measure and γ a finite measure on $\mathcal{X}$. Let σ be a real nonnegative function on $\mathbf{A}$ and $\{a_N\}_{N=1}^\infty$ a nondecreasing real sequence diverging to infinity.

A weighted sample $\{\chi^{N,i}, v^{N,i}\}_{i=1}^N$ is said to be asymptotically normal for $(\mu, A, W, \sigma, \gamma, \{a_N\})$ if

$$a_N \Omega_N^{-1} \sum_{i=1}^N v^{N,i} (f(\chi^{N,i}) - \mu(f)) \xrightarrow{D} \mathcal{N}(0, \sigma^2(f)), \quad (9a)$$

$$a_N^2 \Omega_N^{-2} \sum_{i=1}^N (v^{N,i})^2 g(\chi^{N,i}) \xrightarrow{P} \gamma(g), \quad (9b)$$

$$a_N \Omega_N^{-1} \max_{1 \leq i \leq N} v^{N,i} \xrightarrow{P} 0, \quad (9c)$$

as $N \rightarrow \infty$, for any $f \in A$ and any $g \in W$, where $\Omega_N = \sum_{j=1}^N v^{N,j}$.

In the following two theorems (slight modifications of what has previously been given by Douc and Moulines (2008)) we claim asymptotic normality for the weighted particle systems generated by the SPF and the RBPF, respectively.

Theorem 1. (Asymptotic normality of the SPF) Assume that the initial particle system $\{x_1^{N,i}, w_1^{N,i}\}_{i=1}^N$ is asymptotically normal for $(\phi_1, A_1, W_1, \sigma_1, \phi_1, \{\sqrt{N}\})$. Define recursively the sets

$$A_t \triangleq \{f \in L^2(X^t, \phi_t) : R_{t-1}(\cdot, W_{t-1}f) \in A_{t-1}, \\ R_{t-1}(\cdot, W_{t-1}^2 f^2) \in W_{t-1}\},$$

$$W_t \triangleq \{f \in L^1(X^t, \phi_t) : R_{t-1}(\cdot, W_{t-1}^2 |f|) \in W_{t-1}\}.$$

Assume that, for any $t \geq 1$, $R_t(\cdot, W_t^2) \in W_t$. Then, for any $t \geq 1$, the weighted particle system $\{x_{1:t}^{N,i}, w_t^{N,i}\}_{i=1}^N$ generated by the SPF is asymptotically normal for $(\phi_t, A_t, W_t, \sigma_t, \phi_t, \{\sqrt{N}\})$. The asymptotic variance is, for $f \in A_t$, given by

$$\sigma_t^2(f) = \sigma_{t-1}^2(R_{t-1}(\cdot, W_{t-1}\bar{f})) \\ + \phi_{t-1}[R_{t-1}(\cdot, (W_{t-1}\bar{f})^2)], \quad (10)$$

$$\bar{f} = f - \phi_t(f).$$

Proof. See Appendix A.

Theorem 2. (Asymptotic normality of the RBPF) Under analogous conditions and definitions as in Theorem 1, for any $t \geq 1$ the particle system $\{\xi_{1:t}^{N,i}, \omega_t^{N,i}\}_{i=1}^N$ generated by the RBPF is asymptotically normal for $(\phi_t^m, A_t^m, W_t^m, \tau_t, \phi_t^m, \{\sqrt{N}\})$. The asymptotic variance is, for $g \in A_t^m$, given by

$$\tau_t^2(g) = \tau_{t-1}^2(R_{t-1}^m(\cdot, W_{t-1}^m \bar{g})) \\ + \phi_{t-1}^m[R_{t-1}^m(\cdot, (W_{t-1}^m \bar{g})^2)], \quad (11)$$

$$\bar{g} = g - \phi_t^m(g).$$

Proof. See Appendix A.

Recall from Remark 2 that the SPF and the RBPF are really just two particle filters, targeting different distributions, hence the similarity between the two theorems above. Actually, we could have given one, more general, theorem applicable to both filters. The reason for why we have chosen to present them separately is for clarity and to introduce all the required notation.

As previously pointed out, the RBPF will intuitively produce better estimates than the SPF, i.e. we expect $\tau_t^2(\phi_t^c(f)) \leq \sigma_t^2(f)$. Let us therefore define the variance difference

$$\Delta_t(f) \triangleq \sigma_t^2(f) - \tau_t^2(\phi_t^c(f)). \quad (12)$$

The problem that we are concerned with in this paper is to find an explicit expression for this quantity. This will be provided in the next section.

4. THE MAIN RESULT

To analyse the variance difference (12) we shall need the following assumption (similar to what is used by Chopin (2004)).

Assumption 1. For each $\tilde{\xi}_{1:t-1} \in X_\xi^{t-1}$, the two measures

$$\int_{X_\xi^{t-1}} \phi_{t-1}^c(d\tilde{z}_{1:t-1} | \tilde{\xi}_{1:t-1}) R_{t-1}(\{\tilde{\xi}_{1:t-1}, \tilde{z}_{1:t-1}\}, dx_{1:t}) \quad (13a)$$

and

$$a_t(\xi_{1:t}) R_{t-1}^m(\tilde{\xi}_{1:t-1}, d\xi_{1:t}) \pi_t^c(dz_{1:t} | \xi_{1:t}) \quad (13b)$$

agree on X^t , for some positive function $a_t : X_\xi^t \rightarrow \mathbb{R}$ and some transition kernel π_t^c from X_ξ^t to X_z^t , for which $\phi_t^c(\cdot | \xi_{1:t}) \ll \pi_t^c(\cdot | \xi_{1:t})$.

The basic meaning of this assumption is to create a connection between the proposal kernels R_{t-1} and R_{t-1}^m . It is natural that we need some kind of connection. Otherwise the asymptotic variance expressions (10) and (11) would be completely decoupled, and it would not be possible to draw any conclusions from a comparison. Still, as we shall see in the next section, Assumption 1 is fairly weak.

We are now ready to state the main result of this paper.

Theorem 3. Under Assumption 1, and using the definitions from Theorem 1 and Theorem 2, for any $f \in \hat{A}_t$,

$$\Delta_t(f) = \Delta_{t-1}(R_{t-1}(\cdot, W_{t-1}\bar{f})) \\ + \phi_{t-1}^m \left[R_{t-1}^m \left(\cdot, \left(\frac{1-a_t}{a_t} \right) (W_{t-1}^m \bar{\psi})^2 + a_t \text{Var}_{\pi_t^c}(W_{t-1}\bar{f}) \right) \right], \quad (14)$$

where

$$\bar{\psi} = \phi_t^c(f) - \phi_t(f), \quad (15a)$$

$$\hat{A}_t \triangleq \{f \in \mathbb{F}(X^t) : \phi_t^c(f) \in A_t^m\} \cap A_t. \quad (15b)$$

Proof. See Appendix A.

5. RELATIONSHIP BETWEEN THE PROPOSALS KERNELS

To understand the results given in the previous section, we shall have a closer look at the relationship between the proposal kernels imposed by Assumption 1. We shall do this for a certain family of proposal kernels. More precisely, assume that the kernels can be written

$$R_{t-1}(\tilde{x}_{1:t-1}, dx_{1:t}) = r_{t-1}(dx_t | x_{1:t-1}) \delta_{\tilde{x}_{1:t-1}}(dx_{1:t-1}), \quad (16a)$$

$$R_{t-1}^m(\tilde{\xi}_{1:t-1}, d\xi_{1:t}) = r_{t-1}^m(d\xi_t | \xi_{1:t-1}) \delta_{\tilde{\xi}_{1:t-1}}(d\xi_{1:t-1}). \quad (16b)$$

Informally, this means that when a trajectory $(x_{1:t}^{N,i}$ or $\xi_{1:t}^{N,i})$ is sampled at time t , we keep the “old” trajectory up to time $t-1$ and simply append a sample from time t . This is the case for most PFs (when targeting the joint smoothing distribution), but not all, see e.g. the resample-move algorithm by Gilks and Berzuini (2001).

Furthermore, let r_{t-1} be factorised as

$$r_{t-1}(dx_t | x_{1:t-1}) = q_{t-1}^c(dz_t | \xi_{1:t}, z_{1:t-1}) q_{t-1}^m(d\xi_t | \xi_{1:t-1}, z_{1:t-1}). \quad (16c)$$

Assume that $q_{t-1}^m \ll r_{t-1}^m$ and define the kernel

$$\nu_t(dz_{1:t} | \xi_{1:t}) \triangleq \frac{dq_{t-1}^m(\cdot | \xi_{1:t-1}, z_{1:t-1})}{dr_{t-1}^m(\cdot | \xi_{1:t-1})}(\xi_t) \times \phi_{t-1}^c(dz_{1:t-1} | \xi_{1:t-1}) q_{t-1}^c(dz_t | \xi_{1:t}, z_{1:t-1}). \quad (17)$$

It can now be verified that the choice

$$a_t(\xi_{1:t}) = \int_{\mathcal{X}_z^t} \nu_t(dz_{1:t} | \xi_{1:t}), \quad (18)$$

$$\pi_t^c(dz_{1:t} | \xi_{1:t}) = \frac{\nu_t(dz_{1:t} | \xi_{1:t})}{a_t(\xi_{1:t})}, \quad (19)$$

satisfies Assumption 1, given that $\phi_t^c(\cdot | \xi_{1:t}) \ll \pi_t^c(\cdot | \xi_{1:t})$.

Hence, the function a_t takes the role of a normalisation of the kernel ν_t to obtain a transition kernel π_t^c . One interesting fact is that, from (14), we cannot guarantee that $\Delta_t(f)$ is nonnegative for arbitrary functions a_t . At first this might seem counterintuitive, since it would mean that the variance is higher for the RBPF than for the SPF. The explanation lies in that Assumption 1, relating the proposal kernels in the two filters, is fairly weak. In other words, we have not assumed that the proposal kernels are “equally good”. For instance, say that the optimal proposal kernel is used in the SPF, whereas the RBPF uses a poor kernel. It is then no longer clear that the RBPF will outperform the SPF. In the next section we shall see that if both filters use their respective bootstrap proposal kernel, a case when the term “equally good” makes sense, then $\Delta_t(f)$ will indeed be nonnegative. However, for other proposal kernels, it is not clear that there is an analogue between the SPF and the RBPF in the same sense.

5.1 Example: Bootstrap kernels

Let $Q(dx_t | x_{t-1})$ be the Markov transition kernel of the process X . In the bootstrap SPF we choose the proposal kernel according to (16a) with

$$r_{t-1}(dx_t | x_{1:t-1}) = Q(dx_t | x_{t-1}), \quad (20)$$

where, for $A \in \mathcal{X}$,

$$\begin{aligned} Q(A | X_{t-1}) &= P(X_t \in A | X_{t-1}) \\ &= P(X_t \in A | X_{1:t-1}, Y_{1:t-1}). \end{aligned} \quad (21)$$

The second equality follows from the Markov property of the process. In the RBPF, the analogue of the bootstrap proposal kernel is to use (16b) with

$$r_{t-1}^m(A | \Xi_{1:t-1}) = P(\Xi_t \in A | \Xi_{1:t-1}, Y_{1:t-1}), \quad (22)$$

for $A \in \mathcal{X}_\xi$.

It can be verified that these choices fulfill Assumption 1 with

$$a_t \equiv 1, \quad (23)$$

and

$$\pi_t^c(A | \Xi_{1:t}) = P(Z_{1:t} \in A | \Xi_{1:t}, Y_{1:t-1}), \quad (24)$$

for $A \in \mathcal{X}_z^t$. Hence, π_t^c is indeed the predictive distribution of $Z_{1:t}$ conditioned on $\Xi_{1:t}$ and based on the measurements up to time $t-1$. In this case we can also write $\pi_t(dx_{1:t}) = \pi_{t-1}^m(d\xi_{1:t})\pi_t^c(dz_{1:t} | \xi_{1:t})$, which highlights the connection between the predictive distributions in the two filters. In

this case, due to (23), the variance difference (14) can be simplified to

$$\begin{aligned} \Delta_t(f) &= \Delta_{t-1}(R_{t-1}(\cdot, W_{t-1}\bar{f})) \\ &\quad + \phi_{t-1}^m[R_{t-1}^m(\cdot, \text{Var}_{\pi_t^c}(W_{t-1}\bar{f}))]. \end{aligned} \quad (25)$$

Hence, $\Delta_t(f)$ can be written as a sum (though, we have expressed it in a recursive form here) in which each term is an expectation of a conditional variance. It is thus ensured to be nonnegative.

6. DISCUSSION

In Theorem 3 we gave an explicit expression for the difference in asymptotic variance between the SPF and the RBPF. This expression can be used as a guideline for when it is beneficial to apply Rao-Blackwellisation, and when it is not. The variance expressions given in this paper are asymptotic. Consequently, they do not apply exactly to the variances of the estimators (6) and (7), for a finite number of particles. Still, a reasonable approximation of the accuracy of the estimator (6) is

$$\text{Var}(\hat{f}_{\text{SPF}}^N) \approx \frac{\sigma_t^2(f)}{N}, \quad (26)$$

and similarly for (7)

$$\text{Var}(\hat{f}_{\text{RBPF}}^N) \approx \frac{\tau_t^2(\phi_t^c(f))}{N}. \quad (27)$$

Now, assume that the computational effort required by the RBPF, using M particles, equals that required by the SPF, using N particles (thus, $M < N$ since, in general, the RBPF is more computationally demanding than the SPF per particle). We then have

$$\frac{\text{Var}(\hat{f}_{\text{SPF}}^N)}{\text{Var}(\hat{f}_{\text{RBPF}}^M)} \approx \frac{M}{N} \left(1 + \frac{\Delta_t(f)}{\tau_t^2(\phi_t^c(f))} \right). \quad (28)$$

Whether or not this quantity is greater than one tells us if it is beneficial to use Rao-Blackwellisation. The crucial point is then to compute the ratio $\Delta_t(f)/\tau_t^2(\phi_t^c(f))$, which in itself is a challenging problem. One option is to apply an RBPF to estimate this ratio, but to sort out the details of how this can be done is a topic for future work.

As a final remark, for the special case discussed in Section 5.1, the variance difference (25) resembles the last term in (8). They are both composed of an expectation of a conditional variance. One important difference though, is that the dependence on the weight function W_{t-1} is visible in (25). As an example, if the test function is restricted to $f \in \mathcal{L}^1(\mathcal{X}_\xi^t, \phi_t^m)$ the gain in variance indicated by (8) would be zero (since $\text{Var}(f(\Xi_{1:t}) | \Xi_{1:t}) \equiv 0$), but this is not the case for the actual gain (25).

7. CONCLUSIONS

We have analysed the Rao-Blackwellised particle filter in a fairly general setting, and provide an explicit expression for the reduction of asymptotic variance obtained from Rao-Blackwellisation. This expression is expected to be of practical use, since it can serve as an indicator for when it is beneficial to apply Rao-Blackwellisation, and when it is not. We are currently investigating efficient methods, based on the analytical expression, for estimating the obtained variance reduction.

Appendix A. PROOFS

Proof of Theorem 1. In Theorem 10 in Douc and Moulines (2008), take

$$L_{t-1}(\tilde{x}_{1:t-1}, dx_{1:t}) = W_{t-1}(x_{1:t})R_{t-1}(\tilde{x}_{1:t-1}, dx_{1:t}), \quad (\text{A.1})$$

which satisfies the conditions of the hypothesis. Furthermore, $\phi_{t-1}L_{t-1}(X^t) = 1$. Now, the results follow for the choice $\kappa = 0$. ■

In Theorem 10 in Douc and Moulines (2008), the asymptotic normality of a particle system obtained after resampling is considered. Compared to Theorem 1 of this paper, they thus obtain an additional term in the expression for the asymptotic variance.

Proof of Theorem 2. As the previous proof, with

$$L_{t-1}^m(\tilde{\xi}_{1:t-1}, d\xi_{1:t}) = W_{t-1}^m(\xi_{1:t})R_{t-1}^m(\tilde{\xi}_{1:t-1}, d\xi_{1:t}). \quad (\text{A.2})$$

Proof of Theorem 3. Let Assumption 1 be satisfied. Consider

$$\begin{aligned} \phi_t(dx_{1:t}) &= \int \phi_{t-1}(d\tilde{x}_{1:t-1})W_{t-1}(x_{1:t})R_{t-1}(\tilde{x}_{1:t-1}, dx_{1:t}) \\ &= a_t(\xi_{1:t})W_{t-1}(x_{1:t})\pi_t^m(d\xi_{1:t})\pi_t^c(dz_{1:t} | \xi_{1:t}), \end{aligned} \quad (\text{A.3})$$

where we have used (3) and (4) for the first equality, and (5) and Assumption 1 for the second equality (recall that π_t^m is defined analogously to (3)).

However, we may also write

$$\phi_t = \frac{d\phi_t^m}{d\pi_t^m} \frac{d\phi_t^c}{d\pi_t^c} \pi_t^m \pi_t^c. \quad (\text{A.4})$$

Hence, we have two candidates for the Radon-Nikodym derivative of ϕ_t with respect to $\pi_t^m \pi_t^c$ which, $\pi_t^m \pi_t^c$ -almost surely, implies

$$a_t(\xi_{1:t})W_{t-1}(x_{1:t}) = W_{t-1}^m(\xi_{1:t}) \frac{d\phi_t^c(\cdot | \xi_{1:t})}{d\pi_t^c(\cdot | \xi_{1:t})}(z_{1:t}). \quad (\text{A.5})$$

Consider arbitrary $\varphi \in \tilde{\mathbf{A}}_t$. Using (5) and Assumption 1 we may write

$$\phi_{t-1}[R_{t-1}(\cdot, \varphi)] = \phi_{t-1}^m[R_{t-1}^m(\cdot, a_t\pi_t^c(\varphi))]. \quad (\text{A.6})$$

Comparing (A.6) and (10), we see that we can let φ take the role of $(W_{t-1}\bar{f})^2$. Hence, consider

$$\pi_t^c((W_{t-1}\bar{f})^2) = (\pi_t^c(W_{t-1}\bar{f}))^2 + \text{Var}_{\pi_t^c}(W_{t-1}\bar{f}), \quad (\text{A.7})$$

where, using (A.5) we have π_t^m -almost surely,

$$\begin{aligned} \pi_t^c(W_{t-1}\bar{f}) &= \int \frac{W_{t-1}^m(\xi_{1:t})}{a_t(\xi_{1:t})} \frac{d\phi_t^c(\cdot | \xi_{1:t})}{d\pi_t^c(\cdot | \xi_{1:t})}(z_{1:t}) \\ &\quad \times \bar{f}(\{\xi_{1:t}, z_{1:t}\})\pi_t^c(dz_{1:t} | \xi_{1:t}) \\ &= \frac{W_{t-1}^m(\xi_{1:t})}{a_t(\xi_{1:t})} \phi_t^c(\bar{f}) = \frac{W_{t-1}^m(\xi_{1:t})}{a_t(\xi_{1:t})} \bar{\psi}(\xi_{1:t}). \end{aligned} \quad (\text{A.8})$$

Combining (A.7) and (A.8) we get, π_t^m -almost surely,

$$\begin{aligned} a_t(\xi_{1:t})\pi_t^c((W_{t-1}\bar{f})^2) &= \frac{(W_{t-1}^m(\xi_{1:t})\bar{\psi}(\xi_{1:t}))^2}{a_t(\xi_{1:t})} + a_t(\xi_{1:t})\text{Var}_{\pi_t^c}(W_{t-1}\bar{f}). \end{aligned} \quad (\text{A.9})$$

Let L_{t-1} and L_{t-1}^m be defined as in (A.1) and (A.2), respectively. Then, $R_{t-1}(\cdot, W_{t-1}\bar{f}) = L_{t-1}(\cdot, \bar{f})$ and $R_{t-1}^m(\cdot, W_{t-1}^m\bar{\psi}) = L_{t-1}^m(\cdot, \bar{\psi})$.

Using (12), (10), (11) and the above results, the difference in asymptotic variance can be expressed as

$$\begin{aligned} \Delta_t(f) &= \sigma_{t-1}^2(L_{t-1}(\cdot, \bar{f})) - \tau_{t-1}^2(L_{t-1}^m(\cdot, \bar{\psi})) \\ &\quad + \phi_{t-1}[R_{t-1}(\cdot, (W_{t-1}\bar{f})^2)] - \phi_{t-1}^m[R_{t-1}^m(\cdot, (W_{t-1}^m\bar{\psi})^2)] \\ &= \sigma_{t-1}^2(L_{t-1}(\cdot, \bar{f})) - \tau_{t-1}^2(L_{t-1}^m(\cdot, \bar{\psi})) \\ &\quad + \phi_{t-1}^m\left[R_{t-1}^m\left(\cdot, \left(\frac{1}{a_t} - 1\right)(W_{t-1}^m\bar{\psi})^2 + a_t\text{Var}_{\pi_t^c}(W_{t-1}\bar{f})\right)\right] \end{aligned} \quad (\text{A.10})$$

(recall that $\pi_t^m = \phi_{t-1}^m R_{t-1}^m$ which ensures that we, due to the expectation w.r.t. $\phi_{t-1}^m R_{t-1}^m$ in (A.10), can make use of the equality in (A.9)).

Finally, by straightforward, but somewhat tedious manipulations

$$\phi_{t-1}^c(L_{t-1}(\cdot, \bar{f})) = L_{t-1}^m(\cdot, \bar{\psi}), \quad \pi_t^m\text{-almost surely.} \quad (\text{A.11})$$

Hence,

$$\sigma_{t-1}^2(L_{t-1}(\cdot, \bar{f})) - \tau_{t-1}^2(L_{t-1}^m(\cdot, \bar{\psi})) = \Delta_{t-1}(L_{t-1}(\cdot, \bar{f})), \quad (\text{A.12})$$

which completes the proof. ■

REFERENCES

- Chopin, N. (2004). Central limit theorem for sequential Monte Carlo methods and its application to Bayesian inference. *The Annals of Statistics*, 32(6), 2385–2411.
- Douc, R. and Moulines, E. (2008). Limit theorems for weighted samples with applications to sequential Monte Carlo. *The Annals of Statistics*, 36(5), 2344–2376.
- Doucet, A., de Freitas, N., Murphy, K., and Russell, S. (2000a). Rao-Blackwellised particle filtering for dynamic Bayesian networks. In *Proceedings of the Sixteenth Conference on Uncertainty in Artificial Intelligence*, 176–183. Stanford, USA.
- Doucet, A., Godsill, S.J., and Andrieu, C. (2000b). On sequential Monte Carlo sampling methods for Bayesian filtering. *Statistics and Computing*, 10(3), 197–208.
- Doucet, A. and Johansen, A. (2011). A tutorial on particle filtering and smoothing: Fifteen years later. In D. Crisan and B. Rozovsky (eds.), *The Oxford Handbook of Nonlinear Filtering*. Oxford University Press.
- Gilks, W.R. and Berzuini, C. (2001). Following a moving target – Monte Carlo inference for dynamic Bayesian models. *Journal of the Royal Statistical Society. Series B (Statistical Methodology)*, 63(1), 127–146.
- Gordon, N.J., Salmond, D.J., and Smith, A.F.M. (1993). Novel approach to nonlinear/non-Gaussian Bayesian state estimation. *Radar and Signal Processing, IEE Proceedings F*, 140(2), 107–113.
- Karlsson, R., Schön, T.B., and Gustafsson, F. (2005). Complexity analysis of the marginalized particle filter. *IEEE Transactions on Signal Processing*, 53(11), 4408–4411.
- Pitt, M.K. and Shephard, N. (1999). Filtering via simulation: Auxiliary particle filters. *Journal of the American Statistical Association*, 94(446), 590–599.
- Schön, T., Gustafsson, F., and Nordlund, P.J. (2005). Marginalized particle filters for mixed linear/nonlinear state-space models. *IEEE Transactions on Signal Processing*, 53(7), 2279–2289.
- Uglanov, A.V. (1991). Fubini's theorem for vector-valued measures. *Math. USSR Sbornik*, 69(2), 453–463.