

* Overview

- Graphical Models & Markov Properties

- Inference Problems: given $P_G(x)$ find $\begin{cases} P(x_i) \\ \arg \max_x P_G(x) \\ \text{Compute } Z \end{cases}$

- Belief Propagation

- Variational methods

- Gibbs Sampling

- Learning Graphical model

given samples : $x^{(1)}, x^{(2)}, \dots, x^{(N)} \in \mathcal{X}^n$

↳ Learn the structure of the graph G

↳ Learn the parameters of factor.

* Structural Learning

- $X \in \mathbb{R}^{n \times n}$ Random Vector
- Directed Graphical Model: $|E| = \binom{n}{2} = \frac{n(n-1)}{2}$

$$\# \text{ of possible graphs} \leq 3^{|E|} \approx 3^{\frac{n(n-1)}{2}}$$

- Given $x^{(1)}, \dots, x^{(N)}$ indep. samples from unknown $P(x)$
 - How do we "score" each graph?
 - How do we find the graph with highest score?
- There are 2 ways to approach such statistical problems.

Frequentist

- Assume graph G and $P = \{P(x_i | x_{\pi_i})\}_{i=1}^n$ are deterministic but unknown

- Maximum likelihood (ML) estimate find (G, P) that maximizes:

$$\max_{G, P} \underbrace{\sum_{j=1}^N \log P_{G, P}(x^{(j)})}_{\text{SCORE}(G, P)}$$

Bayesian

- Assume (G, P) is randomly drawn from known dist.

$$Q_{G, P}$$

- Maximum A Posterior (MAP) estimate, maximizes:

$$\max_{G, P} P(G, P | x^{(1)}, \dots, x^{(N)})$$

* Frequentist: Structural Learning

$$\hat{G} \triangleq \arg \max_G \left\{ \max_P \underbrace{\frac{1}{N} \sum_{j=1}^N \log P_{G, P}(x^{(j)})}_{\text{SCORE}(G)} \right\}$$

* simple case with $n=2$, $X = (X_1, X_2) \in \{0,1\}^2$

Samples $\begin{matrix} x_1 & x_2 \\ \downarrow & \downarrow \end{matrix}$ (0,0) (0,1) (1,0) (1,1)

$\hat{P}_{x_1, x_2} =$

	x_1	x_2
0	.1	.3
1	.2	.4

Case 1: $G_1 = (\textcircled{1} \textcircled{2})$ independent $\rightarrow P_{12}(x_1, x_2) = P_1(x_1) \cdot P_2(x_2)$

ML estimate:

$$\max_{P_1} \frac{1}{N} \sum_{j=1}^N \log P_1(x_1^{(j)}) + \max_{P_2} \frac{1}{N} \sum_{j=1}^N \log P_2(x_2^{(j)})$$

def. of empirical distribution

$$= \max_{P_1 = (P_1(0), P_1(1))} \hat{P}_1(0) \cdot \log P_1(0) + \hat{P}_1(1) \cdot \log P_1(1) + \dots$$

$$= \max_{P_1} \sum_{x_1} \hat{P}_1(x_1) \cdot \log P_1(x_1) + \dots$$

$$= \max_{P_1} \sum_{x_1} \hat{P}_1(x_1) \log \hat{P}_1(x_1) + \sum_{x_1} \hat{P}_1(x_1) \log \frac{P_1(x_1)}{\hat{P}_1(x_1)} + \dots$$

$$\boxed{-H(\hat{P}_1)}$$

Does not depend P_1

$$-D_{KL}(\hat{P}_1 \parallel P_1) \leq 0$$

max achieved when

$$P_1 = \hat{P}_1, P_2 = \hat{P}_2$$

$$\text{for } G_1 (\textcircled{1} \textcircled{2}) \max_{P_1 \times P_2} \frac{1}{N} \sum_{j=1}^N \log P_1(x_1^{(j)}) \cdot P_2(x_2^{(j)})$$

$$\text{SCORE}(G_1) = -H(\hat{P}_1) - H(\hat{P}_2)$$

Case 2: $G_2 = (\textcircled{1} \rightarrow \textcircled{2})$

$$\max_{P_{12}(x_1, x_2)} \frac{1}{N} \sum_{j=1}^N P_{12}(x_{12}^{(j)}) = -H(\hat{P}_{12}) - \underbrace{D_{KL}(\hat{P}_{12} \parallel P_{12})}_{\text{max achieved at}}$$

$$P_{12} = \hat{P}_{12}$$

$$\text{SCORE}(G_2) = -H(\hat{P}_{12})$$

$$\mathcal{L}(\textcircled{1} \textcircled{2}) = -H(\hat{\beta}_1) - H(\hat{\beta}_2)$$

$$\mathcal{L}(\textcircled{1} \rightarrow \textcircled{2}) = -H(\hat{\beta}_{1,2})$$

follows from Independence
makes entropy ↑

Remark 1: $-H(\hat{\beta}_{1,2}) \geq -H(\hat{\beta}_1) - H(\hat{\beta}_2)$

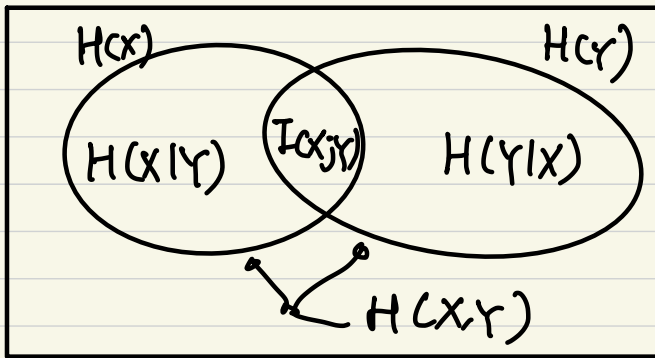
$$H(\beta_{1,2}) \leq H(\beta_1) + H(\beta_2)$$

Refresher on H, I . X, Y random vectors from joint $P_{XY}(x, y)$

$$I(X; Y) \triangleq \sum_{x, y} P(x, y) \cdot \log \frac{P(x, y)}{P(x) \cdot P(y)}$$

$$H(X) \triangleq \sum_x -P(x) \log P(x)$$

$$H(Y|X) \triangleq \sum_{x, y} -P(x, y) \log P(y|x)$$



$$H(Y) = H(Y|X) + I(X; Y)$$

$$H(X) + H(Y) = H(X, Y) + I(X; Y)$$

$$\mathcal{L}(\textcircled{1} \rightarrow \textcircled{2}) \geq \mathcal{L}(\textcircled{1} \textcircled{2})$$

$$-H(\hat{\beta}_{1,2}) \geq -H(\hat{\beta}_1) - H(\hat{\beta}_2)$$

always get the densest graph = Overfitting

↑ even if the truth is $\textcircled{1} \textcircled{2}$,
you choose $\textcircled{1} \rightarrow \textcircled{2}$

Remark 2. depending on N and target false positive rate β
decision is made by:

$$\mathcal{L}(\textcircled{1} \rightarrow \textcircled{2}) - \mathcal{L}(\textcircled{1} \textcircled{2}) = -H(\hat{\beta}_{1,2}) + H(\hat{\beta}_1) + H(\hat{\beta}_2)$$

$$= +I_{\hat{\beta}_{1,2}}(X_1; X_2)$$

$$\text{output} \begin{cases} \textcircled{1} \textcircled{2} & \text{if } I_{\hat{\beta}_{1,2}}(X_1; X_2) > \frac{\epsilon \beta}{N} \\ \textcircled{1} \rightarrow \textcircled{2} & \text{else} \end{cases}$$

* Maximum Likelihood Approach for a DAG

$$G^* \in \arg\max_G \left\{ \max_{\{P(X_i | X_{\pi_i})\}_{i=1}^n} \frac{1}{N} \sum_{j=1}^N \log \prod_{i=1}^n P_i(x_i^{(j)} | x_{\pi_i}^{(j)}) \right\}$$

$$P_i(x_i | X_{\pi_i}) = \hat{p}_i(x_i | X_{\pi_i})$$

SCORE(G)

$$\Rightarrow \frac{1}{N} \sum_{j=1}^N \log \left(\prod_{i=1}^n \hat{p}_i(x_i^{(j)} | x_{\pi_i}^{(j)}) \right)$$

$$= \sum_{i=1}^n \frac{1}{N} \sum_{j=1}^N \log \hat{p}_i(x_i^{(j)} | x_{\pi_i}^{(j)})$$

def empirical distribution $\rightarrow$

$$= \sum_{i=1}^n \sum_{x_i, x_{\pi_i}} \hat{p}_i(x_i, x_{\pi_i}) \cdot \log \hat{p}_i(x_i | x_{\pi_i})$$

$$= \sum_{i=1}^n \{ -H_{\hat{p}}(X_i | X_{\pi_i}) \} \quad \leftarrow H(X|Y) = H(X) - I(X;Y)$$

$$= \sum_{i=1}^n \underbrace{I_{\hat{p}}(X_i | X_{\pi_i})}_{\substack{\uparrow \\ \text{depends on} \\ \text{Graph}}} - \underbrace{H_{\hat{p}}(X_i)}_{\text{does not depend on } G}$$

* Remark: this gives a "SCORE" = likelihood for any given DAG G.

Search over all DAGs.

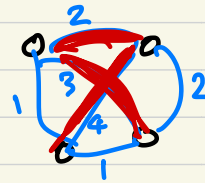
$$I_{\hat{p}}(X_i | X_{\pi_i}) \geq I_{\hat{p}}(X_i | X_{\pi'_i}) \quad \text{if } \pi_i \supset \pi'_i$$

$\Rightarrow$ we need to restrict the family of graphs we are searching over. = set of trees

* Chow-Liu algorithm: searches over all Trees, efficiently, exactly.

Step 1: Create a complete graph (Undirected)
over $V = \{1, 2, \dots, n\}$
give edge weight $w_{ij} = I_{\hat{p}}(x_i; x_j)$

$$W = \begin{bmatrix} w_{ij} \end{bmatrix}$$

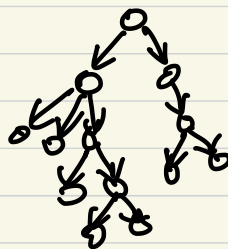


Step 2: find the max-weight spanning Tree, for example,
using Kruskal's algorithm.

Claim: Chow-Liu finds the optimal tree that maximizes

$$\max_{G \in \text{Trees}} \sum_{i=1}^n I_{\hat{p}}(x_i; x_{\pi_i})$$

↑ single node.



* Impractical approach for undirected graphical model learning
 consider learning an Ising Model $(G, \{\theta_i\}_{i=1}^n, \{\theta_{ij}\}_{i,j=1}^n)$

$$P_{G,\theta}(x) = \frac{1}{Z_\theta} \cdot \prod_{i=1}^n e^{\theta_i x_i} \cdot \prod_{(i,j) \in E} e^{\theta_{ij} x_i x_j}$$

$x_i \in \{-1, +1\}$
 $\theta = \begin{bmatrix} \theta_1 & & \theta_{1j} \\ & \ddots & \\ \theta_{ij} & & \theta_n \end{bmatrix}$

samples $\{x^{(1)}, x^{(2)}, \dots, x^{(N)}\} = D$

The likelihood

$$P_{G,\theta}(D) = \prod_{l=1}^N P_{G,\theta}(x^{(l)})$$

$$= \prod_{l=1}^N \frac{1}{Z_\theta} \prod_{i \in V} e^{\theta_i x_i^{(l)}} \cdot \prod_{(i,j) \in E} e^{\theta_{ij} x_i^{(l)} x_j^{(l)}}$$

$$= \exp \left\{ -N \log Z_\theta + \sum_{i \in V} N \cdot \hat{M}_{ii} \cdot \theta_i + \sum_{(i,j) \in E} N \cdot \hat{M}_{ij} \cdot \theta_{ij} \right\}$$

$$\hat{M}_{ii} \triangleq \frac{1}{N} \sum_{l=1}^N x_i^{(l)}$$

$$\hat{M}_{ij} \triangleq \frac{1}{N} \sum_{l=1}^N x_i^{(l)} x_j^{(l)}$$

the log-likelihood is

$$\min L(G, \theta, D) = -\frac{1}{N} \log P_{G,\theta}(D)$$

$$= \underbrace{\frac{\Phi(\theta)}{\log Z_\theta}}_{\text{convex}}$$

$$- \langle \hat{M}, \theta \rangle \cdot N$$

$$\begin{bmatrix} \hat{M}_{ii} & \hat{M}_{ij} \end{bmatrix} \begin{bmatrix} \theta_i & \theta_{ij} \end{bmatrix}$$

linear in θ

Remark:

involves exp-time

* learning is easier if inference is easier.