

d d a d a
c h a c h

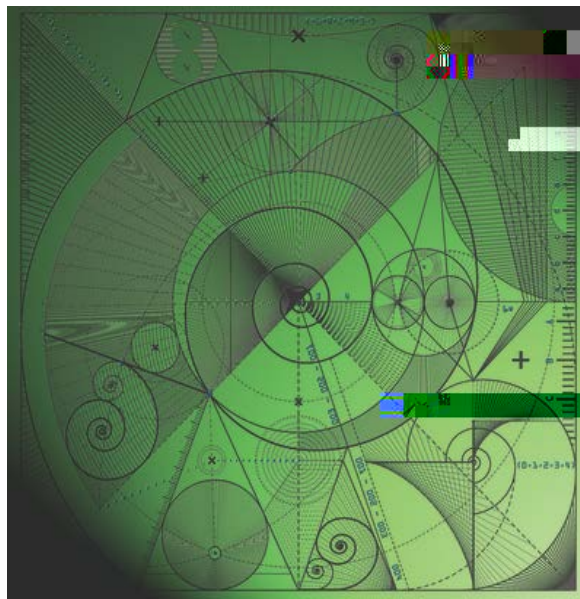
e g c d h c h h

g e g c

Lazy ABC

by

c c h g c a



1 Introduction

Approximate Bayesian computation (ABC) algorithms are a popular method of inference for a wide class of otherwise intractable probability models in applications such as population genetics, ecology, and systems biology (Beaumont, 2010; Marin et al., 2012). They select parameter vectors for which datasets simulated from the model of interest are sufficiently close to the observations. A bottleneck is the computational cost of producing the large quantity of model simulations needed, which becomes increasingly severe for more detailed models. However, it is often clear during a simulation that it is unpromising. For example it is likely to produce a poor match or to require excessive computation time. This paper presents lazy ABC, an importance sampling method which abandons some such simulations, a step referred to as early stopping exploiting information from incomplete simulations to save time. The result me. inflation17 16(infl)-267(ar16(inf)-285(8im)2415)-376(inf(detail(infscal(prob

The final likelihood estimator is based on X and Y

predictions given θ forming X in the notation above) to gain speed benefits without incurring additional approximation errors. More generally, there has been much interest over the past decade in Bayesian inference algorithms with random weights (e.g. Beaumont, 2003; Andrieu and Roberts, 2009; Fearnhead et al., 2010; Tran et al., 2014). A novelty of lazy ABC is that it introduces a random factor to the weights to reduce computation time, rather than to deal with intractability.

Rejection control in sequential Bayesian algorithms (Liu et al., 1998) uses a similar idea to lazy ABC. Here after the first t stages of the sequential analysis, a proposal (typically a sequence of latent states at times $2, \dots, t$) with weight w is allowed to continue with probability $\alpha = \min(1, w/c)$ for some constant c . On continuation the weight is updated to $w = cw$ and otherwise a new proposal is generated. Novelties of the current work are finding an optimal form for α and allowing it depend on information other than w .

The remainder of the paper is structured as follows. Section 2 contains background material. To help later developments this presents ABC within the framework of random weight importance sampling. Section 3 gives the lazy ABC algorithm and proves it targets the correct distribution. Section 4 presents theory and practical methods for tuning the algorithm. A related result is also given on the optimal importance distribution for standard ABC importance sampling. Section 5 contains the application to spatial extremes and Section 6 is a discussion. Appendices contain proofs and material on lazy ABC with multiple stopping decisions.

2 Importance sampling

Consider analysing data y_{obs} under a probability model with density $(y|j)$ and parameters θ . The likelihood is defined as $L(\theta) = (y_{\text{obs}}|j)$. Bayesian inference introduces a prior distribution with density $\pi(\theta)$ and aims to find the posterior distribution $\pi(\theta|y_{\text{obs}}) = \pi(\theta)L(\theta)/(y_{\text{obs}})$, where $(y_{\text{obs}}) = \int \pi(\theta)L(\theta)d\theta$, or at least to estimate the posterior expect-

tation $E[h(\boldsymbol{\theta})y_{\text{obs}}]$ of a generic function $h(\boldsymbol{\theta})$. Importance sampling is a simple method to do this. Parameter values $\boldsymbol{\theta}_{1:N}$ are simulated independently from an importance density $q(\boldsymbol{\theta})$ and given weights $w_i = L(\boldsymbol{\theta}_i | y_{\text{obs}}) / q(\boldsymbol{\theta}_i)$ (n.b. $\boldsymbol{\theta}_{1:N}$ represents the sequence $\boldsymbol{\theta}_1, \dots, \boldsymbol{\theta}_N$. Similar notation is used later.) It is assumed throughout that $q(\boldsymbol{\theta}) > 0$ whenever $L(\boldsymbol{\theta}) > 0$. Each of the $(\boldsymbol{\theta}_i; w_i)$ pairs can be computed in parallel, allowing for efficient implementation.

A Monte Carlo estimate of $E[h(\boldsymbol{\theta})y_{\text{obs}}]$ is $\hat{h} = \frac{\sum_{i=1}^N h(\boldsymbol{\theta}_i) w_i}{\sum_{i=1}^N w_i}$. Two properties of importance sampling estimates are

$$\hat{h} \xrightarrow{P} E[h(\boldsymbol{\theta})y_{\text{obs}}] \text{ almost surely as } N \rightarrow \infty; \quad (1)$$

$$E\left[\frac{1}{N} \sum_{i=1}^N w_i\right] = E(y_{\text{obs}}); \quad (2)$$

See Geyer (1989) for further details.

be seen by noting that Algorithm 1a is equivalent to a deterministic weight importance

Input (general):

Prior density $p(\cdot)$ and importance density $g(\cdot)$.

Number of iterations to perform N .

Input (RW-IS):

Likelihood estimator $\hat{L}$.

Input (ABC):

Observed data y_{obs} .

Summary statistics $S(\cdot)$, distance function $d(\cdot)$;

choice of $S(\cdot)$ involves a trade-off between low dimension and informativeness. For further background details on all aspects of ABC see the review articles of Beaumont (2010) and Marin et al. (2012).

3 Lazy ABC

This section defines lazy ABC and shows it produces valid results.

Definition 1. Lazy ABC is Algorithm 1a, using a likelihood estimator of the form (3) under conditions C1-C3.

$$\hat{\ell}_{\text{lazy}} = \begin{cases} \hat{\ell}_{\text{ABC}} = \ell(\theta; X) & \text{with probability } \ell(\theta; X) \\ 0 & \text{otherwise} \end{cases} \quad (3)$$

C1 $\ell(\theta; x)$ is a function with codomain $[0, 1]$.

C2 $\ell(\theta; x) > 0$ whenever $\Pr(\hat{\ell}_{\text{ABC}} > 0 | \theta; x) > 0$.

C3 The random variable X is such that both X_j and $Y_j | x$ can be simulated from.

The following theorem shows that the estimator $\hat{\ell}_{\text{lazy}}$ can be used in a RW-IS algorithm, or a pseudo-marginal MCMC algorithm, and give valid results.

Theorem 1. Conditional on θ , $\hat{\ell}_{\text{lazy}}$ is a non-negative unbiased estimator of $\ell_{\text{ABC}}(\theta)$.

Proof. Non-negativity is immediate. For unbiasedness first observe that $E(\hat{\ell}_{\text{lazy}} | \theta; x)$ equals zero when $\ell(\theta; x) = 0$ and $E(\hat{\ell}_{\text{ABC}} | \theta; x)$ otherwise. By C2 if $\ell(\theta; x) = 0$ then $\Pr(\hat{\ell}_{\text{ABC}} > 0 | \theta; x) = 0$ and so $E(\hat{\ell}_{\text{ABC}} | \theta; x) = 0$. Hence $E(\hat{\ell}_{\text{lazy}} | \theta; X) = E(\hat{\ell}_{\text{ABC}} | \theta; X)$. Taking expectations over

the continuation simulation stage It is often useful later to have $(\theta; x) = (\theta(\theta; x))$ where $(\theta; x)$ is referred to as the decision statistics

Notation is now introduced for expected CPU times: $T_1(\theta)$ is for steps 1 and 2 above conditional on θ , $T_2(\theta; \theta)$ is for step 3 conditional on $(\theta; \theta)$ and $T(\theta)$ is for simulation from $\hat{\mathcal{L}}_{ABC}$ conditional on θ . The first two are roughly the times of the initial simulation and continuation stages, but are defined to cover all steps involved in simulating from $\hat{\mathcal{L}}_{lazy}$.

It is assumed that

$$T(\theta) = T_1(\theta) + E[T_2(\theta; \theta) | \theta] \quad (4)$$

Roughly speaking this states that drawing from $\hat{\mathcal{L}}_{lazy}$ conditional on no early stopping takes at least as long as drawing from $\hat{\mathcal{L}}_{ABC}$. The difference is due to computational overheads of considering stopping. It is also convenient to define $\bar{T}_1 = E[T_1(\theta)]$ and $\bar{T}_2(\theta) = E[T_2(\theta; \theta) | \theta]$.

3.1 Lazy importance sampling

The above approach can be generalised to non-ABC situations to give lazy importance sampling (LIS). This is Algorithm 1a using a likelihood estimator of the form:

$$\hat{\mathcal{L}}_{lazy} = \begin{cases} \hat{\mathcal{L}} = (\theta; X) & \text{with probability } (\theta; X) \\ 0 & \text{otherwise} \end{cases}$$

In addition to conditions C1 and C2 above assume:

- C4 The distribution $(X; \hat{\mathcal{L}})^j$ is such that $\hat{\mathcal{L}}^j$ is a non-negative unbiased estimator of $\mathcal{L}(\theta)$, and both X^j and $\hat{\mathcal{L}}^j; x$ can be simulated from.

This framework can be used when $\hat{\mathcal{L}}$ is an expensive unbiased estimator. It also allows cases where either or both of X and $\hat{\mathcal{L}}$ are non-random. For example X may be a deterministic approximation of the likelihood and $\hat{\mathcal{L}}^j$ may be a point mass at $\mathcal{L}(\theta)$.

All theorems and proofs of this paper also hold for lazy importance sampling, replacing $L_{ABC}(\cdot)$ and $\hat{L}_{ABC}$ with $L(\cdot)$ and $\hat{L}$, and making other small modifications noted in the text. In particular Theorem 1 shows that given conditions C1, C2 and C4, $\hat{L}_{lazy,j}$ is a non-negative unbiased estimator of $L(\cdot)$. However the practical application of lazy importance sampling is challenging as discussed in Section 6.

4 Tuning

There is considerable freedom to tune lazy ABC through the choice of δ (when to consider stopping) and γ (the function assigning continuation probabilities). Section 4.1 proves a result on the most efficient choice of δ . This theory is used in Section 4.2 to motivate practical tuning methods.

Note that the case where $\delta(\cdot)$ is based on χ^2_A for discrete A does not require the theoretical results below. Here $\delta(\cdot)$ values can be selected by numerical optimisation of an estimate of the algorithm's efficiency based on pilot simulations. The methods that follow detail construction of such an estimate.

4.1 Theory

A commonly used tool for the analysis of importance sampling algorithms is the effective sample size (ESS). Liu (1996) argued that typically the variance of the importance sampling estimator is roughly equal to that of N_e independent samples where

$$N_e = N / E(W)^2 = E(W)^{-2};$$

and the random variable W is the weight generated in an iteration of importance sampling. The argument of Liu generalises immediately to RW-IS algorithms through the interpretation of them as importance sampling algorithms on an augmented parameter space given in Section 2.2.

Define efficiency as $N_e = T$ where T is the CPU time of the algorithm (i.e. ignoring any execution time savings due to parallelisation.) Assume that $\hat{W}$ follows a central limit theorem in N . Then the delta method gives that for large N efficiency asymptotically equals

$$\frac{E(W)^2 = E(W^2)}{E(T) = N}.$$

Theorem 2. Fix some decision statistics $(\hat{W}; x)$. Amongst continuation probability functions of the form $g(\hat{W}; x) = g(\hat{W}(\hat{W}; x))$, asymptotic efficiency is maximised by the following expression for some $\gamma > 0$,

$$g(\hat{W}) = \min_{\gamma} \left\{ 1; \frac{E[\hat{L}_{ABC}^2 \frac{(\hat{W})^2}{g(\hat{W})^2} j]}{T_2(\hat{W})} \right\}^{\frac{3}{5} \gamma^{1/2}} \quad (5)$$

Proof. See Appendix B.1. □

Remark 1. Suppose $g(\hat{W}) = u(\hat{W})$ i.e. this fraction is completely determined by $\hat{W}$. For example this is the case in ABC rejection sampling where $g(\hat{W}) = u(\hat{W})$. Then (5) becomes

$$g(\hat{W}) = \min \left\{ 1; u(\hat{W}) \frac{(\hat{W})^2}{T_2(\hat{W})} \right\}^{\frac{3}{5} \gamma^{1/2}} \quad (6)$$

where $(\hat{W}) = E[\hat{L}_{ABC} j] = \Pr(d(S(Y); S(y_{obs})) \leq j)$.

Remark 2. Theorem 2 and Remark 1 hold for LIS with $\hat{L}_{ABC}$ replaced by $\hat{L}^2$.

A simple closed form expression for $\hat{W}$ does not appear possible. In the practical tuning methods below $\hat{W}$ is found numerically, and the behaviour of this numerical estimate investigated by simulation study (see Figure 1B).

By viewing ABC-IS as a special case of lazy ABC, Theorem 2 can be applied to find the optimal choice of $g(\hat{W})$ for ABC-IS.

Corollary 1. The asymptotic efficiency of ABC-IS is maximised by $g(\hat{W}) = \frac{(\hat{W})^2}{T(\hat{W})}$, where $(\hat{W}) = E(\hat{L}_{ABC} j)$.

Proof. See Appendix B.2. □

Remark 3. A corresponding result to Corollary 1 holds for RW-IS with $\gamma(\cdot) = E(\hat{L}^2_j)$.

Remark 4. The special case of Corollary 1 with $\gamma(\cdot)$ constant matches the result of Appendix A in Fearnhead and Prangle (2012).

Note that it is not clear what the optimal choice of $\gamma(\cdot)$ is for lazy ABC. The examples later use typical choices from the ABC literature, but a better choice may improve lazy ABC performance further.

4.2 Methods

Theorem 2 motivates choosing γ by estimating (6). This section details a method to implement this approach. Its effectiveness is discussed in Section 6.

Tuning begins with a pilot run of N^0 iterations of ABC-IS. This is used to estimate $\gamma(\cdot)$ and $T_2(\cdot)$ for various choices of X and γ , considering only γ is such that Remark 1 can be applied. Under each of these choices, γ is found by numerically maximising an estimate of efficiency. The optimal choice of X and γ is then made. Note that estimation of $\gamma(\cdot)$ and $T_2(\cdot)$ is challenging if γ is high dimensional e.g. for $\gamma = (\gamma; x)$. Therefore a low dimensional γ is recommended. To ensure Remark 1 applies $\gamma(\cdot) = g(\cdot)$ can form one component of γ if necessary. Following tuning N iterations of lazy ABC are performed (unless the estimated efficiency gains are judged inadequate). Detailed comments on several aspects of this method follow.

4.2.1 Estimation of $T_2(\cdot)$

It may often suffice to treat $T_2(\cdot)$ as constant and estimate it as the mean CPU time of the

shows $T_2(\cdot)$ varies little relative to $(\cdot)$. Alternatively, statistical methods such as regression can be used for estimation, which is straightforward when $\mathbf{X}$ is low dimensional.

4.2.2 Estimation of $(\cdot)$

Estimation of $(\cdot)$ is more difficult. Two approaches are suggested: the "standard" approach, producing $\hat{\lambda}^{(1)}$, attempts accurate estimation but involves strong assumptions; the "conservative" approach, producing $\hat{\lambda}^{(2)}$, is stronger.

4.2.3 Estimating efficiency

The tuning method outlined above requires the use of N^0 pilot run iterations to estimate the efficiency of lazy ABC under various choices of tuning details (in particular X , β and γ). It is sufficient to estimate $[E(W^2)E(T)]^{-1}$, as this equals efficiency up to a constant of proportionality. This can be used to estimate efficiency relative to ABC-IS, which is a particularly interpretable form of the results as it shows the efficiency improvement of using lazy ABC.

Assume that for a particular choice of tuning details the following are available for $i = 1, \dots, N^0$: $t_i^{(1)}$ - initial simulation stage time; $t_i^{(2)}$ - continuation simulation stage time; α_i - continuation probability; $\hat{\mu}_i$ - estimate of $E(\hat{L}_{ABC} | j_i)$; u_i - ratio $\frac{\hat{\mu}_i}{\mu} = g(\beta)$. An estimate up to proportionality of efficiency is then $[W^2 \hat{T}]^{-1}$ where $W^2 = N^{0-1} \sum_{i=1}^{N^0} u_i^2 \alpha_i = \hat{\mu}$ and $\hat{T} = \sum_{i=1}^{N^0} t_i^{(1)} + \sum_{i=1}^{N^0} \alpha_i t_i^{(2)}$. An estimate of efficiency of ABC-IS is formed by taking $\frac{1}{\hat{\mu}}$. Note that this typically overestimates T due to the overheads of considering stopping (see (4)). A more precise estimate would be possible using further pilot simulations of standard ABC.

4.2.4 Combining pilot and main run output

To make efficient use of the pilot run, it can be used in the final output as well as for tuning. This is done by appending the pilot sequence of (w, v) pairs to that from the main algorithm. Loosely speaking, since each individual sequence targets the same distribution, so does the combined sequence. More technically, it is straightforward to see that ABC versions of relations (1) and (2) are roughly true for the combined sequence when N and N^0 are large, and are exactly true as $N \rightarrow \infty$ regardless of N^0 . Also note that on appending the sequences, gains in efficiency are possible by multiplying the weights of one sequence by a constant, but this is not implemented here as little improvement was observed in the application later.

4.2.5 Choice of ϵ

In ABC-IS, an appropriate value ϵ is often unknown a priori and is instead chosen based on the simulated $d(S(Y); s(y_{\text{obs}}))$ values. For lazy ABC in this situation one can use the pilot run to select a preliminary conservative choice of ϵ_1 as in Section 4.2.2 and perform lazy ABC with $\epsilon = \epsilon_1$. Alternative values of ϵ can then be investigated by updating the realised $\hat{L}_{\text{ABC}}$ values in the weight calculations. For $\epsilon < \epsilon_1$ this simply reduces the number of non-zero weights. However $\epsilon > \epsilon_1$ is not recommended as this may introduce large weights and destabilise the importance sampling approximation.

Erhardt and Smith focus on the Whittle-Matern correlation function with zero nugget

$$\rho(h; c; \nu) = \frac{2^{1-\nu}}{\Gamma(\nu)} \left(\frac{h}{c}\right)^{\nu} K_{\nu}\left(\frac{h}{c}\right);$$

where Γ is the gamma function and K_{ν} is the modified Bessel function of the third kind with order ν . This has two parameters: range $c > 0$ and smoothness $\nu > 0$.

A density function for the Schlather process is not available for $D > 2$, making inference difficult. Schlather (2002) provides a near-exact algorithm to simulate from the process based on only a finite number of copies of Y_i , motivating the use of ABC by Erhardt and Smith. They applied ABC rejection and importance sampling with a uniform prior on $[0, 10]^2$ and investigated several choices of summary statistics. The analysis here focuses on the choice they found most successful, based on tripletwise extremal coefficient estimators. Given a triple of 3 locations, i, j, k , this estimator is

$$\hat{e}_{ijk} = P_{t=1}^T \frac{1}{1 + \max(y_{t,i}; y_{t,j}; y_{t,k})}.$$

There are $O(D^3)$ such summaries, so Erhardt and Smith calculate a vector of mean values within 100 clusters of triples, and use these as summary statistics. Their clustering process finds triples of similar shapes, ignoring differences of location and rotation. The ABC distance function between two vectors m_1 and m_2 of cluster means is

$$d(m_1; m_2) = \sqrt{\sum_{i=1}^D (m_{1i} - m_{2i})^2}. \quad (7)$$

Although applying dimension reduction techniques to such high dimensional summaries has been shown to often improve ABC results (Fearnhead and Prangle, 2012), this is not investigated here as the aim is to investigate the efficiency improvements of lazy ABC.

of the approach of Erhardt and Smith.

5.2 Methods

Exploratory investigation of ABC code with $D = 20$ and $T = 100$ showed that the majority of time was spent simulating the data (7.1ms/iteration) and calculating extremal coefficient estimates (17.9ms/iteration), with the remaining steps being brief (3.1ms/iteration). The time costs of the first two of these scaled with D as roughly proportional to D and D^3 respectively, so the latter is expected to dominate for large D . Furthermore, interrupting and then resuming operations during the calculation of extremal coefficient estimates is much simpler to implement than during simulation of data. Therefore the initial simulation stage of the lazy ABC analysis was chosen to be simulating the data at all locations, and extremal coefficient estimates at a subset of locations $\mathcal{L}$. The continuation simulation stage was to calculate the remaining extremal coefficient estimates.

The decision statistic $\hat{d}$ was constructed as follows. Let m_{1i} be the i th cluster mean for the observed data. Let m_{2i} be the i th cluster mean for the simulated data using only extremal coefficient estimates available at the initial simulation stage, and $\mathcal{B}$ be the set of clusters for which any such estimates are available. Then define
$$\hat{d} = \sqrt{\frac{1}{|\mathcal{B}|} \sum_{i \in \mathcal{B}} (m_{1i} - m_{2i})^2}.$$
 This is an estimate of the ABC distance $d_{m_{2i}}$.

simulations with nearby values of $\hat{d}$. This was done for several $\hat{d}$ values and interpolated estimates elsewhere formed $q(\hat{d})$. For the importance sampling case, $\log$ was also included in each regression so that a number of functions mapping to estimates of $\Pr(d = j\hat{d}; u)$ for various $\hat{d}$ values were produced which were used for interpolation. Calculation of $\hat{z}^*$ was as described in Section 4.2.2, taking n to give 100 acceptances in the pilot run. Given estimates of T_2 and $\hat{z}^*$, tuning was performed as described in Section 4.2, with optimisation over possible choices of d by backwards selection.

Three simulation studies were performed. The first replicated the rejection sampling analysis of Erhardt and Smith on several simulated datasets. These used $D = 20$; $T = 100$ and true parameter values shown in Table 1. Each dataset used a different set of observation locations with integer coordinates sampled from $[0, 10]^2$. The first analysis was a replication of the standard ABC analysis, using n values corresponding to 200 acceptances. Then lazy ABC was performed on the same datasets under each method of estimating T_2 to compare the methods fairly, lazy ABC used the same value as standard ABC and reused its random seeds so that the sequence of $(X; Y)$ realisations is also the same.

The second simulation study investigated rejection sampling for a single larger simulated dataset with $D = 35$, $c = 0.5$ and $\sigma = 1$. Locations were chosen as before. As in a real application σ was not assumed to be known in advance and the approach of Section 4.2.5 was used to select this post-hoc. A complication for this dataset was that the simulation of Gaussian processes was difficult when both parameters were large: the default "direct method", based on Choleski decomposition, sometimes produced numerical errors. Simulation was possible via the turning bands method (TBM) but much slower (roughly 150 times the CPU time). A two stage simulation method was implemented. First the direct method was attempted and if this failed TBM was used. To save time lazy ABC was implemented with multiple stopping decisions, the first taking place after attempting the direct method. This has a binary decision statistic indicating success or failure. The second stopping decision is as described earlier. Tuning was performed as described in Appendix A.1.2, using $\hat{z}^*$ tuned as described

above by either the standard or conservative tuning method. The standard method used to give 30 acceptances in the pilot run. As before all analyses reused the same random seeds.

Finally an importance sampling analysis was performed on the larger dataset. A sample of 10^4 log parameter values was taken from simulations of the preceding standard ABC analysis with distances below the 0.3 quantile. A Gaussian mixture distribution was constructed with locations given by this sample and variances equal to twice the empirical variance of the sample. After truncation to the prior support, this was used to give $q(\theta)$, where θ now represents the log parameters. This choice follows the suggestions of Beaumont et al. (2009), noting that using the log scale produced a better fit to the sample and that the subsample was used to avoid slow density calculations. The preceding $n = 35$ analysis was then repeated.

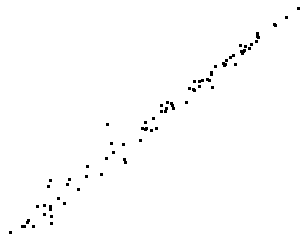


Figure 1: Details of a simulation study applying lazy ABC to spatial extremes corresponding to the first row of Table 1. Panels A-C concentrate on the standard tuning approach. Panel A Pilot run values of $\hat{d}$ and d . The dashed line shows the value of d . Panel B Estimated efficiency for different values of L . The dashed line shows the realised efficiency. Panel C Estimated efficiency for the best choices of L of various lengths output by backwards selection. The dashed line shows the realised efficiency. Panel D Values of $\hat{d}$ and d from non-pilot simulations under standard (solid line) and conservative (dashed line) tuning. The marks on the horizontal axis indicate the simulations which resulted in positive weights. (For this panel conservative tuning was performed using d as selected by standard tuning.)

Range	Smooth	Standard Time (10 ³ s)	Time (10 ³ s)	Lazy Sample size	ESS	Relative e Estimated	ciency Actual
0.5	1	32.0	8.0 (11.6)	196 (199)	196.0 (198.7)	4.08 (3.28)	4.00 (2.79)
1	1	31.3	7.3 (9.8)	200 (200)	199.9 (200.0)	4.34 (4.31)	4.35 (3.25)
1	3	31.3	8.2 (11.2)	194 (198)	182.5 (196.5)	3.77 (3.43)	3.51 (2.79)
3	1	31.2	7.7 (11.1)	194 (200)	189.9 (200.0)	4.18 (3.56)	3.89 (2.86)
3	3	31.2	7.4 (11.0)	192 (199)	175.8 (199.0)	4.43 (3.65)	3.79 (2.87)
5	3	31.3	8.3 (11.1)	200 (200)	200.0 (200.0)	3.73 (3.49)	3.85 (2.87)

		Standard			Lazy			Relative efficiency
		Time (10 ³ s)	Sample size	ESS	Time (10 ³ s)	Sample size	ESS	
RS standard	2.61	241.4	210	210	22.8	200	139.6	7.0
RS conservative	2.61	241.4	207	207	25.5	200	197.7	9.0
IS standard	2.33	136.4	209	168	61.9	200	165	2.2
IS conservative	2.33	136.4	209	168	51.0	200	162	2.6

Table 2: Simulation study on a spatial extremes dataset with $D = 35$. Results are shown for rejection and importance sampling with standard and conservative tuning. The rejection sampling output was used to create the importance density. The final choice of δ is shown. For IS the two δ values are equal but there is a small difference for RS. The lazy ABC output includes the pilot run and the tuning time.

and weighting the accepted simulations accordingly, the algorithm targets exactly the same distribution as standard ABC, in the sense that Monte Carlo estimates of functions $\mathbb{E}(f)$ and of the model evidence converge to unchanged values.

Results have been provided on the optimal tuning of the lazy ABC stopping rule and used to motivate a practical tuning method. This has been demonstrated for a computationally challenging application where it has produced improvements in efficiency (ESS/CPU time) over standard ABC of up to 8 times. One case of this application involved multiple stopping decisions. This illustrated two potential uses of lazy ABC: firstly to consider stopping every simulation based on whether it appears promising, secondly to consider stopping after particular events which are suspected a priori to indicate unpromising results.

The tuning method is based on estimating the optimal choice of δ , (6). The most difficult part was estimating $\pi(\delta) = \Pr(\delta(S(Y); S(y_{\text{obs}})) \leq j)$ from pilot run data. Two approaches to this were described, a standard approach of direct estimation and a conservative approach of estimation using a larger value than is of interest for ABC. The latter approach improves robustness and make estimation simpler at the cost of some inefficiency. Both approaches performed well in the simulation studies but some improvements are desirable. Firstly, estimation of $\pi(\delta)$ involves extrapolation which may produce inaccurate results. Secondly, several choices by the user are required, especially for the standard approach. A more automated approach would be useful for lazy versions of ABC SMC algorithms, where a new choice of δ would be needed for each value, or alternatively for lazy ABC algo-

rithms which adapt as more simulations become available. It would be of interest to find suboptimal but robust choices of addressing these issues.

Lazy ABC with multiple stopping decisions is an extension to the framework of the main paper and is described in Appendix A. A tuning method is given when the decision statistics for all stopping decisions are discrete, and also some cases where one decision statistic is continuous. For more complex cases tuning results are not available. For now it is recommended to discretise most decision statistics to avoid this difficulty.

Also, Section 3.1 showed that a generalisation to the non-ABC setting, lazy importance sampling, is possible, and the theoretical results of the paper carry over to this. However exploratory analysis suggests tuning this in practice is more challenging than lazy ABC. This is because it is necessary to estimate $\psi(\theta) = E(\hat{L}^2 j)$ (see Remark 2), and this expectation can be strongly influenced by the upper tail of $\hat{L} j$ which is hard to estimate from pilot run output. For lazy ABC, $\hat{L}_{ABC} j$ is Bernoulli avoiding this difficulty. A related point is that lazy ABC can be generalised to allow a non-uniform ABC kernel. This gives $\hat{L}_{ABC}$ with a known upper bound so that estimation of $\psi(\theta)$ seems feasible.

This paper has concentrated on importance sampling, which is widely used by ABC practitioners, but the lazy ABC approach can be extended to ABC versions of MCMC and SMC, which are more efficient algorithms. The tuning results are applicable to SMC algorithms, but further practical methods are needed, as mentioned above. Further theory on optimal tuning is necessary for MCMC, although good performance may be possible with ad-hoc tuning. Examining the connections between lazy ABC and rejection control (Liu et al., 1998) may also be fruitful, especially to design algorithms in which partial

A Multiple stopping decisions

The lazy ABC framework of Section 3 allows multiple stopping decisions, as follows. As in that section assume Y is a deterministic transformation of a latent vector $X_{1:p}$.

Example A1: Multiple stopping decisions Let $X = X_{1:p}$ and $(\cdot; x) = \prod_{i=1}^Q (\cdot; x_{1:t_i})$.

Thus, for each $1 \leq i \leq s$, once simulation of $X_{1:t_i}$ has been performed then $\hat{L}_{\text{lazy}}$ is set to zero with a certain probability, in which case no further simulation is necessary. It is often be useful to let $(i)(\cdot; x_{1:t_i}) = (i)(\cdot_i(\cdot; x_{1:t_i}))$. That is, each stopping decision has associated decision statistics $\cdot_i$.

Example A2: Multiple random stopping times As for Example A1 but with each t_i replaced with a random stopping time variable T_i . This permits stopping to be considered when various random events occur, without imposing a fixed order of occurrence.

The following alternative characterisation of these examples is useful below.

Lemma 1. For any $1 \leq i \leq s$, Examples A1 and A2 can be represented as a lazy importance sampling algorithm with continuation probability $(i)(\cdot_i)$ and

$$\hat{L} = \begin{cases} \hat{L}_{\text{ABC}} = \cdot_i(\cdot; X) & \text{with probability } (i)(\cdot_i) \\ 0 & \text{otherwise} \end{cases}$$

where $\cdot_i(\cdot; x) = \prod_{j \in i} (\cdot_j)(\cdot_j)$.

Proof. The likelihood estimator stated can easily be verified to have the same distribution as $\hat{L}_{\text{lazy}}$. □

It is also helpful to define $T_{2i}(\cdot; \cdot_{1:s})$ as the expected time remaining from the calculation of $\cdot_i$ until the likelihood estimate is computed conditional on $\cdot$ and $\cdot_{1:s}$, and $T_{2i}(\cdot_i)$

A.1 Tuning

The efficiency estimate of Section 4.2.3 can be used in a multiple stopping decision setting given a choice of $\hat{\tau}$. It is necessary to update the estimator $\hat{\tau}$ given there which is usually a straightforward task. Sections A.1.1 and A.1.2 describe situations of practical interest where the optimal form of $\hat{\tau}$ can be derived. However in general the problem is challenging, as illustrated by Section A.1.3.

A.1.1 Discrete decision statistics

Suppose $(x, y) \in [0, 1] \times [0, 1]$ is a point in the unit square. Let (x_1, y_1) and (x_2, y_2) be the two points in the unit square such that $(x_1, y_1) \leq (x, y) \leq (x_2, y_2)$ and $(x_1, y_1) \neq (x_2, y_2)$. Let (x_3, y_3) and (x_4, y_4) be the two points in the unit square such that $(x_3, y_3) \leq (x, y) \leq (x_4, y_4)$ and $(x_3, y_3) \neq (x_4, y_4)$. Let (x_5, y_5) and (x_6, y_6) be the two points in the unit square such that $(x_5, y_5) \leq (x, y) \leq (x_6, y_6)$ and $(x_5, y_5) \neq (x_6, y_6)$. Let (x_7, y_7) and (x_8, y_8) be the two points in the unit square such that $(x_7, y_7) \leq (x, y) \leq (x_8, y_8)$ and $(x_7, y_7) \neq (x_8, y_8)$. Let (x_9, y_9) and (x_{10}, y_{10}) be the two points in the unit square such that $(x_9, y_9) \leq (x, y) \leq (x_{10}, y_{10})$ and $(x_9, y_9) \neq (x_{10}, y_{10})$. Let (x_{11}, y_{11}) and (x_{12}, y_{12}) be the two points in the unit square such that $(x_{11}, y_{11}) \leq (x, y) \leq (x_{12}, y_{12})$ and $(x_{11}, y_{11}) \neq (x_{12}, y_{12})$. Let (x_{13}, y_{13}) and (x_{14}, y_{14}) be the two points in the unit square such that $(x_{13}, y_{13}) \leq (x, y) \leq (x_{14}, y_{14})$ and $(x_{13}, y_{13}) \neq (x_{14}, y_{14})$. Let (x_{15}, y_{15}) and (x_{16}, y_{16}) be the two points in the unit square such that $(x_{15}, y_{15}) \leq (x, y) \leq (x_{16}, y_{16})$ and $(x_{15}, y_{15}) \neq (x_{16}, y_{16})$. Let (x_{17}, y_{17}) and (x_{18}, y_{18}) be the two points in the unit square such that $(x_{17}, y_{17}) \leq (x, y) \leq (x_{18}, y_{18})$ and $(x_{17}, y_{17}) \neq (x_{18}, y_{18})$. Let (x_{19}, y_{19}) and (x_{20}, y_{20}) be the two points in the unit square such that $(x_{19}, y_{19}) \leq (x, y) \leq (x_{20}, y_{20})$ and $(x_{19}, y_{19}) \neq (x_{20}, y_{20})$. Let (x_{21}, y_{21}) and (x_{22}, y_{22}) be the two points in the unit square such that $(x_{21}, y_{21}) \leq (x, y) \leq (x_{22}, y_{22})$ and $(x_{21}, y_{21}) \neq (x_{22}, y_{22})$. Let (x_{23}, y_{23}) and (x_{24}, y_{24}) be the two points in the unit square such that $(x_{23}, y_{23}) \leq (x, y) \leq (x_{24}, y_{24})$ and $(x_{23}, y_{23}) \neq (x_{24}, y_{24})$. Let (x_{25}, y_{25}) and (x_{26}, y_{26}) be the two points in the unit square such that $(x_{25}, y_{25}) \leq (x, y) \leq (x_{26}, y_{26})$ and $(x_{25}, y_{25}) \neq (x_{26}, y_{26})$. Let (x_{27}, y_{27}) and (x_{28}, y_{28}) be the two points in the unit square such that $(x_{27}, y_{27}) \leq (x, y) \leq (x_{28}, y_{28})$ and $(x_{27}, y_{27}) \neq (x_{28}, y_{28})$. Let (x_{29}, y_{29}) and (x_{30}, y_{30}) be the two points in the unit square such that $(x_{29}, y_{29}) \leq (x, y) \leq (x_{30}, y_{30})$ and $(x_{29}, y_{29}) \neq (x_{30}, y_{30})$. Let (x_{31}, y_{31}) and (x_{32}, y_{32}) be the two points in the unit square such that $(x_{31}, y_{31}) \leq (x, y) \leq (x_{32}, y_{32})$ and $(x_{31}, y_{31}) \neq (x_{32}, y_{32})$. Let (x_{33}, y_{33}) and (x_{34}, y_{34}) be the two points in the unit square such that $(x_{33}, y_{33}) \leq (x, y) \leq (x_{34}, y_{34})$ and $(x_{33}, y_{33}) \neq (x_{34}, y_{34})$. Let (x_{35}, y_{35}) and (x_{36}, y_{36}) be the two points in the unit square such that $(x_{35}, y_{35}) \leq (x, y) \leq (x_{36}, y_{36})$ and $(x_{35}, y_{35}) \neq (x_{36}, y_{36})$. Let (x_{37}, y_{37}) and (x_{38}, y_{38}) be the two points in the unit square such that $(x_{37}, y_{37}) \leq (x, y) \leq (x_{38}, y_{38})$ and $(x_{37}, y_{37}) \neq (x_{38}, y_{38})$. Let (x_{39}, y_{39}) and (x_{40}, y_{40}) be the two points in the unit square such that $(x_{39}, y_{39}) \leq (x, y) \leq (x_{40}, y_{40})$ and $(x_{39}, y_{39}) \neq (x_{40}, y_{40})$. Let (x_{41}, y_{41}) and (x_{42}, y_{42}) be the two points in the unit square such that $(x_{41}, y_{41}) \leq (x, y) \leq (x_{42}, y_{42})$ and $(x_{41}, y_{41}) \neq (x_{42}, y_{42})$. Let (x_{43}, y_{43}) and (x_{44}, y_{44}) be the two points in the unit square such that $(x_{43}, y_{43}) \leq (x, y) \leq (x_{44}, y_{44})$ and $(x_{43}, y_{43}) \neq (x_{44}, y_{44})$. Let (x_{45}, y_{45}) and (x_{46}, y_{46}) be the two points in the unit square such that $(x_{45}, y_{45}) \leq (x, y) \leq (x_{46}, y_{46})$ and $(x_{45}, y_{45}) \neq (x_{46}, y_{46})$. Let (x_{47}, y_{47}) and (x_{48}, y_{48}) be the two points in the unit square such that $(x_{47}, y_{47}) \leq (x, y) \leq (x_{48}, y_{48})$ and $(x_{47}, y_{47}) \neq (x_{48}, y_{48})$. Let (x_{49}, y_{49}) and (x_{50}, y_{50}) be the two points in the unit square such that $(x_{49}, y_{49}) \leq (x, y) \leq (x_{50}, y_{50})$ and $(x_{49}, y_{49}) \neq (x_{50}, y_{50})$. Let (x_{51}, y_{51}) and (x_{52}, y_{52}) be the two points in the unit square such that $(x_{51}, y_{51}) \leq (x, y) \leq (x_{52}, y_{52})$ and $(x_{51}, y_{51}) \neq (x_{52}, y_{52})$. Let (x_{53}, y_{53}) and (x_{54}, y_{54}) be the two points in the unit square such that $(x_{53}, y_{53}) \leq (x, y) \leq (x_{54}, y_{54})$ and $(x_{53}, y_{53}) \neq (x_{54}, y_{54})$. Let (x_{55}, y_{55}) and (x_{56}, y_{56}) be the two points in the unit square such that $(x_{55}, y_{55}) \leq (x, y) \leq (x_{56}, y_{56})$ and $(x_{55}, y_{55}) \neq (x_{56}, y_{56})$. Let (x_{57}, y_{57}) and (x_{58}, y_{58}) be the two points in the unit square such that $(x_{57}, y_{57}) \leq (x, y) \leq (x_{58}, y_{58})$ and $(x_{57}, y_{57}) \neq (x_{58}, y_{58})$. Let (x_{59}, y_{59}) and (x_{60}, y_{60}) be the two points in the unit square such that $(x_{59}, y_{59}) \leq (x, y) \leq (x_{60}, y_{60})$ and $(x_{59}, y_{59}) \neq (x_{60}, y_{60})$. Let (x_{61}, y_{61}) and (x_{62}, y_{62}) be the two points in the unit square such that $(x_{61}, y_{61}) \leq (x, y) \leq (x_{62}, y_{62})$ and $(x_{61}, y_{61}) \neq (x_{62}, y_{62})$. Let (x_{63}, y_{63}) and (x_{64}, y_{64}) be the two points in the unit square such that $(x_{63}, y_{63}) \leq (x, y) \leq (x_{64}, y_{64})$ and $(x_{63}, y_{63}) \neq (x_{64}, y_{64})$. Let (x_{65}, y_{65}) and (x_{66}, y_{66}) be the two points in the unit square such that $(x_{65}, y_{65}) \leq (x, y) \leq (x_{66}, y_{66})$ and $(x_{65}, y_{65}) \neq (x_{66}, y_{66})$. Let (x_{67}, y_{67}) and (x_{68}, y_{68}) be the two points in the unit square such that $(x_{67}, y_{67}) \leq (x, y) \leq (x_{68}, y_{68})$ and $(x_{67}, y_{67}) \neq (x_{68}, y_{68})$. Let (x_{69}, y_{69}) and (x_{70}, y_{70}) be the two points in the unit square such that $(x_{69}, y_{69}) \leq (x, y) \leq (x_{70}, y_{70})$ and $(x_{69}, y_{69}) \neq (x_{70}, y_{70})$. Let (x_{71}, y_{71}) and (x_{72}, y_{72}) be the two points in the unit square such that $(x_{71}, y_{71}) \leq (x, y) \leq (x_{72}, y_{72})$ and $(x_{71}, y_{71}) \neq (x_{72}, y_{72})$. Let (x_{73}, y_{73}) and (x_{74}, y_{74}) be the two points in the unit square such that $(x_{73}, y_{73}) \leq (x$

A.1.3 Multiple continuous decision statistics

Consider the setting of A.1.2 with the modification that every $u_i(\cdot; x)$ is continuous and there exists a corresponding function $u_i(\cdot) = g(\cdot)$. The same approach as above gives equations of the form

$$u_i^{(1)}(\cdot) = \min \left(1, u_i(\cdot) \frac{u_i(\cdot)^{1/2}}{T_{2i}(\cdot)} \right);$$

for $i = 1, \dots, s$. The definition of u_i involves $u_j^{(1)}$ for all $j \neq i$, and T_{2i} will also involve many of these terms. Thus deriving the optimal $u_i^{(1)}$ functions involves solving a complicated system of non-linear implicit equations.

B Tuning proofs

All results are proved for the general case of LIS as described in Section 3.1. For lazy ABC replace $\hat{L}$ with $\hat{L}_{ABC}$.

B.1 Proof of Theorem 2

In LIS the importance sampling weight W equals $\frac{\hat{L}(\cdot)}{(\cdot)g(\cdot)}$ with probability $(\cdot)$ and zero otherwise. Hence:

$$E(W^2) = \int \frac{E[\hat{L}^2(\cdot; y)] (\cdot)^2}{(\cdot)g(\cdot)^2} (\cdot; y) g(\cdot) d\cdot dy = \int \frac{(\cdot)}{(\cdot)} g(\cdot) d\cdot; \quad (9)$$

where $(\cdot) = E \hat{L}^2 \frac{(\cdot)}{g(\cdot)}^2$ (which equals $E \hat{L}_{ABC}^2 \frac{(\cdot)}{g(\cdot)}^2$ in the ABC case) and $g(\cdot) = \int (\cdot; j) g(\cdot) d\cdot$.

The expected time of a single iteration of the LIS algorithm is

$$E(T) = N = T_1 + \int (\cdot) T_2(\cdot; \cdot) (\cdot; j) g(\cdot) d\cdot = T_1 + \int (\cdot) T_2(\cdot) g(\cdot) d\cdot; \quad (10)$$

Note that $E(W)$ is a constant, so choosing (λ) to maximise the expression for asymptotic efficiency in Section 4.1 is equivalent to minimising $E(W^2)E(T)=N$. Call this problem P . Consider also the problem $\mathcal{P}(\lambda)$, minimising $E(W^2)$ under the constraint $E(T)=N$ and $P(\lambda; \lambda)$, minimising $E(W^2) + \lambda [E(T)-N]$, or equivalently

$$Z = \frac{E(W^2)}{E(T)}$$

solution:

$$g(\theta) = \frac{q(\theta)}{T(\theta)}^{1/2} : \quad (13)$$

(Recall the LIS definition of $q(\theta)$ from Remark 3 and note that the lazy ABC definition in Corollary 1 can be derived from this.) Various choices of q , such as q_1 , give a solution which also meets the constraint on T . These all give algorithms which are equivalent to RW-IS with $g(\theta) / q(\theta) = \frac{q(\theta)}{T(\theta)}^{1/2}$ as claimed.

References

- Andrieu, C. and Roberts, G. O. (2009). The pseudo-marginal approach for efficient Monte Carlo computations. *The Annals of Statistics* pages 697{725.
- Beaumont, M. A. (2003). Estimation of population growth or decline in genetically monitored populations. *Genetics* 164(3):1139{1160.
- Beaumont, M. A. (2010). Approximate Bayesian computation in evolution and ecology. *Annual Review of Ecology, Evolution and Systematics* 41:379{406.
- Beaumont, M. A., Cornuet, J. M., Marin, J.-M., and Robert, C. P. (2009). Adaptive approximate Bayesian computation. *Biometrika*, 96(4):983{990.
- Buzbas, E. O. and Rosenberg, N. A. (2013). AABC: approximate approximate Bayesian computation when simulating a large number of data sets is computationally infeasible. Preprint. Available at <http://www.arxiv.org/abs/1301.6282>.
- Erhardt, R. J. and Smith, R. L. (2012). Approximate Bayesian computing for spatial extremes. *Computational Statistics & Data Analysis* 56(6):1468{1481.
- Fearnhead, P., Papaspiliopoulos, O., Roberts, G. O., and Stuart, A. (2010). Random-weight particle filtering of continuous time processes. *Journal of the Royal Statistical Society: Series B* 72(4):497{512.

- Fearnhead, P. and Prangle, D. (2012). Constructing summary statistics for approximate Bayesian computation: Semi-automatic ABC. *Journal of the Royal Statistical Society, Series B* 74:419{474.
- Geweke, J. (1989). Bayesian inference in econometric models using Monte Carlo integration. *Econometrica* 57(6):1317{1339.
- Liu, J. S. (1996). Metropolized independent sampling with comparisons to rejection sampling and importance sampling. *Statistics and Computing* 6(2):113{119.
- Liu, J. S., Chen, R., and Wong, W. H. (1998). Rejection control and sequential importance sampling. *Journal of the American Statistical Association*, 93(443):1022{1031.
- Marin, J.-M., Pudlo, P., Robert, C. P., and Ryder, R. J. (2012). Approximate Bayesian computational methods. *Statistics and Computing* 22(6):1167{1180.
- Meeds, E. and Welling, M. (2014). GPS-ABC: Gaussian process surrogate approximate Bayesian computation. Preprint. Available at <http://arxiv.org/abs/1404.4344>.

Wood, S. (2011). Fast stable restricted maximum likelihood and marginal likelihood estimation of semiparametric generalized linear models. *Journal of the Royal Statistical Society: Series B* 73(1):3{36.