



Computer Intensive Statistics

Andi Wang
andi.wang@warwick.ac.uk

20th–24th July, 2026

Compiled: July 20, 2026

Part 1

Introduction, Motivation & Basics

What is Computer Intensive Statistics

Computer, *n.* A device or machine for performing or facilitating calculation.

Compare Middle French computeur person who makes calculations (1578).

Intensive, *adj.* Of very high degree or force, vehement.

French intensif, -ive (14–15th cent. in Hatzfeld & Darmesteter).

Statistics, *n.* The systematic collection and arrangement of numerical facts or data of any kind; (also) the branch of science or mathematics concerned with the analysis and interpretation of numerical data and appropriate ways of gathering such data.

In early use after French statistique and German Statistik.

What Makes Statistics Computer Intensive?

Some *good* reasons for using computer-intensive methods:

Complexity Complex models cannot often be dealt with analytically.

Intractability Models which are not available analytically.

Laziness Computer time is cheap; human time isn't.

Scale Large data sets bring fresh challenges.

We won't address the *bad* reasons here. . .

Vevox.app 124-467-150

What is your familiarity with Computer Intensive Statistics?

What Makes Statistics Computer Intensive?

Some *good* reasons for using computer-intensive methods:

Complexity Complex models cannot often be dealt with analytically.

Intractability Models which are not available analytically.

Laziness Computer time is cheap; human time isn't.

Scale Large data sets bring fresh challenges.

We won't address the *bad* reasons here. . .

Vevox.app 124-467-150

What is your familiarity with Computer Intensive Statistics?

Part 1— Section 1

Motivation

Problems

Motivating Problem: Decision making under uncertainty

Decision making under uncertainty is at the heart of science / industry / society.

Modelling the world

There is some underlying state of the world $\theta \in \Theta$, which we have access to through observations $\mathbf{y}$, connected via a model $f(\mathbf{y}|\theta)$.

Suppose we have a set of actions $\mathcal{A}$ and a loss function ℓ : we would like to choose the 'best' action $a \in \mathcal{A}$, to minimize

$$\mathcal{L}(a) = \mathbb{E}[\ell(\theta, a)]. \quad (1)$$

Here the $\mathbb{E}$ represents the uncertainty regarding the 'true' θ .

What is the appropriate distribution for θ , given our observations $\mathbf{y}$? How do we compute (1)?

Motivating Problem: Decision making under uncertainty

Decision making under uncertainty is at the heart of science / industry / society.

Modelling the world

There is some underlying state of the world $\theta \in \Theta$, which we have access to through observations $\mathbf{y}$, connected via a model $f(\mathbf{y}|\theta)$.

Suppose we have a set of actions $\mathcal{A}$ and a loss function ℓ : we would like to choose the 'best' action $a \in \mathcal{A}$, to minimize

$$\mathcal{L}(a) = \mathbb{E}[\ell(\theta, a)]. \quad (1)$$

Here the $\mathbb{E}$ represents the uncertainty regarding the 'true' θ .

What is the appropriate distribution for θ , given our observations $\mathbf{y}$? How do we compute (1)?

Problems

Motivating Problem: Decision making under uncertainty

Decision making under uncertainty is at the heart of science / industry / society.

Modelling the world

There is some underlying state of the world $\boldsymbol{\theta} \in \Theta$, which we have access to through observations $\mathbf{y}$, connected via a model $f(\mathbf{y}|\boldsymbol{\theta})$.

Suppose we have a set of actions $\mathcal{A}$ and a loss function ℓ : we would like to choose the 'best' action $a \in \mathcal{A}$, to minimize

$$\mathcal{L}(a) = \mathbb{E}[\ell(\boldsymbol{\theta}, a)]. \quad (1)$$

Here the $\mathbb{E}$ represents the uncertainty regarding the 'true' $\boldsymbol{\theta}$.

What is the appropriate distribution for $\boldsymbol{\theta}$, given our observations $\mathbf{y}$? How do we compute (1)?

Motivating Problem: Bayesian Inference

Bayesian statistics

- Data $\mathbf{y}_1, \dots, \mathbf{y}_n$ and model $f(\mathbf{y}_i|\boldsymbol{\theta})$ where $\boldsymbol{\theta}$ is some parameter of interest.

$$\text{Likelihood } L(\boldsymbol{\theta}; \mathbf{y}_1, \dots, \mathbf{y}_n) = \prod_{i=1}^n f(\mathbf{y}_i|\boldsymbol{\theta})$$

- In the Bayesian framework $\boldsymbol{\theta}$ is a random variable with prior distribution $f^{\text{prior}}(\boldsymbol{\theta})$. After observing $\mathbf{y}_1, \dots, \mathbf{y}_n$, the posterior density of $\boldsymbol{\theta}$ is

$$\begin{aligned} f^{\text{post}}(\boldsymbol{\theta}) &= f(\boldsymbol{\theta}|\mathbf{y}_1, \dots, \mathbf{y}_n) \\ &= \frac{f^{\text{prior}}(\boldsymbol{\theta})L(\boldsymbol{\theta}; \mathbf{y}_1, \dots, \mathbf{y}_n)}{\int_{\Theta} f^{\text{prior}}(\boldsymbol{\vartheta})L(\boldsymbol{\vartheta}; \mathbf{y}_1, \dots, \mathbf{y}_n) d\boldsymbol{\vartheta}} \end{aligned}$$

- Often this is intractable—we need an approximation.

Motivating Problem: Hypothesis Testing

Testing Example: Chi-Squared Test of goodness of fit

- $T = \sum_{k=1}^K \frac{(O_k - E_k)^2}{E_k}$
- Asymptotic argument: $T \overset{d}{\approx} \chi_{K-1}^2$ under regularity conditions.

What if we *don't* have many observations of every category?

What if we don't know the form of their distributions?

Motivating Problem: Medians

Testing Medians

- Two populations $X_1, X_2, \dots, X_{n_X} \stackrel{\text{iid}}{\sim} f_X$ and $Y_1, Y_2, \dots, Y_{n_Y} \stackrel{\text{iid}}{\sim} f_Y$.
- $X_{[1]} \leq X_{[2]} \leq \dots \leq X_{[n_X]}$ are the associated order statistics.
- $T_X = X_{[(n+1)/2]}$ is the sample median; T_Y defined analogously.
- We want to know whether the *medians* of the two populations are *significantly different*.

What if we don't know the form of their distributions f_X, f_Y ?

We don't know the distributions of T_X, T_Y under H_0 and can't sample from it!

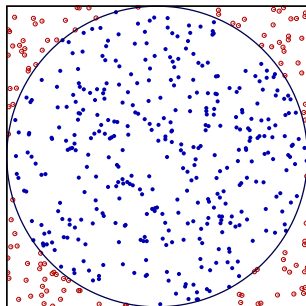
Simulation-based Methods

- Doing statistics backwards:

*Representing the solution of a problem as a parameter of a hypothetical population, and using a random sequence of numbers to construct a sample of the population, from which statistical estimates of the parameter (*p values, confidence intervals, or other quantities of interest*) can be obtained.*

Preliminary Example: Raindrop experiment for π

- Consider “uniform rain” on the square $[-1, 1] \times [-1, 1]$, i.e. the two coordinates $X, Y \stackrel{\text{iid}}{\sim} \text{U}[-1, 1]$.
- Probability that a rain drop falls in the circle is



$$\begin{aligned}
 \mathbb{P}(\text{drop within circle}) &= \frac{\text{area of the unit circle}}{\text{area of the square}} \\
 &= \frac{\int\int_{\{x^2+y^2 \leq 1\}} 1 \, dx dy}{\int\int_{\{-1 \leq x, y \leq 1\}} 1 \, dx dy} = \frac{\pi}{2 \cdot 2} = \frac{\pi}{4}.
 \end{aligned}$$

Preliminary Example: Raindrop experiment for π

- Given π , we can compute $\mathbb{P}(\text{drop within circle}) = \frac{\pi}{4}$.
- Given n independent raindrops, the number of rain drops falling in the circle, Z_n is a binomial random variable:

$$Z_n \sim \text{Bin} \left(n, p = \frac{\pi}{4} \right).$$

- So we can estimate p with

$$\hat{p} = \frac{Z_n}{n},$$

- and π by

$$\hat{\pi} = 4\hat{p} = 4 \cdot \frac{Z_n}{n}.$$

Preliminary Example: Raindrop experiment for π

- Result obtained for $n = 100$ raindrops:
77 points inside the circle.
- Resulting estimate of π is

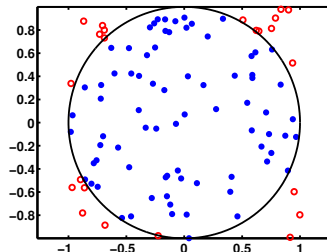
$$\hat{\pi} = \frac{4 \cdot Z_n}{n} = \frac{4 \cdot 77}{100} = 3.08,$$

(rather poor estimate).

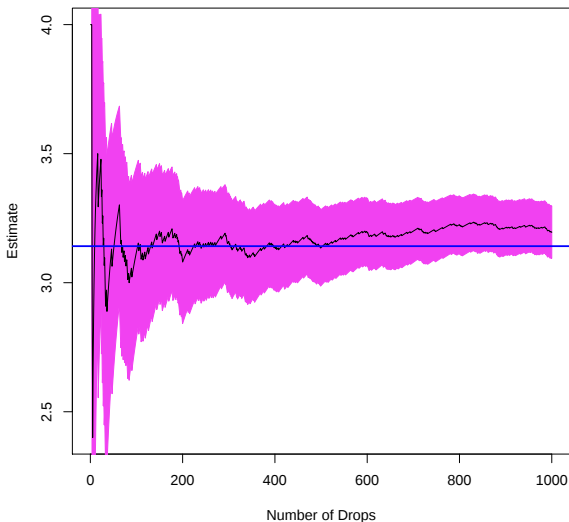
- However: the *law of large numbers* **guarantees** that

$$\hat{\pi}_n = \frac{4 \cdot Z_n}{n} \rightarrow \pi$$

almost surely for $n \rightarrow \infty$.



Preliminary Example: Raindrop experiment for π



Preliminary Example: Raindrop experiment for π

- How fast does $\hat{\pi}$ converge to π ?
Central limit theorem gives the answer.

- $(1 - 2\alpha)$ confidence interval for p ($\hat{p}_n = Z_n/n$):

$$\left[\hat{p}_n - z_{1-\alpha} \sqrt{\frac{\hat{p}_n(1 - \hat{p}_n)}{n}}, \hat{p}_n + z_{1-\alpha} \sqrt{\frac{\hat{p}_n(1 - \hat{p}_n)}{n}} \right]$$

- $(1 - 2\alpha)$ confidence interval for π ($\hat{\pi}_n = 4\hat{p}_n$):

$$\left[\hat{\pi}_n - z_{1-\alpha} \sqrt{\frac{\hat{\pi}_n(4 - \hat{\pi}_n)}{n}}, \hat{\pi}_n + z_{1-\alpha} \sqrt{\frac{\hat{\pi}_n(4 - \hat{\pi}_n)}{n}} \right]$$

- Width of the interval is $O(n^{-1/2})$, thus speed of convergence $O_{\mathbb{P}}(n^{-1/2})$.

Preliminary Example: Raindrop experiment for π

Recall the two core elements of this example:

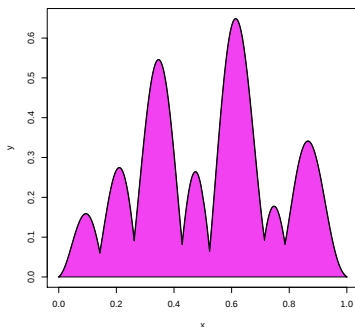
- 1 Write the quantity of interest (here π) as an expectation:

$$\pi = 4\mathbb{P}(\text{drop within circle}) = \mathbb{E}\left(4 \cdot \mathbb{I}_{\{\text{drop within circle}\}}\right)$$

- 2 Replace this algebraic representation with a sample approximation.
 - SLLN guarantees that the sample approximation converges to the algebraic representation.
 - CLT gives information about the speed of convergence.

The Generalisation to Monte Carlo Integration

$$f : [0, 1] \rightarrow [0, 1]$$



$$\int_0^1 f(x) dx = \int_0^1 \int_0^{f(x)} 1 dt dx = \iint_{\{(x,t): t \leq f(x)\}} 1 dt dx = \frac{\iint_{\{(x,t): t \leq f(x)\}} 1 dt dx}{\iint_{\{0 \leq x, t \leq 1\}} 1 dt dx}.$$

Comparison of the speed of convergence

- Monte Carlo integration is $O_{\mathbb{P}}(n^{-1/2})$.
- Numerical integration of a *one-dimensional* function by Riemann sums is $O(n^{-1})$.
- Monte Carlo does not compare favourably for one-dimensional problems.
- However:
 - Monte Carlo estimates are often *unbiased*.
 - Order of convergence of Monte Carlo integration is *independent* of dimension.
 - Order of convergence of numerical integration techniques deteriorates with increasing dimension.

Monte Carlo methods can be a good choice for high-dimensional integrals.

Views of Simulation-based Inference

Direct approximation of a quantity of interest.

- Careful construction of random experiment for particular task at hand.
- Justify with a dedicated argument in each case.

Approximation of *integrals* of interest.

- Represent quantity of interest as expectation w.r.t. some f .
- Use sample average to approximate expectation.
- Appeal to SLLN and CLT.

Approximation of *distributions* of interest.

- Represent quantity of interest as a function of distribution f .
- Use empirical measure of sample to approximate f .
- Appeal to Glivenko–Cantelli theorem.

Theoretical Motivation of Sample Approximation

Theorem (Strong Law of Large Numbers)

Let $X_1, X_2, \dots \stackrel{iid}{\sim} f$, and let $\varphi : E \rightarrow \mathbb{R}$ with $\mathbb{E} [|\varphi(X_1)|] < \infty$.
Then:

$$\frac{1}{n} \sum_{i=1}^n \varphi(X_i) \xrightarrow{a.s.} \mathbb{E} [\varphi(X_1)] .$$

Theorem (Central Limit Theorem)

Let $X_1, \dots \stackrel{iid}{\sim} f_X$ and let $\varphi : E \rightarrow \mathbb{R}^k$ with $\Sigma = \text{Var} [\varphi(X)] < \infty$.
Then as $n \rightarrow \infty$:

$$\sqrt{n} \left[\frac{1}{n} \sum_{i=1}^n \varphi(X_i) - \mathbb{E} [\varphi(X_1)] \right] \xrightarrow{\mathcal{D}} N(\mathbf{0}, \Sigma) .$$

Theoretical Motivation of Sample Approximation

Theorem (Glivenko–Cantelli)

Let $X_1, \dots \stackrel{iid}{\sim} f_X$ have cdf F_X .

Let

$$F_n(x) = \frac{1}{n} \sum_{i=1}^n \mathbb{I}_{(-\infty, x]}(X_i).$$

Then as $n \rightarrow \infty$:

$$\sup_x |F_n(x) - F(x)| \xrightarrow{a.s.} 0.$$

Part 1— Section 2

Randomized Testing

Randomized Testing

- One simple example of computer intensive statistics.
- We'll revisit *how* we can implement these things later.
- Art of testing: find a set R_α such that

$$\mathbb{P}(T \in R_\alpha; H_0) = \alpha$$

and

$$\mathbb{P}(T \in R_\alpha; H_1) > \alpha.$$

- What if we don't know the distribution of the test statistic, f_T ?

Is a Die Fair?

- Given n rolls of a die, we want to establish whether it's fair.
- Canonical example of a χ^2 -test. . .
- Compute

$$T = \sum_{k=1}^K \frac{(O_k - E_k)^2}{E_k}$$

- $T \stackrel{\text{approx}}{\sim} \chi_{K-1}^2$ by asymptotic arguments.
- What if the asymptotics don't hold?

A Randomized Goodness of Fit Test

- Imagine we have 9 measured rolls (and can't easily obtain more):

Value	1	2	3	4	5	6
Count	0	1	0	2	2	4

- If the die is fair we *expect* 1.5 observations of each value.
- The test statistic is:

$$T = \frac{1.5^2 + 0.5^2 + 1.5^2 + 0.5^2 + 0.5^2 + 2.5^2}{1.5} = 7\frac{2}{3}$$

- The asymptotics *certainly* don't hold:

$$(O_k - E_k)^2 \in \{0.5^2, 1.5^2, 2.5^2, 3.5^2, 4.5^2, 5.5^2, 6.5^2, 7.5^2\}.$$

- But we can *simulate* from H_0 .

An R Implementation

Randomized Goodness of Fit Testing: Setup

```
p <- 1/6 * c(1,1,1,1,1,1)
n <- 9
r <- 10000
ob <- rmultinom(r,n,p)
ex <- n*p
T <- colSums((ob - ex)^2/ex)
```

How many elements in T are larger than the observed value?

Randomized Goodness of Fit Testing: Comparison

```
t <- 23/3
m <- sum(T >= (t - 1E-9)) #T discrete
print(m/r)
```

Randomized testing: results

Does this look fair? Vote!

Vevox.app 124-467-150

Value	1	2	3	4	5	6
Count	0	1	0	2	2	4

Randomized Tests

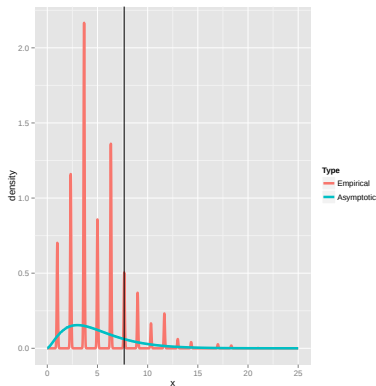
Randomized testing: results

Empirical p -value:

0.1848

Asymptotic p -value:

0.1860



Randomized Test in General

- Given a hypothesis, H_0 and an alternative, H_1 , and data $\mathbf{x}$ which realises $\mathbf{X}$ under H_0 :
 - Obtain a realisation $\mathbf{u}$ of $\mathbf{U}$
 ($\mathbf{U}|\mathbf{X} \sim f_{\mathbf{U}|\mathbf{X}}$ from some known distribution).
 - Compute R_α such that $\mathbb{P}((\mathbf{X}, \mathbf{U}) \in R_\alpha; H_0) = \alpha$.
 - Reject H_0 if $(\mathbf{x}, \mathbf{u}) \in R_\alpha$.

Goodness of Fit Test in General Form

- Let $f_{\mathbf{U}|\mathbf{X}}(\mathbf{u}|\mathbf{x}) = \prod_{i=1}^r f_{T(\mathbf{X})}(u_i; H_0)$.
 In practice: sample $\mathbf{Z}_i \stackrel{\text{iid}}{\sim} f_{\mathbf{X}}(\cdot; H_0)$ and set $U_i = T(\mathbf{Z}_i)$, where $T(\mathbf{X})$ is a real-valued summary of $\mathbf{X}$.
- Let $R_\alpha = \{(\mathbf{x}, \mathbf{u}) : T(\mathbf{x}) > u_{[r(1-\alpha)]}\}$, where $u_{[i]}$ is the i^{th} order statistic.

Are Those Medians Different (Part I)?

- Consider testing for different medians:

$$H_0 : X_1, \dots, X_{n_X} \stackrel{\text{iid}}{\sim} f_X(\cdot; m) \quad Y_1, \dots, Y_{n_Y} \stackrel{\text{iid}}{\sim} f_Y(\cdot; m)$$

$$H_1 : X_1, \dots, X_{n_X} \stackrel{\text{iid}}{\sim} f_X(\cdot; m) \quad Y_1, \dots, Y_{n_Y} \stackrel{\text{iid}}{\sim} f_Y(\cdot; m')$$

- And we'll assume a particular example for the form of the two distributions:

$$f_X(x; m) = f_Y(x; m) = \frac{1}{2} \exp(-|x - m|)$$

- Letting $\tilde{X} = X_{[(n_X+1)/2]}$ and $\tilde{Y} = Y_{[(n_Y+1)/2]}$:

$$\begin{aligned} \tilde{X} - \tilde{Y} &= (\tilde{X} - m) - (\tilde{Y} - m) \\ &= (X - m)_{[(n_X+1)/2]} - (Y - m)_{[(n_Y+1)/2]} \end{aligned}$$

- So the distribution of $\tilde{X} - \tilde{Y}$ is *independent* of $m | H_0$.

Randomized Tests

- A Randomized test:
 - Let $T = \tilde{X} - \tilde{Y}$.
 - Draw $i = 1, \dots, r$ copies of $\mathbf{X}$ and $\mathbf{Y}$ with $m = 0$:

$$X'_{1,\dots,n_X,i} \stackrel{\text{iid}}{\sim} f_X(\cdot; 0),$$

$$Y'_{1,\dots,n_Y,i} \stackrel{\text{iid}}{\sim} f_Y(\cdot; 0).$$

- Compute the difference between their medians:

$$i = 1, \dots, r : \quad T'_i = X'_{[(n_X+1)/2],i} - Y'_{[(n_Y+1)/2],i}.$$

- Let $p = (1 + |\{i : T'_i \geq T\}|)/(r + 1)$.
- Reject H_0 if $p < \alpha$ (a one-sided test; $H_1 : m' < m$).

But surely this is cheating: what if we *don't* know so much (like f_X and f_Y)?

Permutation Tests

- Consider the hypotheses:

$$H_0 : \quad X_1, \dots, X_{n_X} \stackrel{\text{iid}}{\sim} f_X(\cdot) \quad Y_1, \dots, Y_{n_Y} \stackrel{\text{iid}}{\sim} f_Y(\cdot) \\ f_X = f_Y \text{ (so in particular } F_X^{-1}(0.5) = F_Y^{-1}(0.5))$$

$$H_1 : \quad X_1, \dots, X_{n_X} \stackrel{\text{iid}}{\sim} f_X(\cdot) \quad Y_1, \dots, Y_{n_Y} \stackrel{\text{iid}}{\sim} f_Y(\cdot) \\ F_X^{-1}(0.5) \neq F_Y^{-1}(0.5)$$

where f_X and f_Y are *unknown*.

- Here, F_X^{-1} and F_Y^{-1} are assumed to exist.
- Sample medians are natural test statistics, but:
 - We don't know their distribution under H_0 .
 - And can't sample from that distribution.
- What can we do?

Permutation Tests

- Let $\mathbf{Z} = (X_1, \dots, X_{n_X}, Y_1, \dots, Y_{n_Y})$ be an $n = n_X + n_Y$ vector.
- Now let

$$T(\mathbf{Z}) = \text{median}(Z_1, \dots, Z_{n_X}) - \text{median}(Z_{n_X+1}, \dots, Z_n)$$

- And let $\pi \in \mathcal{P} \subseteq \{1, \dots, n\}^n$ denote a permutation, writing:

$$\pi \mathbf{Z} := (Z_{\pi_1}, Z_{\pi_2}, \dots, Z_{\pi_n})$$

- Now, under H_0 :

$$\forall \pi \in \mathcal{P} : \quad T(\pi \mathbf{Z}) \stackrel{\mathcal{D}}{=} T(\mathbf{Z})$$

- So we can construct a valid p -value by checking if $T(\mathbf{Z}) > T(\pi \mathbf{Z})$ for each π .
- We *just* need to compute $T(\pi \mathbf{Z})$ for every $\pi \in \mathcal{P} \dots$

A Randomized Permutation Test

- We can sample elements uniformly from $\mathcal{P}$:
 - Sample $\pi_1 \sim U(1, \dots, n)$.
 - Sample $\pi_2 \sim U(\{1, \dots, n\} \setminus \{\pi_1\})$.
 - $\vdots$
 - Sample $\pi_n \sim U(\{1, \dots, n\} \setminus \{\pi_1, \dots, \pi_{n-1}\})$.
- We can do this many times to approximate the law of $T(\pi \mathbf{z})$ when $\pi \sim U(\mathcal{P})$:
 - Sample $\boldsymbol{\pi}_1, \dots, \boldsymbol{\pi}_k \stackrel{\text{iid}}{\sim} U(\mathcal{P})$.
 - Compute $T_1 = T(\boldsymbol{\pi}_1 \mathbf{z}), \dots, T_k = T(\boldsymbol{\pi}_k \mathbf{z})$.
 - Use the empirical distribution of $(T_1, \dots, T_k)$ to approximate the law of $T(\boldsymbol{\pi} \mathbf{z})$.
- This provides a general strategy for nonparametric testing.

Summary of Part 1

- Motivation: decision making, Bayesian inference, hypothesis testing, . . .
- Towards simulation-based inference (see later).
- Randomized Tests
- Permutation Tests
- Corenflos, A. (2024). Computationally efficient statistical inference in Markovian models [PhD thesis, Aalto University].
- Young, G. A. (1994) Bootstrap: More than a stab in the dark? Statistical Science, 9, 382–395.

Part 2

Simulation and the Monte Carlo Method

Simulation

- We've seen *motivation* of simulation for inference.
- We've seen *examples* of simulation-based methods.
- Now we need methods for simulation.

Part 2— Section 3

The Monte Carlo Method

Monte Carlo Method

- A generic scheme for approximating expectations.
- To approximate $I = \mathbb{E}_f [\varphi(X)]$,
 - E.g. recall $\mathcal{L}(a) = \mathbb{E}_f[\ell(\boldsymbol{\theta}, a)]$
- Draw $X_1, \dots, X_n \stackrel{\text{iid}}{\sim} f$,
- Use $\hat{I}_{\text{mc}} = \frac{1}{n} \sum_{i=1}^n \varphi(X_i)$.
- Convergence follows from SLLN, CLT, ...

Recall: The Three Views of the Monte Carlo Method

Direct Approximation Design an experiment such that:

$$\varphi(X) \sim f_{\varphi(X)}$$

constructed such that it has the expectation of interest.

Integral Approximation We're interested in

$$\mathbb{E}_f [\varphi(X)]$$

and know how to approximate such.

Distributional Approximation We're interested in

$$\mathbb{E}_f [\varphi(X)]$$

so obtain an approximation of f with respect to which we can compute expectations.

Contrasting Views of Monte Carlo

- Usual explanation of the Monte Carlo Method, with $X_1, \dots \stackrel{\text{iid}}{\sim} f$ approximating the integral:

$$\frac{1}{n} \sum_{i=1}^n \varphi(X_i) \xrightarrow{a.s.} \mathbb{E}_f [\varphi(X)]$$

- Another perspective, approximate the distribution:
 - let $\hat{f}^n = \frac{1}{n} \sum_{i=1}^n \delta_{X_i}$
 - if $\hat{f}^n \Rightarrow f$
 - then we automatically have that

$$\mathbb{E}_{\hat{f}^n} [\varphi(X)] \rightarrow \mathbb{E}_f [\varphi(X)]$$

for every continuous bounded φ .

Part 2— Section 4

PRNGs

Problem: (how) can computers produce random numbers?

von Neumann's perspective

Any one who considers arithmetical methods of reproducing random digits is, of course, in a state of sin. . . there is no such thing as a random number—there are only methods of producing random numbers, and a strict arithmetic procedure is of course not such a method.

As in so many other areas, von Neumann was completely correct.



Three Resolutions of this Philosophical Paradox

- 1 Use Exogeneous Randomness (TRNGs)
See www.random.org or
http://en.wikipedia.org/wiki/Hardware_random_number_generator.
- 2 Pseudorandom Number Generators (PRNGs; c.f. *Statistical Computing* module)
Sacrifice randomness whilst mimicking its *relevant statistical properties*.
- 3 Quasirandom Number Sequences (QRNSs)
Sacrifice randomness in exchange for *minimising discrepancy*.

All have advantages and disadvantages; we'll focus on PRNGs.

Part 2— Section 5

Sampling From Distributions

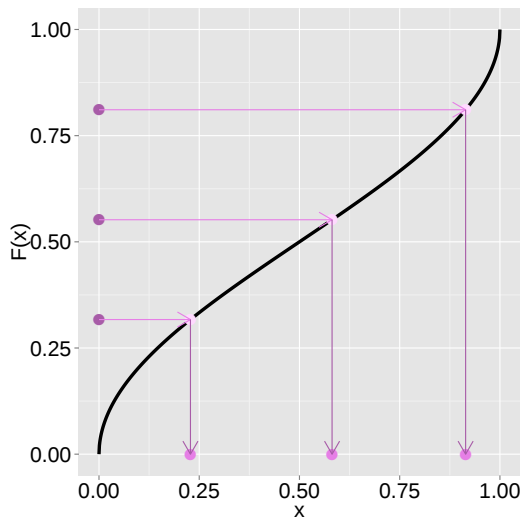
Transformation Methods

- Assume we have a *good* PRNG.
- How can we obtain (pseudo)samples from other distributions?
- General framework:
 - Treat output of PRNG as a stream of iid $U[0, 1]$ RVs.
 - Use laws of probability to transform these to obtain RVs with other distributions.
 - Treat transformed PRNG output as RVs of the target distribution.
- But, how?

Inversion Sampling

The Inversion method

Let $U \sim U[0, 1]$ and
let F be an invertible CDF.
Then $F^{-1}(U)$ has the CDF F .



Inversion Sampling

The Inversion method

Let $U \sim U[0, 1]$ and F be an invertible CDF.

Then $F^{-1}(U)$ has the CDF F .

Inversion Sampling: A simple algorithm for drawing $X \sim F$

- 1 Draw $U \sim U[0, 1]$.
- 2 Set $X = F^{-1}(U)$.

Example: Exponential distribution

The exponential distribution with rate $\lambda > 0$ has the CDF ($x \geq 0$)

$$\begin{aligned}F_{\lambda}(x) &= 1 - \exp(-\lambda x) \\F_{\lambda}^{-1}(u) &= -\log(1 - u)/\lambda.\end{aligned}$$

So we have a simple algorithm for drawing $X \sim \text{Exp}(\lambda)$:

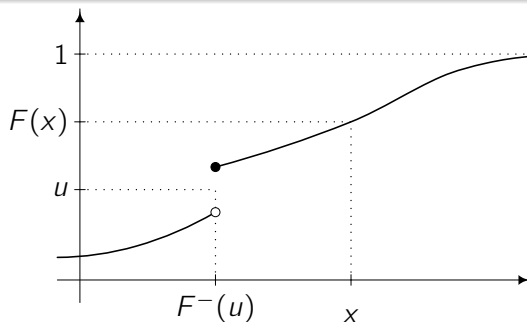
- 1 Draw $U \sim U[0, 1]$.
- 2 Set $X = -\frac{\log(1 - U)}{\lambda}$.

Actually, setting $X = -\frac{\log(U)}{\lambda}$ makes more sense.

The Generalised Inverse of the CDF

Generalised inverse of the CDF

$$F^{-}(u) := \inf\{x : F(x) \geq u\}$$



Replacing F^{-1} with F^{-} yields a generally-applicable inversion sampling algorithm — key is $F^{-}(u) \leq x \Leftrightarrow u \leq F(x)$.

Box–Muller: Fast Normally-Distributed Random Variables

- Consider (X_1, X_2) their polar representation (R, θ) :

$$X_1 = R \cdot \cos(\theta), \quad X_2 = R \cdot \sin(\theta)$$

- The following equivalence holds (with θ, R independent):

$$X_1, X_2 \stackrel{\text{iid}}{\sim} \mathbf{N}(0, 1) \iff \theta \sim \mathbf{U}[0, 2\pi] \text{ and } R^2 \sim \mathbf{Exp}(1/2)$$

- Given $U_1, U_2 \stackrel{\text{iid}}{\sim} \mathbf{U}[0, 1]$ set

$$R = \sqrt{-2 \log(U_1)}, \quad \theta = 2\pi U_2.$$

- By substitution

$$X_1 = \sqrt{-2 \log(U_1)} \cdot \cos(2\pi U_2),$$

$$X_2 = \sqrt{-2 \log(U_1)} \cdot \sin(2\pi U_2).$$

Box-Muller: Algorithm

Box-Muller method

- 1 Draw

$$U_1, U_2 \stackrel{\text{iid}}{\sim} U[0, 1].$$

- 2 Set

$$X_1 = \sqrt{-2 \log(U_1)} \cdot \cos(2\pi U_2),$$

$$X_2 = \sqrt{-2 \log(U_1)} \cdot \sin(2\pi U_2).$$

- 3 Output $X_1, X_2 \stackrel{\text{iid}}{\sim} N(0, 1)$.

The Limitations of Simple Transformations...

- When F^{-} is available and cheap to evaluate, inversion sampling is very efficient. But:
 - We often don't have access to F ;
 - even if we do, F^{-} may be difficult/impossible to obtain.
 - The multivariate case can be even harder.
- Clever custom transformations:
 - are costly to develop,
 - require considerable ingenuity,
 - are completely infeasible in complicated scenarios.
- We need alternatives.

The Fundamental Theorem of simulation

Fundamental Theorem of Simulation

Sampling from a density f is equivalent to sampling uniformly from the area between f and the ordinal axes and discarding the “vertical” component.

- Follows from the identity

$$f(x) = \int_0^{f(x)} 1 \, du = \int_0^\infty \underbrace{1_{0 < u < f(x)}}_{=f(x,u)} du.$$

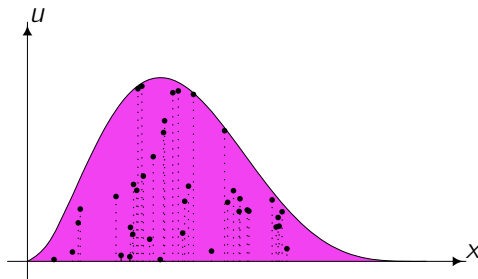
- i.e. $f(x)$ can be interpreted as the marginal density of a uniform distribution on the area under the density $f(x)$:

$$\{(x, u) : 0 \leq u \leq f(x)\}.$$

Rejection

First element of rejection sampling

- We can sample from f by sampling from the area under the density.



- If $(X, U) \sim U(\{(x, u) : 0 \leq u \leq f(x)\})$ then $X \sim f$.

Second Element of Rejection Sampling

- Generally $\mathcal{G} = \{(x, u) : 0 \leq u \leq f(x)\}$ is complicated: we can't sample uniformly from it—at least not directly.
- Idea: Instead:
 - Sample from some $\mathcal{A} \supseteq \mathcal{G}$.
 - Keep only those points which lie within $\mathcal{G}$.
 - *Reject* the rest.

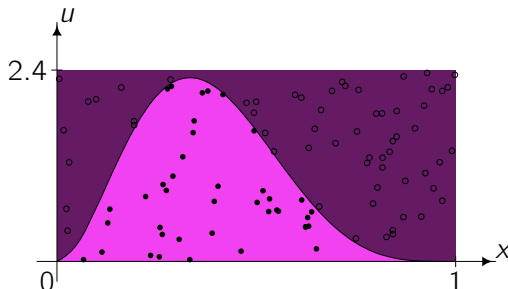
Rejection

Example: Sampling from a Beta(3, 5) distribution (1)

- 1 Draw (X, U) from the dark rectangle, i.e.:

$$X \sim U(0, 1) \quad U \sim U(0, 2.4) \quad X \perp U.$$

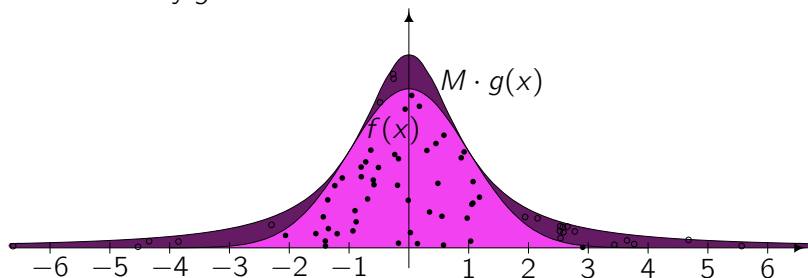
- 2 Accept X as a sample from f if (X, U) lies under the density.



Step 2 is equivalent to: Accept X if $U \leq f(X)$,
 i.e. accept X with probability $\mathbb{P}(U \leq f(X)|X = x) = f(X)/2.4$.

Example: Sampling from a Beta(3, 5) distribution (2)

- Algorithm:
 - 1 Draw $X \sim U(0, 1)$.
 - 2 Accept X as a sample from Beta(3, 5) w.p. $f(X)/2.4$.
- Not every density can be bounded by a box.
- Natural generalisation: replace M times $U[0, 1]$ with M times another density g .



A General Algorithm

Algorithm: Rejection sampling

Given two densities f, g with $f(x) \leq M \cdot g(x)$ for all x , we can generate a sample from f by

1. Draw $X \sim g$.
2. Accept X as a sample from f with probability

$$\frac{f(X)}{M \cdot g(X)},$$

otherwise go back to step 1.

For $f(x) \leq M \cdot g(x)$ to hold for all x , f *cannot* have heavier tails than g .

A Useful Trick

Avoiding Unknown Constants

If we know only $\tilde{f}(x)$ and $\tilde{g}(x)$, where $f(x) = C \cdot \tilde{f}(x)$, and $g(x) = D \cdot \tilde{g}(x)$, we can carry out rejection sampling using acceptance probability

$$\frac{\tilde{f}(X)}{M \cdot \tilde{g}(X)}$$

provided $\tilde{f}(x) \leq M \cdot \tilde{g}(x)$ for all x .

Can be useful in Bayesian statistics:

$$\begin{aligned} f^{\text{post}}(\theta) &= \frac{f^{\text{prior}}(\theta)L(\theta; \mathbf{y}_1, \dots, \mathbf{y}_n)}{\int_{\Theta} f^{\text{prior}}(\vartheta)L(\vartheta; \mathbf{y}_1, \dots, \mathbf{y}_n) d\vartheta} \\ &= C \cdot f^{\text{prior}}(\theta)L(\theta; \mathbf{y}_1, \dots, \mathbf{y}_n). \end{aligned}$$

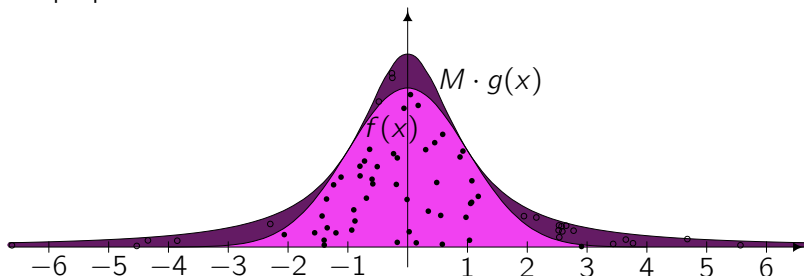
Rejection

Example: Sampling from $N(0, 1)$

- Recall the $N(0, 1)$ and Cauchy densities:

$$f(x) = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{x^2}{2}\right), \quad g(x) = \frac{1}{\pi(1+x^2)}.$$

- For $M = \sqrt{2\pi} \cdot \exp(-1/2)$ we have that $f(x) \leq Mg(x)$. So we can use rejection sampling targeting f using g as proposal.



Non-example: Sampling from a Cauchy Distribution

- We cannot sample the other way round: from a Cauchy distribution using a Normal as proposal distribution.
- The Cauchy distribution has heavier tails than the Normal distribution: there is no $M \in \mathbb{R}$ such that

$$\frac{1}{\pi(1+x^2)} \leq M \cdot \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{x^2}{2}\right).$$

Vevox.app 124-467-150

How would you sample from a Cauchy distribution?

Non-example: Sampling from a Cauchy Distribution

- We cannot sample the other way round: from a Cauchy distribution using a Normal as proposal distribution.
- The Cauchy distribution has heavier tails than the Normal distribution: there is no $M \in \mathbb{R}$ such that

$$\frac{1}{\pi(1+x^2)} \leq M \cdot \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{x^2}{2}\right).$$

Vevox.app 124-467-150

How would you sample from a Cauchy distribution?

An Alternative to Rejection

- Rejection sampling discards many samples.
- This seems wasteful.
- Couldn't we, instead, *weight* samples based on the acceptance probability?

The fundamental identities behind importance sampling

Assume that $g(x) > 0$ for (almost) all x with $f(x) > 0$:

$$\mathbb{P}(X \in \mathcal{X}) = \int_{\mathcal{X}} f(x) \, dx = \int_{\mathcal{X}} g(x) \underbrace{\frac{f(x)}{g(x)}}_{=: w(x)} \, dx = \int_{\mathcal{X}} g(x) w(x) \, dx.$$

Assume that $g(x) > 0$ for (almost) all x with $f(x) \cdot \varphi(x) \neq 0$

$$\begin{aligned} \mathbb{E}_f(\varphi(X)) &= \int f(x) \varphi(x) \, dx = \int g(x) \underbrace{\frac{f(x)}{g(x)}}_{=: w(x)} \varphi(x) \, dx \\ &= \int g(x) w(x) \varphi(x) \, dx = \mathbb{E}_g(w(X) \cdot \varphi(X)). \end{aligned}$$

The fundamental identities behind importance sampling

- Consider $X_1, \dots, X_n \sim g$ and $\mathbb{E}_g |w(X) \cdot \varphi(X)| < \infty$. Then

$$\begin{aligned} \frac{1}{n} \sum_{i=1}^n w(X_i) \varphi(X_i) &\xrightarrow[n \rightarrow \infty]{a.s.} \mathbb{E}_g(w(X) \cdot \varphi(X)) \\ \implies \frac{1}{n} \sum_{i=1}^n w(X_i) \varphi(X_i) &\xrightarrow[n \rightarrow \infty]{a.s.} \mathbb{E}_f(\varphi(X)). \end{aligned}$$

- Thus we can estimate $\mu := \mathbb{E}_f(\varphi(X))$ by

- Sample $X_1, \dots, X_n \sim g$,
- $\tilde{\mu} := \frac{1}{n} \sum_{i=1}^n w(X_i) \varphi(X_i)$.

The importance sampling algorithm

Algorithm: Importance Sampling

Choose g such that $\text{supp}(g) \supseteq \text{supp}(f \cdot \varphi)$.

① For $i = 1, \dots, n$:

① Generate $X_i \sim g$.

② Set $w(X_i) = \frac{f(X_i)}{g(X_i)}$.

② Return

$$\tilde{\mu} = \frac{\sum_{i=1}^n w(X_i) \varphi(X_i)}{n}$$

as an estimate of $\mathbb{E}_f(\varphi(X))$.

- Importance sampling does not yield realisations from f ,
but a *weighted sample* (X_i, W_i) ,
which can be used for estimating expectations $\mathbb{E}_f(\varphi(X))$,
or approximating f itself.

Basic properties of the importance sampling estimate

- We have already seen that $\tilde{\mu}$ is consistent if $\text{supp}(g) \supseteq \text{supp}(f \cdot \varphi)$ and $\mathbb{E}_g|w(X) \cdot \varphi(X)| < \infty$, as

$$\tilde{\mu} := \frac{1}{n} \sum_{i=1}^n w(X_i) \varphi(X_i) \xrightarrow[n \rightarrow \infty]{a.s.} \mathbb{E}_f(\varphi(X))$$

- The expected value of the weights is $\mathbb{E}_g(w(X)) = 1$.
- $\tilde{\mu}$ is unbiased (see theorem below)

Theorem 2.2: Bias and Variance of Importance Sampling

$$\begin{aligned} \mathbb{E}_g(\tilde{\mu}) &= \mu, \\ \text{Var}_g(\tilde{\mu}) &= \frac{\text{Var}_g(w(X) \cdot \varphi(X))}{n}. \end{aligned}$$

Optimal proposals

Theorem (Optimal proposal)

The proposal distribution g that minimises the variance of $\tilde{\mu}$ is

$$g^*(x) = \frac{|\varphi(x)|f(x)}{\int |\varphi(t)|f(t) dt}.$$

- Theorem of little practical use: the optimal proposal involves $\int |\varphi(t)|f(t) dt$, which is the integral we want to estimate!
- Practical relevance:
Choose g such that it is close to $|\varphi(x)| \cdot f(x)$.

Super-efficiency of importance sampling

- For the optimal g^* we have that

$$\text{Var}_f \left(\frac{\varphi(X_1) + \cdots + \varphi(X_n)}{n} \right) > \text{Var}_{g^*}(\tilde{\mu}),$$

if φ is not almost surely constant.

Superefficiency of importance sampling

The variance of the importance sampling estimate can be *less* than the variance obtained by sampling directly from the target f .

- Intuition: Importance sampling allows us to choose a g that focuses on areas which contribute most to $\int \varphi(x)f(x) dx$.
- Even sub-optimal proposals can be super-efficient.

Importance Sampling Example 1: Setup

Compute $\mathbb{E}_f|X|$ for $X \sim t_3$ by ...

- (a) sampling directly from t_3 .
- (b) using a t_1 distribution as proposal distribution.
- (c) using a $N(0, 1)$ distribution as proposal distribution.

Vevox.app 124-467-150

Which of these methods is best?

Reminder:

$$g_{t_3}(x) = \frac{2}{\pi\sqrt{3}} \cdot \frac{1}{\left(1 + \frac{x^2}{3}\right)^2}, \quad g_{t_1}(x) = \frac{1}{\pi} \cdot \frac{1}{1 + x^2}.$$

Importance Sampling Example 1: Setup

Compute $\mathbb{E}_f|X|$ for $X \sim t_3$ by ...

- (a) sampling directly from t_3 .
- (b) using a t_1 distribution as proposal distribution.
- (c) using a $N(0, 1)$ distribution as proposal distribution.

Vevox.app 124-467-150

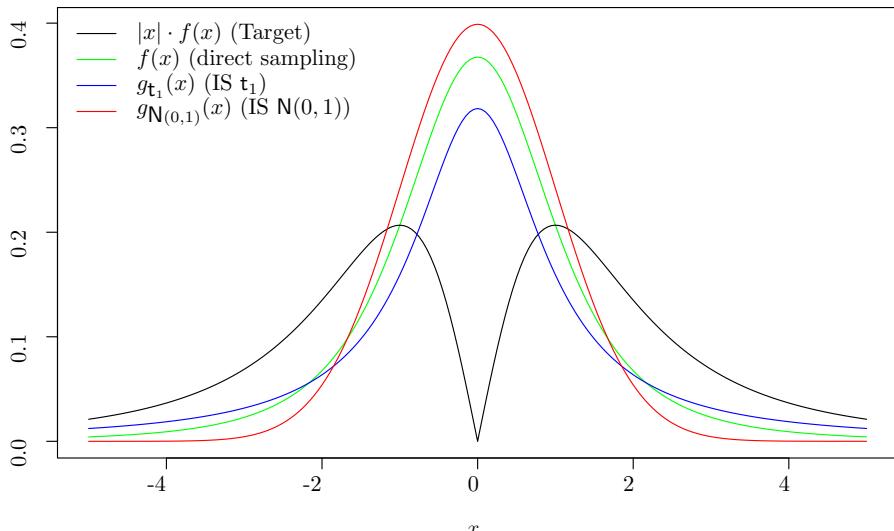
Which of these methods is best?

Reminder:

$$g_{t_3}(x) = \frac{2}{\pi\sqrt{3}} \cdot \frac{1}{\left(1 + \frac{x^2}{3}\right)^2}, \quad g_{t_1}(x) = \frac{1}{\pi} \cdot \frac{1}{1 + x^2}.$$

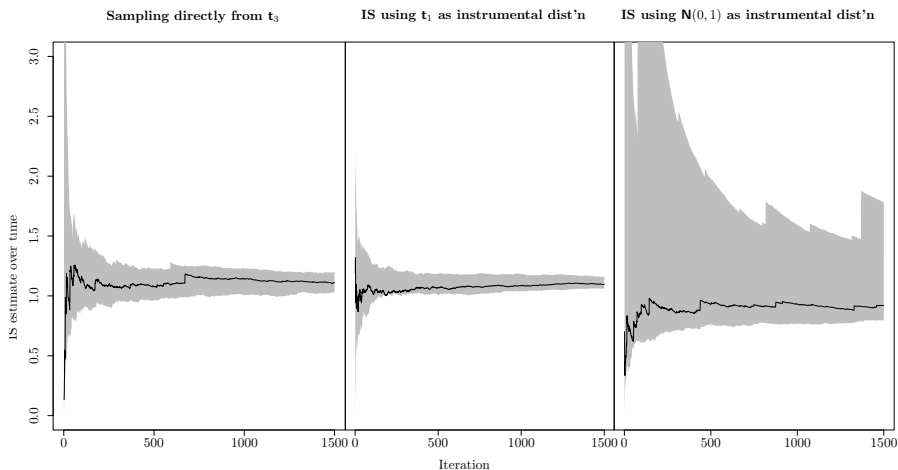
Importance Sampling

IS Example: Densities



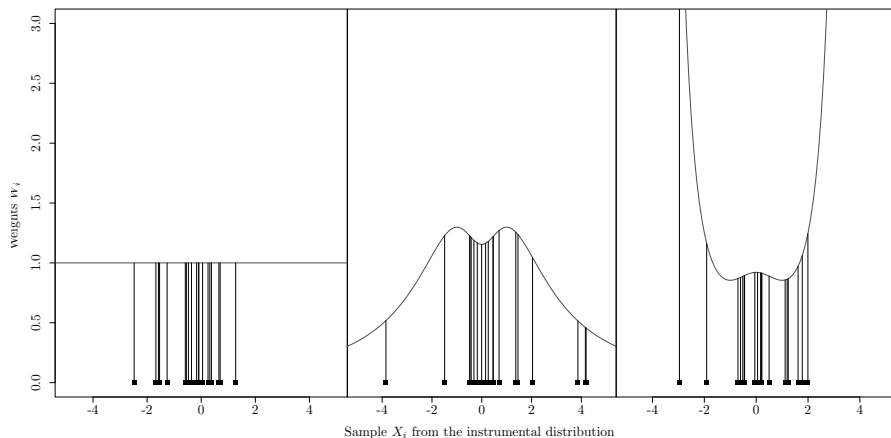
Importance Sampling

IS Example: Estimates obtained



Importance Sampling

IS Example: Weights

Sampling directly from t_3 IS using t_1 as instrumental dist'nIS using $N(0, 1)$ as instrumental dist'n

Importance Sampling

Another Example: Rare Events (1)

Consider

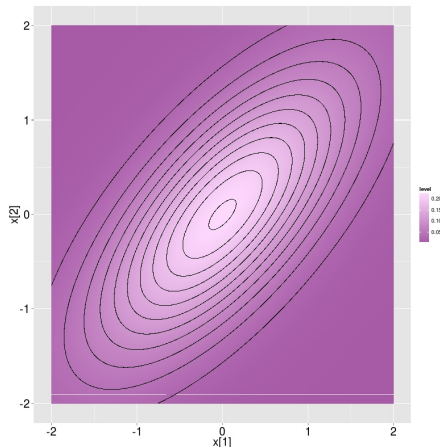
$$f(x, y) = \mathcal{N} \left(\begin{pmatrix} x \\ y \end{pmatrix}; \mu, \Sigma \right),$$

where

$$\mu = \begin{pmatrix} 0 \\ 0 \end{pmatrix}, \quad \Sigma = \begin{bmatrix} 1 & 0.7 \\ 0.7 & 1 \end{bmatrix}.$$

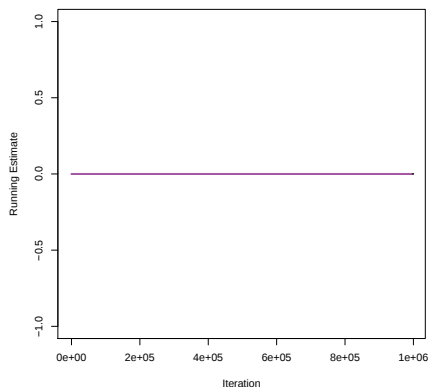
Consider

$$\varphi(x, y) = \mathbb{I}_{[4, \infty)}(x) \mathbb{I}_{[4, \infty)}(y).$$



Another Example: Rare Events (2)

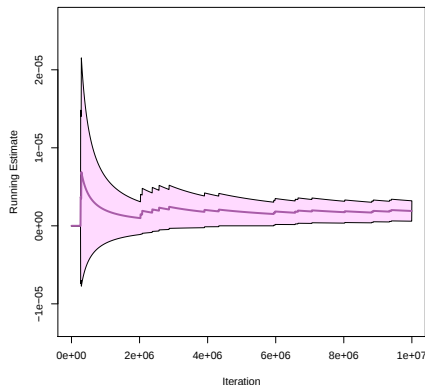
Using simple Monte Carlo with 1,000,000 samples from f :



shaded region shows *estimated* 99.7% confidence interval.

Another Example: Rare Events (3)

Using simple Monte Carlo with 10,000,000 samples from f :

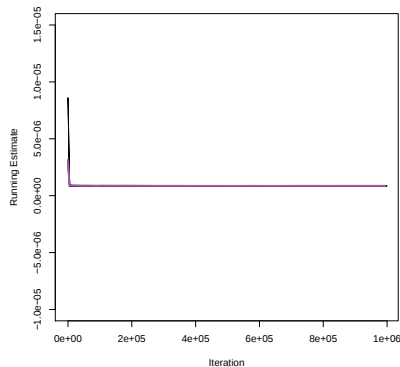


shaded region shows *estimated* 99.7% confidence interval.

Another Example: Rare Events (4)

Using importance sampling with 1,000,000 samples from

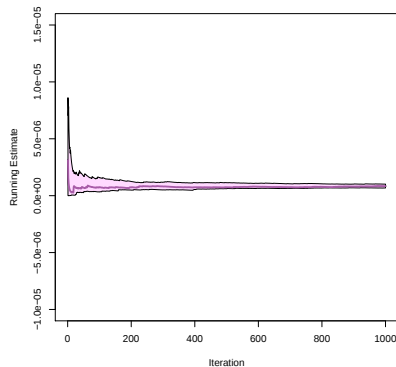
$$g(x, y) = \exp(-(x - 4) - (y - 4))\mathbb{I}_{x \geq 4}\mathbb{I}_{y \geq 4}:$$



shaded region shows range of 100 replications.

Another Example: Rare Events (5)

Using importance sampling with 1,000 samples from $g(x, y) = \exp(-(x - 4) - (y - 4))\mathbb{I}_{x \geq 4}\mathbb{I}_{y \geq 4}$:

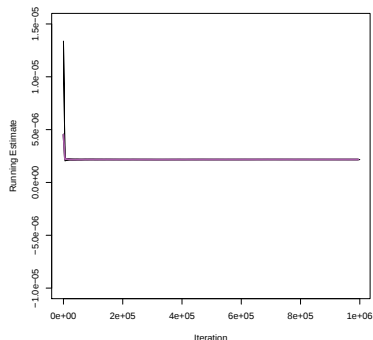


shaded region shows range of 100 replications.

Another Example: Rare Events (6)

Using importance sampling with 1,000,000 samples from

$$g(x, y) = N \left(\begin{pmatrix} x \\ y \end{pmatrix}; \begin{pmatrix} 4 \\ 4 \end{pmatrix}, \Sigma \mid x \geq 4, y \geq 4 \right) :$$

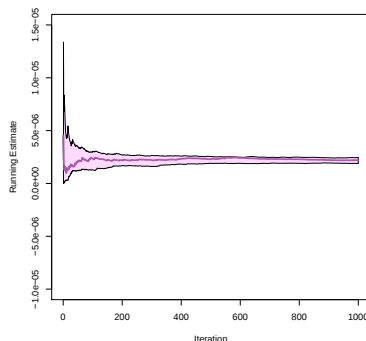


shaded region shows range of 100 replications

Another Example: Rare Events (7)

Using importance sampling with 1,000 samples from

$$g(x, y) = N \left(\begin{pmatrix} x \\ y \end{pmatrix}; \begin{pmatrix} 4 \\ 4 \end{pmatrix}, \Sigma \mid x \geq 4, y \geq 4 \right) :$$



shaded region shows range of 100 replications

Importance Sampling

We only need f up to a multiplicative constant.

- Assume $f(x) = C\tilde{f}(x)$. Then

$$\tilde{\mu} = \frac{1}{n} \sum_{i=1}^n w(X_i) \varphi(X_i) = \frac{1}{n} \sum_{i=1}^n \frac{C\tilde{f}(X_i)}{g(X_i)} \varphi(X_i)$$

C does not cancel out. Knowing $\tilde{f}(\cdot)$ is not enough.

- Idea: Estimate C using the sample, via $\sum_{i=1}^n w(X_i)$, i.e. consider the *self-normalised estimator*

$$\hat{\mu} = \frac{\frac{1}{n} \sum_{i=1}^n w(X_i) \varphi(X_i)}{\frac{1}{n} \sum_{j=1}^n w(X_j) 1}$$

- Now we have that $\hat{\mu}$ does not depend on C :

$$\hat{\mu} = \frac{\sum_{i=1}^n w(X_i) \varphi(X_i)}{\sum_{i=1}^n w(X_i)} = \frac{\sum_{i=1}^n \frac{\tilde{f}(X_i)}{g(X_i)} \varphi(X_i)}{\sum_{i=1}^n \frac{\tilde{f}(X_i)}{g(X_i)}}$$

The importance sampling algorithm (2)

Algorithm: Importance Sampling using self-normalised weights

Choose g such that $\text{supp}(g) \supseteq \text{supp}(f)$.

- 1 For $i = 1, \dots, n$:
 - 1 Generate $X_i \sim g$.
 - 2 Set $w(X_i) = \frac{f(X_i)}{g(X_i)}$.
- 2 Return

$$\hat{\mu} = \frac{\sum_{i=1}^n w(X_i) \varphi(X_i)}{\sum_{i=1}^n w(X_i)}$$

as an estimate of $\mathbb{E}_f(\varphi(X))$.

Basic properties of the self-normalised estimate

- $\hat{\mu}$ is consistent as

$$\hat{\mu} = \underbrace{\frac{\sum_{i=1}^n w(X_i) \varphi(X_i)}{n}}_{=\tilde{\mu} \rightarrow \mathbb{E}_f(\varphi(X))} \underbrace{\frac{n}{\sum_{i=1}^n w(X_i)}}_{\rightarrow 1} \xrightarrow[n \rightarrow \infty]{a.s.} \mathbb{E}_f(\varphi(X)),$$

(provided $\text{supp}(g) \supseteq \text{supp}(f)$ and $\mathbb{E}_g |w(X) \cdot \varphi(X)| < \infty$).

Theorem: Bias and Variance (ctd.)

$$\begin{aligned} \mathbb{E}_g(\hat{\mu}) &= \mu + \frac{\mu \text{Var}_g(w(X)) - \text{Cov}_g[w(X), w(X) \cdot \varphi(X)]}{n} + O(n^{-2}) \\ \text{Var}_g(\hat{\mu}) &= \frac{\text{Var}_g(w(X) \cdot \varphi(X)) - 2\mu \text{Cov}_g[w(X), w(X) \cdot \varphi(X)]}{n} \\ &\quad + \frac{\mu^2 \text{Var}_g(w(X))}{n} + O(n^{-2}) \end{aligned}$$

Finite variance estimators

- Importance sampling estimates are consistent for many choices of g .
- More important in practice: we want *finite variance estimators*:

$$\text{Var}(\tilde{\mu}) = \text{Var}\left(\frac{\sum_{i=1}^n w(X_i)\varphi(X_i)}{n}\right) < \infty$$

- Sufficient (albeit restrictive) conditions for finite variance of $\tilde{\mu}$:
 - $f(x) \leq M \cdot g(x)$ and $\text{Var}_f(\varphi(X)) < \infty$, or
 - E is compact, f is bounded above on E , and g is bounded below on E .
- Note: If f has heavier tails than g , then the weights may have *infinite* variance!

Summary of Part 2

- Transformation: Inversion sampling
- Transformation: Case-specific methods such as Box–Muller
- Rejection Sampling
- Importance Sampling

Part 3

Markov chain Monte Carlo

Part 3— Section 6

Motivation and Basics

Motivating MCMC

Why do we need other, more complicated methods?

- Transformation's great when it works.
- Rejection sampling's good when M is small.
- Importance sampling works well with good proposals.
- What do we do when we can't meet any of these requirements?

One Approach

Markov Chain Monte Carlo methods (MCMC)

- Key idea: Create a *dependent* sample, i.e. $X^{(t)}$ depends on the previous value $X^{(t-1)}$.
 Allows for “local” updates.
- Yields an “approximate sample” from the target distribution.
- More mathematically speaking: yields a Markov chain with the target distribution f as stationary distribution.
- Under conditions, the realised chain provides approximations of $\mathbb{E}_f [\varphi(X)]$ and of f itself.

Markov Chains

Markov Chain (N.B. Terminology varies)

A *discrete time* Markov process taking values in a *general space*:

$$X^{(0)} \sim \mu_0 \quad \text{Initial Dist.}$$

$$X^{(t)} | \left(X^{(0)} = x^{(0)}, \dots, X^{(t-1)} = x^{(t-1)} \right) \sim K(x^{(t-1)}, \cdot) \quad \text{Kernel}$$

Stationary Distribution

f is a *stationary* or *invariant* distribution for a Markov Chain on E with kernel K if for all measurable sets A ,

$$\int_A \int_E f(x) K(x, y) dx dy = \int_A f(y) dy.$$

Will later see alternative perspective based on *categorical probability*

Motivating MCMC

Heuristically Motivating MCMC

- If $X^{(0)}, \dots$ is an f -invariant Markov chain and $X^{(t)} \sim f$ for some t then $X^{(t+s)} \sim f \quad \forall s \in \mathbb{N}$.
- So if $X^{(t)}$ is “approximately independent” of $X^{(t+s)}$ for large enough s then
 - $X^{(t)}, X^{(t+s)}, \dots, X^{(t+ks)}, \dots$ is approximately $\overset{\text{iid}}{\sim} f$,
 - $X^{(t+1)}, X^{(t+s+1)}, \dots, X^{(t+ks+1)}, \dots$ is approximately $\overset{\text{iid}}{\sim} f$,
 - $\vdots$
 - $X^{(t+s-1)}, X^{(t+2s-1)}, \dots, X^{(t+ks-1)}, \dots$ is approximately $\overset{\text{iid}}{\sim} f$.
- We might conjecture that for such a chain, for some large s :

$$\frac{1}{n} \sum_{k=1}^n \varphi(X^{(t+ks)}) \rightarrow \mathbb{E}_f [\varphi(X)] \quad \text{and} \quad \frac{1}{n} \sum_{k=1}^n \varphi(X^{(k)}) \rightarrow \mathbb{E}_f [\varphi(X)] .$$

Some Questions to Answer

- Can we formalise this heuristic argument?
~→ ergodic theory
- How can we construct f -invariant Markov kernels?
~→ various types of sampler
- How do we initialise the chain?
~→ transient phases and burn-in
- How do we know if it's working?
~→ ergodic theory and convergence diagnostics

A Motivating Convergence Result

Theorem (A Simple Ergodic Theorem)

If $(X_i)_{i \in \mathbb{N}}$ is an f -irreducible, f -invariant, recurrent $\mathbb{R}^d$ -valued Markov chain, then the following strong law of large numbers holds for any integrable function $\varphi : \mathbb{R}^d \rightarrow \mathbb{R}$:

$$\lim_{t \rightarrow \infty} \frac{1}{t} \sum_{i=1}^t \varphi(X_i) \stackrel{a.s.}{=} \int \varphi(x) f(x) dx.$$

for almost every starting value x .

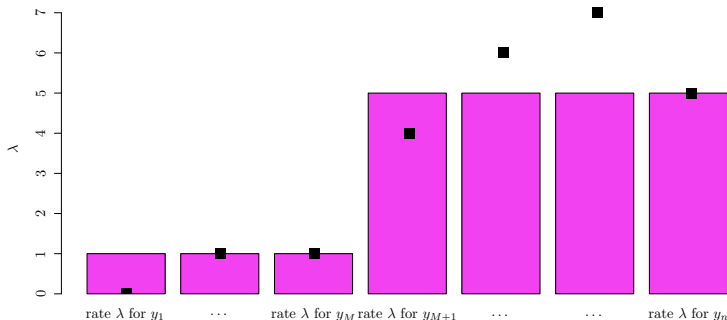
Note: this gives no *rate* of convergence.

Part 3— Section 7

The Gibbs Sampler

A Motivating Example

Example: Poisson change point model I



$$Y_i \sim \text{Poi}(\lambda_1) \quad \text{for} \quad i = 1, \dots, M,$$

$$Y_i \sim \text{Poi}(\lambda_2) \quad \text{for} \quad i = M + 1, \dots, n.$$

A Motivating Example

Example: Poisson change point model II

Objective: (Bayesian) inference about the parameters λ_1 , λ_2 , and M given observed data $y_1, \dots, y_n$.

- Prior distributions: $\lambda_j \sim \text{Gamma}(\alpha_j, \beta_j)$ ($j = 1, 2$), i.e.

$$f(\lambda_j) = \frac{1}{\Gamma(\alpha_j)} \lambda_j^{\alpha_j-1} \beta_j^{\alpha_j} \exp(-\beta_j \lambda_j).$$

(discrete uniform prior on M , i.e. $p(M) \propto 1$).

- Likelihood: $L(\lambda_1, \lambda_2, M; y_1, \dots, y_n)$

$$= \left(\prod_{i=1}^M \frac{\exp(-\lambda_1) \lambda_1^{y_i}}{y_i!} \right) \cdot \left(\prod_{i=M+1}^n \frac{\exp(-\lambda_2) \lambda_2^{y_i}}{y_i!} \right).$$

A Motivating Example

Example: Poisson change point model III

- Joint distribution $f(y_1, \dots, y_n, \lambda_1, \lambda_2, M)$

$$\begin{aligned} &= L(\lambda_1, \lambda_2, M; y_1, \dots, y_n) \cdot f(\lambda_1) \cdot f(\lambda_2) \cdot p(M) \\ &\propto \left(\prod_{i=1}^M \frac{\exp(-\lambda_1) \lambda_1^{y_i}}{y_i!} \right) \cdot \left(\prod_{i=M+1}^n \frac{\exp(-\lambda_2) \lambda_2^{y_i}}{y_i!} \right) \\ &\quad \cdot \frac{1}{\Gamma(\alpha_1)} \lambda_1^{\alpha_1-1} \beta_1^{\alpha_1} \exp(-\beta_1 \lambda_1) \cdot \frac{1}{\Gamma(\alpha_2)} \lambda_2^{\alpha_2-1} \beta_2^{\alpha_2} \exp(-\beta_2 \lambda_2) \end{aligned}$$

- Joint posterior distribution $f(\lambda_1, \lambda_2, M | y_1, \dots, y_n)$

$$\begin{aligned} &\propto \lambda_1^{\alpha_1-1+\sum_{i=1}^M y_i} \exp(-(\beta_1 + M) \lambda_1) \\ &\quad \cdot \lambda_2^{\alpha_2-1+\sum_{i=M+1}^n y_i} \exp(-(\beta_2 + n - M) \lambda_2) \end{aligned}$$

A Motivating Example

Example: Poisson change point model IV

- Conditional on M (i.e. if M was known) we have

$$f(\lambda_1 | y_1, \dots, y_n, M) \propto \lambda_1^{\alpha_1 - 1 + \sum_{i=1}^M y_i} \exp(-(\beta_1 + M)\lambda_1),$$

i.e.

$$\lambda_1 | Y_1, \dots, Y_n, M \sim \text{Gamma} \left(\alpha_1 + \sum_{i=1}^M y_i, \beta_1 + M \right),$$

$$\lambda_2 | Y_1, \dots, Y_n, M \sim \text{Gamma} \left(\alpha_2 + \sum_{i=M+1}^n y_i, \beta_2 + n - M \right).$$

- $p(M | \dots) \propto \lambda_1^{\sum_{i=1}^M y_i} \cdot \lambda_2^{\sum_{i=M+1}^n y_i} \cdot \exp((\lambda_2 - \lambda_1) \cdot M).$

A Motivating Example

Example: Poisson change point model V

This suggests an iterative algorithm:

- ① Draw λ_1 from $\lambda_1|Y_1, \dots, Y_n, M$, i.e. draw

$$\lambda_1 \sim \text{Gamma} \left(\alpha_1 + \sum_{i=1}^M y_i, \beta_1 + M \right).$$

- ② Draw λ_2 from $\lambda_2|Y_1, \dots, Y_n, M$, i.e. draw

$$\lambda_2 \sim \text{Gamma} \left(\alpha_2 + \sum_{i=M+1}^n y_i, \beta_2 + n - M \right).$$

- ③ Draw M from $M|Y_1, \dots, Y_n, \lambda_1, \lambda_2$, i.e. draw

$$p(M) \propto \lambda_1^{\sum_{i=1}^M y_i} \cdot \lambda_2^{\sum_{i=M+1}^n y_i} \cdot \exp((\lambda_2 - \lambda_1) \cdot M).$$

The systematic scan Gibbs sampler

Algorithm: (Systematic scan) Gibbs sampler

Starting with $(X_1^{(0)}, \dots, X_p^{(0)})$ iterate for $t = 1, 2, \dots$

1. Draw $X_1^{(t)} \sim f_{X_1|X_{-1}}(\cdot | X_2^{(t-1)}, \dots, X_p^{(t-1)})$.

...

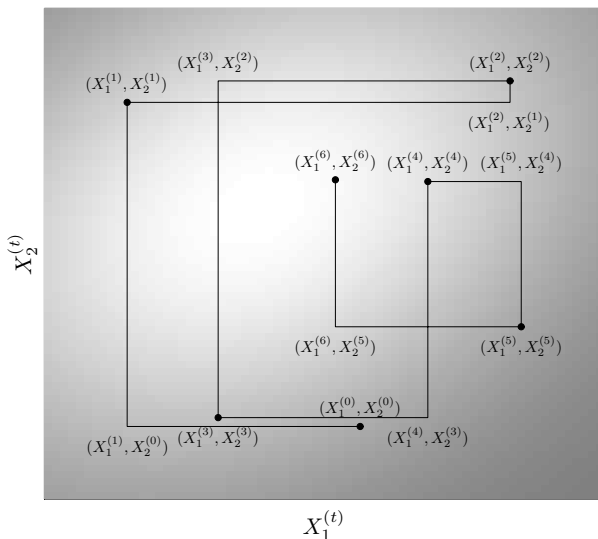
j. Draw $X_j^{(t)} \sim f_{X_j|X_{-j}}(\cdot | X_1^{(t)}, \dots, X_{j-1}^{(t)}, X_{j+1}^{(t-1)}, \dots, X_p^{(t-1)})$.

...

p. Draw $X_p^{(t)} \sim f_{X_p|X_{-p}}(\cdot | X_1^{(t)}, \dots, X_{p-1}^{(t)})$.

The Algorithm

Illustration of the systematic scan Gibbs sampler



The random scan Gibbs sampler

Algorithm: (Random scan) Gibbs sampler

Starting with $(X_1^{(0)}, \dots, X_p^{(0)})$ iterate for $t = 1, 2, \dots$

- 1 Draw an index j from a distribution on $\{1, \dots, p\}$ (e.g. uniform).
- 2 Draw $X_j^{(t)} \sim f_{X_j|X_{-j}}(\cdot | X_1^{(t-1)}, \dots, X_{j-1}^{(t-1)}, X_{j+1}^{(t-1)}, \dots, X_p^{(t-1)})$, and set $X_\iota^{(t)} := X_\iota^{(t-1)}$ for all $\iota \neq j$.

Later on, we will show (using categorical probability) that indeed this Gibbs sampler (and the previous systematic scan version) have the joint distribution $f(x_1, \dots, x_p)$ as their invariant distribution.

The random scan Gibbs sampler

Algorithm: (Random scan) Gibbs sampler

Starting with $(X_1^{(0)}, \dots, X_p^{(0)})$ iterate for $t = 1, 2, \dots$

- 1 Draw an index j from a distribution on $\{1, \dots, p\}$ (e.g. uniform).
- 2 Draw $X_j^{(t)} \sim f_{X_j|X_{-j}}(\cdot | X_1^{(t-1)}, \dots, X_{j-1}^{(t-1)}, X_{j+1}^{(t-1)}, \dots, X_p^{(t-1)})$, and set $X_\iota^{(t)} := X_\iota^{(t-1)}$ for all $\iota \neq j$.

Later on, we will show (using categorical probability) that indeed this Gibbs sampler (and the previous systematic scan version) have the joint distribution $f(x_1, \dots, x_p)$ as their invariant distribution.

Examples

Recall our Poisson Changepoint Model

- Joint posterior distribution $f(\lambda_1, \lambda_2, M | y_1, \dots, y_n)$

$$\begin{aligned} \propto & \lambda_1^{\alpha_1 - 1 + \sum_{i=1}^M y_i} \exp(-(\beta_1 + M)\lambda_1) \\ & \cdot \lambda_2^{\alpha_2 - 1 + \sum_{i=M+1}^n y_i} \exp(-(\beta_2 + n - M)\lambda_2) \end{aligned}$$

- Full Posterior Distributions

$$\lambda_1 | Y_1, \dots, Y_n, M \sim \text{Gamma} \left(\alpha_1 + \sum_{i=1}^M y_i, \beta_1 + M \right),$$

$$\lambda_2 | Y_1, \dots, Y_n, M \sim \text{Gamma} \left(\alpha_2 + \sum_{i=M+1}^n y_i, \beta_2 + n - M \right).$$

- and $p(M | \dots) \propto \lambda_1^{\sum_{i=1}^M y_i} \cdot \lambda_2^{\sum_{i=M+1}^n y_i} \cdot \exp((\lambda_2 - \lambda_1) \cdot M)$.

Examples

An R Implementation

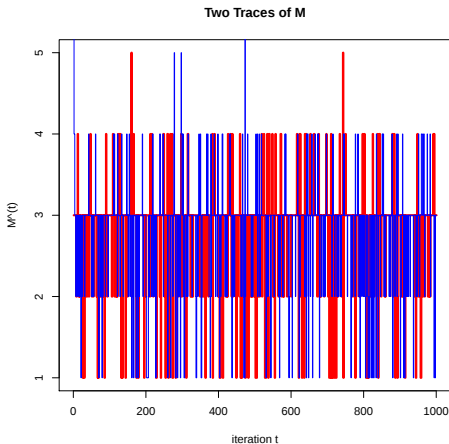
```

cdist.M <- function(lambda1,lambda2) {
  dist.M.log <- cumsum(y[1:n-1]) * log(lambda1) +
    (sum(y)-cumsum(y[1:n-1]))*log(lambda2) +
    (lambda2-lambda1) * (1:(n-1))
  dist.M <- exp(dist.M.log - mean(dist.M.log))
  dist.M <- dist.M / sum(dist.M)
}

pmix.gibbs <- function(M,lambda1,lambda2,t) {
  r <- array(NA,c(t+1,3))
  r[1,] <- c(M,lambda1,lambda2)
  for (i in 1:t) {
    #lambda1
    r[i+1,2] <- rgamma(1,a1+sum(y[1:r[i,1]]), b1+r[i,1])
    #lambda2
    r[i+1,3] <- rgamma(1,a2+sum(y[(r[i,1]+1):n]), b2+n-r[i,1])
    #M
    r[i+1,1] <- sample.int(n-1,1,prob=cdist.M(r[i+1,2],r[i+1,3]))
  }
  r
}

```

Examples

Traces and Estimates: M 

Consider two
differently-initialised chains.

Chain 1:
 $(M, \lambda_1, \lambda_2)^{(0)} = (3, 1, 2)$

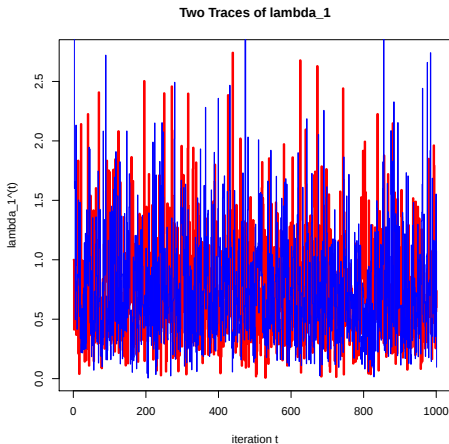
Chain 2:
 $(M, \lambda_1, \lambda_2)^{(0)} = (6, 4, \frac{1}{2})$

Estimated Posterior *Modes*:

Chain 1: 3

Chain 2: 3

Examples

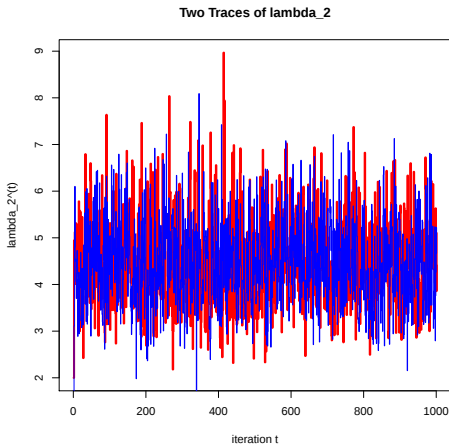
Traces and Estimates: λ_1 

Estimated Posterior *Means*:

Chain 1: 0.76

Chain 2: 0.78

Examples

Traces and Estimates: λ_2 Estimated Posterior *Means*:

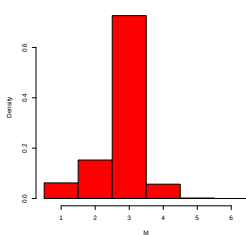
Chain 1: 4.51

Chain 2: 4.47

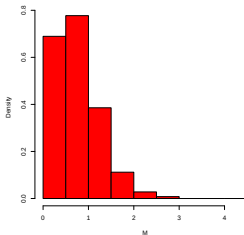
Examples

Histograms: Approximations of the Posterior

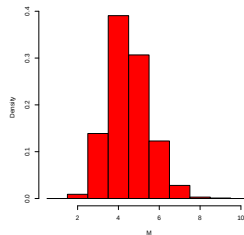
Histogram of M from chain 1



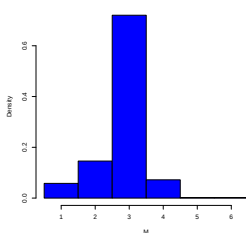
Histogram of lambda_1 from chain 1



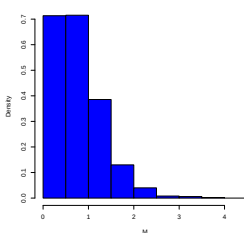
Histogram of lambda_2 from chain 1



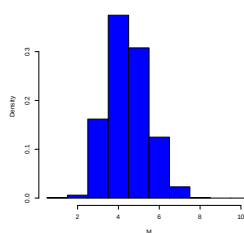
Histogram of M from chain 2



Histogram of lambda_1 from chain 2



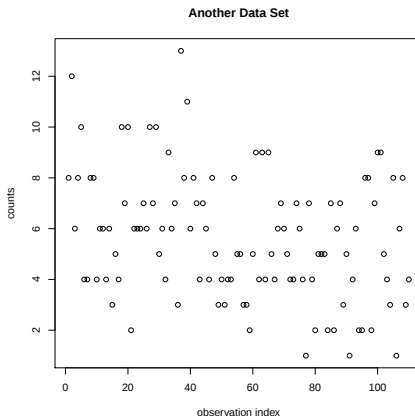
Histogram of lambda_2 from chain 2



Examples

Poisson Change-Point Model: More Challenging Data I

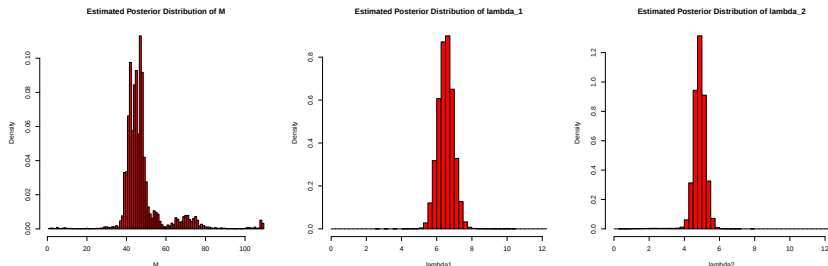
Consider the more realistic data:



Examples

Poisson Change-Point Model: More Challenging Data II

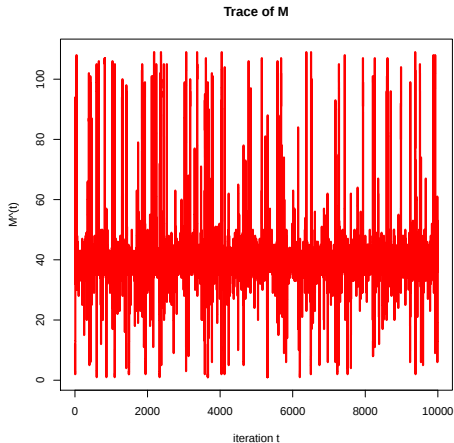
From a chain of length 100,000 we obtain the following histograms:



Data was generated with: `y <- c(rpois(40,7),rpois(70,5))`

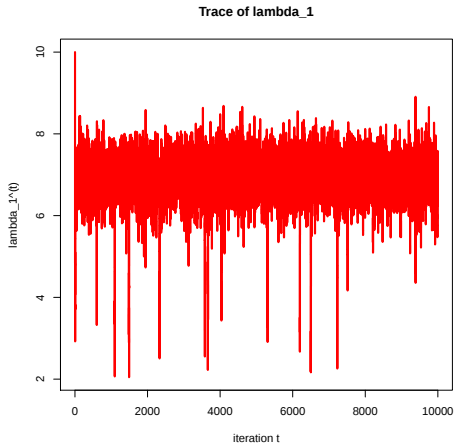
Examples

Poisson Change-Point Model: More Challenging Data III



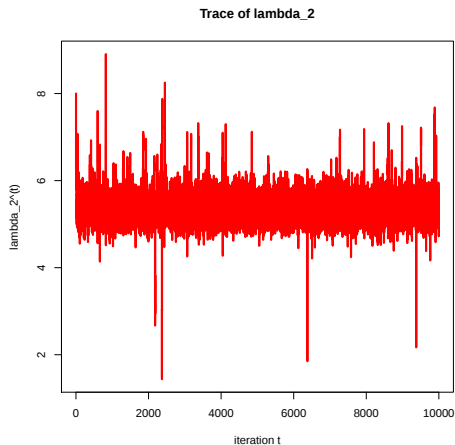
Examples

Poisson Change-Point Model: More Challenging Data IV



Examples

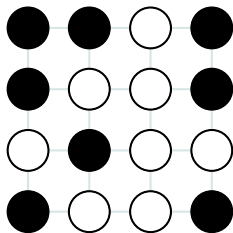
Poisson Change-Point Model: More Challenging Data V



Examples

Example: The Ising Model

The Ising model on $(\mathcal{V}, \mathcal{E})$ each $v_i \in \mathcal{V}$ has an associated $x_i \in \{-1, +1\}$:



$$\pi(x_1, \dots, x_m)$$

$$= \frac{1}{Z} \exp \left(J \sum_{(i,j) \in \mathcal{E}} x_i \cdot x_j \right)$$

$$= \frac{1}{Z} \exp(-J|\mathcal{E}|) \exp \left(2J \sum_{(i,j) \in \mathcal{E}} \mathbb{I}(x_i = x_j) \right)$$

$$= \frac{1}{Z'} \exp \left(2J \sum_{(i,j) \in \mathcal{E}} \mathbb{I}(x_i = x_j) \right).$$

$$\pi(x_j | x_{-j}) = \exp \left(J \sum_{i:(i,j) \in \mathcal{E}} x_i x_j \right) / \left[\exp \left(-J \sum_{i:(i,j) \in \mathcal{E}} x_i \right) + \exp \left(J \sum_{i:(i,j) \in \mathcal{E}} x_i \right) \right].$$

Examples

The Core Logic in R

```

tr <- list()
tr[[1]] <- x <- array(0,c(m,n))

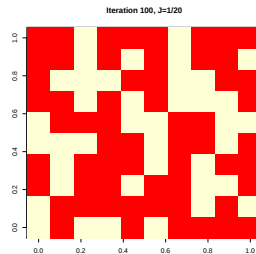
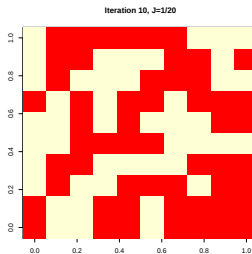
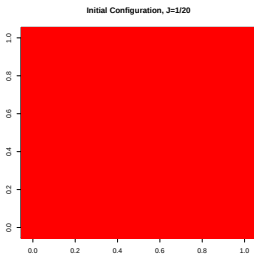
for (t in 1:100) {
  for(i in 1:m) {
    for(j in 1:n) {
      ns <- neighbours(m,n,i,j)
      p1 <- 0
      for(k in 1:length(ns)) {
        p1 <- p1 + x[(ns[[k]])[1] , (ns[[k]])[2]]
      }
      p0 <- length(ns) - p1
      pp <- c(exp(J*p0),exp(J*p1))
      pp <- pp / sum(pp)
      x[i,j] <- sample(c(0,1),1,prob=pp)
    }
  }
  tr[[t+1]] <- x
}

```

Examples

The Gibbs Sampler for Ising Models I

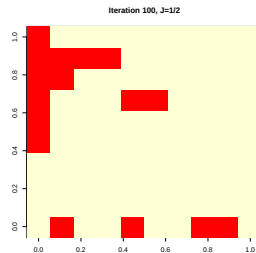
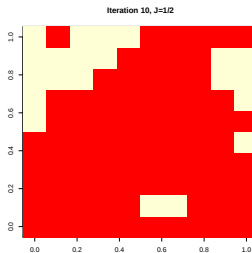
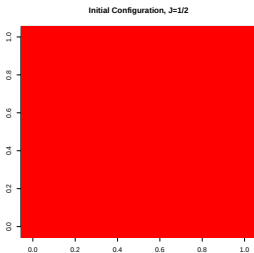
Samples 1, 10, and 100 with $J = 0.05$:



Examples

The Gibbs Sampler for Ising Models II

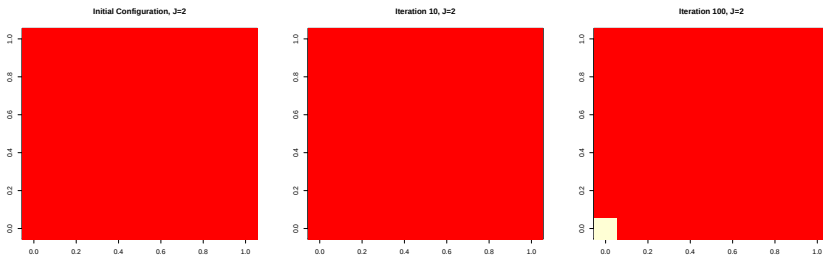
Samples 1, 10, and 100 with $J = 0.50$:



Examples

The Gibbs Sampler for Ising Models III

Samples 1, 10, and 100 with $J = 1.00$:



Solutions include the *Swendsen-Wang* algorithm (c.f. assessment) or *perfect simulation*...

Examples

The Ising Model and Image Reconstruction

The Ising Model is widely used in statistics as a prior distribution.

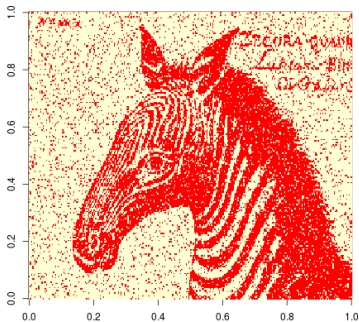
- Consider image denoising: x an $m \times n$ image on $\mathcal{V} \subseteq \mathbb{Z}^2$ with obvious neighbourhood structure $\mathcal{E}$:
- Observe y where $y_v = x_v$ wp $1 - \epsilon$.
- Prior: $X \sim \text{Ising}(J, \mathcal{V}, \mathcal{E})$.
- Likelihood:
$$L(x; y) = \prod_{v \in \mathcal{V}} [(1 - \epsilon)\mathbb{I}\{y_v = x_v\} + \epsilon\mathbb{I}\{y_v \neq x_v\}].$$
- Posterior:

$$p(x|y) \propto \exp \left(2J \sum_{(i,j) \in \mathcal{E}} \mathbb{I}(x_i = x_j) \right) \cdot \prod_{v \in \mathcal{V}} [(1 - \epsilon)\mathbb{I}\{y_v = x_v\} + \epsilon\mathbb{I}\{y_v \neq x_v\}]$$

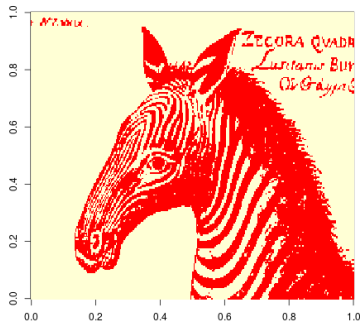
Examples

Ludolphus' Zebra

<https://upload.wikimedia.org/wikipedia/commons/a/af/ZebraLudolphus.jpg>



Noisy Image / Samples



Ground Truth

Examples

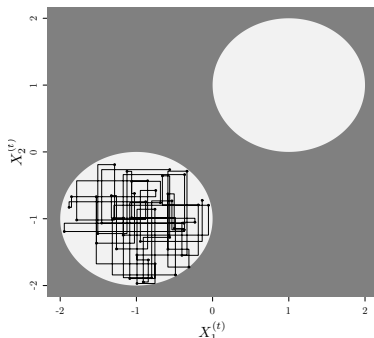
A Pathological Example: The Reducible Gibbs sampler

Consider Gibbs sampling from the uniform distribution

$$f(x_1, x_2) = \frac{1}{2\pi} \mathbb{I}_{C_1 \cup C_2}(x_1, x_2),$$

$$C_1 := \{(x_1, x_2) : \|(x_1, x_2) - (1, 1)\| \leq 1\}$$

$$C_2 := \{(x_1, x_2) : \|(x_1, x_2) + (1, 1)\| \leq 1\}$$



The resulting Markov chain is *reducible*:

It stays forever in either C_1 or C_2 .

Part 3— Section 8

The Metropolis–Hastings Algorithm

The Metropolis–Hastings algorithm

Algorithm: Metropolis–Hastings

Starting with $\mathbf{X}^{(0)} := (X_1^{(0)}, \dots, X_p^{(0)})$ iterate for $t = 1, 2, \dots$

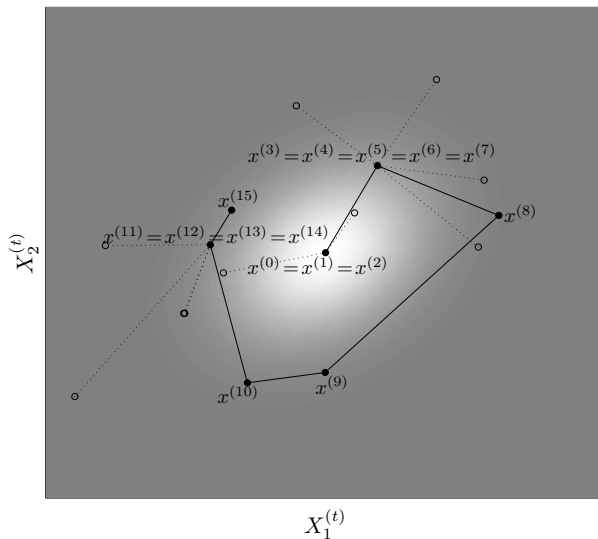
- 1 Draw $\mathbf{X} \sim q(\cdot | \mathbf{X}^{(t-1)})$.
- 2 Compute

$$\alpha(\mathbf{X} | \mathbf{X}^{(t-1)}) = \min \left\{ 1, \frac{f(\mathbf{X}) \cdot q(\mathbf{X}^{(t-1)} | \mathbf{X})}{f(\mathbf{X}^{(t-1)}) \cdot q(\mathbf{X} | \mathbf{X}^{(t-1)})} \right\}.$$

- 3 With probability $\alpha(\mathbf{X} | \mathbf{X}^{(t-1)})$ set $\mathbf{X}^{(t)} = \mathbf{X}$, otherwise set $\mathbf{X}^{(t)} = \mathbf{X}^{(t-1)}$.

The Algorithm

Illustration of the Metropolis–Hastings method



Basic properties of the Metropolis–Hastings algorithm

- The probability that a newly proposed value is accepted given $\mathbf{X}^{(t-1)} = \mathbf{x}^{(t-1)}$ is

$$a(\mathbf{x}^{(t-1)}) = \int \alpha(\mathbf{x}|\mathbf{x}^{(t-1)})q(\mathbf{x}|\mathbf{x}^{(t-1)}) d\mathbf{x}.$$

- The probability of remaining in state $\mathbf{X}^{(t-1)}$ is

$$\mathbb{P}(\mathbf{X}^{(t)} = \mathbf{X}^{(t-1)} | \mathbf{X}^{(t-1)} = \mathbf{x}^{(t-1)}) = 1 - a(\mathbf{x}^{(t-1)}).$$

- The probability of acceptance does not depend on the normalisation constant: If $f(\mathbf{x}) = C \cdot \tilde{f}(\mathbf{x})$, then

$$\alpha(\mathbf{X}|\mathbf{X}^{(t-1)}) = \min \left(1, \frac{\tilde{f}(\mathbf{X}) \cdot q(\mathbf{X}^{(t-1)}|\mathbf{X})}{\tilde{f}(\mathbf{X}^{(t-1)}) \cdot q(\mathbf{X}|\mathbf{X}^{(t-1)})} \right)$$

Transition Kernel

Lemma (Transition Kernel of Metropolis–Hastings)

The transition kernel of the Metropolis–Hastings algorithm is

$$K(\mathbf{x}^{(t-1)}, \mathbf{x}^{(t)}) = \alpha(\mathbf{x}^{(t)} | \mathbf{x}^{(t-1)}) q(\mathbf{x}^{(t)} | \mathbf{x}^{(t-1)}) \\ + (1 - a(\mathbf{x}^{(t-1)})) \delta_{\mathbf{x}^{(t-1)}}(\mathbf{x}^{(t)}),$$

Lemma (Detailed Balance and Metropolis Hastings)

The Metropolis–Hastings kernel satisfies the detailed balance condition, and hence is f -invariant;

$$K(\mathbf{x}^{(t-1)}, \mathbf{x}^{(t)}) f(\mathbf{x}^{(t-1)}) = K(\mathbf{x}^{(t)}, \mathbf{x}^{(t-1)}) f(\mathbf{x}^{(t)}).$$

We will prove this later using categorical probability.

Random-walk Metropolis: Idea

- In the Metropolis–Hastings algorithm the proposal is from $\mathbf{X} \sim q(\cdot | \mathbf{X}^{(t-1)})$.
- A popular choice for the proposal is $q(\mathbf{x} | \mathbf{x}^{(t-1)}) = g(\mathbf{x} - \mathbf{x}^{(t-1)})$ with g symmetric, thus

$$\mathbf{X} = \mathbf{X}^{(t-1)} + \varepsilon, \quad \varepsilon \sim g.$$

- Probability of acceptance becomes

$$\min \left\{ 1, \frac{f(\mathbf{X}) \cdot g(\mathbf{X} - \mathbf{X}^{(t-1)})}{f(\mathbf{X}^{(t-1)}) \cdot g(\mathbf{X}^{(t-1)} - \mathbf{X})} \right\} = \min \left\{ 1, \frac{f(\mathbf{X})}{f(\mathbf{X}^{(t-1)})} \right\}.$$

- We accept ...
 - every move to a more probable state with probability 1.
 - moves to less probable states with a probability $f(\mathbf{X})/f(\mathbf{x}^{(t-1)}) < 1$.

Random-walk Metropolis: Algorithm

Random-Walk Metropolis

Starting with $\mathbf{X}^{(0)} := (X_1^{(0)}, \dots, X_p^{(0)})$ and using a symmetric random walk proposal g , iterate for $t = 1, 2, \dots$

- 1 Draw $\boldsymbol{\varepsilon} \sim g$ and set $\mathbf{X} = \mathbf{X}^{(t-1)} + \boldsymbol{\varepsilon}$.
- 2 Compute

$$\alpha(\mathbf{X}|\mathbf{X}^{(t-1)}) = \min \left\{ 1, \frac{f(\mathbf{X})}{f(\mathbf{X}^{(t-1)})} \right\}.$$

- 3 With probability $\alpha(\mathbf{X}|\mathbf{X}^{(t-1)})$ set $\mathbf{X}^{(t)} = \mathbf{X}$, otherwise set $\mathbf{X}^{(t)} = \mathbf{X}^{(t-1)}$.

Popular choices for g are (multivariate) Gaussians or t-distributions (the latter having heavier tails)

Example 3.4: Bayesian probit model (1)

- Medical study on infections resulting from birth by Cæsarean section.
- 3 influence factors:
 - indicator whether the Cæsarian was planned or not (z_{i1}),
 - indicator of whether additional risk factors were present at the time of birth (z_{i2}), and
 - indicator of whether antibiotics were given as a prophylaxis (z_{i3}).
- Response variable: number of infections Y_i that were observed amongst n_i patients having the same covariates.

# births		planned	risk factors	antibiotics
infection	total			
y_i	n_i	z_{i1}	z_{i2}	z_{i3}
11	98	1	1	1
1	18	0	1	1
0	2	0	0	1
23	26	1	1	0
28	58	0	1	0
0	9	1	0	0
8	40	0	0	0

Example 3.4: Bayesian probit model (2)

- Model for Y_i :

$$Y_i \sim \text{Bin}(n_i, \pi_i), \quad \pi = \Phi(\mathbf{z}'_i \boldsymbol{\beta}),$$

where $\mathbf{z}_i = (1, z_{i1}, z_{i2}, z_{i3})$ and $\Phi(\cdot)$ being the CDF of a $N(0, 1)$.

- Prior on the parameter of interest $\boldsymbol{\beta}$: $\boldsymbol{\beta} \sim N(\mathbf{0}, \mathbb{I}/\lambda)$.
- The posterior density of $\boldsymbol{\beta}$ is

$$f(\boldsymbol{\beta} | y_1, \dots, y_n) \propto \left(\prod_{i=1}^N \Phi(\mathbf{z}'_i \boldsymbol{\beta})^{y_i} \cdot (1 - \Phi(\mathbf{z}'_i \boldsymbol{\beta}))^{n_i - y_i} \right) \cdot \exp \left(-\frac{\lambda}{2} \sum_{j=0}^3 \beta_j^2 \right)$$

Example 3.4: Bayesian probit model (3)

Use the following “random walk Metropolis” algorithm.
Starting with any $\beta^{(0)}$ iterate for $t = 1, 2, \dots$:

1. Draw $\epsilon \sim N(\mathbf{0}, \Sigma)$ and set $\beta = \beta^{(t-1)} + \epsilon$.
2. Compute

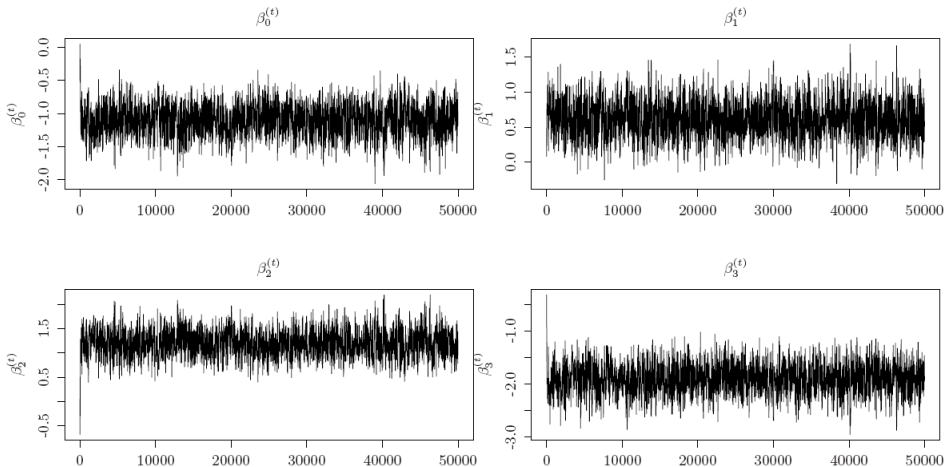
$$\alpha(\beta|\beta^{(t-1)}) = \min \left\{ 1, \frac{f(\beta|Y_1, \dots, Y_n)}{f(\beta^{(t-1)}|Y_1, \dots, Y_n)} \right\}.$$

3. With probability $\alpha(\beta|\beta^{(t-1)})$ set $\beta^{(t)} = \beta$, otherwise set $\beta^{(t)} = \beta^{(t-1)}$.

(for the moment we use $\Sigma = 0.08 \cdot \mathbb{I}$, and $\lambda = 10$).

Random-walk Metropolis with Examples

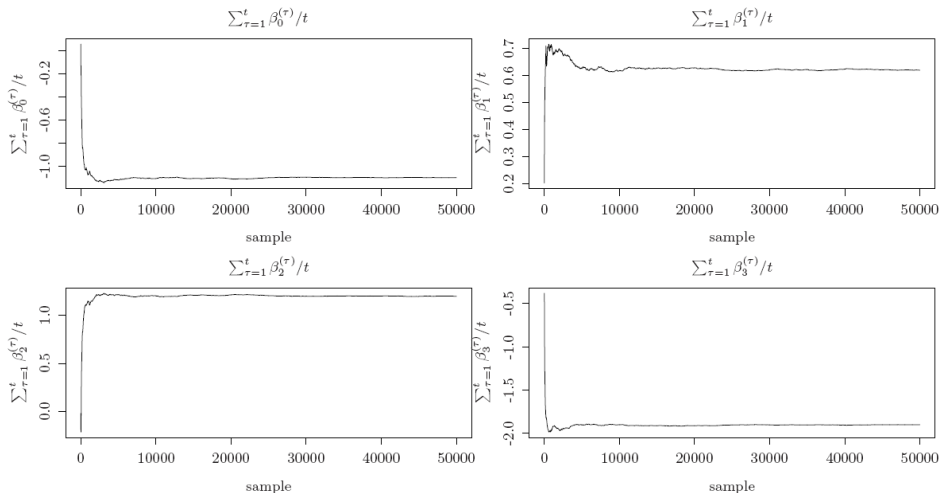
Example 3.4: Bayesian probit model (4)



Convergence of the $\beta_j^{(t)}$ is to a distribution, not a value!

Random-walk Metropolis with Examples

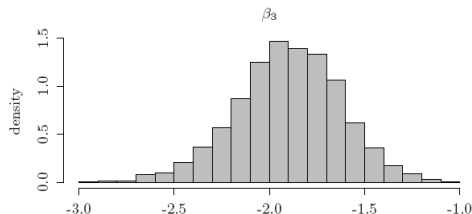
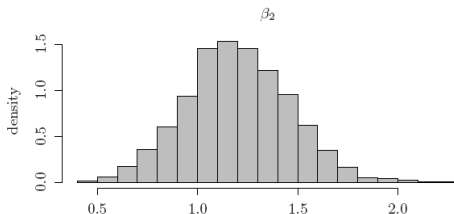
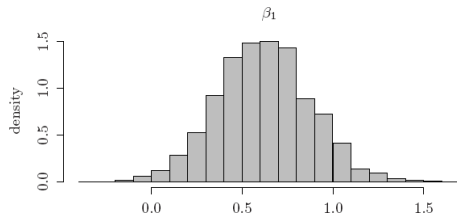
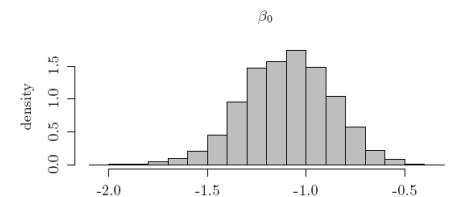
Example 3.4: Bayesian probit model (5)



Convergence of cumulative averages $\sum_{\tau=1}^t \beta_j^{(\tau)} / t$ is to a value.

Random-walk Metropolis with Examples

Example 3.4: Bayesian probit model (6)



Example 3.4: Bayesian probit model (7)

		Posterior mean	95% credible interval	
intercept	β_0	-1.0952	-1.4646	-0.7333
planned	β_1	0.6201	0.2029	1.0413
risk factors	β_2	1.2000	0.7783	1.6296
antibiotics	β_3	-1.8993	-2.3636	-1.471

Choosing a good proposal distribution

- Ideally: Markov chain with small correlations $\rho(\mathbf{X}^{(t-1)}, \mathbf{X}^{(t)})$. Yields fast exploration of the support of the target f .
- Two sources for this correlation:
 - correlation between current state $\mathbf{X}^{(t-1)}$ and newly proposed value $\mathbf{X} \sim q(\cdot | \mathbf{X}^{(t-1)})$
(can be reduced using a proposal with high variance),
 - correlation introduced by retaining a value $\mathbf{X}^{(t)} = \mathbf{X}^{(t-1)}$ because the proposal $\mathbf{X}$ has been rejected
(can be reduced using a proposal with small variance).
- Trade-off for finding compromise between:
 - fast exploration of the space (good mixing behaviour),
 - obtaining a large probability of acceptance.
- For multivariate distributions: covariance of proposal should reflect the covariance structure of the target.

Example: Choice of proposal (1)

- Target distribution: $N(0, 1)$ (i.e. $f(\cdot) = \phi_{(0,1)}(\cdot)$).
- We want to use a random walk Metropolis algorithm with

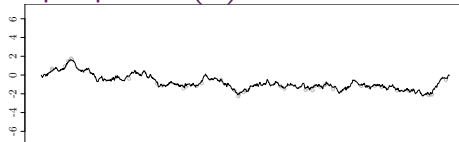
$$\varepsilon \sim N(0, \sigma^2).$$

- What is the optimal choice of σ^2 ?
- We consider four choices $\sigma^2 = 0.01, 1, 5, 100$.

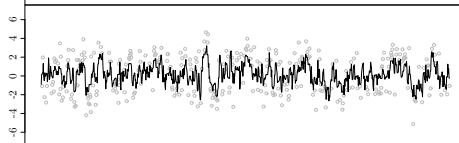
Random-walk Metropolis with Examples

Example 5.3: Choice of proposal (2)

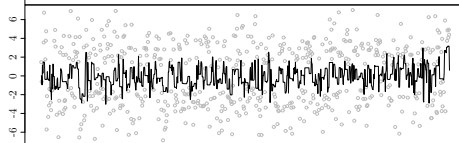
$$\sigma^2 = 0.01$$



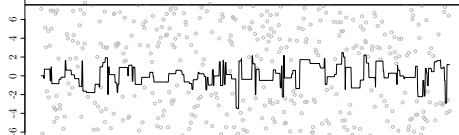
$$\sigma^2 = 1$$



$$\sigma^2 = 5$$



$$\sigma^2 = 100$$



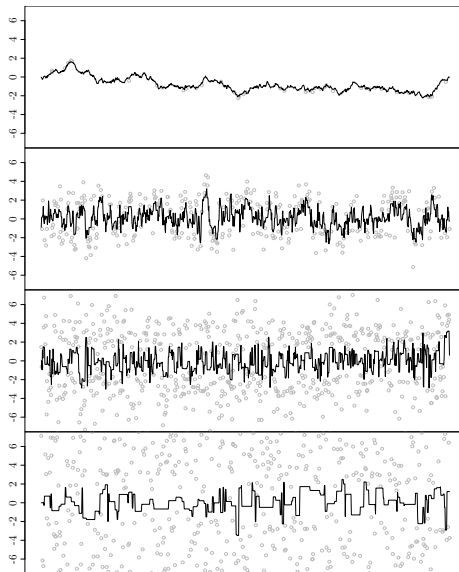
Random-walk Metropolis with Examples

$$\sigma^2 = 0.01$$

$$\sigma^2 = 1$$

$$\sigma^2 = 5$$

$$\sigma^2 = 100$$



Which proposal looks best?

Vevox.app 124-467-150

Example 5.3: Choice of proposal (4)

	Autocorrelation $\rho(X^{(t-1)}, X^{(t)})$		Probability of acceptance $\alpha(X, X^{(t-1)})$	
	Mean	95% CI	Mean	95% CI
$\sigma^2 = 0.1^2$	0.9901	(0.9891, 0.9910)	0.9694	(0.9677, 0.9710)
$\sigma^2 = 1$	0.7733	(0.7676, 0.7791)	0.7038	(0.7014, 0.7061)
$\sigma^2 = 2.38^2$	0.6225	(0.6162, 0.6289)	0.4426	(0.4401, 0.4452)
$\sigma^2 = 10^2$	0.8360	(0.8303, 0.8418)	0.1255	(0.1237, 0.1274)

Suggests: Optimal choice is $\sigma^2 = 2.38^2 = 5.66 > 1$.

Example 5.4: Bayesian probit model (revisited)

- So far we used: $\text{Var}(\epsilon) = 0.08 \cdot \mathbb{I}$.
- Better choice: Let $\text{Var}(\epsilon)$ reflect the covariance structure
- Frequentist asymptotic theory: $\text{Var}(\hat{\beta}^{\text{m.l.e}}) = (\mathbf{Z}'\mathbf{D}\mathbf{Z})^{-1}$, $\mathbf{D}$ is a suitable diagonal matrix.
- Better choice: $\text{Var}(\epsilon) = 2 \cdot (\mathbf{Z}'\mathbf{D}\mathbf{Z})^{-1}$.
- Increases rate of acceptance from 13.9% to 20.0% and reduces autocorrelation:

$\Sigma = 0.08 \cdot \mathbf{I}$	β_0	β_1	β_2	β_3
Autocorrelation $\rho(\beta_j^{(t-1)}, \beta_j^{(t)})$	0.9496	0.9503	0.9562	0.9532
$\Sigma = 2 \cdot (\mathbf{Z}'\mathbf{D}\mathbf{Z})^{-1}$	β_0	β_1	β_2	β_3
Autocorrelation $\rho(\beta_j^{(t-1)}, \beta_j^{(t)})$	0.8726	0.8765	0.8741	0.8792

(In this example $\det(0.08 \cdot \mathbb{I}) = \det(2 \cdot (\mathbf{Z}'\mathbf{D}\mathbf{Z})^{-1})$.)

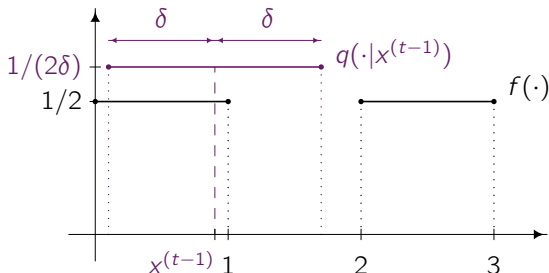
Pathological Example: Reducible Metropolis–Hastings

Consider the target distribution

$$f(x) = (\mathbb{I}_{[0,1]}(x) + \mathbb{I}_{[2,3]}(x))/2.$$

and the proposal distribution $q(\cdot|x^{(t-1)})$:

$$X|X^{(t-1)} = x^{(t-1)} \sim U[x^{(t-1)} - \delta, x^{(t-1)} + \delta]$$



Reducible if $\delta \leq 1$: the chain stays either in $[0, 1]$ or $[2, 3]$.

The Metropolised Independence Sampler

Independent proposals: choose $q(\cdot|x) = q(\cdot)$.

Algorithm 5.3 The Independence Sampler

Starting with $\mathbf{X}^{(0)} := (X_1^{(0)}, \dots, X_p^{(0)})$ iterate for $t = 1, 2, \dots$

1. Draw $\mathbf{X} \sim q(\cdot)$.
2. Compute

$$\alpha(\mathbf{X}|\mathbf{X}^{(t-1)}) = \min \left\{ 1, \frac{f(\mathbf{X}) \cdot q(\mathbf{X}^{(t-1)})}{f(\mathbf{X}^{(t-1)}) \cdot q(\mathbf{X})} \right\}.$$

3. With probability $\alpha(\mathbf{X}|\mathbf{X}^{(t-1)})$ set $\mathbf{X}^{(t)} = \mathbf{X}$, otherwise set $\mathbf{X}^{(t)} = \mathbf{X}^{(t-1)}$.

Acceptance Rate

Proposition (Acceptance Rate of Independence Sampler)

If $f(\mathbf{x})/q(\mathbf{x}) \leq M < \infty$ the acceptance rate of the independence sampler is at least as high as that of the corresponding rejection sampler.

Gibbs Samplers Revisited

What about full conditionals as MH proposals?

- For $\mathbf{X} = (X_1, \dots, X_p)$:
- Consider $q(\mathbf{X}|\mathbf{x}^{(t-1)}) = \delta_{X_{-p}^{(t-1)}}(X_{-p})f_{X_p|X_{-p}}(X_p|X_{-p})$.

Remark

A Gibbs sampler step is a special case of the Metropolis–Hastings algorithm.

Part 3— Section 9

The Exchange Algorithm

Moving beyond vanilla MH

The Metropolis–Hastings algorithm is an extremely general construction (more like a ‘*meta-algorithm*’) and it can be extended significantly.

We will see this later when we discuss the *involution* approach of Andrieu, Lee, Livingstone (2020).

For now, as a taster: an extension to the case of *doubly-intractable models*.

Double Intractable Models

Recall we have a likelihood

$$L(\boldsymbol{\theta}; \mathbf{y}^{\text{obs}}) = f(\mathbf{y}^{\text{obs}} | \boldsymbol{\theta}) / Z(\boldsymbol{\theta}),$$

where we have made explicit the normalising constant $Z(\boldsymbol{\theta})$ which may depend on the parameters (known as the *partition function* in physics).

In some cases, this $Z(\boldsymbol{\theta})$ is intractable - these are referred to as *doubly-intractable models*.

This causes serious problems since we will need to compute (e.g. for random walk proposal)

$$1 \wedge \frac{f_{\text{prior}}(\boldsymbol{\theta}) L(\boldsymbol{\theta}; \mathbf{y}^{\text{obs}})}{f_{\text{prior}}(\boldsymbol{\vartheta}) L(\boldsymbol{\vartheta}; \mathbf{y}^{\text{obs}})} = 1 \wedge \frac{f_{\text{prior}}(\boldsymbol{\theta}) f(\mathbf{y}^{\text{obs}} | \boldsymbol{\theta})}{f_{\text{prior}}(\boldsymbol{\vartheta}) f(\mathbf{y}^{\text{obs}} | \boldsymbol{\vartheta})} \cdot \frac{Z(\boldsymbol{\vartheta})}{Z(\boldsymbol{\theta})},$$

for different values $\boldsymbol{\theta} \neq \boldsymbol{\vartheta}$: these Z factors **do not cancel out!**

Double Intractable Models

Recall we have a likelihood

$$L(\boldsymbol{\theta}; \mathbf{y}^{\text{obs}}) = f(\mathbf{y}^{\text{obs}} | \boldsymbol{\theta}) / Z(\boldsymbol{\theta}),$$

where we have made explicit the normalising constant $Z(\boldsymbol{\theta})$ which may depend on the parameters (known as the *partition function* in physics).

In some cases, this $Z(\boldsymbol{\theta})$ is intractable - these are referred to as *doubly-intractable models*.

This causes serious problems since we will need to compute (e.g. for random walk proposal)

$$1 \wedge \frac{f^{\text{prior}}(\boldsymbol{\theta}) L(\boldsymbol{\theta}; \mathbf{y}^{\text{obs}})}{f^{\text{prior}}(\boldsymbol{\vartheta}) L(\boldsymbol{\vartheta}; \mathbf{y}^{\text{obs}})} = 1 \wedge \frac{f^{\text{prior}}(\boldsymbol{\theta}) f(\mathbf{y}^{\text{obs}} | \boldsymbol{\theta})}{f^{\text{prior}}(\boldsymbol{\vartheta}) f(\mathbf{y}^{\text{obs}} | \boldsymbol{\vartheta})} \cdot \frac{Z(\boldsymbol{\vartheta})}{Z(\boldsymbol{\theta})},$$

for different values $\boldsymbol{\theta} \neq \boldsymbol{\vartheta}$: these Z factors **do not cancel out!**

Example

Recall the Ising model from before:

$$\pi(x_1, \dots, x_m | \theta) = \frac{1}{Z(\theta)} \exp \left(\theta \sum_{(i,j) \in \mathcal{E}} x_i \cdot x_j \right),$$

but now suppose $x_1, \dots, x_m$ are fixed and we want to infer $\theta > 0$; then

$$Z(\theta) = \sum_{(x_1, \dots, x_m) \in \{-1, +1\}^m} \exp \left(\theta \sum_{(i,j) \in \mathcal{E}} x_i \cdot x_j \right),$$

which could be a *very, very* big sum.

Is Bayesian inference still possible in such models?

The Exchange Algorithm

The Exchange Algorithm¹ is one ingenious solution to this problem.

It is an *auxiliary variable approach*: we suppose we are able to sample for any θ , new ‘data’

$$\mathbf{y} \sim f(\mathbf{y}|\theta)/Z(\theta).$$

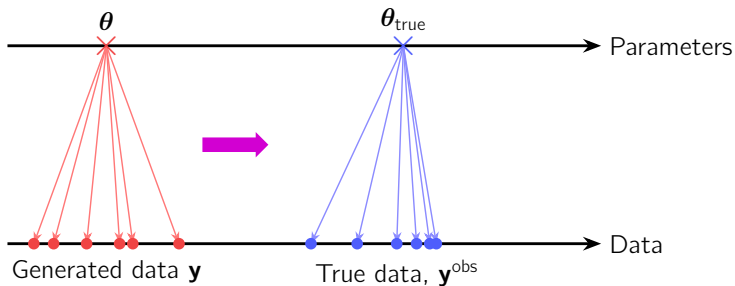
Then at each MCMC iteration, we will sample new ‘data’ and use this auxiliary data to cancel out the intractable $Z(\theta)$!

¹Murray, I., Ghahramani, Z., MacKay, D. (2012). MCMC for doubly-intractable distributions. Proceedings of the 22nd Conference on Uncertainty in Artificial Intelligence, UAI 2006, 359-366.

Doubly Intractable Models

Exchange algorithm intuition

Suppose we are at some θ - we generate synthetic data $\mathbf{y} \sim f(\mathbf{y}|\theta)$.



If this synthetic data isn't similar to the actual data $\mathbf{y}^{\text{obs}}$, that suggests it's a poor value of the parameter!

The Exchange Algorithm

Starting with $\theta^{(0)}$ iterate for $t = 1, 2, \dots$

- 1 Draw $\vartheta \sim q(\cdot | \theta^{(t-1)})$.
- 2 Draw auxiliary data $\mathbf{y} \sim f(\mathbf{y} | \vartheta) / Z(\vartheta)$.
- 3 Compute $\alpha(\vartheta, \mathbf{y} | \theta^{(t-1)})$ given by

$$\begin{aligned} 1 \wedge \frac{f^{\text{pr}}(\vartheta) q(\theta^{(t-1)} | \vartheta) f(\mathbf{y}^{\text{obs}} | \vartheta) / Z(\vartheta) \cdot f(\mathbf{y} | \theta^{(t-1)}) / Z(\theta^{(t-1)})}{f^{\text{pr}}(\theta^{(t-1)}) q(\vartheta | \theta^{(t-1)}) f(\mathbf{y}^{\text{obs}} | \theta^{(t-1)}) / Z(\theta^{(t-1)}) \cdot f(\mathbf{y} | \vartheta) / Z(\vartheta)} \\ = 1 \wedge \frac{f^{\text{pr}}(\vartheta) q(\theta^{(t-1)} | \vartheta) f(\mathbf{y}^{\text{obs}} | \vartheta) f(\mathbf{y} | \theta^{(t-1)})}{f^{\text{pr}}(\theta^{(t-1)}) q(\vartheta | \theta^{(t-1)}) f(\mathbf{y}^{\text{obs}} | \theta^{(t-1)}) f(\mathbf{y} | \vartheta)} \end{aligned}$$

- 4 With probability $\alpha(\vartheta, \mathbf{y} | \theta^{(t-1)})$ set $\theta^{(t)} = \vartheta$, otherwise set $\theta^{(t)} = \theta^{(t-1)}$.

We will derive this acceptance probability and show reversibility!

Doubly Intractable Models

Summary of Part 3

- Motivation
- MCMC
- Gibbs Samplers
- Metropolis–Hastings-type Algorithms
- The Exchange Algorithm

Part 4

Theory and Practice

Part 4— Section 10

Interlude: Introduction to Categorical Probability

Categorical probability: motivation

Statistics is built on probability theory, which in turn is built on measure theory.

Technically speaking, we need to construct an appropriate probability space $(\Omega, \mathcal{F}, \mathbb{P})$, and then all random variables we need should be constructed as measurable mappings $X : \Omega \rightarrow \mathbb{R}$.

Needless to say, *nobody bothers with this in practice!*

In practice, we just call functions like *sample*, *rnorm*, *runif* etc. and simply generate whatever variables we need on the fly.

Categorical probability

Categorical probability² takes a different approach: instead of focussing on technical details such as sample spaces and σ -algebras, we take a user-oriented approach and focus on the transformations, that is, *kernels* (morphisms) and *composition*.

For this course, you do not need to know any category theory!

For us, we will simply make use of *string diagrams* as a way of streamlining proofs regarding Monte Carlo algorithms.

²Fritz, T. (2020). A synthetic approach to Markov kernels, conditional independence and theorems on sufficient statistics. *Advances in Mathematics*, 370, 107239. Cho, K., Jacobs, B. (2019). Disintegration and Bayesian inversion via string diagrams. *Mathematical Structures in Computer Science*, 29(7), 938-971.

Categorical probability

Categorical probability² takes a different approach: instead of focussing on technical details such as sample spaces and σ -algebras, we take a user-oriented approach and focus on the transformations, that is, *kernels* (morphisms) and *composition*.

For this course, you do not need to know any category theory!

For us, we will simply make use of *string diagrams* as a way of streamlining proofs regarding Monte Carlo algorithms.

²Fritz, T. (2020). A synthetic approach to Markov kernels, conditional independence and theorems on sufficient statistics. *Advances in Mathematics*, 370, 107239, Cho, K., Jacobs, B. (2019). Disintegration and Bayesian inversion via string diagrams. *Mathematical Structures in Computer Science*, 29(7), 938-971.

Categorical probability

Categorical probability² takes a different approach: instead of focussing on technical details such as sample spaces and σ -algebras, we take a user-oriented approach and focus on the transformations, that is, *kernels* (morphisms) and *composition*.

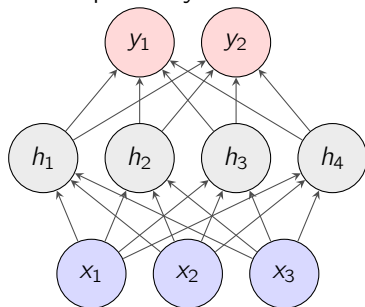
For this course, you do not need to know any category theory!

For us, we will simply make use of *string diagrams* as a way of streamlining proofs regarding Monte Carlo algorithms.

²Fritz, T. (2020). A synthetic approach to Markov kernels, conditional independence and theorems on sufficient statistics. *Advances in Mathematics*, 370, 107239, Cho, K., Jacobs, B. (2019). Disintegration and Bayesian inversion via string diagrams. *Mathematical Structures in Computer Science*, 29(7), 938-971.

String diagrams - neural network diagrams

You have probably all seen this diagram (or something similar)...

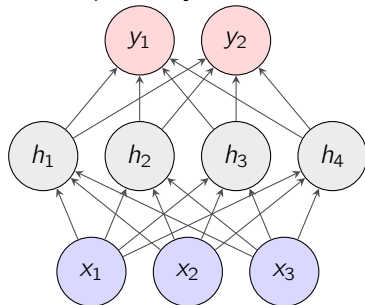


It represents a 3-layer neural network with 1 hidden layer, 3 input neurons and 2 output neurons...

We are going to formalise this kind of diagram and use it as a computational tool!

String diagrams - neural network diagrams

You have probably all seen this diagram (or something similar)...



It represents a 3-layer neural network with 1 hidden layer, 3 input neurons and 2 output neurons...

We are going to formalise this kind of diagram and use it as a computational tool!

Kernels as random functions

We have already seen Markov kernels $K(x, dy)$, $q(\cdot|\mathbf{X})$.

Intuitively, think of a kernel $K(x, dy)$ as a *function with a random output*: $K : \mathcal{X} \rightarrow \mathcal{Y}$ takes input $x \in \mathcal{X}$ and outputs a *random* $y \in \mathcal{Y}$.

My first kernels

```
add_Gaussian <- function(x){
  z = rnorm(1)
  return x+z
}

mult_Gaussian <- function(x){
  w = rnorm(1)
  return w*x
}
```

Kernels as random functions

We have already seen Markov kernels $K(x, dy)$, $q(\cdot|\mathbf{X})$.

Intuitively, think of a kernel $K(x, dy)$ as a *function with a random output*: $K : \mathcal{X} \rightarrow \mathcal{Y}$ takes input $x \in \mathcal{X}$ and outputs a *random* $y \in \mathcal{Y}$.

My first kernels

```
add_Gaussian <- function(x){
  z = rnorm(1)
  return x+z
}

mult_Gaussian <- function(x){
  w = rnorm(1)
  return w*x
}
```

Kernels as random functions

We have already seen Markov kernels $K(x, dy)$, $q(\cdot|\mathbf{X})$.

Intuitively, think of a kernel $K(x, dy)$ as a *function with a random output*: $K : \mathcal{X} \rightarrow \mathcal{Y}$ takes input $x \in \mathcal{X}$ and outputs a *random* $y \in \mathcal{Y}$.

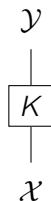
My first kernels

```
add_Gaussian <- function(x){
  z = rnorm(1)
  return x+z
}

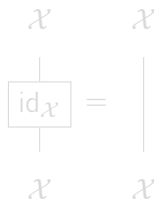
mult_Gaussian <- function(x){
  w = rnorm(1)
  return w*x
}
```

First string diagrams

Given a kernel $K : \mathcal{X} \rightarrow \mathcal{Y}$, we will depict it visually as follows:

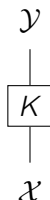


For each space $\mathcal{X}$, there is an *identity kernel* $\text{id}_{\mathcal{X}} : \mathcal{X} \rightarrow \mathcal{X}$ given by $x \mapsto \delta_x$ which ‘does nothing’. This kernel is simply not depicted:

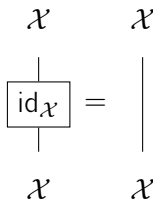


First string diagrams

Given a kernel $K : \mathcal{X} \rightarrow \mathcal{Y}$, we will depict it visually as follows:



For each space $\mathcal{X}$, there is an *identity kernel* $\text{id}_{\mathcal{X}} : \mathcal{X} \rightarrow \mathcal{X}$ given by $x \mapsto \delta_x$ which ‘does nothing’. This kernel is simply not depicted:



Composition of kernels

Chapman–Kolmogorov

Given two kernels

$$K_1 : \mathcal{X} \rightarrow \mathcal{Y}, \quad K_2 : \mathcal{Y} \rightarrow \mathcal{Z},$$

we can form the composite kernel $K_2 \circ K_1 : \mathcal{X} \rightarrow \mathcal{Z}$ through the *Chapman–Kolmogorov* equation:

$$K_2 \circ K_1(x, dz) = \int_{\mathcal{Y}} K_1(x, dy) K_2(y, dz).$$

The identity kernel $\text{id}_{\mathcal{X}} : \mathcal{X} \rightarrow \mathcal{X}$ given by $x \mapsto \delta_x$ has the property that

$$P \circ \text{id}_{\mathcal{X}} = P, \quad \text{id}_{\mathcal{X}} \circ P' = P'.$$

Composition of kernels

Chapman–Kolmogorov

Given two kernels

$$K_1 : \mathcal{X} \rightarrow \mathcal{Y}, \quad K_2 : \mathcal{Y} \rightarrow \mathcal{Z},$$

we can form the composite kernel $K_2 \circ K_1 : \mathcal{X} \rightarrow \mathcal{Z}$ through the *Chapman–Kolmogorov* equation:

$$K_2 \circ K_1(x, dz) = \int_{\mathcal{Y}} K_1(x, dy) K_2(y, dz).$$

The identity kernel $\text{id}_{\mathcal{X}} : \mathcal{X} \rightarrow \mathcal{X}$ given by $x \mapsto \delta_x$ has the property that

$$P \circ \text{id}_{\mathcal{X}} = P, \quad \text{id}_{\mathcal{X}} \circ P' = P'.$$

Chapman–Kolmogorov pseudo-code

The corresponding pseudo-code is extremely simple:

Kernel composition

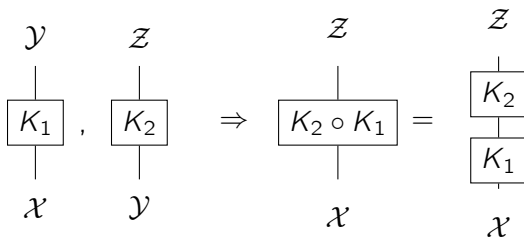
```
kernel1 <- function(x){           # K_1(x, dy)
  y = [some complex stochastic operation]
  return y
}

kernel2 <- function(y){           # K_2(y, dz)
  z = [another complex stochastic operation]
  return z
}

k2_after_k1 = function(x){        # K_2 \circ K_1 (x, dz)
  y = kernel1(x)
  z = kernel2(y)
  return z
}
```

Chapman-Kolmogorov string diagram

Just as in our neural network diagram, composition is simply connecting the boxes:



Parallel composition and swap

Parallel composition

Given any two kernels

$$K : \mathcal{X} \rightarrow \mathcal{Y}, \quad K' : \mathcal{X}' \rightarrow \mathcal{Y}',$$

we can also compose these in parallel via the *independent coupling*:

$$K \otimes K' : \mathcal{X} \otimes \mathcal{X}' \rightarrow \mathcal{Y} \otimes \mathcal{Y}',$$

$$K \otimes K'(x, x'; dy, dy') = K(x, dy)K'(x', dy').$$

Here $\mathcal{X} \otimes \mathcal{X}'$ denotes the *product space* (equipped with product σ -algebra).

In addition, we have swap operation $(x, y) \mapsto \delta_{(y, x)}$.

Parallel composition and swap

Parallel composition

Given any two kernels

$$K : \mathcal{X} \rightarrow \mathcal{Y}, \quad K' : \mathcal{X}' \rightarrow \mathcal{Y}',$$

we can also compose these in parallel via the *independent coupling*:

$$K \otimes K' : \mathcal{X} \otimes \mathcal{X}' \rightarrow \mathcal{Y} \otimes \mathcal{Y}',$$

$$K \otimes K'(x, x'; dy, dy') = K(x, dy)K'(x', dy').$$

Here $\mathcal{X} \otimes \mathcal{X}'$ denotes the *product space* (equipped with product σ -algebra).

In addition, we have swap operation $(x, y) \mapsto \delta_{(y, x)}$.

Parallel composition pseudo-code

Parallel composition

```

kernel <- function(x){                                     # K(x, dy)
  y = [something complex]
  return y
}

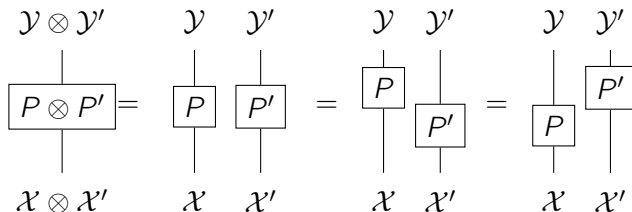
kernelp <- function(xp){                                   # K'(x', dy')
  yp = [something else complex]
  return yp
}

k_otimes_kp = function(x,xp){                             # K \otimes K' (x,x'; dy,dy')
  y = kernel(x)
  yp = kernelp(xp)
  return c(y,yp)
}

```

More string diagrams

Parallel composition is simply concatenation of diagrams:



Morphisms side-by-side are (conditionally) independent.

Key fact

String diagrams are *purely topological*; you can freely deform them (as long as you don't cut wires or rewire).

More string diagrams

Parallel composition is simply concatenation of diagrams:

$$\begin{array}{c}
 \mathcal{Y} \otimes \mathcal{Y}' \\
 \hline
 \boxed{P \otimes P'} \\
 \hline
 \mathcal{X} \otimes \mathcal{X}'
 \end{array}
 =
 \begin{array}{cc}
 \mathcal{Y} & \mathcal{Y}' \\
 \hline
 \boxed{P} & \boxed{P'} \\
 \hline
 \mathcal{X} & \mathcal{X}'
 \end{array}
 =
 \begin{array}{cc}
 \mathcal{Y} & \mathcal{Y}' \\
 \hline
 \boxed{P} & \hline
 \boxed{P'} \\
 \hline
 \mathcal{X} & \mathcal{X}'
 \end{array}
 =
 \begin{array}{cc}
 \mathcal{Y} & \mathcal{Y}' \\
 \hline
 \boxed{P} & \boxed{P'} \\
 \hline
 \mathcal{X} & \mathcal{X}'
 \end{array}$$

Morphisms side-by-side are (conditionally) independent.

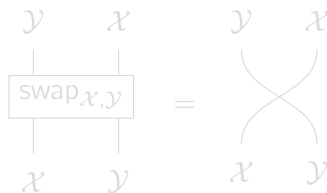
Key fact

String diagrams are *purely topological*; you can freely deform them (as long as you don't cut wires or rewire).

String diagrams

The compositional structure we have - sequential and parallel composition - is a special categorical structure known a *symmetric monoidal category*.

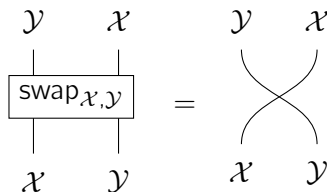
We also have a special swap isomorphism $(x, y) \mapsto \delta_{y,x}$:



String diagrams

The compositional structure we have - sequential and parallel composition - is a special categorical structure known a *symmetric monoidal category*.

We also have a special swap isomorphism $(x, y) \mapsto \delta_{y,x}$:



Individual probability measures

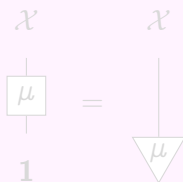
Everything so far has been a kernel $P : \mathcal{X} \rightarrow \mathcal{Y}$. But what if we want to have a single probability measure μ on $\mathcal{X}$?

Individual measure

An individual measure μ can be viewed as a kernel

$$\mu : \mathbf{1} \rightarrow \mathcal{X},$$

where $\mathbf{1} := \{*\}$ is the *one-point space* (with its trivial σ -algebra).



Individual probability measures

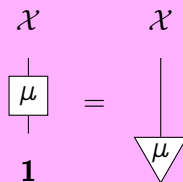
Everything so far has been a kernel $P : \mathcal{X} \rightarrow \mathcal{Y}$. But what if we want to have a single probability measure μ on $\mathcal{X}$?

Individual measure

An individual measure μ can be viewed as a kernel

$$\mu : \mathbf{1} \rightarrow \mathcal{X},$$

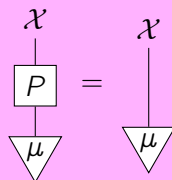
where $\mathbf{1} := \{*\}$ is the *one-point space* (with its trivial σ -algebra).



Invariance of kernels

Invariance

If we have $\mu : \mathbf{1} \rightarrow \mathcal{X}$ and $P : \mathcal{X} \rightarrow \mathcal{X}$, we say P is μ -invariant if



This exactly corresponds to our previous definition

$$\mu(dx) = \int \mu(dx') P(x', dx).$$

Markov categories

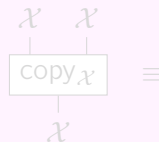
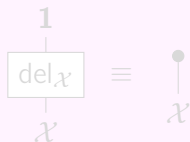
When coding we often want to use the same variable multiple times (as in the neural network example), and we can discard variables we no longer need.

Copy and delete kernels

Explicitly, we have the following kernels:

$$\text{copy}_{\mathcal{X}} : \mathcal{X} \rightarrow \mathcal{X} \otimes \mathcal{X}, \quad x \mapsto \delta_x \otimes \delta_x$$

$$\text{del}_{\mathcal{X}} : \mathcal{X} \rightarrow \mathbf{1}, \quad x \mapsto \delta_*$$



Markov categories

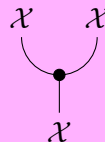
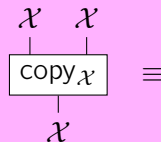
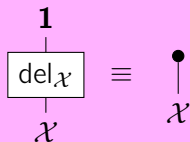
When coding we often want to use the same variable multiple times (as in the neural network example), and we can discard variables we no longer need.

Copy and delete kernels

Explicitly, we have the following kernels:

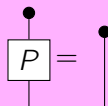
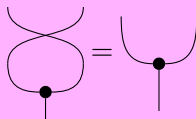
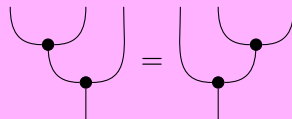
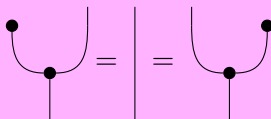
$$\text{copy}_{\mathcal{X}} : \mathcal{X} \rightarrow \mathcal{X} \otimes \mathcal{X}, \quad x \mapsto \delta_x \otimes \delta_x$$

$$\text{del}_{\mathcal{X}} : \mathcal{X} \rightarrow \mathbf{1}, \quad x \mapsto \delta_*$$



copy and del properties

Together copy and delete satisfy some natural properties:

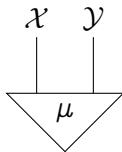


Where the last equation holds for any $P : \mathcal{X} \rightarrow \mathcal{Y}$, and reflects *normalisation*:

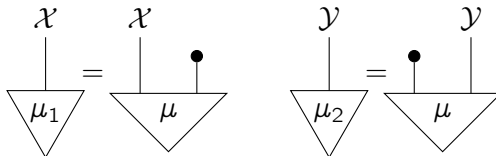
$$P(x, \mathcal{Y}) = 1, \quad \forall x \in \mathcal{X}.$$

del is marginalisation

The delete morphism precisely corresponds to *marginalisation*:
given a bivariate distribution μ on product space $\mathcal{X} \otimes \mathcal{Y}$,



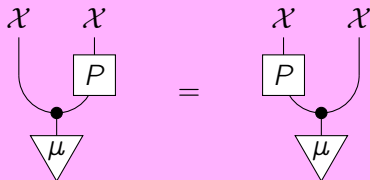
Its marginals can be represented as



Reversibility

Reversibility in string diagrams

Given $\mu : \mathbf{1} \rightarrow \mathcal{X}$ and $P : \mathcal{X} \rightarrow \mathcal{X}$, we say that P is μ -reversible if



In equations, this reduces to the *detailed balance* condition

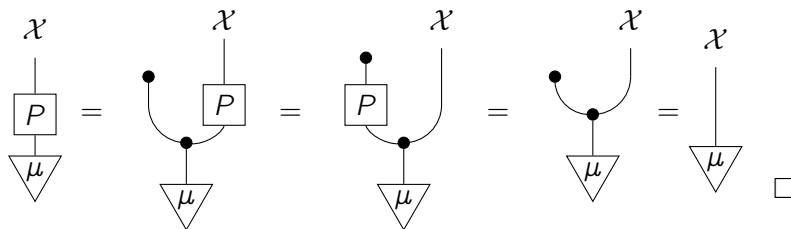
$$\mu(dx)P(x, dx') = \mu(dx')P(x', dx).$$

Reversibility implies invariance

Lemma

Suppose $P : \mathcal{X} \rightarrow \mathcal{X}$ is μ -reversible. Then it is μ -invariant.

Proof:



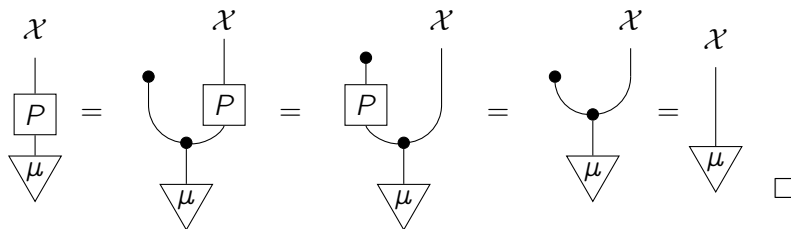
Exercise: instantiate this proof in equations and check it aligns with the 'standard' proof from Markov chain theory.

Reversibility implies invariance

Lemma

Suppose $P : \mathcal{X} \rightarrow \mathcal{X}$ is μ -reversible. Then it is μ -invariant.

Proof:



Exercise: instantiate this proof in equations and check it aligns with the 'standard' proof from Markov chain theory.

Part 4— Section 11

Theoretical Considerations and Convergence Results

Invariance of Gibbs sampler

Gibbs sampler recap

Recall we have a joint density $f(X_1, \dots, X_p)$ on

$$\mathcal{X}_{1:p} := \mathcal{X}_1 \otimes \dots \mathcal{X}_p.$$

For systematic scan Gibbs, we iteratively sample from the *full conditional distributions*

$$X_j \sim f_{X_j|X_{-j}}(\cdot | X_1, \dots, X_{j-1}, X_{j+1}, \dots, X_p).$$

We can view these full conditionals as kernels

$$f_{X_j|X_{-j}} : \mathcal{X}_{1:(j-1)} \otimes \mathcal{X}_{(j+1):p} \rightarrow \mathcal{X}_j.$$

Invariance of Gibbs sampler

Gibbs sampler recap

Recall we have a joint density $f(X_1, \dots, X_p)$ on

$$\mathcal{X}_{1:p} := \mathcal{X}_1 \otimes \dots \mathcal{X}_p.$$

For systematic scan Gibbs, we iteratively sample from the *full conditional distributions*

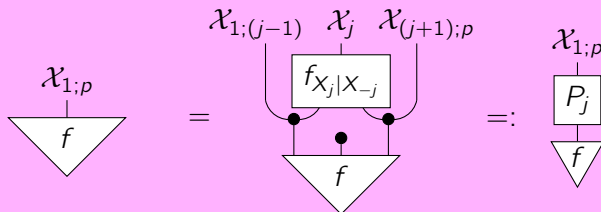
$$X_j \sim f_{X_j|X_{-j}}(\cdot | X_1, \dots, X_{j-1}, X_{j+1}, \dots, X_p).$$

We can view these full conditionals as kernels

$$f_{X_j|X_{-j}} : \mathcal{X}_{1:(j-1)} \otimes \mathcal{X}_{(j+1):p} \rightarrow \mathcal{X}_j.$$

Conditionals as kernels

Viewing $f_{X_j|X_{-j}} : \mathcal{X}_{1:(j-1)} \otimes \mathcal{X}_{(j+1):p} \rightarrow \mathcal{X}_j$ as a kernel, its defining property is:

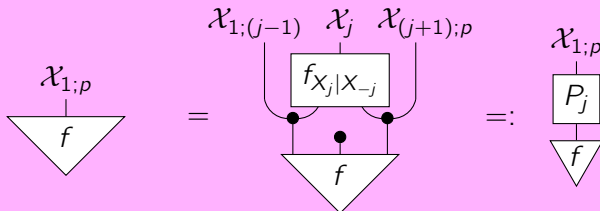


Where we have defined $P_j : \mathcal{X}_{1:p} \rightarrow \mathcal{X}_{1:p}$.
By construction, P_j is f -invariant.

P_j precisely corresponds to step j in Gibbs sampling: we refresh coordinate X_j from the full conditional given all the others.

Conditionals as kernels

Viewing $f_{X_j|X_{-j}} : \mathcal{X}_{1:(j-1)} \otimes \mathcal{X}_{(j+1):p} \rightarrow \mathcal{X}_j$ as a kernel, its defining property is:



Where we have defined $P_j : \mathcal{X}_{1:p} \rightarrow \mathcal{X}_{1:p}$.
By construction, P_j is f -invariant.

P_j precisely corresponds to step j in Gibbs sampling: we refresh coordinate X_j from the full conditional given all the others.

Invariance of Gibbs samplers

Proposition

The systematic scan Gibbs sampler and the random scan Gibbs sampler both possess f as invariant distribution.

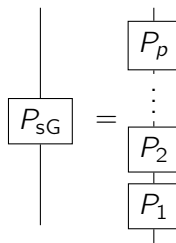
Proof. The random scan Gibbs sampler can be represented as the mixture

$$P_{\text{rG}} = \frac{1}{p} \sum_{j=1}^p P_j,$$

and this is similarly f -invariant by f -invariance of each P_j .

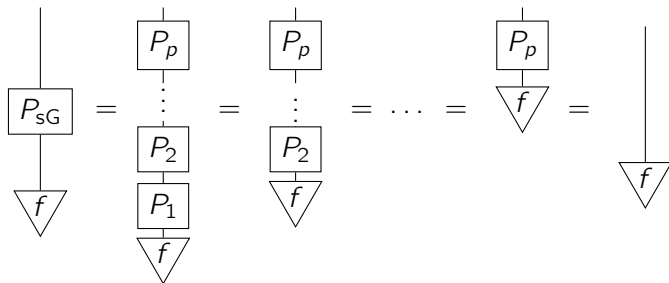
Systematic scan Gibbs

The systematic scan Gibbs sampler can be represented as:



Systematic scan Gibbs II

Since each P_j is f -invariant,



Ergodic theorem

Theorem (Ergodicity of the Gibbs Sampler)

If the Markov chain generated by the Gibbs sampler is irreducible and recurrent (which is e.g. the case when the positivity condition holds), then for any integrable function $\varphi : E \rightarrow \mathbb{R}$

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{t=1}^n \varphi(\mathbf{X}^{(t)}) \stackrel{a.s.}{=} \mathbb{E}_f(\varphi(\mathbf{X}))$$

for almost every starting value $\mathbf{X}^{(0)}$.

Thus we can approximate expectations $\mathbb{E}_f(\varphi(\mathbf{X}))$ by their empirical counterparts using *a single* Markov chain.

A Simple Example: bivariate Gaussian model

$$\begin{pmatrix} X_1 \\ X_2 \end{pmatrix} \sim N_2 \left(\begin{pmatrix} \mu_1 \\ \mu_2 \end{pmatrix}, \begin{pmatrix} \sigma_1^2 & \sigma_{12} \\ \sigma_{12} & \sigma_2^2 \end{pmatrix} \right)$$

- Associated marginal distributions

$$X_1 \sim N(\mu_1, \sigma_1^2),$$

$$X_2 \sim N(\mu_2, \sigma_2^2)$$

- Associated full conditionals

$$(X_1 | X_2 = x_2) \sim N(\mu_1 + \sigma_{12}/\sigma_2^2(x_2 - \mu_2), \sigma_1^2 - (\sigma_{12})^2/\sigma_2^2)$$

$$(X_2 | X_1 = x_1) \sim N(\mu_2 + \sigma_{12}/\sigma_1^2(x_1 - \mu_1), \sigma_2^2 - (\sigma_{12})^2/\sigma_1^2)$$

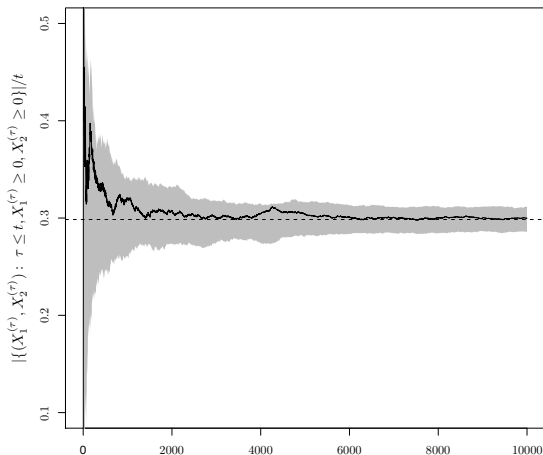
- Gibbs sampler consists of iterating for $t = 1, 2, \dots$

1. Draw $X_1^{(t)} \sim N(\mu_1 + \sigma_{12}/\sigma_2^2(X_2^{(t-1)} - \mu_2), \sigma_1^2 - (\sigma_{12})^2/\sigma_2^2)$.

2. Draw $X_2^{(t)} \sim N(\mu_2 + \sigma_{12}/\sigma_1^2(X_1^{(t)} - \mu_1), \sigma_2^2 - (\sigma_{12})^2/\sigma_1^2)$.

Results for Gibbs Samplers

Using the ergodic theorem we can estimate $\mathbb{P}(X_1 \geq 0, X_2 \geq 0)$ by the proportion of samples $(X_1^{(t)}, X_2^{(t)})$ with $X_1^{(t)} \geq 0$ and $X_2^{(t)} \geq 0$:



Reversibility of Metropolis–Hastings

Recall Metropolis–Hastings is a two-step algorithm: we make a *proposal* (possibly with auxiliary variables) followed by an *accept-reject* step.

Random walk Metropolis / Independence Sampler: we directly propose the next state which is accepted or rejected.

Exchange Algorithm: we propose the next state *along with an auxiliary variable*: both of these (along with the current state) are used in the accept-reject step.

In order to handle both algorithms (and many more!) we will make use of an approach based on involutions.³

³Andrieu, C., Lee, A., Livingstone, S. (2020). A general perspective on the Metropolis-Hastings kernel. <http://arxiv.org/abs/2012.14881>

Reversibility of Metropolis–Hastings

Recall Metropolis–Hastings is a two-step algorithm: we make a *proposal* (possibly with auxiliary variables) followed by an *accept-reject* step.

Random walk Metropolis / Independence Sampler: we directly propose the next state which is accepted or rejected.

Exchange Algorithm: we propose the next state *along with an auxiliary variable*: both of these (along with the current state) are used in the accept-reject step.

In order to handle both algorithms (and many more!) we will make use of an approach based on involutions.³

³Andrieu, C., Lee, A., Livingstone, S. (2020). A general perspective on the Metropolis-Hastings kernel. <http://arxiv.org/abs/2012.14881>

Reversibility of Metropolis–Hastings

Recall Metropolis–Hastings is a two-step algorithm: we make a *proposal* (possibly with auxiliary variables) followed by an *accept-reject* step.

Random walk Metropolis / Independence Sampler: we directly propose the next state which is accepted or rejected.

Exchange Algorithm: we propose the next state *along with an auxiliary variable*: both of these (along with the current state) are used in the accept-reject step.

In order to handle both algorithms (and many more!) we will make use of an approach based on involutions.³

³Andrieu, C., Lee, A., Livingstone, S. (2020). A general perspective on the Metropolis-Hastings kernel. <http://arxiv.org/abs/2012.14881>

Results for Metropolis–Hastings Algorithms

ALL (2020) framework

The key idea of ALL (2020) is that virtually any MH-type algorithm can be described using three building blocks:

An **augmented state space** $\mathcal{E}$: typically of the form $\mathcal{E} = \mathcal{X} \otimes \mathcal{Z}$,
and a measure μ on $\mathcal{E}$, typically of the form

$$d\mu(x, z) = f(dx)Q(x, dz).$$

An **involution** $\phi : \mathcal{E} \rightarrow \mathcal{E}$. This means that ϕ is its own inverse:

$$\phi \circ \phi = \text{id}_{\mathcal{E}}.$$

An **acceptance probability** $\alpha : \mathcal{E} \rightarrow [0, 1]$.

The core of the algorithm is then of the form

$$\Pi_{\text{MH}}(\cdot | \xi) = \alpha(\xi)\delta_{\phi(\xi)} + (1 - \alpha(\xi))\delta_{\xi}.$$

ALL (2020) framework

The key idea of ALL (2020) is that virtually any MH-type algorithm can be described using three building blocks:

An augmented state space $\mathcal{E}$: typically of the form $\mathcal{E} = \mathcal{X} \otimes \mathcal{Z}$, and a measure μ on $\mathcal{E}$, typically of the form

$$d\mu(x, z) = f(dx)Q(x, dz).$$

An involution $\phi : \mathcal{E} \rightarrow \mathcal{E}$. This means that ϕ is its own inverse:

$$\phi \circ \phi = \text{id}_{\mathcal{E}}.$$

An acceptance probability $\alpha : \mathcal{E} \rightarrow [0, 1]$.

The core of the algorithm is then of the form

$$\Pi_{\text{MH}}(\cdot | \xi) = \alpha(\xi) \delta_{\phi(\xi)} + (1 - \alpha(\xi)) \delta_{\xi}.$$

ALL (2020) framework

The key idea of ALL (2020) is that virtually any MH-type algorithm can be described using three building blocks:

An augmented state space $\mathcal{E}$: typically of the form $\mathcal{E} = \mathcal{X} \otimes \mathcal{Z}$, and a measure μ on $\mathcal{E}$, typically of the form

$$d\mu(x, z) = f(dx)Q(x, dz).$$

An involution $\phi : \mathcal{E} \rightarrow \mathcal{E}$. This means that ϕ is its own inverse:

$$\phi \circ \phi = \text{id}_{\mathcal{E}}.$$

An acceptance probability $\alpha : \mathcal{E} \rightarrow [0, 1]$.

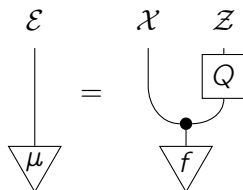
The core of the algorithm is then of the form

$$\Pi_{\text{MH}}(\cdot | \xi) = \alpha(\xi) \delta_{\phi(\xi)} + (1 - \alpha(\xi)) \delta_{\xi}.$$

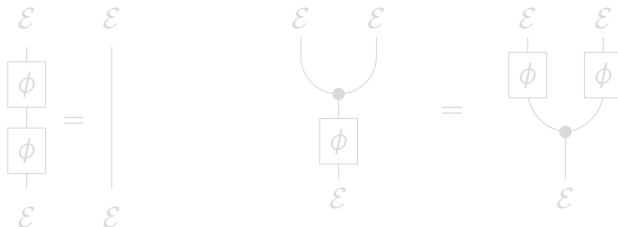
Results for Metropolis–Hastings Algorithms

String diagrammatically

$$d\mu(x, z) = f(dx)Q(x, dz) \text{ on } \mathcal{E} = \mathcal{X} \otimes \mathcal{Z}:$$



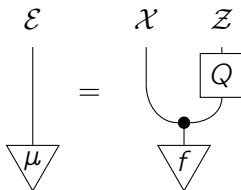
$\phi : \mathcal{E} \rightarrow \mathcal{E}$ is a deterministic involution:



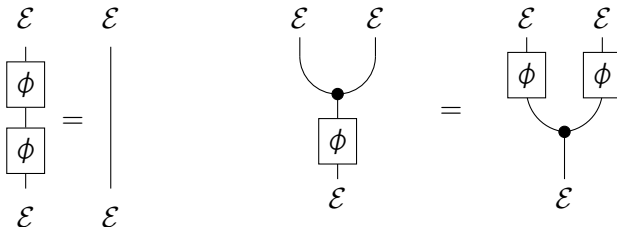
Results for Metropolis–Hastings Algorithms

String diagrammatically

$d\mu(x, z) = f(dx)Q(x, dz)$ on $\mathcal{E} = \mathcal{X} \otimes \mathcal{Z}$:



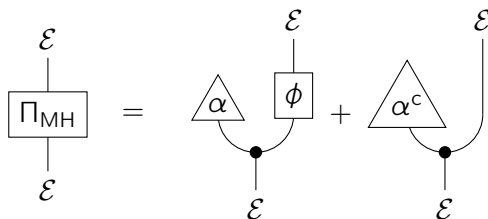
$\phi : \mathcal{E} \rightarrow \mathcal{E}$ is a deterministic involution:



Kernel Π_{MH} on augmented space

$\alpha : \mathcal{E} \rightarrow [0, 1]$ acceptance probability;

$$\Pi_{\text{MH}}(\cdot | \xi) = \alpha(\xi) \delta_{\phi(\xi)} + (1 - \alpha(\xi)) \delta_{\xi}.$$



Read upwards triangle as ‘do this with probability’ and define $\alpha^c := 1 - \alpha$.

Augmentation and involution

Have $\mu(d\xi) = f(x)Q(x, dz)$ on $\mathcal{E} = \mathcal{X} \otimes \mathcal{Z}$; involution $\phi : \mathcal{E} \rightarrow \mathcal{E}$ and acceptance probability $\alpha : \mathcal{E} \rightarrow [0, 1]$.

$$\Pi_{\text{MH}}(\cdot|\xi) = \alpha(\xi)\delta_{\phi(\xi)} + (1 - \alpha(\xi))\delta_{\xi}.$$

We now recover $P_{\text{MH}} : \mathcal{X} \rightarrow \mathcal{X}$.

Metropolis–Hastings P_{MH} via involutions

Suppose we are at $x \in \mathcal{X}$

- ① Sample $z \sim Q(x, dz)$ to get $\xi := (x, z)$.
- ② Sample $\xi' = (x', z') \sim \Pi_{\text{MH}}(\cdot|\xi)$: with probability $\alpha(\xi)$ return $\xi' = \phi(\xi)$ else $\xi' = \xi$
- ③ Discard z' and return x' .

Augmentation and involution

Have $\mu(d\xi) = f(x)Q(x, dz)$ on $\mathcal{E} = \mathcal{X} \otimes \mathcal{Z}$; involution $\phi : \mathcal{E} \rightarrow \mathcal{E}$ and acceptance probability $\alpha : \mathcal{E} \rightarrow [0, 1]$.

$$\Pi_{\text{MH}}(\cdot|\xi) = \alpha(\xi)\delta_{\phi(\xi)} + (1 - \alpha(\xi))\delta_{\xi}.$$

We now recover $P_{\text{MH}} : \mathcal{X} \rightarrow \mathcal{X}$.

Metropolis–Hastings P_{MH} via involutions

Suppose we are at $x \in \mathcal{X}$

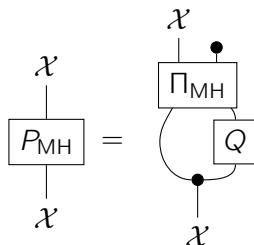
- ① Sample $z \sim Q(x, dz)$ to get $\xi := (x, z)$.
- ② Sample $\xi' = (x', z') \sim \Pi_{\text{MH}}(\cdot|\xi)$: with probability $\alpha(\xi)$ return $\xi' = \phi(\xi)$ else $\xi' = \xi$
- ③ Discard z' and return x' .

Augmentation and involution string diagrams

Metropolis–Hastings P_{MH} via involutions

Suppose we are at $x \in \mathcal{X}$

- ① Sample $z \sim Q(x, dz)$ to get $\xi := (x, z)$.
- ② Sample $\xi' = (x', z') \sim \Pi_{MH}$: with probability $\alpha(\xi)$ return $\xi' = \phi(\xi)$ else $\xi' = \xi$
- ③ Discard z' and return x' .



Example: Random Walk Metropolis

Suppose we are at $x \in \mathcal{X}$

- ① Sample $z \sim Q(x, dz)$ to get $\xi := (x, z)$.
- ② Sample $\xi' = (x', z') \sim \Pi_{\text{MH}}$: with probability $\alpha(\xi)$ return $\xi' = \phi(\xi)$ else $\xi' = \xi$
- ③ Discard z' and return x' .

Have Lebesgue density f on $\mathcal{X} = \mathbb{R}^d$ and take $\mathcal{Z} = \mathbb{R}^d$.

Define proposal Q , involution ϕ (check!), acceptance prob α :

$$Q(x, dz) = \mathcal{N}(dz; 0, \sigma^2 I_d), \quad \phi(x, z) = (x + z, -z),$$

$$\alpha(x, z) = 1 \wedge \frac{f(x + z)}{f(x)}.$$

Example: Independence Sampler

Suppose we are at $x \in \mathcal{X}$

- ① Sample $z \sim Q(x, dz)$ to get $\xi := (x, z)$.
- ② Sample $\xi' = (x', z') \sim \Pi_{MH}$: with probability $\alpha(\xi)$ return $\xi' = \phi(\xi)$ else $\xi' = \xi$
- ③ Discard z' and return x' .

Have density f on general space $\mathcal{X}$ and proposal measure μ on $\mathcal{X}$;
set $\mathcal{Z} = \mathcal{X}$.

Define

$$Q(x, dz) = \mu(dz), \quad \phi(x, z) = (z, x),$$

$$\alpha(x, z) = 1 \wedge \frac{f(z)\mu(x)}{f(x)\mu(z)}.$$

Example: Exchange Algorithm

Suppose we are at $x \in \mathcal{X}$

- ① Sample $z \sim Q(x, dz)$ to get $\xi := (x, z)$.
- ② Sample $\xi' = (x', z') \sim \Pi_{\text{MH}}$: with probability $\alpha(\xi)$ return $\xi' = \phi(\xi)$ else $\xi' = \xi$
- ③ Discard z' and return x' .

Have density $f = f^{\text{prior}}(\boldsymbol{\theta})f(\mathbf{y}^{\text{obs}}|\boldsymbol{\theta})/Z(\boldsymbol{\theta})$ on $\mathcal{X}$ for data $\mathbf{y}^{\text{obs}} \in \mathcal{Y}$.
Set $\mathcal{Z} = \mathcal{Y} \otimes \mathcal{X}$.

$$Q(\boldsymbol{\theta}; d\mathbf{y}, d\boldsymbol{\vartheta}) = q(d\boldsymbol{\vartheta}|\boldsymbol{\theta}) \cdot \frac{f(d\mathbf{y}|\boldsymbol{\vartheta})}{Z(\boldsymbol{\vartheta})},$$

$$\phi(\boldsymbol{\theta}, \mathbf{y}, \boldsymbol{\vartheta}) = (\boldsymbol{\vartheta}, \mathbf{y}, \boldsymbol{\theta}).$$

Acceptance prob $\alpha(\boldsymbol{\theta}, \mathbf{y}, \boldsymbol{\vartheta})$ as before.

Proving reversibility

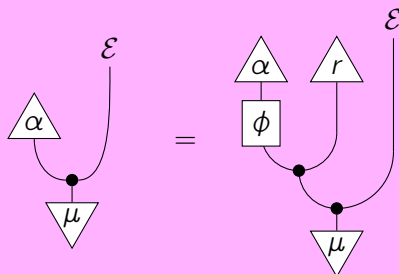
The goal is to show that this general procedure is reversible in *general* (for any proper choices of μ, Q, α).

- 1 We first show that the kernel $\Pi_{\text{MH}} : \mathcal{E} \rightarrow \mathcal{E}$ is μ -reversible.
- 2 We show that the augmentation - discard technique preserves reversibility.
- 3 We conclude that P_{MH} is f -reversible.

Step 1: reversibility of Π_{MH} ⁴

Theorem

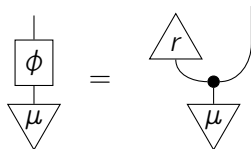
Given involution $\phi : \mathcal{E} \rightarrow \mathcal{E}$, probability measure μ on $\mathcal{E}$, let $r := \frac{d(\phi \circ \mu)}{d\mu}$ be the Radon–Nikodym derivative. Then Π_{MH} is μ -reversible if and only if



⁴Cornish, Wang (2026). A categorical account of the Metropolis–Hastings algorithm. Proceedings of the Forty-First Annual ACM/IEEE Symposium on Logic in Computer Science (LICS 2026)

Radon–Nikodym condition

The Radon–Nikodym condition means that



that is, when all densities exist

$$r(\xi) = \frac{\mu(\phi(\xi))}{\mu(\xi)} \cdot |J(\xi)|,$$

where $J = \det \left(\frac{d\phi}{d\xi}(\xi) \right)$ is the Jacobian (often just equal to 1).

Condition on α

The necessary and sufficient condition can be expressed as

$$\alpha(\xi) = r(\xi) \cdot \alpha \circ \phi(\xi), \quad \text{for } \mu\text{-a.e. } \xi.$$

A sufficient condition for this is:

$$\alpha(\xi) = a(r(\xi)),$$

where $a : [0, \infty) \rightarrow [0, 1]$ is a *balancing function*:

$$a(t) = \begin{cases} 0 & t = 0; \\ t \cdot a(1/t) & t > 0. \end{cases}$$

Balancing function examples

Examples: Metropolis–Hastings and Barker's:

$$a(t) = 1 \wedge t, \quad a(t) = \frac{t}{1+t}.$$

It can be shown that the Metropolis choice is optimal in the sense that it maximises acceptance rates, but Barkers choice is at most a constant factor 2 worse (and there are good reasons to use it!)⁵.

⁵Gonçalves, F. B., Latuszynski, K., Roberts, G. O. (2017). Barker's algorithm for Bayesian inference with intractable likelihoods. *Braz. J. Probab. Stat.*, 31(4), 732-745.

Balancing function examples

Examples: Metropolis–Hastings and Barker’s:

$$a(t) = 1 \wedge t, \quad a(t) = \frac{t}{1+t}.$$

It can be shown that the Metropolis choice is optimal in the sense that it maximises acceptance rates, but Barkers choice is at most a constant factor 2 worse (and there are good reasons to use it!)⁵.

⁵Gonçalves, F. B., Latuszynski, K., Roberts, G. O. (2017). Barker’s algorithm for Bayesian inference with intractable likelihoods. *Braz. J. Probab. Stat.*, 31(4), 732-745.

Random Walk example

Recall $\mu(dx, dz) = f(dx)\mathcal{N}(dz; 0, \sigma I_d)$ and $\phi(x, z) = (x + z, -z)$.

$$r(x, z) = \frac{d\mu^\phi}{d\mu}(x, z) = \frac{f(x + z)\mathcal{N}(-z; 0, \sigma I_d)}{f(x)\mathcal{N}(z; 0, \sigma I_d)} = \frac{f(x + z)}{f(x)}.$$

Choosing

$$a(t) = 1 \wedge t$$

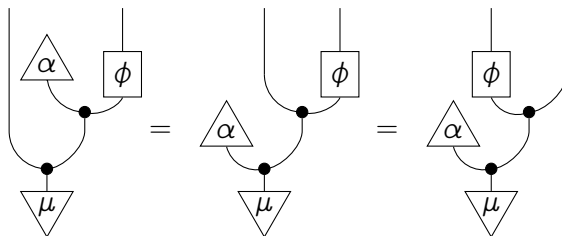
and setting $\alpha(x, z) = a(r(x, z))$ recovers the RWM acceptance probability,

$$1 \wedge \frac{f(x + z)}{f(x)}.$$

Results for Metropolis–Hastings Algorithms

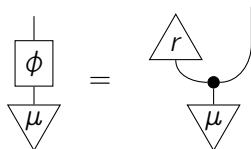
Proof

The desired reversibility is the following:

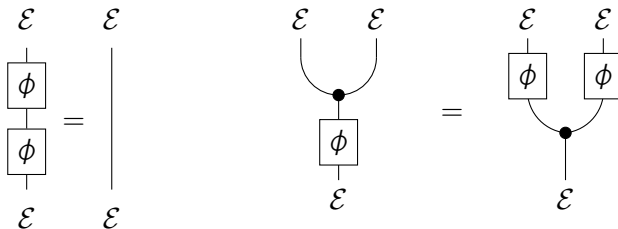


Results for Metropolis–Hastings Algorithms

Tools at our disposal



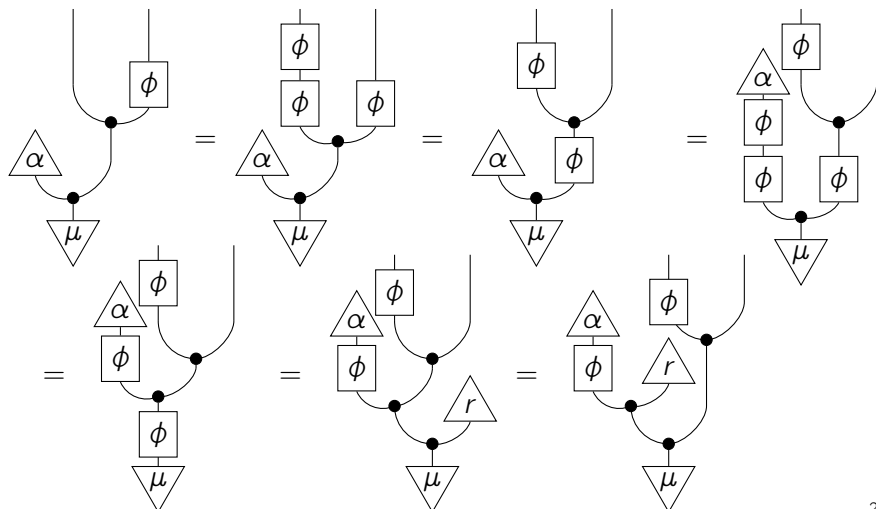
$\phi : \mathcal{E} \rightarrow \mathcal{E}$ is a deterministic involution:



We will make use of these repeatedly to prove the theorem.

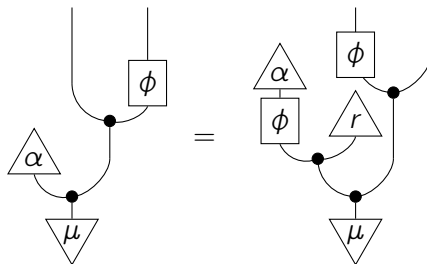
Results for Metropolis–Hastings Algorithms

Proof of theorem



Proof of theorem II

We have established that

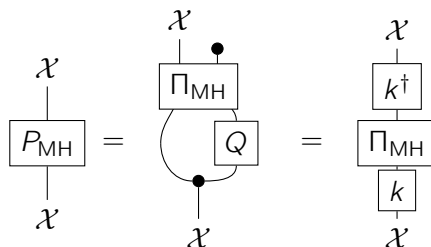


It then follows that if the condition holds, we immediately have reversibility.

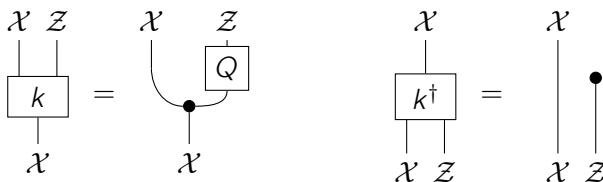
Conversely if we are reversible, we can simply delete the left output string and deduce the condition holds. $\square$

Results for Metropolis–Hastings Algorithms

Auxiliary variables and discarding



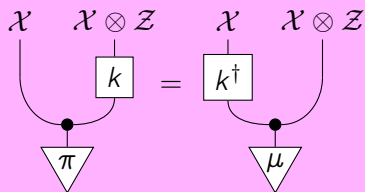
We have proven that the inner Π_{MH} is μ -reversible.



Auxiliary variables and discarding

Bayesian inverses

It can be shown (simple exercise!) that $k : \mathcal{X} \rightarrow \mathcal{X} \otimes \mathcal{Z}$ and $k^\dagger : \mathcal{X} \otimes \mathcal{Z} \rightarrow \mathcal{X}$ are actually *Bayesian inverses*, meaning that:



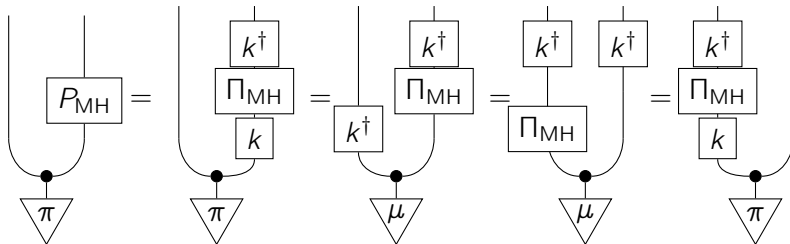
(Recall $\mu = k \circ \pi$.)

Completing the proof

Theorem: Reversibility of generalised MH

The kernel $P_{\text{MH}} : \mathcal{X} \rightarrow \mathcal{X}$ is π -reversible.

Proof.



□

Theoretical properties of Metropolis–Hastings

- The Markov chain $(\mathbf{X}^{(0)}, \mathbf{X}^{(1)}, \dots)$ is (strongly) *irreducible* if $q(\mathbf{x}|\mathbf{x}^{(t-1)}) > 0$ for all $\mathbf{x}, \mathbf{x}^{(t-1)} \in \text{supp}(f)$.
(See, e.g., Roberts & Tweedie, 1996, for weaker conditions.)
- Such a chain is *recurrent* if it is irreducible.
(See e.g., Tierney, 1994.)
- The chain is *aperiodic* if there is positive probability that the chain remains in the current state, i.e. $\mathbb{P}(\mathbf{X}^{(t)} = \mathbf{X}^{(t-1)}) > 0$ (for a suitable group of “current states”).

Virtually all MCMC chains used in practice satisfy these conditions.

Theorem (A Simple Ergodic Theorem)

If $(X_i)_{i \in \mathbb{N}}$ is an f -irreducible, f -invariant, recurrent $\mathbb{R}^d$ -valued Markov chain then the following strong law of large numbers holds for any integrable function $\varphi : \mathbb{R}^d \rightarrow \mathbb{R}$:

$$\lim_{t \rightarrow \infty} \frac{1}{t} \sum_{i=1}^t \varphi(X_i) \stackrel{\text{a.s.}}{=} \int \varphi(x) f(x) dx.$$

for almost every starting value x .

Theorem (A Central Limit Theorem)

Under technical regularity conditions the following CLT holds for a recurrent, f -invariant Markov chain, and a function $\varphi : E \rightarrow \mathbb{R}$ which has at least two finite moments:

$$\lim_{t \rightarrow \infty} \sqrt{t} \left[\frac{1}{t} \sum_{i=1}^t \varphi(X_i) - \int \varphi(x) f(x) dx \right] \stackrel{\mathcal{D}}{=} N(0, \sigma^2(\varphi)),$$

$$\sigma^2(\varphi) = \mathbb{E} \left[(f(X_1) - \bar{\varphi})^2 \right] + 2 \sum_{k=2}^{\infty} \mathbb{E} \left[(\varphi(X_1) - \bar{\varphi})(\varphi(X_k) - \bar{\varphi}) \right],$$

where $\bar{\varphi} = \int \varphi(x) f(x) dx$.

Optimal Scaling

Much effort has gone into determining optimal scaling rules:

Diffusion Limits Under strong assumptions:

$$\lim_{p \rightarrow \infty} \frac{X_1^{(\lfloor tp \rfloor)}}{\sqrt{p}} \xrightarrow{d} \text{Diffusion}$$

where p is *dimension* and the *speed* of the diffusion depends upon proposal scale.

ESJD Seek to maximise:

$$\int f(x) K(x, y; \theta) (y - x)^2 dx dy$$

Rule of Thumb Optimal RWM Scaling depends upon dimension:

$p = 1$ Acceptance rate of around 0.44.

$p \geq 5$ Acceptance rate of around 0.234.

Part 4— Section 12

Convergence Diagnostics

Motivation: The Need for Convergence Diagnostics

The need for convergence diagnostics

- Theory guarantees (under certain conditions) the convergence of the Markov chain $\mathbf{X}^{(t)}$ to the desired distribution.
- This does not imply that a *finite* sample from such a chain yields a good approximation to the target distribution.
- Validity of the approximation must be confirmed in practice.
- Convergence diagnostics help answering this question.
- Convergence diagnostics are *not* perfect and should be treated with a good amount of scepticism.

Motivation: The Need for Convergence Diagnostics

Different diagnostic tasks

Convergence to the target distribution Does $\mathbf{X}^{(t)}$ yield a sample from the target distribution?

- Has reached $(\mathbf{X}^{(t)})_t$ a stationary regime?
- Does $(\mathbf{X}^{(t)})_t$ cover the support of the target distribution?

Convergence of averages Is $\sum_{t=1}^T \varphi(\mathbf{X}^{(t)})/T \approx \mathbb{E}_f(\varphi(\mathbf{X}))$?

Comparison to i.i.d. sampling How much information is contained in the sample from the Markov chain compared to an i.i.d. sample?

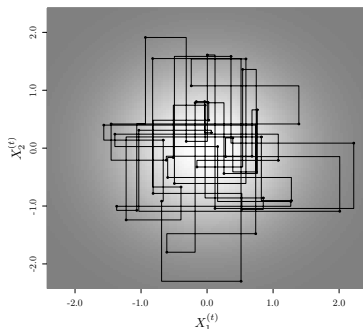
Motivation: The Need for Convergence Diagnostics

Pathological example 1: potentially slowly mixing

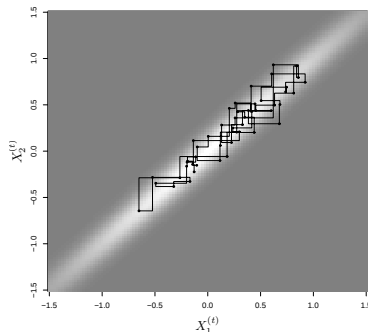
Gibbs sampler from a bivariate Gaussian with correlation

$$\rho(X_1, X_2)$$

$$\rho(X_1, X_2) = 0.3$$



$$\rho(X_1, X_2) = 0.99$$

For correlations $\rho(X_1, X_2)$ close to ± 1 the chain mixes poorly.

Motivation: The Need for Convergence Diagnostics

Pathological example 2: no central limit theorem

The following MCMC algorithm has the $\text{Beta}(\alpha, 1)$ distribution as stationary distribution:

Starting with any $X^{(0)}$ iterate for $t = 1, 2, \dots$

1. With probability $1 - X^{(t-1)}$, set $X^{(t)} = X^{(t-1)}$.
2. Otherwise draw $X^{(t)} \sim \text{Beta}(\alpha + 1, 1)$.

Markov chain converges very slowly (no central limit theorem applies).

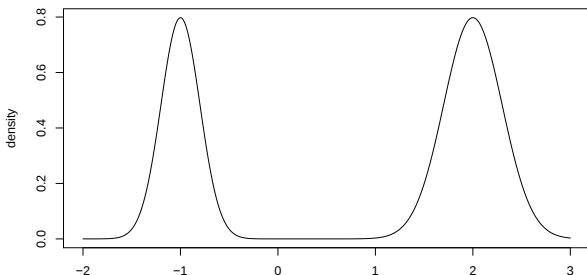
Motivation: The Need for Convergence Diagnostics

Pathological example 3: nearly reducible chain

Metropolis–Hastings sample from a mixture of two well-separated Gaussians, i.e. the target is

$$f(x) = 0.4 \cdot \phi_{(-1, 0.2^2)}(x) + 0.6 \cdot \phi_{(2, 0.3^2)}(x).$$

If the variance of the proposal is too small, the chain cannot move from one population to the other.



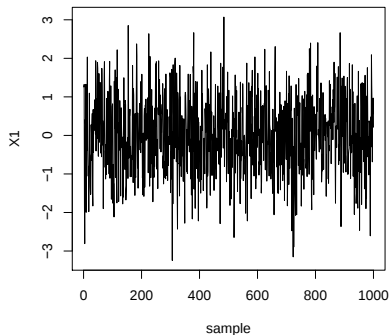
Basic plots

- Plot the sample paths $(X_j^{(t)})_t$.
should be oscillating very fast and show very little structure.
- Plot the cumulative averages $(\sum_{\tau=1}^t \varphi(X_j^{(\tau)})/t)_t$.
should be converging to a value.
- Only very obvious problems visible in these plots.
- Difficult to assess multivariate distributions from univariate projections.

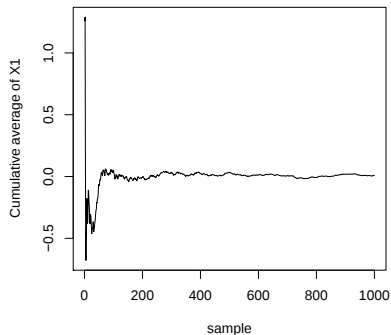
Elementary Techniques for Assessing Convergence

Plots for pathological example 1 ($\rho(X_1, X_2) = 0.3$)

Sample paths



Cumulative averages

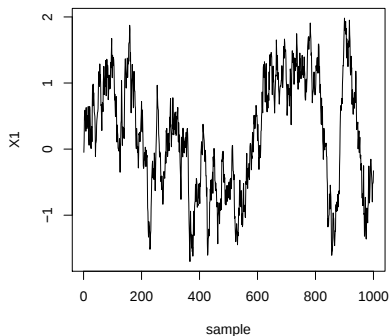


Looks OK.

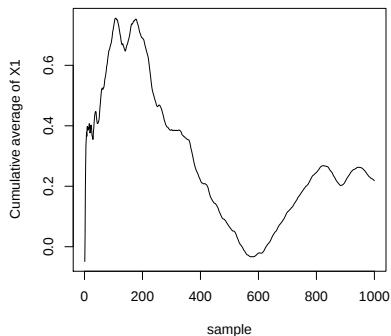
Elementary Techniques for Assessing Convergence

Plots for pathological example 1 ($\rho(X_1, X_2) = 0.99$)

Sample paths



Cumulative averages

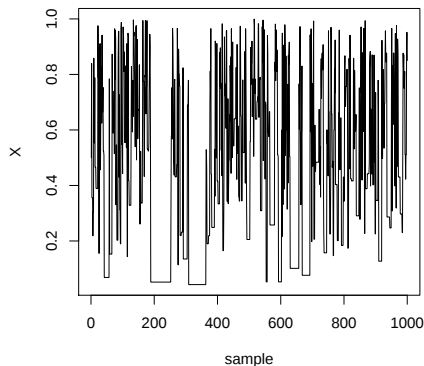


Slow mixing speed can be detected.

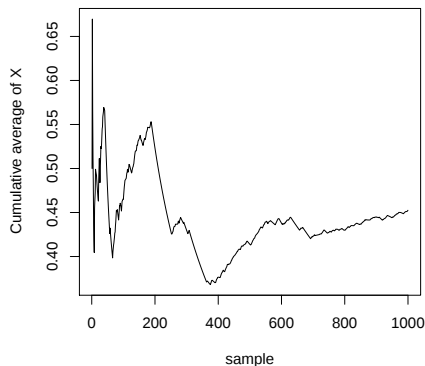
Elementary Techniques for Assessing Convergence

Plots for pathological example 2

Sample paths



Cumulative averages

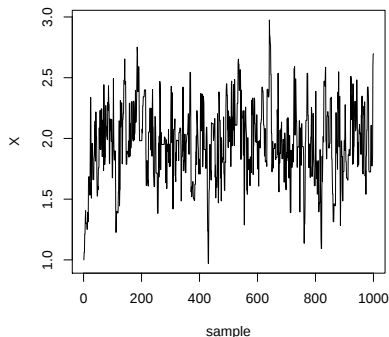


Slow convergence of the mean can be detected.

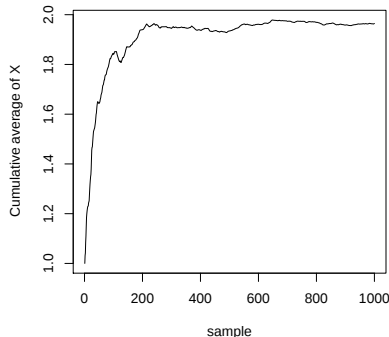
Elementary Techniques for Assessing Convergence

Plots for pathological example 3

Sample paths



Cumulative averages



We *cannot* detect that the sample only covers one part of the distribution.

(“you’ve only seen where you’ve been”)

Further Convergence Diagnostics

Comparing multiple chains

- Compare $L > 1$ chains $(\mathbf{X}^{(1,t)})_t, \dots, (\mathbf{X}^{(L,t)})_t$.
- Initialised using overdispersed values $\mathbf{X}^{(1,0)}, \dots, \mathbf{X}^{(L,0)}$.
- Idea: Variance and range of each chain $(\mathbf{X}^{(l,t)})_t$ should equal the range and variance of all chains pooled together.
- Compare basic plots for the different chains.
- Quantitative measure:
 - Compute distance $\delta_\alpha^{(l)}$ between α and $(1 - \alpha)$ quantile of $(X_k^{(l,t)})_t$.
 - Compute distance $\delta_\alpha^{(\cdot)}$ between α and $(1 - \alpha)$ quantile of the pooled data.
 - The ratio $\hat{S}_\alpha^{\text{interval}} = \frac{\sum_{l=1}^L \delta_\alpha^{(l)} / L}{\delta_\alpha^{(\cdot)}}$ should be around 1.
- Alternative: compare variance within each chain with the pooled variance estimate.
- Choosing suitable initial values $\mathbf{X}^{(1,0)}, \dots, \mathbf{X}^{(L,0)}$ difficult.

oooooooooooooooo
ooooo

ooooooooo
oooooooooooooooooooooooooooo
ooo
o

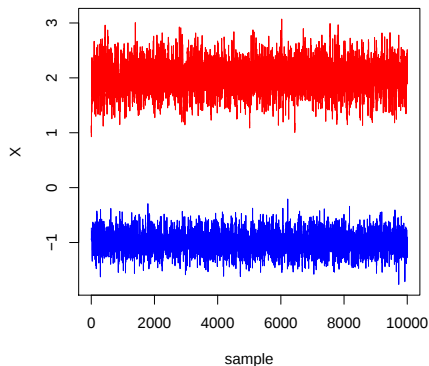
ooooo
ooooo
●oooooooo

oooooo

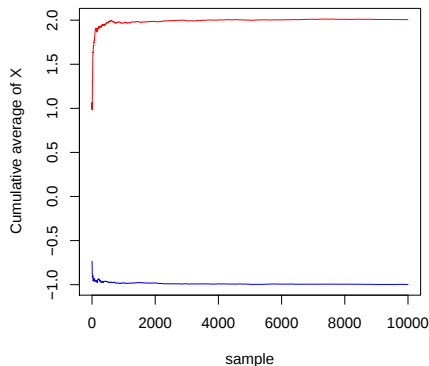
Further Convergence Diagnostics

Comparing multiple chains plots for pathological example 3

Sample paths



Cumulative averages



$\hat{\Sigma}_{\alpha}^{\text{interval}} = 0.2703 \ll 1$; we can detect that the sample only covers one part of the distribution.

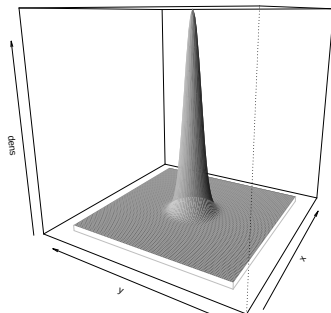
Further Convergence Diagnostics

Comparing multiple chains: A warning

- Consider the **Witch's hat** distribution:

$$f(x_1, x_2) \propto \begin{cases} (1 - \delta)\phi_{(\boldsymbol{\mu}, \sigma^2 \mathbb{I})}(x_1, x_2) + \delta & \text{if } x_1, x_2 \in (0, 1) \\ 0 & \text{otherwise.} \end{cases}$$

- Assume we want to estimate $\mathbb{P}(0.49 < X_1, X_2 \leq 0.51)$ for $\delta = 10^{-3}$, $\boldsymbol{\mu} = (0.5, 0.5)'$, and $\sigma = 10^{-5}$.



Further Convergence Diagnostics

Comparing multiple chains: A warning (II)

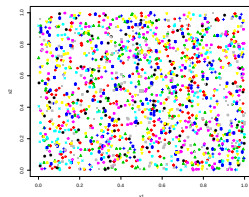
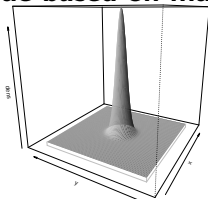
- We can use a Gibbs sampler. Conditional distribution:

$$f(x_1|x_2) \propto \begin{cases} (1 - \delta)\phi_{(\mu, \sigma^2 \cdot \mathbb{I})}(x_1, x_2) + \delta & \text{for } x_1 \in (0, 1) \\ 0 & \text{otherwise.} \end{cases}$$

- But on average only 0.04% of the sampled values lie in $(0.49, 0.51) \times (0.49, 0.51)$ yielding an estimate of:

$$\hat{\mathbb{P}}(0.49 < X_1, X_2 \leq 0.51) = 0.0004.$$

- It is close to impossible to detect this problem with any technique based on multiple initialisations.**



Further Convergence Diagnostics

Riemann sums and control variates

- Consider order statistic $X^{[1]} \leq \dots \leq X^{[T]}$.
- Provided $(X^{[t]})_t = 1 \dots, T$ covers the support of the target, the Riemann sum

$$\sum_{t=2}^T (X^{[t]} - X^{[t-1]}) f(X^{[t]})$$

converges to

$$\int f(x) dx = 1.$$

- Thus if $\sum_{t=2}^T (X^{[t]} - X^{[t-1]}) f(X^{[t]}) \ll 1$, the Markov chain has failed to explore all the support of the target.
- Requires that target density f is available inclusive of normalisation constants.
- Only effective in 1D.

Further Convergence Diagnostics

Riemann sums for pathological example 3

For the chain stuck in the population with mean 2 we obtain

$$\sum_{t=2}^T (X^{[t]} - X^{[t-1]}) f(X^{[t]}) = 0.598 \ll 1,$$

so we can detect that we have not explored the whole distribution.

Further Convergence Diagnostics

Effective sample size

- MCMC algorithms yield a positively correlated sample $(\mathbf{X}^{(t)})_{t=1,\dots,T}$.
- How much less useful is an MCMC sample of size T than an i.i.d. sample of size T ?
- Approximate $(\varphi(\mathbf{X}^{(t)}))_{t=1,\dots,T}$ by an $AR(1)$ process, i.e.:

$$\rho(\varphi(\mathbf{X}^{(t)}), \varphi(\mathbf{X}^{(t+\tau)})) = \rho^{|\tau|}.$$

- Variance of the estimator is

$$\text{Var} \left(\frac{1}{T} \sum_{t=1}^T \varphi(\mathbf{X}^{(t)}) \right) \approx \frac{1+\rho}{1-\rho} \cdot \frac{1}{T} \text{Var} \left(\varphi(\mathbf{X}^{(t)}) \right)$$

- Same variance as an i.i.d. sample of the size $T \cdot \frac{1-\rho}{1+\rho}$.
- Thus define $T \cdot \frac{1-\rho}{1+\rho}$ as *effective sample size*.

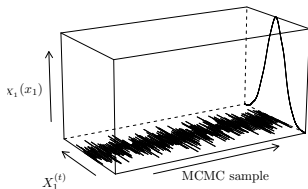
Further Convergence Diagnostics

Effective sample for pathological example 1

Rapidly mixing chain

$$(\rho(X_1, X_2) = 0.3)$$

10,000 samples



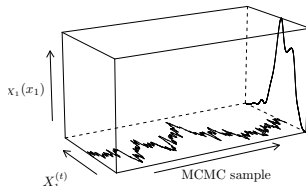
$$\rho(X_1^{(t-1)}, X_1^{(t)}) = 0.078$$

ESS for estimating $\mathbb{E}_f(X_1)$ is
8,547.

Slowly mixing chain

$$(\rho(X_1, X_2) = 0.99)$$

10,000 samples



$$\rho(X_1^{(t-1)}, X_1^{(t)}) = 0.979$$

ESS for estimating $\mathbb{E}_f(X_1)$ is
105.

What Else Can We Do?

- ① More sophisticated convergence diagnostics:
 - Geweke's method based on spectral analysis
 - Raftery's binary-chain method
 - ⋮
- ② Theoretical Computations
 - Convergence rates
 - Mixing times
 - Confidence intervals
- ③ Perfect Simulation
 - Processes with "ordered transitions".
 - Certain spatial processes.

Part 4— Section 13

Practical Considerations

Where do we start?

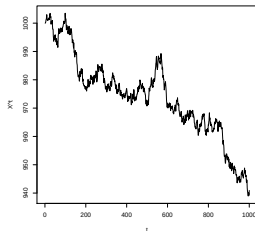
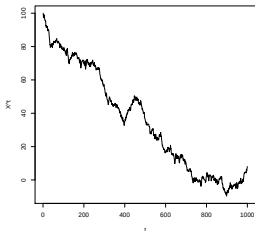
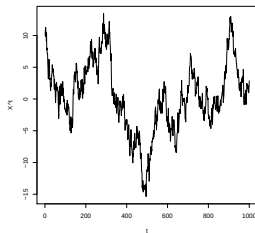
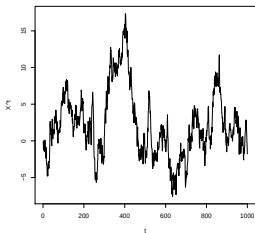
RWM Traces.

Target:

$$f(x) = e^{-|x|/5}/10$$

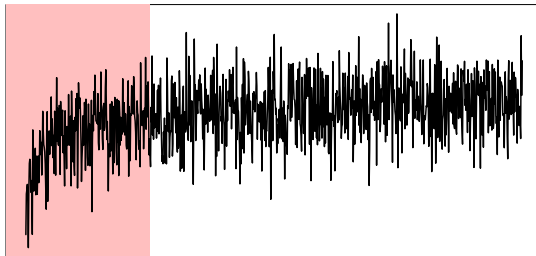
Starting values:

- $X^{(1)} = 0$
- $X^{(1)} = 10$
- $X^{(1)} = 100$
- $X^{(1)} = 1,000$



Practical considerations: Burn-in period

- Theory (ergodic theorems) allows for the use of the entire chain $(\mathbf{X}^{(0)}, \mathbf{X}^{(1)}, \dots)$.
- However distribution of $(\mathbf{X}^{(t)})$ for small t might still be far from the stationary distribution f .
- Can be beneficial to discard the first iterations $\mathbf{X}^{(t)}$, $t = 1, \dots, T_0$ (*burn-in period*).
- Optimal T_0 depends on mixing properties of the chain.



Practical considerations: Multiple Starts?

- Should we use “multiple overdispersed initialisations”?
- Advantages:
 - Exploring different parts of the space.
 - May be useful for assessing convergence.
 - Trivial to parallelize.
- Disadvantages:
 - We need to specify many starting values.
 - What does overdispersed mean, anyway?
 - Every chain needs to reach stationarity.
 - Multiple burn-in periods may be expensive.

```

oooooooooooooooo
ooooo

```

```

oooooooo
oooooooooooooooooooooooooooo
ooo
oo
o

```

```

oooo
oooo
oooooooo

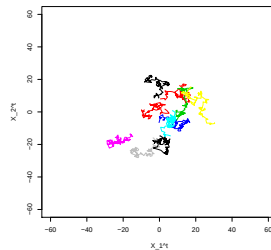
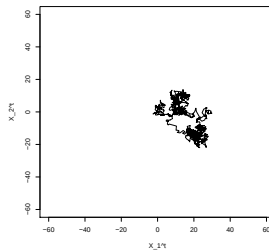
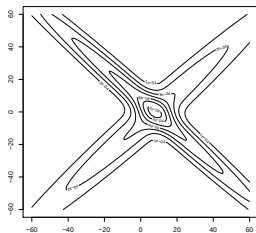
```

```

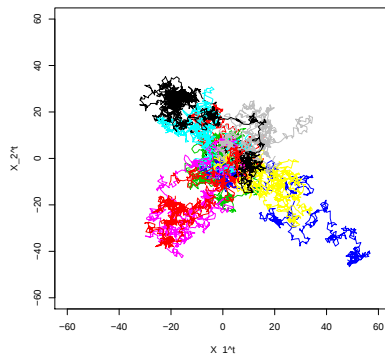
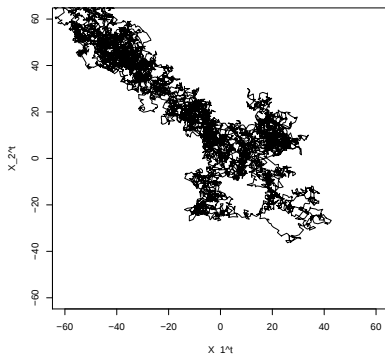
●●oooo

```

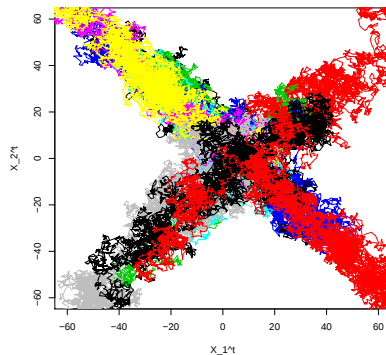
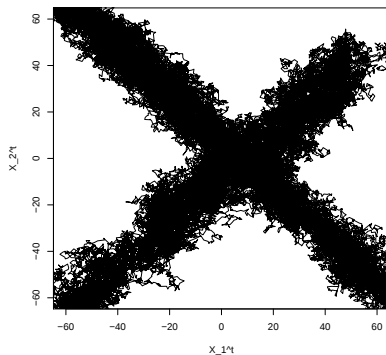
Reducing Correlation

One Chain vs. Many: 1000 or 10×100 

Reducing Correlation

One Chain vs. Many: 10,000 or 10×1000 

Reducing Correlation

One Chain vs. Many: 100,000 or $10 \times 10,000$ 

Practical considerations: Thinning (1)

- MCMC methods typically yield positively correlated chain: $\rho(\mathbf{X}^{(t)}, \mathbf{X}^{(t+\tau)})$ large for small τ .
- Idea: keeping only every m -th value: $(\mathbf{Y}^{(t)})_{t=1, \dots, \lfloor T/m \rfloor}$ with $\mathbf{Y}^{(t)} = \mathbf{X}^{(m \cdot t)}$ instead of $(\mathbf{X}^{(t)})_{t=1, \dots, T}$ (*thinning*).
- $(\mathbf{Y}^{(t)})_t$ exhibits less autocorrelation than $(\mathbf{X}^{(t)})_t$, i.e.

$$\rho(\mathbf{Y}^{(t)}, \mathbf{Y}^{(t+\tau)}) = \rho(\mathbf{X}^{(t)}, \mathbf{X}^{(t+m \cdot \tau)}) < \rho(\mathbf{X}^{(t)}, \mathbf{X}^{(t+\tau)}),$$

if the correlation $\rho(\mathbf{X}^{(t)}, \mathbf{X}^{(t+\tau)})$ decreases monotonically in τ .

- Price: length of $(\mathbf{Y}^{(t)})_{t=1, \dots, \lfloor T/m \rfloor}$ is only $(1/m)$ -th of the length of $(\mathbf{X}^{(t)})_{t=1, \dots, T}$.

Practical considerations: Thinning (2)

- If $\mathbf{X}^{(t)} \sim f$ and corresponding variances exist,

$$\text{Var} \left(\frac{1}{T} \sum_{t=1}^T \varphi(\mathbf{X}^{(t)}) \right) \leq \text{Var} \left(\frac{1}{\lfloor T/m \rfloor} \sum_{t=1}^{\lfloor T/m \rfloor} \varphi(\mathbf{Y}^{(t)}) \right),$$

i.e. thinning cannot be justified when objective is estimating $\mathbb{E}_f(\varphi(\mathbf{X}))$.

- Thinning can be a useful concept
 - if computer has insufficient memory.
 - for convergence diagnostics: $(\mathbf{Y}^{(t)})_{t=1, \dots, \lfloor T/m \rfloor}$ is closer to an i.i.d. sample than $(\mathbf{X}^{(t)})_{t=1, \dots, T}$.

Part 5

Alternative approaches

Part 5— Section 14

Augmentation

Augmentation

- “Making the *space* bigger to make the problem easier.”
- Already saw this when we moved from $P_{\text{MH}} : \mathcal{X} \rightarrow \mathcal{X}$ to $\Pi_{\text{MH}} : \mathcal{E} \rightarrow \mathcal{E}$.
- To target a distribution $f_{\mathbf{X}}(\mathbf{x})$:
 - Construct some $f_{\mathbf{X},\mathbf{Z}}(\mathbf{x}, \mathbf{z})$ on $\mathcal{X} \otimes \mathcal{Z}$
 - such that

$$f_{\mathbf{X}}(\mathbf{x}) = \int_{\mathcal{Z}} f_{\mathbf{X},\mathbf{Z}}(\mathbf{x}, \mathbf{z}) d\mathbf{z}$$

- and $f_{\mathbf{X},\mathbf{Z}}$ is easy to sample from (when $f_{\mathbf{X}}$ is not).
- Versatile technique with many applications.

A Generic Augmentation Scheme

- Given any density $f(\mathbf{x})$, define

$$f(\mathbf{x}, u) := f(\mathbf{x}) \cdot f_{U|\mathbf{X}}(u|\mathbf{x})$$

- with

$$f_{U|\mathbf{X}}(u|\mathbf{x}) = \frac{1}{f(\mathbf{x})} \mathbb{I}_{[0, f(\mathbf{x})]}(u)$$

- Then

$$f(\mathbf{x}, u) = \mathbb{I}_{[0, f(\mathbf{x})]}(u).$$

Rejection Sampling Revisited

Proposition (Rejection Sampling Equivalence)

- Given $f(\mathbf{x})$, define

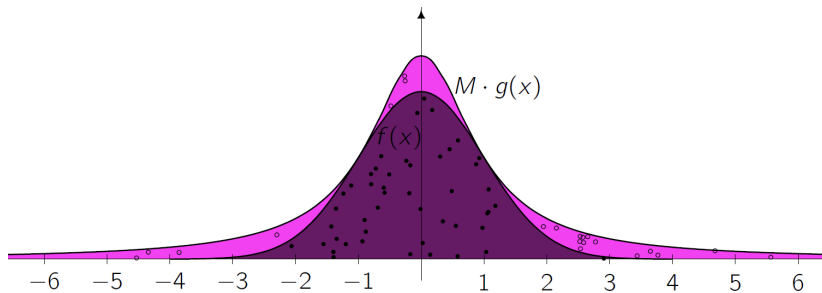
$$f(\mathbf{x}, u) = \mathbb{I}_{[0, f(\mathbf{x})]}(u).$$

- Given proposal $g(\mathbf{x})$ and $M \geq \sup_{\mathbf{x}} f(\mathbf{x})/g(\mathbf{x})$, define

$$g(\mathbf{x}, u) = \frac{1}{M} \mathbb{I}_{[0, M \cdot g(\mathbf{x})]}(u).$$

- Let $w(\mathbf{x}, u) = f(\mathbf{x}, u)/g(\mathbf{x}, u)$
- The associated self-normalised importance sampling estimator of $\mathbb{E}_f[\varphi(\mathbf{X})]$ is the rejection sampling estimator.

Slice sampling



Sample uniformly and weight...

Slice Sampling

- Rejection sampling can be viewed as importance sampling with an extended target distribution. . .
- so can we apply other algorithms to that extended distribution?

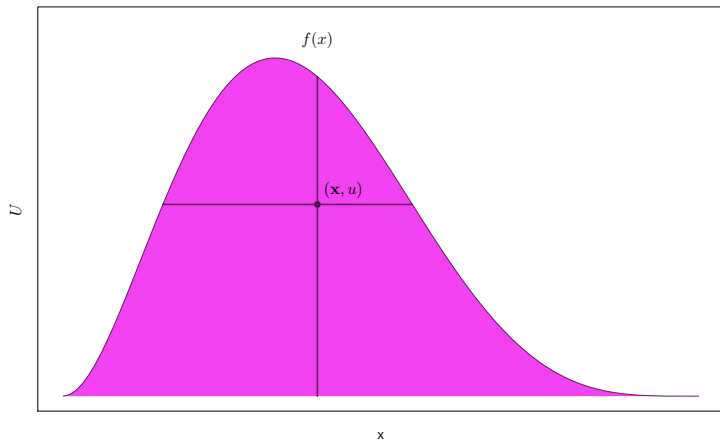
Algorithm: The Slice Sampler

Starting with $(\mathbf{X}^{(0)}, U^{(0)})$ iterate for $t = 1, 2, \dots$

- 1 Draw $\mathbf{X}^{(t)} \sim f_{\mathbf{X}|U}(\cdot | U^{(t-1)})$.
- 2 Draw $U^{(t)} \sim f_{U|\mathbf{X}}(\cdot | \mathbf{X}^{(t)})$.

Slice sampling

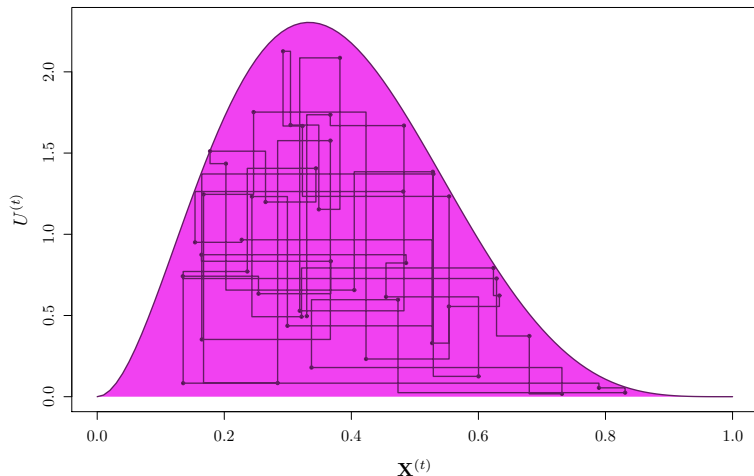
An Illustration of the Conditional Distributions



Slice sampling

A Slice-Sampler Trajectory

Example: Sampling from a **Beta(3,5)** distribution



How Practical Is This?

- Sampling $U \sim U[0, f(\mathbf{X})]$ is *easy*.
- Sampling $\mathbf{X} \sim U(L(U))$ where

$$L(u) := \{\mathbf{x} : f(\mathbf{x}) \geq u\}$$

can be easy. . .

- but it might not be.
- Consider the bivariate density:

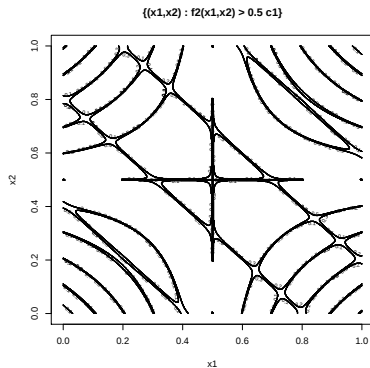
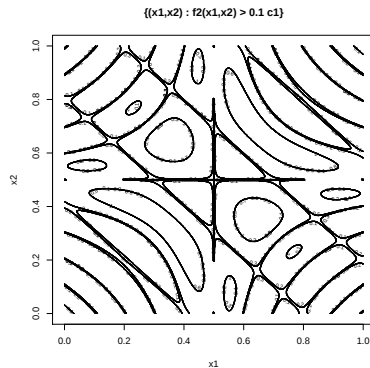
$$f_2(x_1, x_2) = c_1 \cdot \sin^2(x_1 \cdot x_2) \cdot \cos^2(x_1 + x_2) \cdot \exp\left(-\frac{1}{2}(|x_1| + |x_2|)\right).$$

Slice sampling

The Trouble with Slice Sampling

Level sets of:

$$f_2(x_1, x_2) = c_1 \cdot \sin^2(x_1 \cdot x_2) \cdot \cos^2(x_1 + x_2) \cdot \exp\left(-\frac{1}{2}(|x_1| + |x_2|)\right).$$



Here we could use rejection.

Algorithm: The Co-ordinate-wise Slice Sampler

Starting with $(X_1^{(0)}, \dots, X_p^{(0)}, U^{(0)})$ iterate for $t = 1, 2, \dots$

1. Draw $X_1^{(t)} \sim f_{X_1|X_{-1}, U}(\cdot | X_{-1}^{(t-1)}, U^{(t-1)})$.
2. Draw $X_2^{(t)} \sim f_{X_2|X_{-2}, U}(\cdot | X_1^{(t)}, X_3^{(t-1)}, \dots, X_p^{(t-1)}, U^{(t-1)})$.
- $\vdots$
- p. Draw $X_p^{(t)} \sim f_{X_p|X_{-p}, U}(\cdot | X_{-p}^{(t)}, U^{(t-1)})$.
- p+1. Draw $U^{(t)} \sim f_{U|X}(\cdot | \mathbf{X}^{(t)})$.

Algorithm: The Metropolised Slice Sampler

Starting with $(\mathbf{X}^{(0)}, U^{(0)})$ iterate for $t = 1, 2, \dots$

1. Draw $\mathbf{X} \sim q(\cdot | \mathbf{X}^{(t-1)}, U^{(t-1)})$.
2. With probability

$$\min \left(1, \frac{f(\mathbf{X}, U^{(t-1)}) q(\mathbf{X}^{(t-1)} | \mathbf{X}, U^{(t-1)})}{f(\mathbf{X}^{(t-1)}, U^{(t-1)}) q(\mathbf{X} | \mathbf{X}^{(t-1)}, U^{(t-1)})} \right)$$

accept and set $\mathbf{X}^{(t)} = \mathbf{X}$.

Otherwise, set $\mathbf{X}^{(t)} = \mathbf{X}^{(t-1)}$.

2. Draw $U^{(t)} \sim f_{U|\mathbf{X}}(\cdot | \mathbf{X}^{(t)})$.

Data Augmentation I

- *Latent variable models* are common: statistical models with:
 - parameters θ ,
 - observations $\mathbf{y}$, and
 - latent variables, $\mathbf{z}$.
- Typically, the joint distribution, $f_{\mathbf{Y}, \mathbf{Z}, \theta}$, is known,
- but integrating out the latent variables to get $f_{\mathbf{Y}, \theta}$ is not feasible.
- Without $f_{\mathbf{Y}, \theta}$ we can't implement an MCMC algorithm targeting $f_{\theta|\mathbf{Y}}$.
- The basis of data augmentation is to *augment* θ with $\mathbf{z}$ and to run an MCMC algorithm which targets $f_{\theta, \mathbf{Z}|\mathbf{Y}}$.
- This distribution has the correct marginal in θ .

Data Augmentation and Gibbs Samplers

- Gibbs sampling is only feasible when we can sample easily from the full conditionals.
- A technique that can help achieving full conditionals that are easy to sample from is *demarginalisation*:
Introduce a set of auxiliary random variables $Z_1, \dots, Z_r$ such that f is the marginal density of $(X_1, \dots, X_p, Z_1, \dots, Z_r)$, i.e.

$$f(x_1, \dots, x_p) = \int f(x_1, \dots, x_p, z_1, \dots, z_r) d(z_1, \dots, z_r).$$

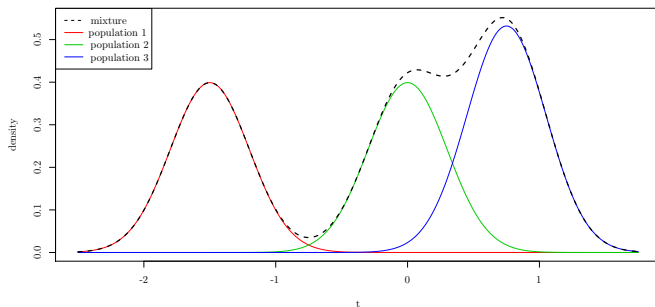
- In many cases there is a “natural choice” of the *completion* $(Z_1, \dots, Z_r)$.

Data Augmentation

Example: Mixture of Gaussians — Model

Consider the following K population mixture model for data $Y_1, \dots, Y_n$:

$$f(y_i) = \sum_{k=1}^K \pi_k \phi(\mu_k, 1/\tau)(y_i)$$



Objective: Bayesian inference for $(\pi_1, \dots, \pi_K, \mu_1, \dots, \mu_K)$.

Example: Mixture of Gaussians — Priors

- The number of components K is assumed to be known.
- The precision parameter τ is assumed to be known.
- $(\pi_1, \dots, \pi_K) \sim \text{Dirichlet}(\alpha_1, \dots, \alpha_K)$, i.e.

$$f_{(\alpha_1, \dots, \alpha_K)}(\pi_1, \dots, \pi_K) = \frac{\Gamma(\sum_{k=1}^K \alpha_k)}{\prod_{k=1}^K \Gamma(\alpha_k)} \prod_{k=1}^K \pi_k^{\alpha_k - 1}.$$

- $(\mu_1, \dots, \mu_K) \sim \text{N}(\mu_0, 1/\tau_0)$, i.e.

$$f_{(\mu_0, \tau_0)}(\mu_k) \propto \exp\left(-\tau_0(\mu_k - \mu_0)^2/2\right).$$

Example: Mixture of Gaussians — Joint distribution

$$f(\mu_1, \dots, \mu_K, \pi_1, \dots, \pi_K, y_1, \dots, y_n) \propto \left(\prod_{k=1}^K \pi_k^{\alpha_k - 1} \right) \cdot \left(\prod_{k=1}^K \exp \left(-\tau_0 (\mu_k - \mu_0)^2 / 2 \right) \right) \cdot \left(\prod_{i=1}^n \sum_{k=1}^K \pi_k \exp \left(-\tau (y_i - \mu_k)^2 / 2 \right) \right)$$

The full conditionals do not seem to come from “nice” distributions.

Use data augmentation: include auxiliary variables $Z_1, \dots, Z_n$ which indicate which population the i -th individual is from, i.e.

$$\mathbb{P}(Z_i = k) = \pi_k \quad \text{and} \quad Y_i | Z_i = k \sim N(\mu_k, 1/\tau).$$

The marginal distribution of Y is as before, so $Z_1, \dots, Z_n$ are indeed a completion.

Example: Mixture of Gaussians — Joint distribution

The joint distribution of the augmented system is

$$\begin{aligned}
 & f(y_1, \dots, y_n, z_1, \dots, z_n, \mu_1, \dots, \mu_K, \pi_1, \dots, \pi_K) \\
 \propto & \left(\prod_{k=1}^K \pi_k^{\alpha_k - 1} \right) \cdot \left(\prod_{k=1}^K \exp \left(-\tau_0 (\mu_k - \mu_0)^2 / 2 \right) \right) \\
 & \cdot \left(\prod_{i=1}^n \pi_{z_i} \exp \left(-\tau (y_i - \mu_{z_i})^2 / 2 \right) \right).
 \end{aligned}$$

The full conditionals now come from “nice” distributions.

Data Augmentation

Example: Mixture of Gaussians — Full conditionals

$$\begin{aligned} \mathbb{P}(Z_i = k | Y_1, \dots, Y_n, \mu_1, \dots, \mu_K, \pi_1, \dots, \pi_K) \\ = \frac{\pi_k \phi_{(\mu_k, 1/\tau)}(y_i)}{\sum_{\ell=1}^K \pi_\ell \phi_{(\mu_\ell, 1/\tau)}(y_i)}, \end{aligned}$$

$$\begin{aligned} \mu_k | Y_1, \dots, Y_n, Z_1, \dots, Z_n, \pi_1, \dots, \pi_K \\ \sim \mathcal{N} \left(\frac{\tau \left(\sum_{i: Z_i=k} Y_i \right) + \tau_0 \mu_0}{|\{i: Z_i = k\}| \tau + \tau_0}, \frac{1}{|\{i: Z_i = k\}| \tau + \tau_0} \right), \end{aligned}$$

$$\begin{aligned} \pi_1, \dots, \pi_K | Y_1, \dots, Y_n, Z_1, \dots, Z_n, \mu_1, \dots, \mu_K \\ \sim \text{Dirichlet}(\alpha_1 + |\{i: Z_i = 1\}|, \dots, \alpha_K + |\{i: Z_i = K\}|). \end{aligned}$$

Example: Mixture of Gaussians — Gibbs sampler

Starting with initial values $\mu_1^{(0)}, \dots, \mu_K^{(0)}, \pi_1^{(0)}, \dots, \pi_K^{(0)}$ iterate for $t = 1, 2, \dots$

1. For $i = 1, \dots, n$:

Draw $Z_i^{(t)}$ from the discrete distribution on $\{1, \dots, K\}$

$$\mathbb{P}(Z_i^{(t)} = k | Y_1, \dots, Y_n, \mu_1^{(t-1)}, \dots, \mu_K^{(t-1)}, \pi_1^{(t-1)}, \dots, \pi_K^{(t-1)}) = \frac{\pi_k \phi_{(\mu_k^{(t-1)}, 1/\tau)}(y_i)}{\sum_{\ell=1}^K \pi_\ell^{(t-1)} \phi_{(\mu_\ell^{(t-1)}, 1/\tau)}(y_i)}.$$

2. For $k = 1, \dots, K$:

Draw

$$\mu_k^{(t)} \sim \mathcal{N} \left(\frac{\tau \left(\sum_{i: Z_i^{(t)} = k} Y_i \right) + \tau_0 \mu_0}{|\{i: Z_i^{(t)} = k\}| \tau + \tau_0}, \frac{1}{|\{i: Z_i^{(t)} = k\}| \tau + \tau_0} \right).$$

3. Draw

$$(\pi_1^{(t)}, \dots, \pi_K^{(t)}) \sim \text{Dirichlet} \left(\alpha_1 + |\{i: Z_i^{(t)} = 1\}|, \dots, \alpha_K + |\{i: Z_i^{(t)} = K\}| \right).$$

- $$\pi(\theta|y) = \frac{f(y|\theta)p(\theta)}{p(y)}.$$

- If both $p(\theta)$ and $f(y|\theta)$ can be evaluated we're done.
- If we *cannot* evaluate $f(y|\cdot)$ even pointwise, then we *can't* directly use the techniques which we've described previously.
- Consider the case in which y is *discrete*.
- We can invoke a clever data augmentation trick which requires only that we can *sample* from $f(\cdot|\theta)$.

ABC and pseudo-marginal methods

- We can define an extended distribution:

$$\pi(\theta, u|y) \propto f(u|\theta)p(\theta)\delta_{y,u}$$

and note that it has, as a marginal distribution, our target:

$$\sum_u \pi(\theta, u|y) \propto \sum_u f(u|\theta)p(\theta)\delta_{y,u} = f(y|\theta)p(\theta).$$

- We can sample $(\theta, u) \sim f(u|\theta)p(\theta)$ and use this as a rejection sampling proposal for our target distribution, keeping samples with probability proportional to

$$\frac{\pi(\theta, u|y)}{f(u|\theta)p(\theta)} \propto \delta_{y,u}.$$

Approximate Bayesian Computation

- When data is not discrete / takes many values, exact matches have no or negligible probability.
- Instead, we keep samples for which $\|u - y\| \leq \epsilon$.
- This leads to a *different* target distribution:

$$\pi_{\theta, u|y}^{\text{ABC}}(\theta, y|u) \propto f(u|\theta)p(\theta)\mathbb{I}_{B(y, \epsilon)}(u),$$

where $B(y, \epsilon) := \{u : |u - y| \leq \epsilon\}$, so

$$\begin{aligned}\pi_{\theta|y}^{\text{ABC}} &\propto \int f(u|\theta)p(\theta)\mathbb{I}_{B(y, \epsilon)}(u)du \\ &\propto p(\theta) \int f(u|\theta)\mathbb{I}_{B(y, \epsilon)}(u)du \\ &\propto p(\theta) \int_{u \in B(y, \epsilon)} f(u|\theta)du.\end{aligned}$$

This approximation amounts to a *smoothing* of the likelihood.

Even More Approximate Bayesian Computation

- Often a further approximation is introduced by considering not the data itself but some low dimensional summary of the data: This leads to a *different* target distribution:

$$\pi_{\theta, u|y}^{\text{ABC}}(\theta, u|y) \propto f(u|\theta)p(\theta)\mathbb{I}_{B(s(y), \epsilon)}(s(u)).$$

- Unless the summary is a sufficient statistic (which it probably isn't) this introduces a difficult to understand approximation.
- Be very careful.

Exact-approximate methods

- Suppose that, for any θ , it is possible to compute an unbiased estimate $\hat{f}(y|\theta)$ of $f(y|\theta)$. Then...

- Using the acceptance probability

$$\alpha(\theta^{(i)}, \theta^*) = \min \left\{ 1, \frac{\hat{f}(y|\theta^*)p(\theta^*)q(\theta^{(i)}|\theta^*)}{\hat{f}(y|\theta^{(i)})p(\theta^{(i)})q(\theta^*|\theta^{(i)})} \right\}$$

yields an MCMC algorithm with target distribution $\pi(\theta|y)$.

- Using the weight

$$w^{(i)} = \frac{\hat{f}(y|\theta^{(i)})p(\theta^{(i)})}{q(\theta^{(i)})}$$

yields an importance sampling algorithm with target distribution $\pi(\theta|y)$.

Beaumont (2003), Andrieu and Roberts (2009), Fearnhead et al. (2010).

ABC and pseudo-marginal methods

Why is this true?

- Write down the joint distribution of *all* of the variables that are being used

$$\hat{f}(y|\theta, u)p(u|\theta)p(\theta)$$

where u are the random variables used to generate the estimate $\hat{f}$.

- An algorithm that simulates from $\pi(\theta, u|y)$ has the correct marginal

$$\begin{aligned} \int_u \pi(\theta, u|y) du &\propto \int_u \hat{f}(y|\theta, u)p(u|\theta)p(\theta) du \\ &= p(\theta) \int_u \hat{f}(y|\theta, u)p(u|\theta) du \\ &= p(\theta) f(y|\theta) \\ &\propto \pi(\theta|y). \end{aligned}$$

Why is this true?

- Using $q\left((\theta^*, u^*) \mid (\theta^{(i)}, u^{(i)})\right) = q(\theta^* \mid \theta^{(i)})p(u^* \mid \theta^*)$ as a proposal within a Metropolis-Hastings algorithm yields the desired acceptance probability.

$$\min \left\{ 1, \frac{\hat{f}(y \mid \theta^*, u^*)p(u^* \mid \theta^*)p(\theta^*)}{\hat{f}(y \mid \theta^{(i)}, u^{(i)})p(u^{(i)} \mid \theta^{(i)})p(\theta^{(i)})} \frac{q(\theta^{(i)} \mid \theta^*)p(u^{(i)} \mid \theta^{(i)})}{q(\theta^* \mid \theta^{(i)})p(u^* \mid \theta^*)} \right\}$$

$$= \min \left\{ 1, \frac{\hat{f}(y \mid \theta^*, u^*)p(\theta^*)}{\hat{f}(y \mid \theta^{(i)}, u^{(i)})p(\theta^{(i)})} \frac{q(\theta^{(i)} \mid \theta^*)}{q(\theta^* \mid \theta^{(i)})} \right\}.$$

- A similar extended space representation may be used in importance sampling.

Part 5— Section 15

Gradient-based methods

The Metropolis-Adjusted Langevin Algorithm

- Based on the Langevin diffusion:

$$d\mathbf{X}_t = -\frac{1}{2}\nabla \log(f(\mathbf{X}_t))dt + d\mathbf{B}_t$$

which is f -invariant *in continuous time*.

- Given target f the MALA proposal mechanism samples:

$$\begin{aligned}\mathbf{X} &\leftarrow \mathbf{X}^{(t-1)} + \epsilon \\ \epsilon &\sim \mathcal{N}\left(-\frac{\sigma^2}{2}\nabla \log f(\mathbf{X}^{(t-1)}), \sigma^2 I_p\right)\end{aligned}$$

at time t .

- Accepts X with the usual MH acceptance probability.

The Metropolis-Adjusted Langevin Algorithm

- Based on the Langevin diffusion:

$$d\mathbf{X}_t = \frac{1}{2} \nabla \log(f(\mathbf{X}_t)) dt + d\mathbf{B}_t$$

which is f -invariant *in continuous time*.

- Given target f the MALA proposal proposes:

$$\mathbf{X} \leftarrow \mathbf{X}^{(t-1)} + \epsilon$$

$$\epsilon \sim \mathcal{N} \left(\frac{\sigma^2}{2} \nabla \log f(\mathbf{X}^{(t-1)}), \sigma^2 I_p \right)$$

at time t .

- Accepts X with the usual MH acceptance probability.
- Optimal acceptance rate (under similar strong conditions) now 0.574.

MALA

MALA Example: Normal (1)

Target $f(x) = \mathcal{N}(0, 1)$

Proposal

$$q(X^{(t-1)}, X) = \mathcal{N}\left(X^{(t-1)} - \frac{\sigma^2 X^{(t-1)}}{2}, \sigma^2\right)$$

Acceptance Probability

$$\begin{aligned}\alpha(X^{(t-1)}, X) &= 1 \wedge \frac{f(X)}{f(X^{(t-1)})} \frac{q(X, X^{(t-1)})}{q(X^{(t-1)}, X)} \\ &= 1 \wedge \exp\left(\frac{1}{2} \left[(X^{(t-1)})^2 - X^2\right]\right) \times \\ &\quad \exp\left(\frac{1}{2\sigma^2} \left[\left\{X - \mu(X^{(t-1)})\right\}^2 - \left\{X^{(t-1)} - \mu(X)\right\}^2\right]\right)\end{aligned}$$

$$\text{where } \mu(x) := x - \frac{x\sigma^2}{2}.$$

Augmentation

oooooooooooo
oooooooooooo
oooooooooooo
oooooooooooo

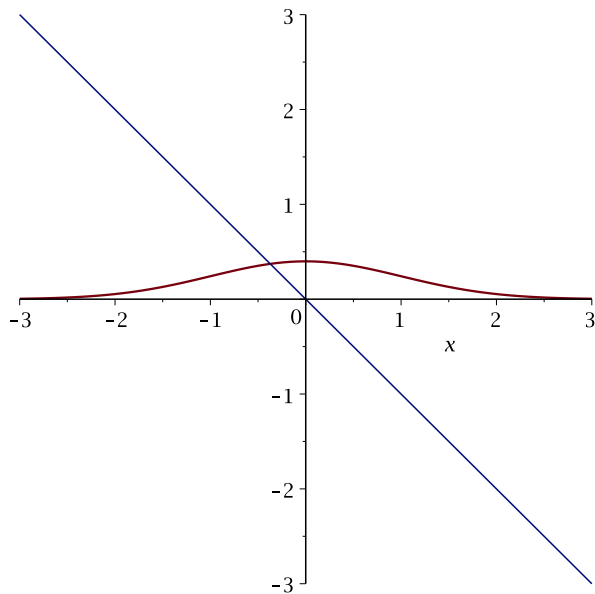
Gradient-based methods

ooo●ooo
oooooooooooooooooooo

Current / future directions

o
oooooo
oooooooooooooooo

MALA



Augmentation

oooooooooooo
oooooooooooo
oooooooooooo
oooooooooooo

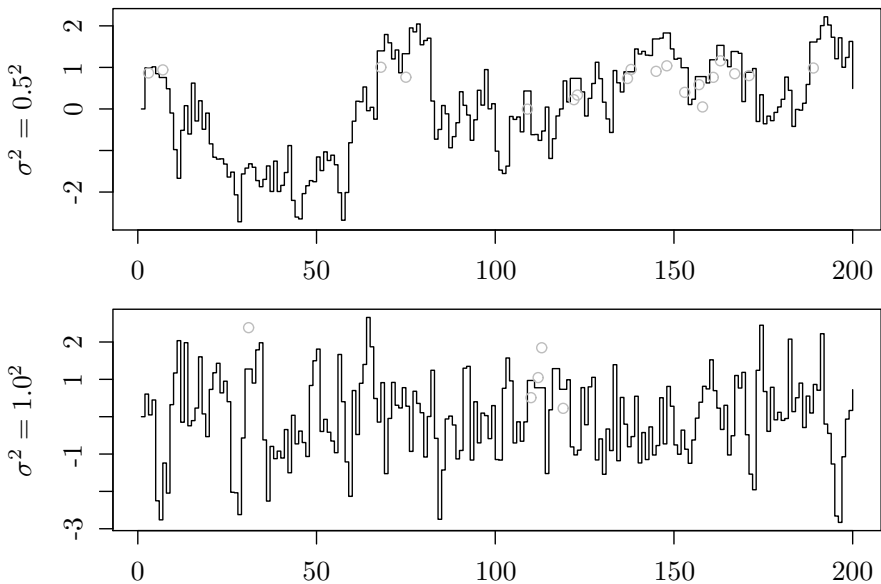
Gradient-based methods

oooo●ooo
oooooooooooooooooooo

Current / future directions

o
ooooooo
oooooooooooo

MALA



Augmentation

oooooooooooo
oooooooooooo
oooooooooooo
oooooooooooo

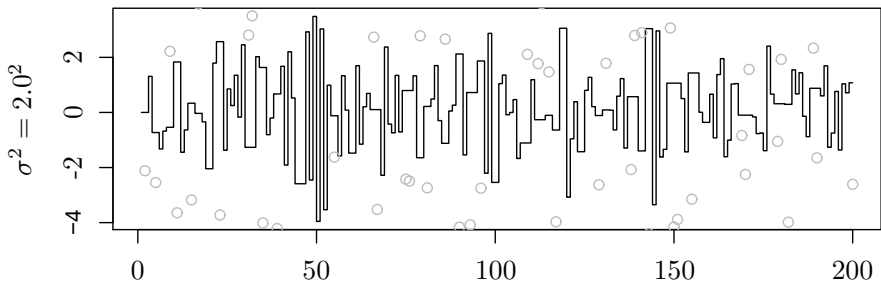
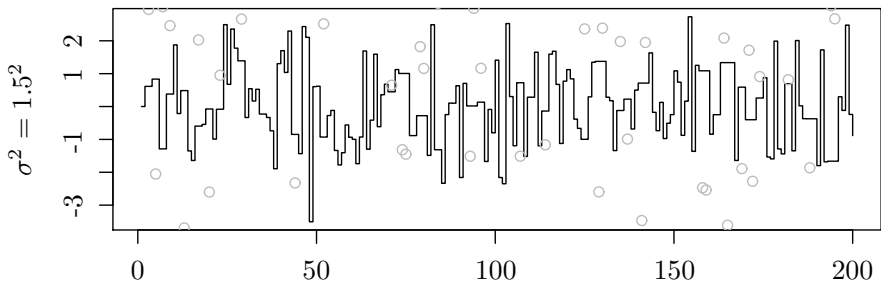
Gradient-based methods

oooo●o
oooooooooooooooooooo

Current / future directions

o
ooooooo
oooooooooooo

MALA



MALA

MALA Example: Normal (2)

RWM	Autocorrelation $\rho(X^{(t-1)}, X^{(t)})$	Probability of acceptance $\alpha(X, X^{(t-1)})$	ESJD
$\sigma^2 = 0.1^2$	0.9901	0.9694	0.010
$\sigma^2 = 1$	0.7733	0.7038	0.448
$\sigma^2 = 2.38^2$	0.6225	0.4426	0.742
$\sigma^2 = 10^2$	0.8360	0.1255	0.337

MALA	Autocorrelation $\rho(X^{(t-1)}, X^{(t)})$	Probability of acceptance $\alpha(X, X^{(t-1)})$	ESJD
$\sigma^2 = 0.5^2$	0.898	0.877	0.246
$\sigma^2 = 1$	0.492	0.961	1.293
$\sigma^2 = 1.5^2$	0.047	0.774	2.137
$\sigma^2 = 2.0^2$	0.011	0.631	4.119

Scaling with dimension

- The number of iterations we must run the following algorithms to obtain one effectively independent point is, as a function of the size of the parameter space d :
 - $O(d)$ for random walk Metropolis-Hastings, which gives an overall computational cost of $O(d^2)$;
 - $O(d^{1/3})$ for the Metropolis-adjusted Langevin algorithm, which gives an overall computational cost of $O(d^{4/3})$.

Hamiltonian / Hybrid Monte Carlo

- Mimics a conservative physical system by introducing momentum.
- Approximate continuous measure-preserving flow using (symplectic) numerical integration.
- Use Metropolis–Hastings accept/reject correction.
- Can mix *much* faster than random walk algorithms.
- Difficulties with multi-modal targets and can be expensive.

c.f. Neal (2011) MCMC using Hamiltonian dynamics. In Brooks et al., 113–162. [Brooks, Gelman, Jones, and Meng (eds.) (2011) Handbook of Markov Chain Monte Carlo. CRC Press.]

Constructing a proposal: dynamics of a ball

- For random walk, we found that we needed to decrease the proposal variance as the dimension increased.
- We would like to have proposals that move a long way, but still have a good probability of acceptance
 - we need a proposal that follows the mass of the distribution.
- Think of the negative log of the target distribution, and consider the idea of setting a ball rolling around this surface
 - someone with a background in physics could describe the dynamics of this ball.
- Idea:
 - give the ball a push in a random direction
 - follow the dynamics of the ball for a while
 - use this as the proposal.

Hamiltonian dynamics

- Hamiltonian mechanics is an abstract formulation of classical mechanics (i.e. equations of motion, etc).
- It describes a system involving two time-evolving vectors θ and v , each of dimension d .
- The “Hamiltonian” $H(\theta, v)$ describes the time evolution of the system, through Hamilton’s equations

$$\frac{d\theta_i}{dt} = \frac{\partial H}{\partial v_i} \quad \frac{dv_i}{dt} = -\frac{\partial H}{\partial \theta_i}$$

for $i = 1, \dots, d$.

- Note that physicists would be very annoyed by the notation here, where the vectors are called q and p instead of θ and v .
- This is very abstract
 - what do these equations mean?

Hamiltonian dynamics: total energy

- In the use of this technique in MCMC, we use these dynamics to describe a frictionless ball rolling around the negative log of the posterior distribution, subject to a gravitational pull.
- The vector θ denotes the position of the ball, and the vector v its momentum
 - recall that momentum is equal to mass times velocity
 - for simplicity we will take the mass of the ball to be 1, which means that momentum equals velocity.
- $H(\theta, v)$ represents the total energy of the ball

$$\underbrace{H(\theta, v)}_{\text{total energy}} = \underbrace{U(\theta)}_{\text{potential energy}} + \underbrace{K(v)}_{\text{kinetic energy}} .$$

Hamiltonian dynamics: potential energy

- Recall from classical mechanics that gravitational potential energy U is equal to mgh , where m is the mass of the ball, g is the gravitational field, and h is the height.
- For simplicity, we simply set m and g to be equal to 1.
- Therefore we simply take $U(\theta)$ to be the height of the ball at θ

$$U(\theta) = -\log \left(\pi(\theta | y) \right).$$

- For example, $U(\theta) = \theta^2$ would correspond to a Gaussian with zero mean.

Hamiltonian dynamics: kinetic energy

- Recall from classical mechanics that kinetic energy K is equal to a half times mass times velocity squared.
- In our case (with $m = 1$, momentum equals velocity). We obtain, in the univariate case, $K = v^2/2$.
- We are looking at the multivariate case, which gives $K(v) = v^T v/2$.

Hamiltonian dynamics: Hamiltonian

- The Hamiltonian in our case is given by

$$H(\theta, v) = -\log(\pi(\theta | y)) + v^T v / 2.$$

- Hamilton's equations in our case are given by

$$\frac{d\theta}{dt} = v \quad \text{and} \quad \frac{dv}{dt} = \nabla \log(\pi(\theta | y)).$$

- These make sense!
 - the rate of change of position is given by the velocity
 - the rate of change of velocity is given by the gradient of the surface.
- To construct a proposal for use in MCMC, we will simply simulate forwards from these dynamics for some time t
 - this simulation defines a deterministic function R_t , mapping $(\theta, v) \mapsto (\theta^*, v^*)$.

Hamiltonian dynamics: properties

- What did we gain from the abstract formulation, rather than simply working out this formulation from classical mechanics?
- Hamiltonian dynamics has some nice mathematical properties, that are particularly useful when constructing MCMC updates (here we follow Neal (2011)).
- **Reversibility.** There is an inverse to R_t , and this can be defined in terms of R_t . We have that R_t^{-1} is given by
 - taking the negative of the velocity (to make the ball go backwards)
 - applying R_t (running the dynamics for time t)
 - taking the negative of the velocity of the result (to make the ball "face" back in the direction it was originally)
 - we need this property for the dynamics to have π as the invariant distribution.

Hamiltonian dynamics: properties

- Conservation of the Hamiltonian.** The dynamics do not change the value of H - the total energy of the ball is conserved.
 - this property is crucial in ensuring that the acceptance probability is high
 - soon we will define the a joint distribution of θ and v in terms of H - the conservation of H under the dynamics will mean that (θ, v) has the same density as (θ^*, v^*) .
- Volume preservation.** Hamiltonian dynamics preserves volume in the space of (θ, v) . This means that no Jacobian is needed when calculating the acceptance probability of a move (as it is in some other methods).

Hamiltonian Monte Carlo

- We now have most of the ingredients needed to define Hamiltonian Monte Carlo.
- We proceed as follows
 - define a joint distribution on (θ, v) such that we can run Hamiltonian dynamics on it in order to obtain points from π
 - describe how to deal with the fact that we cannot simulate Hamiltonian dynamics exactly.

Hamiltonian Monte Carlo: joint distribution

- Define a joint distribution on (θ, v) as follows

$$\begin{aligned}
 \pi_{\theta, v}(\theta, v) &\propto \exp(-H(\theta, v)) \\
 &= \exp(-U(\theta)) \exp(-K(v)) \\
 &= \exp\left(-\left(-\log(\pi(\theta | y))\right)\right) \exp(-v^T v / 2) \\
 &= \pi(\theta | y) \exp(-v^T v / 2).
 \end{aligned}$$

- We see that the joint distribution on (θ, v) has $\pi(\theta | y)$ as its marginal, and that we have a Gaussian distribution on v
 - we could choose a different covariance for this Gaussian distribution on v - this would correspond to using a different mass for the ball in the potential energy.

Using Hamiltonian dynamics as an MCMC move

- Rigorously proving reversibility can also be done using the *involution* framework previously discussed - see Andrieu, Lee, Livingstone (2020): the Metropolis-Hastings acceptance probability for a volume preserving deterministic move T that is an involution is by $\min \left\{ 1, \frac{\pi(T(\theta, v))}{\pi(\theta, v)} \right\}$.
- We define T to be the composition of applying Hamiltonian dynamics $R_t(\theta, v)$, then taking the negative of the velocity component.

Using Hamiltonian dynamics as an MCMC move

- Then, using the conservation of the Hamiltonian, the acceptance probability of applying Hamiltonian dynamics to the joint target is given by

$$\min \left\{ 1, \frac{\pi_{\theta, v}(T(\theta, v))}{\pi_{\theta, v}(\theta, v)} \right\} = \min \left\{ 1, \frac{\exp(-H(T(\theta, v)))}{\exp(-H(\theta, v))} \right\} = 1, \text{ which}$$

means that we would always accept such a move!

- Potentially make very large moves, as long as we choose appropriately the time for which we simulate the dynamics
 - too short, and we will not move far
 - too long, and it is possible that we end up where we started!
- Alternate the dynamics with simulating a new velocity exactly from the target distribution for v , so that we change the direction of the trajectories at different iterations.

Approximating Hamiltonian dynamics

- We cannot simulate Hamiltonian dynamics exactly
 - we must use some solver, just as we did for the Langevin method.
- We use the "leapfrog" method to approximately simulate the dynamics
 - this produces a discretized trajectory that approximates the continuous dynamics
 - the transformation produced using this approach is also reversible and volume preserving.

Approximating Hamiltonian dynamics

- However, the Hamiltonian is not exactly conserved.
 - This means that the acceptance probability is not 1.
- Let T be the transformation given by the leapfrog method, and $(\theta^*, v^*) = T(\theta, v)$. Then, the acceptance probability is

$$\min \left\{ 1, \frac{\pi(T(\theta, v))}{\pi(\theta, v)} \right\} = \min \left\{ 1, \exp \left(-H(\theta^*, v^*) + H(\theta, v) \right) \right\},$$

- Note that, as in standard Metropolis-Hastings, we can use $p(\theta) l(y | \theta)$ in place of $\pi(\theta | y)$, since the normalizing constant $p(y)$ cancels.
- When implementing the leapfrog method, we need $\nabla \log(\pi(\theta | y))$. This is given by $\nabla \log p(\theta_t) + \nabla \log l(y | \theta_t)$ as in the previous lecture.

HMC properties

- Dependence on dimension
 - the optimal τ is proportional to $d^{1/4}$
 - $O(d^{1/4})$ steps are needed to reach a nearly independent point
 - overall cost is $O(d^{5/4})$
 - this beats both random walk and MALA.
- The tuning of HMC makes a big difference to the performance
 - much research is devoted to automating this tuning
 - the "no u-turn sampler" (NUTS), implemented in Stan, is a significant contribution.

Augmentation

oooooooooooo
oooooooooooo
oooooooooooo
oooooooooooo

Gradient-based methods

ooooooo
oooooooooooooooooooo●

Current / future directions

o
ooooooo
oooooooooooooooo

HMC

HMC in action

HMC in action

Part 5— Section 16

Current / future directions

Current / future directions

Everything discussed so far is well-established and arguably 'classical'.

The goal of the remainder of the course is to give you a flavour of some current areas which are popular topics of research.

SBI using deep learning

One of the biggest advances in the last 20 years is the advent of neural networks and deep learning, and a very active research area is to try and utilise advances in deep learning in the context of computer intensive statistics.

To give you a flavour of what current research in this area looks like, we'll discuss one approach termed *Sequential Neural Likelihood*⁶.

⁶Papamakarios, Sterratt, Murray. (2020). Sequential neural likelihood: Fast likelihood-free inference with autoregressive flows. AISTATS 2019, 89.
 Papamakarios, Murray. (2016). Fast e-free inference of simulation models with Bayesian conditional density estimation. Advances in Neural Information Processing Systems, 1036-1044.

SBI using deep learning

One of the biggest advances in the last 20 years is the advent of neural networks and deep learning, and a very active research area is to try and utilise advances in deep learning in the context of computer intensive statistics.

To give you a flavour of what current research in this area looks like, we'll discuss one approach termed *Sequential Neural Likelihood*⁶.

⁶Papamakarios, Sterratt, Murray. (2020). Sequential neural likelihood: Fast likelihood-free inference with autoregressive flows. AISTATS 2019, 89.
 Papamakarios, Murray. (2016). Fast e-free inference of simulation models with Bayesian conditional density estimation. Advances in Neural Information Processing Systems, 1036-1044.

Neural density estimation

One enormous advance in deep learning is *Neural Density Estimation*.

NDE

Suppose we have some joint distribution $p(\mathbf{u}, \mathbf{v})$ and i.i.d. samples $(\mathbf{u}_i, \mathbf{v}_i)_{i=1}^n$.

A *conditional neural density estimator* is a parametric model q_ϕ (it has parameters ϕ) which we train using our samples to approximate the underlying conditional $p(\mathbf{u}|\mathbf{v})$.

For certain classes of estimators (typically based on *normalizing flows*), it can be shown that in the limit $n \rightarrow \infty$, for the optimal parameters ϕ^* , $q_{\phi^*}(\mathbf{u}|\mathbf{v}) = p(\mathbf{u}|\mathbf{v})$.

NDE for CIS

In other words, given (a large number of) samples $(\mathbf{u}_i, \mathbf{v}_i)_{i=1}^n$, it is possible to learn the underlying conditional density $p(\mathbf{u}|\mathbf{v})$ in a black-box manner.

Intractable likelihoods

Recall the *intractable likelihood* setting from before: we have a Bayesian (joint) model $f(\boldsymbol{\theta}, \mathbf{y}) = f(\boldsymbol{\theta})f(\mathbf{y}|\boldsymbol{\theta})$, where $L(\boldsymbol{\theta}; \mathbf{y}) = f(\mathbf{y}|\boldsymbol{\theta})$ *cannot be evaluated...but we assume that for any $\boldsymbol{\theta}$, we can simulate synthetic $\mathbf{y} \sim f(\cdot|\boldsymbol{\theta})$.*

So are able to generate samples $(\boldsymbol{\theta}_i, \mathbf{y}_i) \sim f$ - this is precisely the case where we could use NDE to *learn the intractable likelihood*!

NDE for CIS

In other words, given (a large number of) samples $(\mathbf{u}_i, \mathbf{v}_i)_{i=1}^n$, it is possible to learn the underlying conditional density $p(\mathbf{u}|\mathbf{v})$ in a black-box manner.

Intractable likelihoods

Recall the *intractable likelihood* setting from before: we have a Bayesian (joint) model $f(\boldsymbol{\theta}, \mathbf{y}) = f(\boldsymbol{\theta})f(\mathbf{y}|\boldsymbol{\theta})$, where $L(\boldsymbol{\theta}; \mathbf{y}) = f(\mathbf{y}|\boldsymbol{\theta})$ *cannot be evaluated...but we assume that for any $\boldsymbol{\theta}$, we can simulate synthetic $\mathbf{y} \sim f(\cdot|\boldsymbol{\theta})$.*

So are able to generate samples $(\boldsymbol{\theta}_i, \mathbf{y}_i) \sim f$ - this is precisely the case where we could use NDE to *learn the intractable likelihood!*

NDE for CIS

In other words, given (a large number of) samples $(\mathbf{u}_i, \mathbf{v}_i)_{i=1}^n$, it is possible to learn the underlying conditional density $p(\mathbf{u}|\mathbf{v})$ in a black-box manner.

Intractable likelihoods

Recall the *intractable likelihood* setting from before: we have a Bayesian (joint) model $f(\boldsymbol{\theta}, \mathbf{y}) = f(\boldsymbol{\theta})f(\mathbf{y}|\boldsymbol{\theta})$, where $L(\boldsymbol{\theta}; \mathbf{y}) = f(\mathbf{y}|\boldsymbol{\theta})$ *cannot be evaluated*...but we assume that for any $\boldsymbol{\theta}$, we can *simulate synthetic* $\mathbf{y} \sim f(\cdot|\boldsymbol{\theta})$.

So are able to generate samples $(\boldsymbol{\theta}_i, \mathbf{y}_i) \sim f$ - this is precisely the case where we could use NDE to *learn the intractable likelihood*!

NDE for CIS

In other words, given (a large number of) samples $(\mathbf{u}_i, \mathbf{v}_i)_{i=1}^n$, it is possible to learn the underlying conditional density $p(\mathbf{u}|\mathbf{v})$ in a black-box manner.

Intractable likelihoods

Recall the *intractable likelihood* setting from before: we have a Bayesian (joint) model $f(\boldsymbol{\theta}, \mathbf{y}) = f(\boldsymbol{\theta})f(\mathbf{y}|\boldsymbol{\theta})$, where $L(\boldsymbol{\theta}; \mathbf{y}) = f(\mathbf{y}|\boldsymbol{\theta})$ *cannot be evaluated*...but we assume that for any $\boldsymbol{\theta}$, we can *simulate synthetic* $\mathbf{y} \sim f(\cdot|\boldsymbol{\theta})$.

So are able to generate samples $(\boldsymbol{\theta}_i, \mathbf{y}_i) \sim f$ - this is precisely the case where we could use NDE to *learn the intractable likelihood*!

Neural Likelihood estimation

In practice, we have some given data $\mathbf{y}^{\text{obs}}$ and we want to sample from our posterior $f(\boldsymbol{\theta}|\mathbf{y}^{\text{obs}})$.

Neural Likelihood v1

- 1 Generate a large number of samples $(\boldsymbol{\theta}_i, \mathbf{y}_i) \sim f$ from the joint model.
- 2 Train a neural density estimator $q_{\phi^*}(\mathbf{y}|\boldsymbol{\theta}) \approx f(\mathbf{y}|\boldsymbol{\theta})$ to approximate the intractable likelihood.
- 3 Use this estimator in the case $\mathbf{y} = \mathbf{y}^{\text{obs}}$ to form our posterior approximation

$$f^{\text{NL}}(\boldsymbol{\theta}|\mathbf{y}^{\text{obs}}) \propto f(\boldsymbol{\theta})q_{\phi^*}(\mathbf{y}^{\text{obs}}|\boldsymbol{\theta}).$$

- 4 Run an MCMC scheme targeting $f^{\text{NL}}(\boldsymbol{\theta}|\mathbf{y}^{\text{obs}})$ to generate approximate posterior samples $\boldsymbol{\theta}_1^{\text{NL}}, \dots, \boldsymbol{\theta}_N^{\text{NL}}$.

Neural Likelihood estimation

In practice, we have some given data $\mathbf{y}^{\text{obs}}$ and we want to sample from our posterior $f(\boldsymbol{\theta}|\mathbf{y}^{\text{obs}})$.

Neural Likelihood v1

- ① Generate a large number of samples $(\boldsymbol{\theta}_i, \mathbf{y}_i) \sim f$ from the joint model.
- ② Train a neural density estimator $q_{\phi^*}(\mathbf{y}|\boldsymbol{\theta}) \approx f(\mathbf{y}|\boldsymbol{\theta})$ to approximate the intractable likelihood.
- ③ Use this estimator in the case $\mathbf{y} = \mathbf{y}^{\text{obs}}$ to form our posterior approximation

$$f^{\text{NL}}(\boldsymbol{\theta}|\mathbf{y}^{\text{obs}}) \propto f(\boldsymbol{\theta})q_{\phi^*}(\mathbf{y}^{\text{obs}}|\boldsymbol{\theta}).$$

- ④ Run an MCMC scheme targeting $f^{\text{NL}}(\boldsymbol{\theta}|\mathbf{y}^{\text{obs}})$ to generate approximate posterior samples $\boldsymbol{\theta}_1^{\text{NL}}, \dots, \boldsymbol{\theta}_N^{\text{NL}}$.

- The training samples we generated were from the joint model:

and there is no guarantee that the generated $\mathbf{y}_i$ will be close to the actual $\mathbf{y}^{\text{obs}}$!

- 333

Sequential Neural Likelihood

The idea of Papamakarios, Sterratt, Murray (2020) is to do this in a sequential manner, beginning with the prior:

- 1 Initialise $\hat{f}_0(\boldsymbol{\theta}|\mathbf{y}^{\text{obs}}) = f^{\text{prior}}(\boldsymbol{\theta})$, $\mathcal{D} = \{\}$
- 2 For $r = 1, \dots, R$
 - (a) Obtain (approximate) sample $(\boldsymbol{\theta}_n^r)_{n=1}^N \sim \hat{f}_{r-1}(\boldsymbol{\theta}|\mathbf{y}^{\text{obs}})$ by running an MCMC scheme
 - (b) Draw $\mathbf{y}_n^r \sim f(\mathbf{y}|\boldsymbol{\theta}_n^r)$ for $n = 1, \dots, N$
 - (c) Add $(\boldsymbol{\theta}_n^r, \mathbf{y}_n^r)_{n=1}^N$ to $\mathcal{D}$
 - (d) (Re-)train $q_\phi(\mathbf{y}|\boldsymbol{\theta})$ on $\mathcal{D}$, set

$$\hat{f}_r(\boldsymbol{\theta}|\mathbf{y}^{\text{obs}}) \propto q_\phi(\mathbf{y}^{\text{obs}}|\boldsymbol{\theta})f^{\text{prior}}(\boldsymbol{\theta})$$

- 3 Return $\hat{f}_R(\boldsymbol{\theta}|\mathbf{y}^{\text{obs}})$
- 4 Run MCMC on $\hat{f}_R(\boldsymbol{\theta}|\mathbf{y}^{\text{obs}})$ to get a final sample.

Generative Modelling

So far, the problem has been ‘we are given an (unnormalised) density f , and we want to sample from it’.

The problem

In recent years, a new paradigm as emerged: instead, we are given a (large) dataset $\mathbf{y}_1, \dots, \mathbf{y}_N$, which we assume is i.i.d. from some (unknown) density f .

We would now like to generate new samples from f .

The key challenge is that we *know absolutely nothing about f* , apart from the fact that we have an i.i.d. sample.

Generative Modelling

So far, the problem has been ‘we are given an (unnormalised) density f , and we want to sample from it’.

The problem

In recent years, a new paradigm has emerged: instead, we are given a (large) dataset $\mathbf{y}_1, \dots, \mathbf{y}_N$, which we assume is i.i.d. from some (unknown) density f .

We would now like to generate new samples from f .

The key challenge is that we *know absolutely nothing about f* , apart from the fact that we have an i.i.d. sample.

The problem

In recent years, a new paradigm has emerged: instead, we are given a (large) dataset $\mathbf{y}_1, \dots, \mathbf{y}_N$, which we assume is i.i.d. from some (unknown) density f .

We would now like to generate new samples from f .

The key challenge is that we *know absolutely nothing* about f , apart from the fact that we have an i.i.d. sample.

Example: image generation



Figure: From Kaggle Cat-Faces-Dataset

Given lots of cat pictures, we would like to generate new ones.

Generative modelling

Historically, important advances in generative modelling has been methods such as Generative Adversarial Networks (GANs), Variational Autoencoders (VAEs) and Normalizing Flows.

We will consider what is currently the state of the art: diffusion models and see an application to computer intensive statistics.

Introduction to Diffusion models

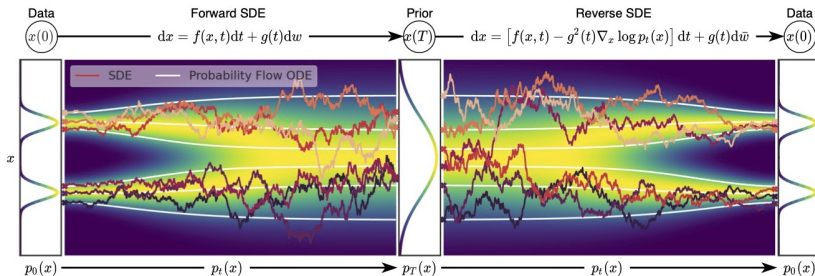


Figure: from Song, Sohl-Dickstein, Kingma, Kumar, Ermon, Poole (2021), Score-Based Generative Modeling through Stochastic Differential Equations, ICLR.

Noising and denoising

We run a diffusion process: a continuous-time stochastic process $(\mathbf{X}_t)_{t \in [0, T]}$ (in practice, you perform a time discretisation). Let $p_t(\mathbf{x})$ be its density at time t .

We initialise the process from $p_0 = f$ - in other words, we begin from one of our data points $\mathbf{X}_0 = \mathbf{y}_i$.

The forwards process is typically very simple: gradually add noise to $\mathbf{X}_0$ so that $\mathbf{X}_T$ is pure Gaussian noise.

$$\mathbf{X}_{t+1} = \alpha \mathbf{X}_t + \sqrt{1 - \alpha} \mathbf{Z}_t,$$

$\mathbf{Z}_t \sim \mathcal{N}(0, I)$ i.i.d., so that $\mathbf{X}_T \approx \mathcal{N}(0, I)$ in distribution for large T .

Noising and denoising

We run a diffusion process: a continuous-time stochastic process $(\mathbf{X}_t)_{t \in [0, T]}$ (in practice, you perform a time discretisation). Let $p_t(\mathbf{x})$ be its density at time t .

We initialise the process from $p_0 = f$ - in other words, we begin from one of our data points $\mathbf{X}_0 = \mathbf{y}_i$.

The forwards process is typically very simple: gradually add noise to $\mathbf{X}_0$ so that $\mathbf{X}_T$ is pure Gaussian noise.

$$\mathbf{X}_{t+1} = \alpha \mathbf{X}_t + \sqrt{1 - \alpha} \mathbf{Z}_t,$$

$\mathbf{Z}_t \sim \mathcal{N}(0, I)$ i.i.d., so that $\mathbf{X}_T \approx \mathcal{N}(0, I)$ in distribution for large T .

Denoising

We thus have a process to transform any data point $\mathbf{X}_0 = \mathbf{y}_i$ into pure Gaussian noise, $\mathbf{X}_T \approx \mathcal{N}(0, I)$ in distribution.

Key idea behind denoising diffusions: *can we do this backwards?*

In other words, we begin with $\mathbf{X}_T \sim \mathcal{N}(0, I)$ (pure noise) and run *backwards in time*,

$$\mathbf{X}_{T-1}, \mathbf{X}_{T-2}, \dots, \mathbf{X}_1, \mathbf{X}_0$$

in such a way that $\mathbf{X}_0$ looks like a new point $\mathbf{X} \sim f$.

Denoising

We thus have a process to transform any data point $\mathbf{X}_0 = \mathbf{y}_i$ into pure Gaussian noise, $\mathbf{X}_T \approx \mathcal{N}(0, I)$ in distribution.

Key idea behind denoising diffusions: *can we do this backwards?*

In other words, we begin with $\mathbf{X}_T \sim \mathcal{N}(0, I)$ (pure noise) and run *backwards in time*,

$$\mathbf{X}_{T-1}, \mathbf{X}_{T-2}, \dots, \mathbf{X}_1, \mathbf{X}_0$$

in such a way that $\mathbf{X}_0$ looks like a new point $\mathbf{X} \sim f$.

Denoising

We thus have a process to transform any data point $\mathbf{X}_0 = \mathbf{y}_i$ into pure Gaussian noise, $\mathbf{X}_T \approx \mathcal{N}(0, I)$ in distribution.

Key idea behind denoising diffusions: *can we do this backwards?*

In other words, we begin with $\mathbf{X}_T \sim \mathcal{N}(0, I)$ (pure noise) and run *backwards in time*,

$$\mathbf{X}_{T-1}, \mathbf{X}_{T-2}, \dots, \mathbf{X}_1, \mathbf{X}_0$$

in such a way that $\mathbf{X}_0$ looks like a new point $\mathbf{X} \sim f$.

Denoising - learning the score

Foundational results from probability theory tell us that the time-reversal of an SDE is also an SDE⁷: $(\mathbf{X}_{T-t})_{t=0}^T$ solves some stochastic differential equation

$$d\mathbf{Y}_t = g(\mathbf{Y}_t, t) dt + h(t) dW_t.$$

So in principle, we could just discretise this SDE and run it...
...Except the drift term $g(\mathbf{Y}_t, t)$ involves the unknown *score function*

$$\nabla \log p_t(\mathbf{X}),$$

where p_t is the density of $\mathbf{X}_t$ at time t .

⁷Anderson (1982). Reverse-time diffusion equation models. *Stoch. Proc. Appl.*, 12(3), 313-326. Haussmann, Pardoux (1986). Time Reversal of Diffusions. *Ann. Probab.*, 14(4), 1188-1205.

Denoising - learning the score

Foundational results from probability theory tell us that the time-reversal of an SDE is also an SDE⁷: $(\mathbf{X}_{T-t})_{t=0}^T$ solves some stochastic differential equation

$$d\mathbf{Y}_t = g(\mathbf{Y}_t, t) dt + h(t) dW_t.$$

So in principle, we could just discretise this SDE and run it...
...Except the drift term $g(\mathbf{Y}_t, t)$ involves the unknown *score function*

$$\nabla \log p_t(\mathbf{X}),$$

where p_t is the density of $\mathbf{X}_t$ at time t .

⁷Anderson (1982). Reverse-time diffusion equation models. *Stoch. Proc. Appl.*, 12(3), 313-326. Haussmann, Pardoux (1986). Time Reversal of Diffusions. *Ann. Probab.*, 14(4), 1188-1205.

Denoising - learning the score

Foundational results from probability theory tell us that the time-reversal of an SDE is also an SDE⁷: $(\mathbf{X}_{T-t})_{t=0}^T$ solves some stochastic differential equation

$$d\mathbf{Y}_t = g(\mathbf{Y}_t, t) dt + h(t) dW_t.$$

So in principle, we could just discretise this SDE and run it...
...Except the drift term $g(\mathbf{Y}_t, t)$ involves the unknown *score function*

$$\nabla \log p_t(\mathbf{X}),$$

where p_t is the density of $\mathbf{X}_t$ at time t .

⁷Anderson (1982). Reverse-time diffusion equation models. *Stoch. Proc. Appl.*, 12(3), 313-326. Haussmann, Pardoux (1986). Time Reversal of Diffusions. *Ann. Probab.*, 14(4), 1188-1205.

Score matching

A huge advance was the technique of *score matching*⁸: assuming we have access to samples $(\mathbf{y}_i) \sim f$, it is possible to approximate the score function using a neural network $s_\phi(\mathbf{X}, t)$:

$$s_\phi(\mathbf{X}, t) \approx \nabla \log p_t(\mathbf{X}).$$

$$d\mathbf{Y}_t = \hat{g}(\mathbf{Y}_t, t) dt + h(t) dW_t,$$

⁸Hyvärinen (2005). Estimation of Non-Normalized Statistical Models by Score Matching. *Journal of Machine Learning Research*, 6(Apr), 695-709.
Song, Ermon (2019). Generative Modeling by Estimating Gradients of the Data Distribution. *Advances in Neural Information Processing Systems*, 32.

Score matching

A huge advance was the technique of *score matching*⁸: assuming we have access to samples $(\mathbf{y}_i) \sim f$, it is possible to approximate the score function using a neural network $s_\phi(\mathbf{X}, t)$:

$$s_\phi(\mathbf{X}, t) \approx \nabla \log p_t(\mathbf{X}).$$

We can then initialise $\mathbf{Y}_0 \sim \mathcal{N}(0, I)$ from pure noise and discretise

$$d\mathbf{Y}_t = \hat{g}(\mathbf{Y}_t, t) dt + h(t) dW_t,$$

where $\hat{g}$ is the same as g but with s_ϕ in place of the true score.

⁸Hyvärinen (2005). Estimation of Non-Normalized Statistical Models by Score Matching. Journal of Machine Learning Research, 6(Apr), 695-709.
Song, Ermon (2019). Generative Modeling by Estimating Gradients of the Data Distribution. Advances in Neural Information Processing Systems, 32.

In other words, given our posterior $f(\boldsymbol{\theta}|\mathbf{y}) \propto f(\boldsymbol{\theta})f(\mathbf{y}|\boldsymbol{\theta})$, is there any way to exploit these diffusion models to generate samples from the posterior?

Conditional score-based diffusion models

We consider a situation now where we want to sample from $f(\boldsymbol{\theta}|\mathbf{y}^{\text{obs}})$.

Similar to before, we imagine initialising $\theta_0 \sim f(\boldsymbol{\theta}|\mathbf{y})$ and then running the forward noising process to generate $\boldsymbol{\theta}_t \sim p_t(\boldsymbol{\theta}_t|\mathbf{y})$ for $t = 1, \dots, T$.

We would like to initialise from pure noise and simulate this process backwards, which would require the *conditional* score

$$s(\boldsymbol{\theta}_t, \mathbf{y}, t) = \nabla_{\boldsymbol{\theta}} \log p_t(\boldsymbol{\theta}_t|\mathbf{y}),$$

which is intractable.

Conditional score-based diffusion models

We consider a situation now where we want to sample from $f(\boldsymbol{\theta}|\mathbf{y}^{\text{obs}})$.

Similar to before, we imagine initialising $\theta_0 \sim f(\boldsymbol{\theta}|\mathbf{y})$ and then running the forward noising process to generate $\boldsymbol{\theta}_t \sim p_t(\boldsymbol{\theta}_t|\mathbf{y})$ for $t = 1, \dots, T$.

We would like to initialise from pure noise and simulate this process backwards, which would require the *conditional* score

$$s(\boldsymbol{\theta}_t, \mathbf{y}, t) = \nabla_{\boldsymbol{\theta}} \log p_t(\boldsymbol{\theta}_t|\mathbf{y}),$$

which is intractable.

Conditional score-based diffusion models

We consider a situation now where we want to sample from $f(\boldsymbol{\theta}|\mathbf{y}^{\text{obs}})$.

Similar to before, we imagine initialising $\theta_0 \sim f(\boldsymbol{\theta}|\mathbf{y})$ and then running the forward noising process to generate $\boldsymbol{\theta}_t \sim p_t(\boldsymbol{\theta}_t|\mathbf{y})$ for $t = 1, \dots, T$.

We would like to initialise from pure noise and simulate this process backwards, which would require the *conditional* score

$$s(\boldsymbol{\theta}_t, \mathbf{y}, t) = \nabla_{\boldsymbol{\theta}} \log p_t(\boldsymbol{\theta}_t|\mathbf{y}),$$

which is intractable.

Conditional score estimation

As before, we approximate this intractable quantity using a neural network

$$s_\phi(\boldsymbol{\theta}_t, \mathbf{y}, t) \approx s(\boldsymbol{\theta}_t, \mathbf{y}, t).$$

This can be done by minimizing⁹ (with respect to ϕ)

$$\int_0^T \mathbb{E}_{p_{t|0}(\boldsymbol{\theta}_t|\boldsymbol{\theta}_0)f(\boldsymbol{\theta}_0,\mathbf{y})} \left[\|s_\phi(\boldsymbol{\theta}_t, \mathbf{y}, t) - \nabla \log p_{t|0}(\boldsymbol{\theta}_t|\boldsymbol{\theta}_0)\|^2 \right] dt.$$

The inner expectation only depends on samples $\boldsymbol{\theta}_0 \sim f(\boldsymbol{\theta})$ (the prior), $\mathbf{y} \sim f(\mathbf{y}|\boldsymbol{\theta})$ from the simulator and $\boldsymbol{\theta}_t \sim p_{t|0}(\boldsymbol{\theta}_t|\boldsymbol{\theta}_0)$ from the forward noising process. We can thus form a Monte Carlo estimator of the objective.

⁹Sharrock, Simons, Liu, Beaumont. (2023). Sequential Neural Score Estimation: Likelihood-Free Inference with Conditional Score Based Diffusion Models. AABI.

Conditional score estimation

As before, we approximate this intractable quantity using a neural network

$$s_{\phi}(\boldsymbol{\theta}_t, \mathbf{y}, t) \approx s(\boldsymbol{\theta}_t, \mathbf{y}, t).$$

This can be done by minimizing⁹ (with respect to ϕ)

$$\int_0^T \mathbb{E}_{p_{t|0}(\boldsymbol{\theta}_t|\boldsymbol{\theta}_0)f(\boldsymbol{\theta}_0,\mathbf{y})} \left[\|s_{\phi}(\boldsymbol{\theta}_t, \mathbf{y}, t) - \nabla \log p_{t|0}(\boldsymbol{\theta}_t|\boldsymbol{\theta}_0)\|^2 \right] dt.$$

The inner expectation only depends on samples $\boldsymbol{\theta}_0 \sim f(\boldsymbol{\theta})$ (the prior), $\mathbf{y} \sim f(\mathbf{y}|\boldsymbol{\theta})$ from the simulator and $\boldsymbol{\theta}_t \sim p_{t|0}(\boldsymbol{\theta}_t|\boldsymbol{\theta}_0)$ from the forward noising process. We can thus form a Monte Carlo estimator of the objective.

⁹Sharrock, Simons, Liu, Beaumont. (2023). Sequential Neural Score Estimation: Likelihood-Free Inference with Conditional Score Based Diffusion Models. AABI.

Conditional score estimation

As before, we approximate this intractable quantity using a neural network

$$s_\phi(\boldsymbol{\theta}_t, \mathbf{y}, t) \approx s(\boldsymbol{\theta}_t, \mathbf{y}, t).$$

This can be done by minimizing⁹ (with respect to ϕ)

$$\int_0^T \mathbb{E}_{p_{t|0}(\boldsymbol{\theta}_t|\boldsymbol{\theta}_0)f(\boldsymbol{\theta}_0,\mathbf{y})} \left[\|s_\phi(\boldsymbol{\theta}_t, \mathbf{y}, t) - \nabla \log p_{t|0}(\boldsymbol{\theta}_t|\boldsymbol{\theta}_0)\|^2 \right] dt.$$

The inner expectation only depends on samples $\boldsymbol{\theta}_0 \sim f(\boldsymbol{\theta})$ (the prior), $\mathbf{y} \sim f(\mathbf{y}|\boldsymbol{\theta})$ from the simulator and $\boldsymbol{\theta}_t \sim p_{t|0}(\boldsymbol{\theta}_t|\boldsymbol{\theta}_0)$ from the forward noising process. We can thus form a Monte Carlo estimator of the objective.

⁹Sharrock, Simons, Liu, Beaumont. (2023). Sequential Neural Score Estimation: Likelihood-Free Inference with Conditional Score Based Diffusion Models. AABI.

Sample generation

We can thus train an estimator of the score $s_{\phi^*}(\boldsymbol{\theta}_t, \mathbf{y}, t)$ and generate samples as follows:

$$\boldsymbol{\theta}_T \sim \mathcal{N}(0, I),$$

recursively simulate

$$\boldsymbol{\theta}_{T-1}, \dots, \boldsymbol{\theta}_1, \boldsymbol{\theta}_0$$

from the reverse time SDE, using $s_{\phi^*}(\boldsymbol{\theta}_t, \mathbf{y}^{\text{obs}}, t)$ in place of the true score.

Augmentation

oooooooooooo
oooooooooooo
oooooooooooo

Gradient-based methods

ooooooo
oooooooooooooooooooo

Current / future directions

o
ooooooo
oooooooooooo●

Generative Modelling

Thank you!