

Bayesian Learning

Video Lecture 0 - Monte Carlo methods

Mattias Villani

Department of Statistics
Stockholm University



mattiasvillani.com



@matvil



@matvil



mattiasvillani

Monte Carlo simulation

- Interest: **expectations of functions** $f(X)$

$$\mathbb{E}_p(f(X)) = \begin{cases} \sum_{\text{all } x} f(x) \cdot p(x) & \text{for discrete } X \\ \int f(x) \cdot p(x) \, dx & \text{for continuous } X \end{cases}$$

- Examples:

- ▶ **Mean** $\mathbb{E}(X)$

$$f(x) = x$$

- ▶ **Variance** $\mathbb{V}(X)$

$$f(x) = (x - \mu)^2$$

- ▶ **Probabilities** $\Pr(a \leq X \leq b)$

$$f(x) = \begin{cases} 1 & \text{if } a \leq x \leq b \\ 0 & \text{otherwise} \end{cases}$$

Properties of sample averages

- **Law of Large Numbers (LLN)**: sample means

$\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i$ converges to population means $\mu = \mathbb{E}_p(X)$

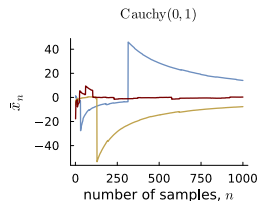
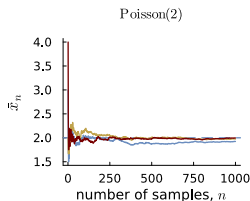
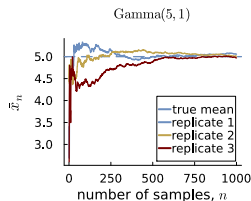
$$\bar{X}_n \xrightarrow{P} \mathbb{E}_p(X)$$

where $X_1, \dots, X_n$ are iid from $p(x)$.

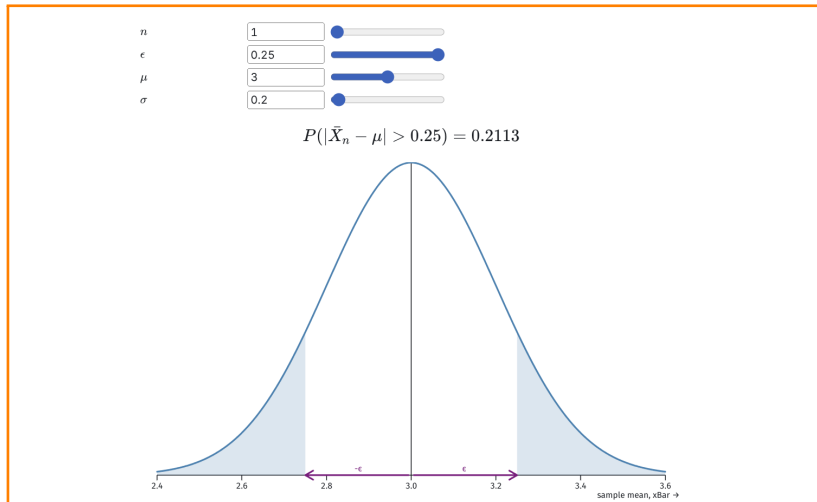
- **Central limit theorem (CLT)**: informally, for large n

$$\bar{X}_n \stackrel{\text{approx}}{\sim} N\left(\mu, \frac{\sigma^2}{n}\right),$$

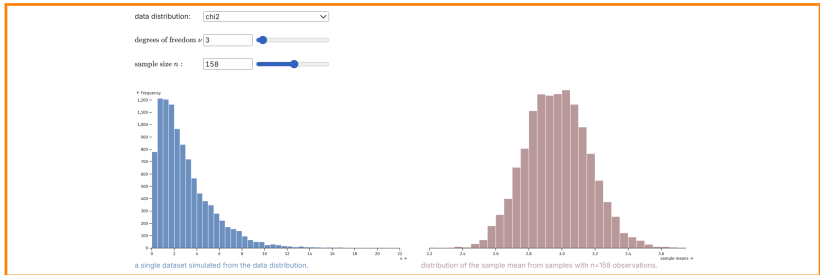
where $\sigma^2 = \mathbb{V}(X)$.



Law of large numbers - widget



Central limit theorem - widget



Sample averages of functions

■ LLN for a function $f(x)$

$$\bar{f}_n = \frac{1}{n} \sum_{i=1}^n f(X_i) \xrightarrow{p} \mathbb{E}_p(f(X))$$

where $X_1, \dots, X_n$ are iid from $p(x)$.

■ CLT for a function $f(x)$

$$\bar{f}_n \stackrel{\text{approx}}{\sim} N\left(\mu_f, \frac{\sigma_f^2}{n}\right),$$

where $\mu_f = \mathbb{E}_p(f(X))$ and $\sigma_f^2 = \mathbb{V}_p(f(X))$.

■ We can estimate σ_f^2 by sample variance of function evaluations

$$\hat{\sigma}_f^2 = \frac{1}{n-1} \sum_{i=1}^n (f(X_i) - \bar{f}_n)^2$$

Monte Carlo simulation

■ Monte Carlo integration

$$\hat{\mu}_f = \frac{1}{m} \sum_{i=1}^m f(x_i), \quad x_i \stackrel{\text{iid}}{\sim} p(x)$$

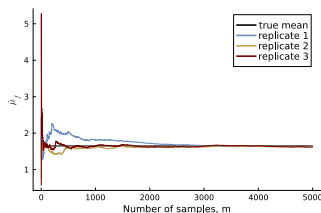
■ **Unbiased** $\mathbb{E}_p(\hat{\mu}_f) = \mu_f$ for all m .

■ **LLN: consistent** $\hat{\mu}_f \xrightarrow{P} \mu_f$ as $m \rightarrow \infty$.

■ **CLT**: quantifies **estimation error** $\hat{\mu}_f - \mu_f$ for finite m .

■ Example: $f(x) = \exp(x)$ and $X \sim N(0, 1)$.

True mean: $\mathbb{E}_p(\exp(X)) = \exp(1/2) \approx 1.648$



Monte Carlo simulation - multivariate

- Two extensions:
 - ▶ Multivariate input: $\mathbf{x}$
 - ▶ Multivariate output: $\mathbf{f}(\mathbf{x})$
- Example: $\mathbf{x} \sim N(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ and $\mathbf{f}(\mathbf{x}) = \mathbf{x}$.
- **Monte Carlo integration**

$$\hat{\boldsymbol{\mu}}_f = \frac{1}{m} \sum_{i=1}^m \mathbf{f}(\mathbf{x}_i), \quad \mathbf{x}_i \stackrel{\text{iid}}{\sim} p(\mathbf{x})$$

- Still unbiased and consistent.
- Multivariate **CLT**:

$$\hat{\boldsymbol{\mu}}_f \stackrel{\text{approx}}{\sim} N\left(\boldsymbol{\mu}_f, \frac{1}{m} \mathbb{V}_p(\mathbf{f}(X))\right),$$

where $\mathbb{V}_p(\mathbf{f}(X))$ is a $d \times d$ covariance matrix.

- Same $\frac{1}{m}$ rate in multivariate case. Works in high dimensions!

Importance sampling

- Aim:

$$\mathbb{E}_p(f(X)) = \int f(x) p(x) dx$$

- **Importance sampling (IS)** samples from a **proposal distribution** $q(x)$ instead of target distribution $p(x)$.

- Based on the trivial identity

$$\int f(x) p(x) dx = \int f(x) \frac{p(x)}{q(x)} q(x) dx = \int f(x) w(x) q(x) dx$$

- **Weight function** (importance function)

$$w(x) = \frac{p(x)}{q(x)}$$

- **Importance sampling estimator**

$$\hat{\mu}_f^{\text{IS}} = \frac{1}{m} \sum_{i=1}^m f(x_i) w(x_i), \quad x_i \stackrel{\text{iid}}{\sim} q(x)$$

Importance sampling

■ Importance sampling estimator

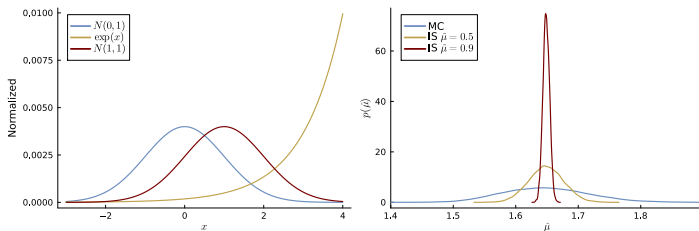
$$\hat{\mu}_f^{\text{IS}} = \frac{1}{m} \sum_{i=1}^m f(x_i) w(x_i), \quad x_i \stackrel{\text{iid}}{\sim} q(x)$$

■ **Unbiased** $\mathbb{E}_q(\hat{\mu}_f^{\text{IS}}) = \mu_f$ and **consistent** $\hat{\mu}_f^{\text{IS}} \xrightarrow{P} \mu_f$.

■ The proposal distribution determines the efficiency (variance).

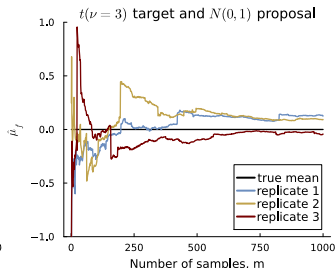
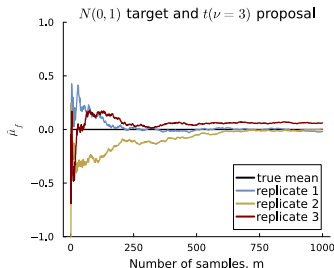
■ Proposal $q(x)$ proportional to integrand $f(x) p(x)$ is ideal.

■ Example: $f(x) = \exp(x)$ and $X \sim N(0, 1)$. Proposal: $N(\tilde{\mu}, 1)$.



The tails of the proposal are important

- **Proposal** $q(x)$ should have **heavier tails** than target $p(x)$.
- Example 1: (OK!):
 - ▶ Target: $N(0, 1)$
 - ▶ Proposal: Student- t with three df.
- Example 2 (not OK!):
 - ▶ Target: Student- t with three df.
 - ▶ Proposal: $N(0, 1)$



Rejection sampling

- Rejection sampling samples from the target distribution $p(x)$.
- Draws from proposal $q(x) \neq p(x)$ and then accepts/rejects draws.
- Need to find a majorization constant

$$p(x) \leq M \cdot q(x) \quad \text{for all } x$$

Rejection sampling:

- 1 Draw a candidate x^* from the proposal distribution $q(x)$.
- 2 Draw $u \sim \text{Uniform}(0, 1)$.
- 3 Accept the candidate x^* if

$$u < \frac{p(x^*)}{M \cdot q(x^*)}$$

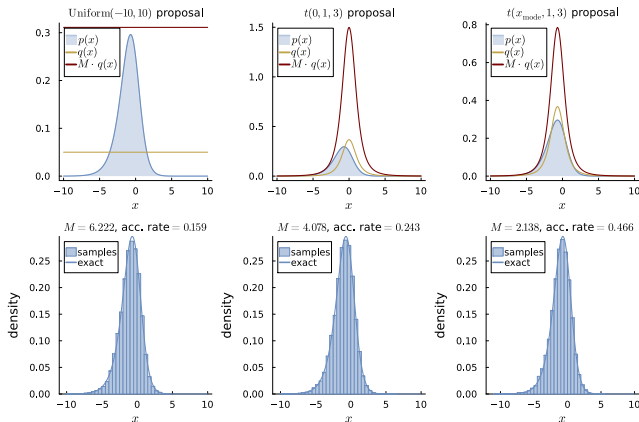
Otherwise, reject the candidate and return to step 1.

Rejection sampling - example

■ Simulation from Logistic-Beta density

$$p(x) = \frac{\sigma(x)^\alpha \sigma(-x)^\beta}{B(\alpha, \beta)}$$

where $\sigma(x) = 1/(1 + e^{-x})$ is the logistic function.



Unnormalized target density

- Both Importance sampling and Rejection sampling works with **unnormalized target density**

$$p(x) = \frac{\tilde{p}(x)}{\int \tilde{p}(x) dx}$$

- The unnormalized density $\tilde{p}(x)$ is known ...
- ... but **normalizing constant** $\int \tilde{p}(x) dx$ cannot be computed.
- Common in **Bayesian learning** where the target is the **posterior distribution**

$$p(\theta|x) = \frac{p(x|\theta)p(\theta)}{\int p(x|\theta)p(\theta) d\theta}$$

- (Much) more later!