

Department of Mathematics and Statistics

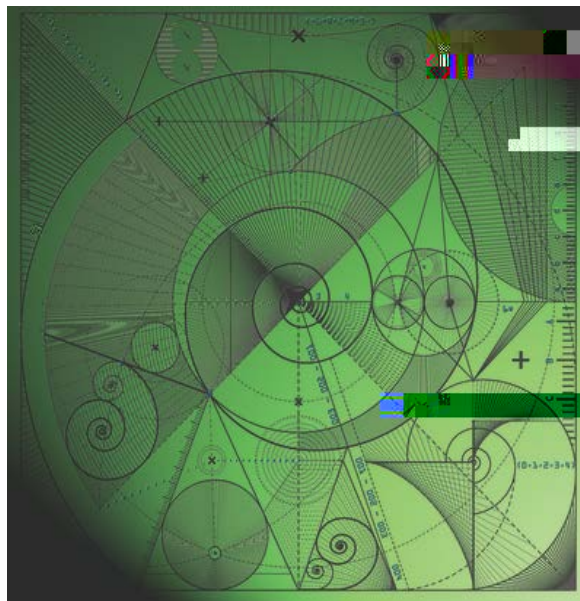
Preprint MPS-2014-11

16 April 2014

# Noisy Monte Carlo: Convergence of Markov chains with approximate transition kernels

by

P. Alquier, N. Friel, R. G. Everitt and A. Boland





1. large datasets: the sample size is too large. This situation is very common nowadays as huge databases can be stored at no cost. For example: in genomics the cost of sequencing has fallen by a factor of  $10^5$  in past decade and a half. This has led to the wide availability of sequence data - the recently announced Personal Genome Project UK aims to sequence 10 human genomes, each consisting of  $3 \times 10^9$  bases.
2. high-dimensional parameter spaces: the sample size might be reasonable, but the number of variables is too large. For example: data assimilation in numerical weather

the target distribution. We then study the special cases of a noisy version of the Exchange algorithm (Murray et al. (2006)), and discretized Langevin Monte Carlo in Section 3. For these noisy algorithms we prove that the total variation distance decreases with the number of iterations,  $N$ , of the randomisation step in the noisy algorithm, and find a bound on this distance in terms of  $N$ . We study in detail an application to intractable likelihood problems in Section 4.

## 2 Noisy MCMC algorithms

It turns out that a useful answer to this question is given by the study of the stability of Markov chains. There have been a long history of research on this topic, we refer the reader to the monograph by Kartashov (1996) and the references therein. Here, we will focus on a more recent method due to Mitrophanov (2005). In order to measure the distance between  $P$  and  $\hat{P}$  recall the definition of the total variation measure between two kernels:

$$\|P - \hat{P}\|_1 := \sup_k |P_k - \hat{P}_k|$$

$$9 N_0 \leq N; 0 < \alpha < 1; L > 0; 8N \leq N_0; \\ \sum_{j=0}^{\infty} V(j) \hat{P}_N(j; d) \leq V(j_0) + L:$$

$$\| \hat{P}_N - P \|_{k_V} \leq \frac{1}{N} \rightarrow 0.$$

Then there exists an  $N_1 \leq N$  such that any  $\hat{P}_N$ , for  $N \geq N_1$ , is geometrically ergodic with limiting distribution  $P$  and  $\| \hat{P}_N - P \|_{k_V} \leq \frac{1}{N} \rightarrow 0$ .

(We refer the reader to (Meyn and Tweedie 1993) for the definition of the  $k_V$  norm). Note that, in contrast to the previous result, we don't know explicitly the rate of convergence of the distance between  $\hat{P}_N$  and  $P$  when  $N$  is fixed. However it is possible to get an estimate of this rate (see Corollary 1 page 189 in (Ferre, Herve and Ledoux 2013)) under stronger assumptions.

## 2.1 Noisy Metropolis-Hastings

The Metropolis-Hastings (M-H) algorithm, sequentially draws candidate observations from a distribution, conditional only upon the last observation, thus inducing a Markov chain. The M-H algorithm is based upon the observation that a Markov chain with transition density  $P(j; i)$  and exhibiting detailed balance for  $\pi$ ,

$$\pi(i)P(j; i) = \pi(j)P(i; j);$$

---

**Algorithm 2** Noisy Metropolis-Hastings algorithm

---

```
for  $n = 0$  to  $I$  do
  Draw  $j^{(n)} \sim h(j | x^{(n)})$ 

  Draw  $y^{(n)} \sim F_{\theta}(y | x^{(n)})$ 

  Set  $x_{n+1} = j^{(n)}$  with probability  $\min(1, \hat{\pi}(j^{(n)} | x^{(n)}; y^{(n)}))$ 

  Otherwise, set  $x_{n+1} = x^{(n)}$ .
end for
```

---

Note that  $\hat{\pi}(j^{(n)} | x^{(n)}; y^{(n)})$  can be thought of as a randomised version of the Metropolis-Hastings acceptance probability.

Obviously, we expect that  $\hat{\lambda}$  is chosen in such a way that  $\frac{\partial}{\partial \lambda} \log \pi(\lambda) = 0$  and so in this case,



that the Langevin algorithm produces a Markov chain and we let  $\tilde{P}$  denote the corresponding transition kernel. Note that, we generally don't have  $(j|y)P = (j|y)$  nor  $P_0 \neq (j|y)$ , however, under some assumptions,  $P_0 \approx (j|y)$  for some  $\epsilon$  close to 0 when  $\epsilon$  is small enough, we discuss this in more detail below.

In practice, it is often the case that  $r \log(\pi_n)$  cannot be computed. Here again, a natural idea is to replace  $r \log(\pi_n)$  by an approximation or an estimate  $\hat{r}^{y_n} \log(\pi_n)$ , possibly using a randomization step  $y_n \sim F_n$ . This yields what we term a noisy Langevin algorithm.

---

#### Algorithm 4 Noisy Langevin algorithm

---

```

for n = 0 to I do
  Draw  $y_n \sim F_n(\cdot)$ .

  Set  $x_{n+1} = x_n + \frac{1}{2} \hat{r}^{y_n} \log(\pi_n | y_n) + C \sim N(0, \sigma^2)$  :
end for

```

---

Note that a similar algorithm has been proposed in (Welling and Teh 2011; Ahn, Korattikara and Welling 2012) in the context of big data situations, where the gradient of the logarithm of the target distribution is estimated using mini-batches of the data.

We let  $\hat{P}$  denote the corresponding transition kernel arising from Algorithm 4. We now prove that the Stochastic gradient Langevin algorithm, (Algorithm 4), will converge to the discrete-time Langevin diffusion with transition kernel resulting from Algorithm 3.

### 2.3 Towards theoretical guarantees for the noisy Langevin algorithm

In this case, the approximation guarantees are not as clear as they are for the noisy Metropolis-Hastings algorithm. To begin, there are two levels of approximation:

the transition kernel  $P$  targets a distribution  $\pi$  that might be far away from  $(j|y)$ .

Moreover, one does not simulate at each step from  $\pi$  but rather from  $\hat{P}$ .

The first point requires one to control the distance between  $\pi$  and  $(j|y)$ . Such an analysis is possible. Here we refer the reader to Proposition 1 in (Dalalyan and Tsybakov 2012) and also to Roberts and Stramer (Roberts and Stramer 2002) for different discretization schemes. It is possible to control  $k\hat{P} - Pk$  as Lemma 2.4 illustrates.

Lemma 2.4

$$kP - \hat{P}k \leq \frac{r}{2}$$

where

$$r = E_{y_n \sim F_n} \exp \frac{1}{2} \left( r \log(\pi_n) - \hat{r}^{y_n} \log(\pi_n) \right)^2 \leq 1 :$$

The paper by Roberts and Tweedie (1996a) contains a complete study of the chain generated by  $P$ . The problem is that it is not uniformly ergodic. So Theorem 2.1 is not the appropriate tool in this situation. However, in some situations, this chain is geometrically ergodic, and in this instance we can use Theorem 2.2 instead (moreover, note that Roberts and Tweedie (1996a) provide the function  $V$  used in the Theorem). We provide an example of such an application in Section 3 below.

## 2.4 Connection with the pseudo-marginal approach

There is a clear connection between this paper and the pseudo-marginal approaches described in (Beaumont 2003) and (Andrieu and Roberts 2009). In both cases a noisy acceptance probability is considered, but in pseudo-marginal approaches this is a consequence of using an estimate of the desired target distribution at each  $i$ , rather than the true value. Before proceeding further, we make precise some of the terminology used in (Beaumont 2003) and (Andrieu and Roberts 2009). These papers describe two alternative algorithms, the "Monte Carlo within Metropolis" (MCWM) approach, and "grouped independence MH" (GIMH). In both cases an unbiased importance sampling estimator  $\hat{\pi}_i$  is used in place of the desired target  $\pi$ , however the overall algorithms proceed slightly differently. The  $(i+1)$ th iteration of the MCWM algorithm is shown in algorithm 5.

---

### Algorithm 5 MCWM

---

for  $n = 0$  to  $I$  do

    Draw  $\theta^0 \sim h(\cdot | j_n)$ .

    Draw  $z^0 \sim G(\cdot | j_n)$ ,  $z \sim G(\cdot | j_n)$ , where  $G$  is an importance proposal and  $\theta^0$  and  $z$  are random vectors of size  $N$ .

    Calculate the acceptance probability,  $\alpha(j_n; j_n)$ , where  $b_z^N$  and  $b_{z^0}^N$  denote the 18 11.9552 Tf 5.03 0 Td

the auxiliary variables. The same argument holds when using any unbiased estimator of the target. As regards our focus in this paper, GIMH is something of a special case, and our framework has more in common with MCWM. We note that despite its exactness, there is no particular reason for estimators from GIMH to be more statistically efficient than those from MCWM.

where  $q(\cdot)$  denotes the prior distribution for  $\theta$ . For example, a naive application of the Metropolis-Hastings algorithm when proposing to move from  $\theta_i$  to  $\theta^0 = h(j, i)$  results in the acceptance probability,

$$\alpha(\theta^0; \theta_i) = \min \left( 1, \frac{q(\theta_i) h(j, i)}{q(\theta^0) h(i, j)} \frac{Z(\theta_i)}{Z(\theta^0)} \right); \quad (4)$$

depending on the intractable ratio  $\frac{Z(\theta_i)}{Z(\theta^0)}$ .

One method to overcome this computational bottleneck is to use an approximation of the likelihood  $f(y_j)$ . A composite likelihood approximation of the true likelihood, such as that of (Besag 1974), is most commonly used. This approximation consists of a product of easily normalised full-conditional distributions. The most basic composite likelihood is the pseudo likelihood which comprised of the product of full-conditional distributions of each,

$$f(y_j) \approx \prod_{i=1}^M f(y_i | y_{-i}; \theta);$$

However this approximation of the true likelihood can give unreliable estimates of (Friel and Pettitt 2004), (Friel et al 2009).

### 3.2 Exchange Algorithm

A more sophisticated approach is to use the Exchange algorithm. Murray et al. (2006) extended the work of Moller et al. (2006) to allow inference on doubly intractable distributions using the exchange algorithm. The algorithm samples from an augmented distribution

$$q(\theta; y^0; j_y) \propto f(y_j) q(\theta) h(j^0, j) f(y^0 | j^0)$$

whose marginal distribution for  $\theta$  is the posterior of interest. Here the auxiliary distribution  $f(y^0 | j^0)$  is the same likelihood model in which  $y$  is defined. By sampling from this augmented distribution, the acceptance formula simplifies, as can be seen in algorithm 6, where the normalising constants arising from the likelihood and auxiliary likelihood cancel. One difficulty of implementing the exchange algorithm is the requirement to sample  $y^0 \sim f(\cdot | j^0)$ , perfect sampling (Propp and Wilson 1996) is often possible for Markov random field models. However when the exchange algorithm is used with MRFs the resultant chains may not mix well. For example, Caimo and Friel (2011) used adaptive direction sampling (Gilks, Roberts and George 1994) to improve the mixing of the exchange algorithm when used with ERGM models.

Murray et al. (2006) proposed the following interpretation of the exchange algorithm. If we compare the acceptance ratios in the M-H and Exchange algorithm, the only difference is that the ratio of the normalising constants in the M-H acceptance probability  $\frac{Z(\theta_i)}{Z(\theta^0)}$  is replaced by  $q(y^0) = q_\theta(y^0)$  in the exchange probability. This ratio of un-normalised likelihoods



---

**Algorithm 7** Noisy Exchange algorithm
 

---

for  $n = 0$  to  $I$  do

    Draw  $j \sim h(\cdot | y_n)$ :

    for  $i = 1$  to  $N$  do

        Draw  $y_i^0 \sim f(\cdot | j^0)$ :

    end for

    Define  $y_0 = f(y_1^0, \dots, y_N^0)$

    Set  $y_{n+1} = j^0$  with probability  $\min(1, \hat{\alpha}(j^0; y_n, y_0))$ , where

$$\hat{\alpha}(j^0; y_n, y_0) = \frac{q_0(y) \prod_{i=1}^N h(j^0 | y_i^0)}{q_n(y) \prod_{i=1}^N h(j^0 | y_n)} \frac{1}{N} \sum_{i=1}^N \frac{q_n(y_i^0)}{q_0(y_i^0)}.$$

    Otherwise, set  $y_{n+1} = y_n$ .

end for

---

(A3) for any  $y$  and  $j^0$  in  $\mathcal{S}$ ,

$$\text{Var}_{y^0 \sim f(y^0 | j^0)} \frac{q_n(y^0)}{q_0(y^0)} < +1 :$$

Note that when (A1) or (A2) is satisfied, we necessarily have that  $\mathcal{S}$  is a bounded set, in this case, we put  $T = \sup_{k \in \mathcal{K}} \|k\|$ . This also means that  $0 < \exp(-TS) \leq q(y) \leq \exp(TS)$  for any  $y$  and  $S$ , we then put  $K := \exp(TS)$ . Also, note that this immediately implies Assumption (A3) because in this case,  $\text{Var}_{y^0 \sim f(y^0 | j^0)} (q_n(y^0) = q_0(y^0)) \leq K$

Note that Liang and Jin (2011) presents a similar algorithm to that above. However in contrast to Lemma 3.1, the results in (Liang and Jin 2011) do not explicitly provide a rate of approximation with respect to  $N$ . Lemma 2.2, page 9 in (Liang and Jin 2011) only states that there exists a  $N$  large enough to reach arbitrarily small accuracy  $> 0$ .

### 3.4 Noisy Langevin algorithm for Gibbs random fields

The discrete-time Langevin approximation (3) is unavailable for Gibbs random fields since the gradient of the log posterior,  $r \log(\pi(y))$  is analytically intractable, in general. However Algorithm 4 can be used using a Monte Carlo estimate of the gradient, as follows.

$$\begin{aligned} \log(\pi(y)) &= \tau s(y) - \log(Z(\tau)) + \log(\pi(y)) \\ r \log(\pi(y)) &= s(y) \frac{Z'(\tau)}{Z(\tau)} + r \log(\pi(y)) \\ &= s(y) \frac{\sum_{j=1}^N s(y_j) \exp(\tau s(y_j))}{\exp(\tau s(y))} + r \log(\pi(y)) \\ &= s(y) E_{y_j} [s(y)] + r \log(\pi(y)) \end{aligned} \quad (7)$$

In practice,  $E_{y_j} [s(y)]$  is usually not known - an exact evaluation of this quantity would require an evaluation of  $Z(\tau)$ . However, it is possible to estimate it through Monte-Carlo simulations. If we simulate  $y = (y_1^0, \dots, y_N^0) \sim \pi(\cdot | j)$ , then  $E_{y_j} [s(y)]$  can be estimated using  $\frac{1}{N} \sum_{i=1}^N s(y_i^0)$ . This gives an estimate of the gradient at  $\tau$  from (7).

$$r \log(\pi(y)) \approx s(y) \frac{1}{N} \sum_{i=1}^N s(y_i^0) + r \log(\pi(y))$$

In turn this yields the following noisy discretized Langevin algorithm.

---

#### Algorithm 8 Noisy discretized Langevin algorithm for Gibbs random fields

---

```

for n = 0 to I do
  for i = 1 to N do
    Draw  $y_i^0 \sim \pi(\cdot | j_n)$ :
  end for
  Define  $y_n = (y_1^0, \dots, y_N^0)$ ,

  Calculate  $r \log(\pi(y_n)) \approx r \log(\pi(y_n)) + s(y_n) \frac{1}{N} \sum_{i=1}^N s(y_i^0)$ :

  Set
     $y_{n+1} = y_n + \frac{\tau}{2} r \log(\pi(y_n)) + \epsilon_n$ ; where  $\epsilon_n$  are i.i.d.  $N(0, \sigma^2)$  :
end for

```

---

Remark that in this case, the bound in Lemma 2.4 can be evaluated.

Lemma 3.3 As soon as  $N > 4kS^2k^{-k^2}$ , the in Lemma 2.4 is finite with

$$= \exp \left( \frac{k \log(N)}{4S^2k^{-k^2}N} \right) \left( 1 + \frac{4k^p - S^k k}{N} \right)^{N+1} \frac{k \log \frac{N}{k}}{4S^2k^{-k^2}N}$$

(where  $k^{-k} = \sup_{f \in \mathcal{F}} \int f(x) dx$ ;  $k \leq 1$ ).

We conclude by an application of Theorem 2.2 that allows to assess the convergence of this scheme when  $N \rightarrow \infty$  when the parameter is real.

Theorem 3.4 Assume that  $\theta \in \mathbb{R}$  and the prior is Gaussian  $\theta \sim N(0; s^2)$



---

Algorithm 9

---

**Algorithm 10**    noisy MALA-exchange
 

---

Initialise; set  $\sigma$ ,  
 for  $i = 1$  to  $N$  do  
     Draw  $y_i \sim f(\cdot | j_0)$ :  
 end for  
 Define  $y_0 = f(y_1; \dots; y_N)g$ ,  
 Calculate  $\log \pi(y_0 | j_0) = r \log \pi(j_0) + s(y_0) - \frac{1}{N} \sum_{i=1}^N s(y_i)$ :

for  $n = 0$  to  $L$  do  
     Draw  $y_0^n = y_0 + \frac{\sigma}{\sqrt{2}} \sqrt{\log \pi(y_n | j_n) + \log \pi(j_n)} N(0; \sigma)$ .

    for  $i = 1$  to  $N$  do  
         Draw  $y_i^0 \sim f(\cdot | j_0)$ :  
     end for  
     define  $y_0^n = f(y_1^0; \dots; y_N^0)g$ .

    Calculate  $\log \pi(y_0^n | j_0^n) = r \log \pi(j_0^n) + s(y_0^n) - \frac{1}{N} \sum_{i=1}^N s(y_i^0)$ :

    Set  $y_{n+1} = y_0^n$  and  $y_{n+1} = y_0$  with probability  $\min(1, \gamma(j_0^n; j_n, y_n))$

    where  $\gamma(j_0^n; j_n, y_n) = \frac{q_0(y_0^n | j_0^n) h(j_n | j_0^n; y_n^0)}{q_n(y_0 | j_n) h(j_0^n | j_n; y_n^0)} \frac{1}{N} \sum_{i=1}^N \frac{q_n(y_i^0 | j_n)}{q_0(y_i^0 | j_0^n)}$

## 4.1 Ising study

The Ising model is defined on a rectangular lattice or grid. It is used to model the spatial distribution of binary variables, taking values  $-1$  and  $1$ . The joint density of the Ising model can be written as

$$f(y) = \frac{1}{Z(\theta)} \exp \left( \sum_{j=1}^M \sum_{i \sim j} \theta_{ij} y_i y_j \right)$$

where  $i \sim j$  denotes that  $i$  and  $j$  are neighbours and  $Z(\theta) = \sum_{y \in \{-1, 1\}^M} \exp \left( \sum_{j=1}^M \sum_{i \sim j} \theta_{ij} y_i y_j \right)$ .

The normalising constant  $Z(\theta)$  is rarely available analytically since this relies on taking the summation over all different possible realisations of the lattice. For a lattice with  $M$  nodes this equates to  $2^{\frac{M(M-1)}{2}}$  different possible lattice formations.

For our study, we simulated 20 grids of size  $16 \times 16$ . This size lattice is sufficiently small enough such that the normalising constant  $Z(\theta)$  can be calculated exactly (36.5 minutes for each graph) using a recursive forward-backward algorithm (Reeves and Pettitt 2004; Friel and Rue 2007), giving a gold standard with which to compare the other algorithms. This is done by calculating the exact density over a fine grid of values,  $\theta_{ij}$ , over the interval  $[-0.4; 0.8]$ , which cover the effective range of values that  $\theta_{ij}$  can take. We normalise  $f(y)$  by numerically integrating over the un-normalised density.

$$\hat{f}(y) = \prod_{i=2}^M \frac{\sum_{j=1}^{i-1} \theta_{ij} y_i y_j}{2} \frac{q_i(y)}{Z(\theta_{i-1})} y_i + \frac{q_{i-1}(y)}{Z(\theta_{i-1})} y_{i-1}; \quad (8)$$

yielding

$$f(y) = \frac{q_i(y)}{Z(\theta_{i-1})} \hat{f}(y).$$

Each of the algorithms was run for 30 seconds on each of the 20 datasets, at each iteration the auxiliary step to draw  $y$

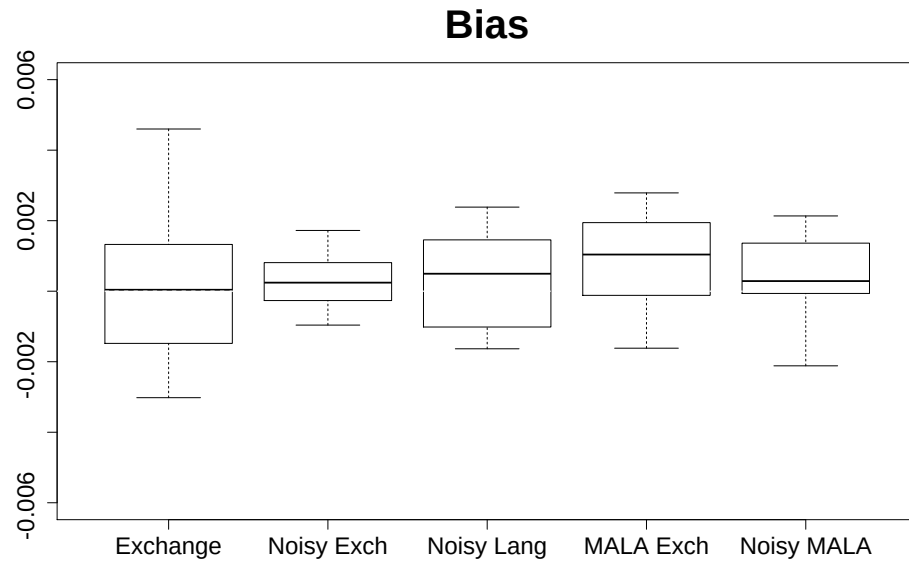


Figure 1: Boxplot of the bias estimate of for 20 datasets corresponding to the exchange, importance sampling exchange, Langevin and MALA algorithms.

Figure 1 shows the bias of the posterior means for each of the algorithms. We see that both

Figure 2: Estimated posterior densities corresponding to the exact and noisy algorithms corresponding to one of the datasets used in the Ising simulation study.

## 4.2 ERGM study

Here we explore how our algorithms may be applied to the exponential random graph model (ERGM) (Robins et al

#### 4.2.1 The Florentine Business dataset

Here, we consider a simple 16 node undirected graph: the Florentine family business graph. This concerns the business relations between some Florentine families in around 1430. The network is displayed in Figure 3. We propose to estimate the following 2-dimensional model.

$$f(y) = \frac{1}{Z(\theta)} \exp(\theta_1 s_1(y) + \theta_2 s_2(y));$$

where  $s_1(y)$  is the number of edges in the graph and  $s_2(y)$  is the number of two-stars.



Figure 3: Florentine family business.

Before we could run the algorithms, certain parameters had to be tuned. We used a Gaussian prior  $N(0;100)$  in all of the algorithms. The Langevin, MALA exchange and noisy MALA exchange algorithms all depend on a stepsize matrix  $\Sigma$ . This matrix determines the scale of proposal values for each of the parameters. This matrix should be set up so that proposed values for  $\theta$  accommodate the different scales of the posterior density of  $\theta$ . In order to have good mixing in the algorithms we chose a  $\Sigma$  which relates to the shape of the posterior density. Our approach was to aim to relate  $\Sigma$  to the covariance of the posterior density. To do this, we equated  $\Sigma$  to an estimate of the inverse of the second derivative of the log posterior at the maximum a posteriori estimate  $\theta^*$ . As the true value of the MAP is unknown, we used a Robbins-Monro algorithm (Robbins and Monro 1951) to estimate this. The Robbins-Monro algorithm takes steps in the direction of the slope of the distribution. It is very similar to Algorithm 8 except without the added noise and follows the stochastic process

$$\theta_{n+1} = \theta_n + \frac{1}{n} \nabla \log f(\theta_n);$$

where  $\sum_{i=0}^N \frac{1}{n} < \infty$  and  $\sum_{i=0}^N \frac{1}{n^2} < \infty$  :

The values of  $\delta$  decrease over time and once the difference between successive values of this process is less than a specified tolerance level, the algorithm is deemed to have converged to the MAP. The second derivative of the log posterior is derived by differentiating (7) yielding

$$r^{-2} \log(\delta_j)$$

MALA exchange but not in the Langevin algorithm. Since our Noisy Langevin algorithm approximates Langevin diffusion we are approximating an approximation. There are two levels of approximations which leaves more room for error.

| Method | Edge | 2-star |
|--------|------|--------|
|--------|------|--------|



Figure 5: Chains, density plot and ACF plot for the 2-star statistic.

#### 4.2.2 The Molecule dataset

The Molecule dataset is a 20 node graph, shown in Figure 6. We consider a four parameter model which includes the number of edges in the graph, the number of two-stars, the number of three-stars and the number of triangles.

$$f(y) = \frac{1}{Z(\theta)} \exp(\theta_1 s_1(y) + \theta_2 s_2(y) + \theta_3 s_3(y) + \theta_4 s_4(y))$$

The  $\theta$  parameter was chosen in a similar fashion to the Florentine business example. The Robbins-Monro algorithm was run for 20,000 iterations to find an estimate of the MAP, 4,000 graphs were then simulated at the estimated MAP and these were used to calculate an estimate of the second derivative using Equation (9). The matrix  $\hat{\Sigma}$  was the ofanthe estimate pa27(

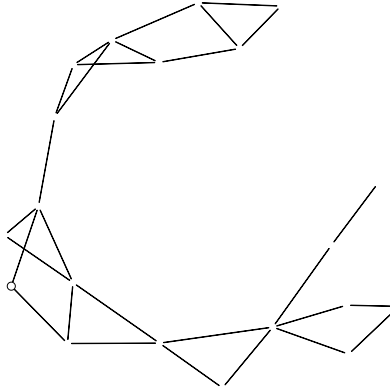


Figure 6: Molecule network

The BERGM algorithm of (Caimo and Friel 2011) was again used as a "ground truth". This algorithm was run for a large number of iterations equating to 4 hours of CPU time. This gave us accurate estimates against which to compare the various algorithms. The various algorithms were each run for 100 seconds of CPU time. Table 2 shows the posterior mean and standard deviations of each of the four parameters for each of the algorithms. The results for the Molecule dataset model are similar to the Florentine business dataset model. In Table 2 we see that the noisy exchange algorithm improved on the standard exchange algorithm. The MALA exchange improved on noisy Langevin and the Noisy MALA improved on the MALA exchange.

Figure 7 and Figure 8 show the densities and the autocorrelation plots of the algorithms. The autocorrelation plots show that the noisy algorithms had less correlation than the exchange algorithm. The densities show that again the algorithms, when run on the Molecule model, performed in the same manner as the Florentine model. The algorithms with the exception of the noisy Langevin algorithm estimated the mode well but underestimated the standard deviation. The noisy Langevin algorithm did not estimate the mean or standard deviations well.

| Method          | Edge  |       | 2-star |       | 3-Star |       | Triangle |       |
|-----------------|-------|-------|--------|-------|--------|-------|----------|-------|
|                 | Mean  | SD    | Mean   | SD    | Mean   | SD    | Mean     | SD    |
| BERGM           | 2.647 | 2.754 | -1.069 | 0.953 | -0.021 | 0.483 | 1.787    | 0.646 |
| Exchange        | 1.889 | 2.142 | -0.797 | 0.744 | -0.138 | 0.385 | 1.593    | 0.519 |
| Noisy Exch      | 1.927 | 2.444 | -0.757 | 0.823 | -0.176 | 0.422 | 1.543    | 0.53  |
| Noisy Lang      | 1.679 | 3.65  | -0.509 | 1.429 | -0.466 | 0.787 | 1.633    | 0.573 |
| MALA Exch       | 2.391 | 2.095 | -0.938 | 0.795 | -0.113 | 0.451 | 1.454    | 0.598 |
| Noisy MALA Exch | 2.731 | 2.749 | -1.054 | 0.886 | -0.041 | 0.417 | 1.519    | 0.492 |

Table 2: Posterior means and standard deviations.

Figure 7: Density plots of the 4 parameters for the molecule example.



result to hold for noisy MCMC algorithms, in which case the effect of this additional variance on top of the aforementioned bias should be a consideration when employing noisy MCMC.

A further area for future work lies in relaxing the requirement for the ideal non-noisy chain to be uniformly ergodic. This property does not hold in many cases: the results in this paper are intended as the first steps towards future work that would obtain results that hold more generally.

### Acknowledgements

The Insight Centre for Data Analytics is supported by Science Foundation Ireland under Grant Number SFI/12/RC/2289. Nial Friel's research was also supported by an Science Foundation Ireland grant: 12/IP/1424.

## References

- Ahn, S., A. Korattikara and M. Welling (2012), Bayesian Posterior Sampling via Stochastic Gradient Fisher Scoring. In Proceedings of the 29th International Conference on Machine Learning
- Andrieu, C. and G. Roberts (2009), The pseudo-marginal approach for efficient Monte-Carlo computations. *The Annals of Statistics* 37(2), 697{725
- Andrieu, C. and M. Vihola (2012), Convergence properties of pseudo-marginal Markov

- Friel, N. and A. N. Pettitt (2004), Likelihood estimation and inference for the autologistic model. *Journal of Computational and Graphical Statistics* 13, 232{246
- Friel, N., A. N. Pettitt, R. Reeves and E. Wit (2009), Bayesian inference in hidden Markov random fields for binary data defined on large lattices. *Journal of Computational and Graphical Statistics* 18, 243{261
- Friel, N. and H. Rue (2007), Recursive computing and simulation-free inference for general factorizable models.



$$\begin{aligned}
&= (d^0) \int_0^Z \int_0^Z dt dy^0 h(tj) F_t(y^0) \min(1, \wedge(\cdot; t; y^0)) \min(1, (\cdot; t)) \\
&\quad + \int_0^Z dy^0 F_o(y^0) h(\cdot^0j) \min(1, (\cdot; \cdot^0)) h(\cdot^0j) \min(1, \wedge(\cdot; \cdot^0j))
\end{aligned}$$

d



Now, note that

$$\begin{aligned} & \int_{-\infty}^{\infty} \frac{1}{2} \exp \left( -\frac{1}{2} k^T \left( \frac{1}{2} (r \log(\cdot) - r^{\wedge y_0} \log(\cdot)) \right) \right) \frac{1}{8} k^{\frac{1}{2} (r \log(\cdot) - r^{\wedge y_0} \log(\cdot))} k^2 dt \\ & = E \int_{-\infty}^{\infty} \exp \left( -\frac{1}{2} a^T X \right) \frac{kak^2}{2} \end{aligned}$$

where  $X \sim N(0, I)$  and  $a = \frac{1}{2} [r \log(\cdot) - r^{\wedge y_0} \log(\cdot)] = 2$ . Then:

$$\begin{aligned} E \int_{-\infty}^{\infty} \exp \left( -\frac{1}{2} a^T X \right) \frac{kak^2}{2} &= \exp \left( -\frac{kak^2}{2} \right) E \exp \left( \frac{1}{2} a^T X \right) \exp \left( \frac{kak^2}{2} \right) \\ &= \exp \left( -\frac{kak^2}{2} \right) E \exp \left( \frac{1}{2} a^T X \right) E \exp \left( \frac{1}{2} a^T X \right) \\ &= \exp \left( -\frac{kak^2}{2} \right) \frac{1}{\sqrt{\text{Var}[\exp(a^T X)]}} \\ &= \exp \left( -\frac{kak^2}{2} \right) \frac{1}{\sqrt{E[\exp(2a^T X)] - E[\exp(a^T X)]^2}} \\ &= \exp \left( -\frac{kak^2}{2} \right) \frac{1}{\sqrt{\exp(2kak^2) - \exp(kak^2)}} \\ &= \frac{1}{\exp(kak^2)} = 1. \end{aligned}$$

So finally,

$$kP = \hat{P}$$

$$p_{\frac{1}{N}}^{\frac{h(j, \theta)}{h(\bar{q})} - \frac{(\theta)q_0(y)}{(\bar{q})q(y)}} \sqrt{\text{Var}_{y_1^0} f(y_{1j}^0, \theta)} \frac{q_n(y_1^0)}{q_0(y_1^0)} :$$

Proof of Theorem 3.2. Under the assumptions of Theorem 3.2, note that (4) leads to

$$(\bar{q}_n, \theta) = \frac{(\theta)q_0(y)Z(\bar{q}_n)h(\bar{q}_n, \theta)}{(\bar{q}_n)q_n(y)Z(\theta)h(\bar{q}_n)} \frac{1}{c^2 c_h^2 K^4} : \tag{10}$$

Let us consider any measurable subset  $B$  of  $\mathbb{R}^2$ . We have

$$P(\cdot; B) = \int_B (d\theta) \frac{1}{Z} \int_Z dth(t, j) \min(1, (\cdot; t)) + \int_B d\theta \frac{1}{Z}$$

with  $m = \frac{\log(1=C)}{\log(\frac{1}{2})}$ .

Proof of Lemma 3.3. Note that

So we have to find an upper bound, uniformly over  $\lambda$ , for

Let us put  $V := \frac{1}{N} \sum_{i=1}^N V^{(i)} := \frac{1}{N} \sum_{i=1}^N \frac{1}{2} f(s(y_i^0) - E_{y^0} f[s(y^0)]g)$  and denote  $V_j$  ( $j = 1; \dots; k$ )

$\exp^k$

and so  $P$  is geometrically ergodic with function  $V$ . We calculate

$$\begin{aligned} \int_{\mathbb{R}} V(x) P^n(x; dx) &= E_{y^0} \left[ \int_{\mathbb{R}} \frac{1}{2} V(x) \exp \left( -\frac{r^{y^0} \log(x/y)}{2} \right) dx \right] \\ &= E_{y^0} \left[ \int_{\mathbb{R}} \frac{1}{2} V(x) dx \right] + \frac{1}{2} (r^{y^0} \log(x/y)) \end{aligned}$$