

today: sampling - approximate inference

Approximate inference - sampling

motivation: NP hard to do exact inference in Ising model $\rightarrow$ need approximation

why sampling? $X = (x_1, \dots, x_p)$

a) simulation $X^{(i)} \sim p$

b) approximate marginal $p(x_i)$

$\rightarrow$ special case of expectations

consider $f: \mathbb{R}^p \rightarrow \mathbb{R}$

we want approximate $\mu = \mathbb{E}_p[f(X)]$

special case: if $f(x) \triangleq \mathbb{1}\{x_A = x_A\} \Rightarrow \mathbb{E}_p[f(X)] = p(x_A = x_A)$

Monte Carlo integration / estimation $\rightarrow$ appears in physics, applied math, ML, statistics

to approximate $\mu = \mathbb{E}_p[f(X)]$

MC estimation: alg:

- n samples $X^{(i)} \stackrel{iid}{\sim} p$
- estimate $\hat{\mu} = \frac{1}{n} \sum_{i=1}^n f(X^{(i)}) = \mathbb{E}_{\hat{p}_n}[f(X)]$

$$\hat{p}_n(x) = \frac{1}{n} \sum_{i=1}^n \mathbb{1}\{x^{(i)} = x\}$$

properties:

1) unbiased estimator $\mathbb{E}[\hat{\mu}] = \mathbb{E}\left[\frac{1}{n} \sum_{i=1}^n f(X^{(i)})\right] = \frac{1}{n} \sum_{i=1}^n \mathbb{E}[f(X^{(i)})] = \frac{1}{n} \sum_{i=1}^n \mu = \mu$

expectation over $(X^{(i)})_{i=1}^n$

μ

this is still true even if $X^{(i)}$'s are dependent

2) expected error (l₂-error) $\mathbb{E}[\|\mu - \hat{\mu}\|_2^2] = \mathbb{E}\left[\left\langle \frac{1}{n} \sum_{i=1}^n f(X^{(i)}) - \mu, \frac{1}{n} \sum_{j=1}^n f(X^{(j)}) - \mu \right\rangle\right]$

$\| \text{tr}(\text{cov}(\hat{\mu}, \hat{\mu})) \|$

$= \mathbb{E}\left[\frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n \langle f(X^{(i)}) - \mu, f(X^{(j)}) - \mu \rangle\right]$

by independence $\Rightarrow$ off diagonal term are zero

i.e. $\mathbb{E}\langle f(X^{(i)}) - \mu, f(X^{(j)}) - \mu \rangle$

$(i \neq j) \text{ terms } \rightarrow \mathbb{E}[f(X^{(i)}) - \mu] \mathbb{E}[f(X^{(j)}) - \mu] = 0$

$$\begin{aligned} & \text{for } (i=j) \text{ terms} \\ & \rightarrow \frac{1}{n^2} \sum_{i,j} \mathbb{E} [\langle f(x^{(i)}) - \mu, f(x^{(j)}) - \mu \rangle] = \frac{A\sigma^2}{n^2} \\ & \quad \mathbb{E} [\|f(x^{(i)}) - \mu\|^2] \triangleq \sigma^2 \\ & \quad \downarrow \\ & \quad \text{tr}(\text{cov}(f(x), f(x))) = \frac{\sigma^2}{n} \end{aligned}$$

$$\mathbb{E} [\|\hat{\mu} - \mu\|^2] = \frac{\sigma^2}{n}$$

note: there is no explicit dimension in rate

(apart σ^2 which could depend implicitly)

e.g. $f(x) = x$
 $x_j \sim \mathcal{N}(0, \sigma^2)$

$$\mathbb{E} [\|f(x) - \mu\|^2] = p\sigma^2$$

How to sample?

- 1) $X \sim \text{Unif}([0,1]) \rightarrow$ pseudo-random generator "rand"
- 2) $X \sim \text{Bernoulli}(p)$ $X = \mathbb{1}\{U \leq p\}$ where $U \sim \text{Unif}([0,1])$
- 3) inverse transform sampling trick

Let F be target c.d.f. of dist p for X $F(x) \triangleq \mathbb{P}\{X \leq x\}$

(first, suppose F is invertible)

Let $X = F^{-1}(U)$ with $U \sim \text{Unif}([0,1])$

claim is that X has cdf $F(x)$ F is monotone

$$\begin{aligned} \text{proof: } \mathbb{P}\{X \leq y\} &= \mathbb{P}\{F^{-1}(U) \leq y\} \stackrel{F \text{ is monotone}}{=} \mathbb{P}\{F(F^{-1}(U)) \leq F(y)\} \\ &\stackrel{F \text{ is invertible}}{=} \mathbb{P}\{U \leq F(y)\} = F(y) \end{aligned}$$

[if F is not invertible, define $X \triangleq \min \{x : F(x) \geq U\}$]

(recall that F is cdf from right)

example:

want $X \sim \text{Exp}(\lambda)$

$$\text{density } p(x) = \lambda \exp(-\lambda x) \mathbb{1}_{\mathbb{R}^+}(x)$$

$$F(x) = 1 - \exp(-\lambda x)$$

$$\text{inverse } F^{-1}(u) = -\frac{1}{\lambda} \log(1-u)$$

multivariate distribution?

can generalise above trick using "chain rule"

$$X_i: p \text{ (dim } p)$$

$X_{1:p}$ (dim p)

from $p(x_{1:p}) = \prod_{i=1}^p p(x_i | x_{1:i-1})$

use cdf for this conditional

"conditional density same" for $X_{1:i-1} = x_{1:i-1}$

$$F_{X_i | X_{1:i-1}}(x_i | x_{1:i-1}) \triangleq \mathbb{P}\{X_i \leq x_i | X_{1:i-1} = x_{1:i-1}\}$$

$$\triangleq \int_{-\infty}^{x_i} p(x_i' | x_{1:i-1}) dx_i'$$

could use $U_1, \dots, U_p \stackrel{iid.}{\sim} \text{Unif}([0,1])$

$$X_1 = F_{X_1}^{-1}(U_1)$$

$$X_2 = F_{X_2 | X_1}^{-1}(U_2 | X_1)$$

$$\vdots$$

$$X_p = F_{X_p | X_{1:p-1}}^{-1}(U_p | X_{1:p-1})$$

inverse of $F_{X_{i+1}}(\cdot | X_i)$

is a very complicated function

(curse of dimensionality)

[aside: "copulas" → model for multivariate data with uniform marginals]

exception is multivariate Gaussian

$$N(\mu, \Sigma) \quad \Sigma = U \Lambda U^T$$

(where $U U^T = I_p$
 Λ is diagonal)

(Cholesky decomposition)
 $\Sigma = L L^T$

generate $V \sim N(0, I_p)$

($V_p \stackrel{iid.}{\sim} N(0,1)$)

$$X = U \Lambda^{1/2} V + \mu$$

$$\mathbb{E}X = \mu$$

$$\text{cov}(X) = U \Lambda^{1/2} \underbrace{\text{cov}(V)}_{I_p} \Lambda^{1/2} U^T = \Sigma$$

Box-Muller transformation to sample (2d) Gaussian

$$R^2 \sim \text{Exp}(1)$$

$$\Rightarrow X \triangleq R \cos \theta$$

$$Y \triangleq R \sin \theta$$

$$\begin{pmatrix} X \\ Y \end{pmatrix} \sim N(0, I)$$

$$\theta \sim \text{Unif}([0, 2\pi])$$

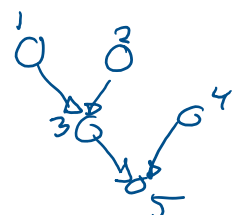
13h34

sampling from a DGM is (relatively) easy e. ancestral sampling

$(x_1, \dots, x_d) \sim p \in \mathcal{S}(G)$ where G is a DAG

$$p(x_1, \dots, x_d) = \prod_{i=1}^d p(x_i | x_{\pi_i})$$

suppose WHOG $1, \dots, d$ is a top-sort of G



ancestral sampling:

for $1, \dots, d$ sample $x_i \sim p(x_i = \cdot | x_{\pi_i})$

these are already sampled by kn. set ancestor

for $i=1, \dots, d$ sample $x_i \sim p(x_i = \cdot | \bar{x}_{-i})$ these are already sampled by bp. ext. property
end

can show by induction $(x_1, \dots, x_d)$ has dist. p

⊛ important side note: when you sample from a joint

you are also sampling from "marginals" by just ignoring joint aspect

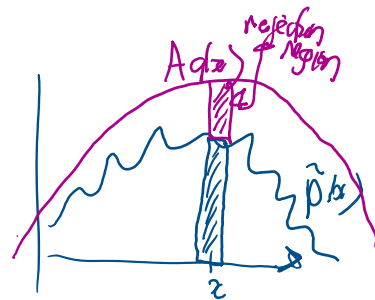
i.e. $(x, y) \sim p(x, y)$ then look at x by itself
 $x \sim p(x)$

rejection sampling:

say $p(x) = \frac{\tilde{p}(x)}{Z_{\tilde{p}}}$; let's say can form a $q(x)$ "proposal" that is easy to sample from

$$Z_{\tilde{p}} \cong \int_x \tilde{p}(x) dx$$

$$s.t. A q(x) \geq \tilde{p}(x) \quad \forall x$$



rule:

- sample $X \sim q(x)$
- Accept with prob. $\frac{\tilde{p}(x)}{A q(x)} \in [0, 1]$
- o.w. reject $\rightarrow$ start again

let's show that accepted samples have correct dist.

(say x discrete)

$$\underbrace{P\{X=x, X \text{ is accepted}\}}_{\text{sampling mechanism}} = \underbrace{P\{X \text{ is accepted} | X=x\}}_{\frac{\tilde{p}(x)}{A q(x)}} \underbrace{P\{X=x\}}_{q(x)} = \frac{\tilde{p}(x)}{A}$$

$$P\{X \text{ is accepted}\} = \sum_x \frac{\tilde{p}(x)}{A} = \frac{Z_{\tilde{p}}}{A}$$

$$P\{X=x | X \text{ is accepted}\} = \frac{\tilde{p}(x)/A}{Z_{\tilde{p}}/A} = \frac{\tilde{p}(x)}{Z_{\tilde{p}}} = p(x)$$

(marginal prob. of acceptance $\rightarrow$ want this to be high)

application to conditioning in a DGM

say want to sample $p(x | \bar{x}_{\mathcal{E}})$

here could use $\tilde{p}(x) = p(x_E, x_{E^c}) \delta(x_E, \bar{x}_E) \Rightarrow p(x) = p(x_E | \bar{x}_E) \delta(x_E, \bar{x}_E)$

let $q(x)$ be original joint in DGM (i.e. $p(x) = q(x)$)

↳ can sample using ancestral sampling

$$q(x) = p_1(x_E, x_{E^c})$$

$$q(x) \geq \tilde{p}(x) \quad \forall x \quad [\text{take } A=1]$$

$$\text{acceptance prob.} : \frac{\tilde{p}(x)}{A q(x)} = \delta(x_E, \bar{x}_E)$$

alg.:

- do ancestral sampling
- accept if $x_E = \bar{x}_E$
- else reject

(rejection sampling for DGM conditionals)

$$P\{\text{accept}\} = \frac{\sum_x \tilde{p}(x)}{A} = p(\bar{x}_E)$$

Importance sampling

in context of computing $\mathbb{E}_p[f(x)] = \mu$ $x \sim p$
 $\leadsto$ can "weight" sample $x^{(i)}$

$$\mathbb{E}_p[f(x)] = \sum_x f(x) p(x) = \sum_x f(x) \frac{p(x)}{q(x)} q(x) \quad \text{for some dist } q$$

$$= \mathbb{E}_q[f(y) \frac{p(y)}{q(y)}] \quad \text{where } y \sim q$$

$$\approx \frac{1}{n} \sum_{i=1}^n g(y^{(i)}) \quad \text{where } y^{(i)} \stackrel{i.i.d.}{\sim} q$$

$$\text{and } g(y) \triangleq f(y) w(y)$$

$$\text{where } w(y) \triangleq \frac{p(y)}{q(y)} \quad \text{"weights"}$$

$$\hat{\mu}_{\text{I.S.}} = \frac{1}{n} \sum_{i=1}^n f(y_i) w_i \quad y_i \stackrel{i.i.d.}{\sim} q$$

$w_i \triangleq \frac{p(y_i)}{q(y_i)}$

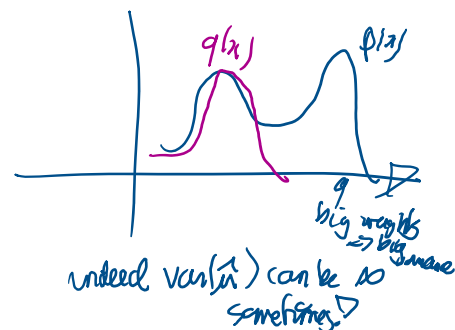
"importance weights"

$$\mathbb{E}[\hat{\mu}_{\text{I.S.}}] = \mu$$

$$\text{Var}[\hat{\mu}] = \frac{1}{n} \left[\mathbb{E}_p[f(x)^2 \frac{p(x)}{q(x)}] - \mu^2 \right]$$

issues when q is small and p is large

intuitively, you want $q(x) \propto f(x) p(x)$



extension to un-normalized dist:

$$p(x) \propto \exp(-\sum_i \phi_i(x_i)) \quad q(x) \propto \exp(-\sum_i \psi_i(x_i))$$

extension to un-normalized dist:

$$p(x) = \frac{\tilde{p}(x)}{Z_p} \quad q(x) = \frac{\tilde{q}(x)}{Z_q}$$

$$\begin{aligned} \mu &= \mathbb{E}_q \left[f(y) \frac{p(y)}{q(y)} \right] \\ &= \mathbb{E}_q \left[f(y) \frac{\tilde{p}(y)}{\tilde{q}(y)} \right] \cdot \frac{Z_q}{Z_p} \end{aligned}$$

$$\begin{aligned} \text{estimate } \frac{Z_p}{Z_q} \text{ with } \hat{Z}_{p/q} &\leq \frac{1}{n} \sum_{i=1}^n \frac{\tilde{p}(y_i)}{\tilde{q}(y_i)} \\ &= \frac{1}{n} \sum_{i=1}^n w_i \end{aligned}$$

$$\hat{\mu}_{UIS} = \frac{\frac{1}{n} \sum_{i=1}^n f(y_i) w_i}{\frac{1}{n} \sum_{i=1}^n w_i} \quad \begin{array}{l} x_i \sim q \\ w_i = \frac{\tilde{p}(y_i)}{\tilde{q}(y_i)} \end{array}$$

note: $\hat{\mu}_{UIS}$ is (slightly biased), but asymptotically unbiased $n \rightarrow \infty$

• this estimator has often lower variance than $\hat{\mu}_{IS}$ even when $Z_p = Z_q = 1$
(normalize 'stabilizes' estimator new weights $\tilde{w}_i = \frac{w_i}{\sum_{j=1}^n w_j} \in [0, 1]$)

see 2017 notes for
• variance reduction (link with SAGA)
• Rao-Blackwellization

Good reference on sampling:
Monte Carlo Statistical Methods, Robert & Casella, 2004