

* Structural Learning Recap.

Approach 1: ML for DAG

$$\text{SCORE}(G) = \sum_{i=1}^n I_{\hat{P}}(X_i; X_{\pi_i}) \leftarrow \{P(X_i | X_{\pi_i})\}$$

[Chow-Liu Algorithm] $\leftrightarrow$ Max-weight Spanning
If search over all trees, then optimization
can be solved efficiently exactly.

Approach 2. ML for Ising models

$$\min_{\theta \in \mathbb{R}^{n \times n}} \log \tilde{Z}_G(\theta) - \langle \hat{M}, \theta \rangle + \lambda \|\theta\|_{L_1}$$

$\sum_{i,j} |\theta_{ij}|$

$\begin{bmatrix} \hat{M}_{11} & \hat{M}_{12} & \dots \end{bmatrix}$
 $\uparrow \quad \uparrow$
 $\frac{1}{N} \sum_{j=1}^N X_1^{(j)} \quad \frac{1}{N} \sum_{j=1}^N X_1^{(j)} X_2^{(j)}$

$\begin{bmatrix} \theta_{11} & \theta_{12} & \dots \\ & \ddots & \ddots \end{bmatrix}$

complex to evaluated

Q. Can we design an algorithm that is efficient & works well in Practice?

Binary random variables $X = \{0, 1\}$

undirected graphical models

$$P(x_1, \dots, x_n) = \frac{1}{Z(\theta)} \exp \left\{ \sum_{i=1}^n \theta_{i1} x_i + \sum_{(i,j) \in E} \theta_{ij} x_i x_j \right\}$$

Strategy is to predict x_i using $X_{-i} \triangleq \{x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_n\}$
subset of this set

Claim: A conditional distribution of x_1 given the rest is

$$(*) \quad \frac{P(x_i=1 | x_2 \dots x_n)}{P(x_i=0 | x_2 \dots x_n)} = \exp \left\{ \theta_{i1} + \sum_{j \in \partial i} \theta_{ij} x_j \right\}$$

this is a logistic regression problem,

to predict x_i from x_2^n , can be solved very efficiently.

$$(*) \rightarrow P(x_i=1 | x_1) = \frac{e^{\theta_{i1} + \sum_j \theta_{ij} x_j}}{1 + e^{\theta_{i1} + \sum_j \theta_{ij} x_j}}$$

$$= \text{Sigmoid} \left(\underbrace{\theta_{i1} + \sum_j \theta_{ij} x_j}_{\text{linear in } \theta_1} \right)$$

Proof (*) >

$$P(x_i=1 | x_2^n) = \frac{P(x_i=1 \wedge x_2^n)}{P(x_2^n)}$$

$$= \frac{1}{Z(\theta)} \cdot \exp \left\{ \theta_{i1} + \sum_{j=2}^n \theta_{ij} x_j + \sum_{\substack{j=2 \\ j \in \partial i}}^n \theta_{ij} x_j \cdot 1 + \sum_{\substack{(i,j) \in E \\ j \neq i}} \theta_{ij} x_i x_j \right\}$$

$$P(x_i=0 | x_2^n) =$$

$$P(X_1=1|X_2^u) = \frac{\exp\{\theta_1 + \sum_j X_j \theta_j\}}{1 + \exp\{\theta_1 + \sum_j X_j \theta_j\}}$$

* Review of logistic regression

• We are given labelled examples of features $\mathbf{z}^{(l)} = (\mathbf{z}_1^{(l)}, \dots, \mathbf{z}_L^{(l)})$

Data = $(\mathbf{z}^{(l)}, y^{(l)}) \}_{l=1}^N$ label $y^{(l)} \in \{0, 1\}$

• Suppose a parametric model with $\mathbf{w} \in \mathbb{R}^L$

$$P(y=1|\mathbf{z}) = \frac{\exp(\sum_{k=1}^L w_k z_k)}{1 + \exp(\sum_{k=1}^L w_k z_k)}$$

$$P(y=0|\mathbf{z}) = \frac{1}{1 + \exp(\sum_{k=1}^L w_k z_k)}$$

• Maximum likelihood

$$\log\text{-likelihood } \mathcal{L}(\{(\mathbf{z}^{(l)}, y^{(l)})\}_{l=1}^N, \mathbf{w} \in \mathbb{R}^L) \triangleq \sum_{l=1}^N \log P_{\mathbf{w}}(y^{(l)}|\mathbf{z}^{(l)})$$

$$= \frac{1}{N} \sum_{l=1}^N \left\{ y^{(l)} \cdot P_{\mathbf{w}}(y=1|\mathbf{z}^{(l)}) + (1-y^{(l)}) \cdot P_{\mathbf{w}}(y=0|\mathbf{z}^{(l)}) \right\}$$

$$= \frac{1}{N} \sum_{l=1}^N \left\{ y^{(l)} \cdot \left(\sum_{k=1}^L w_k z_k^{(l)} \right) - \log \left(1 + \exp \left(\sum_{k=1}^L w_k z_k^{(l)} \right) \right) \right\}$$

utility/objective over $\mathbf{w}$

* this is a concave maximization over $\mathbf{w} \in \mathbb{R}^L$
use gradient ascent to solve efficiently.

* Logistic Regression for neighborhood selection

Structural Learning.

$$P(X_1=1|X_2^u) = \frac{e^{\theta_{11} + \sum_{j \in \partial 1} \theta_{1j} X_j}}{1 + e^{\theta_{11} + \sum_{j \in \partial 1} \theta_{1j} X_j}}$$

$\uparrow$
 G

Logistic regression

$$P(Y=1|Z) = \frac{e^{\sum_k w_k z_k}}{1 + e^{\sum_k w_k z_k}}$$

for node 1, if we know G , then features are $(1, X_{\partial 1})$
label is X_1 .

as we do not know G , suppose ground truths $|\theta_{1j}| \leq 1$,
and degree $\leq K$
this motivates following **Sparse logistic regression**.

$$\begin{aligned} \min_{\substack{\theta_{1\cdot} \\ \cap \\ \mathbb{R}^n \\ (\theta_{11}, \theta_{12}, \dots, \theta_{1n})}} & -\mathcal{L}(\underbrace{\{ (1, x_2^{(1)}, \dots, x_n^{(1)}) \}}_{\text{feature}}, \underbrace{\{ x_1^{(l)} \}_{l=1}^N}_{\text{label}}, \theta_{1\cdot}) \\ \text{s.t.} & \quad \|\theta_{1\cdot}\|_{L_1} \leq K \quad \text{or equivalently} \\ & \quad \sum_j |\theta_{1j}| \leq K \end{aligned}$$

$$(*) \quad \min_{\theta_{1\cdot}} -\mathcal{L}(\cdot, \theta_{1\cdot}) + \lambda \cdot \|\theta_{1\cdot}\|_{L_1}$$

$$\Rightarrow \theta_{1\cdot} \approx (0, \underbrace{\theta_{1,3}, 0, 0, \theta_{1,18}, \dots}_{\text{Graph structure}})$$

Theorem [Klivans, Meka 2017]

If $\max \text{ degree} \leq k$

$P(X_i = \pm 1 | X_{-i}) \in [\delta, 1-\delta]$ for some $\delta > 0$

for some $\varepsilon > 0$

number of samples $N \geq C \cdot \log n \cdot \frac{1}{\varepsilon^2} \cdot k \cdot C_\delta$

Then $(\hat{\theta})$ achieve $\|\hat{\theta} - \theta^*\|_\infty \leq \varepsilon$

$$\max_{i,j} |\hat{\theta}_{ij} - \theta_{ij}^*| \leq \varepsilon$$

with high probability

and runtime $O(n^2 \text{ poly } \frac{1}{\varepsilon})$

if all $|\theta_{ij}^*| > 2\varepsilon$, then threshold $|\hat{\theta}_{ij}| \geq \varepsilon$

learn structure exactly

*

$$P(X^n) = \frac{1}{Z(\theta)} \exp \left\{ \sum_j \theta_{jj} X_j + \sum_{(i,j) \in E} \theta_{ij} X_i X_j \right\}$$

* This principle can be used for more general graphical models

• Gaussian Graphical Models.

$$P(x) = \frac{1}{Z_J} \exp \left\{ -\frac{1}{2} x^T J x \right\}$$

$$J = \Sigma^{-1}$$

v.l. = 0

$\mathcal{X} = \mathbb{R}$

① Compute conditional distribution

$$P(X_1 | X_2^n) = \frac{1}{Z_1} \cdot \exp \left\{ -\frac{1}{2} (J_{11} x_1^2 + 2 \left(\sum_{j \in \mathcal{N}} J_{1j} X_j \right) \cdot x_1) \right\}$$

$$= \frac{1}{\sqrt{2\pi \cdot \frac{1}{J_{11}}}} \cdot \exp \left\{ -\frac{1}{2} \frac{(x_1 - \frac{\sum_{j \in \mathcal{N}} J_{1j} X_j}{J_{11}})^2}{\frac{1}{J_{11}}} \right\}$$

Neighborhood of node 1

② Apply Maximum Likelihood principle to get x_1 from x_2^n

$$\begin{aligned} \min - \frac{1}{N} \sum_{l=1}^N \log P(x_1^{(l)} | x_2^{(l)} \dots x_n^{(l)}) \\ = \frac{\mathcal{J}_{11}}{2N} \cdot \sum_{l=1}^N \left(\underset{\substack{\uparrow \\ \text{label}}}{x_1^{(l)}} - \sum_{j \neq 1} \underbrace{\left(\frac{\mathcal{J}_{1j}}{-\mathcal{J}_{11}} \right)}_{\substack{\uparrow \\ w_j \\ \text{model parameter}}} \cdot \underset{\substack{\uparrow \\ \text{features}}}{x_j^{(l)}} \right)^2 \quad \leftarrow \text{L}_2 \text{-loss} \\ - \frac{1}{2} \log \mathcal{J}_{11} \end{aligned}$$

as we care about structure, we re-parametrize it by $w \in \mathbb{R}^{n-1}$

$$\min_{w \in \mathbb{R}^{n-1}} \frac{1}{N} \sum_{j=1}^N \left(x_1^{(j)} - \sum_{j=2}^n w_j x_j^{(j)} \right)^2 + \lambda \|w\|_1.$$

* Graphical Lasso.

Motivation: — Get $\mathcal{J}$ directly. without separating each neighborhood.

Step 1. Compute the empirical covariance matrix

$$S_{ij} = \frac{1}{N} \sum_{l=1}^N x_i^{(l)} x_j^{(l)}$$

Step 2. maximize likelihood over information matrix $\mathcal{J} \in \mathbb{R}^{m \times m}$

$$\hat{\mathcal{J}} \in \arg \max_{\mathcal{J}} \left\{ \underbrace{\log |\mathcal{J}|}_{\text{determinant}} - \underbrace{\text{Tr}(S \cdot \mathcal{J})}_{\text{Tr}(A) = \sum_i A_{ii}} - \lambda \underbrace{\|\mathcal{J}\|_1}_{\text{regularizer}} \right\}$$

likelihood.

$$\text{s.t. } \underline{\mathcal{J}} \succeq 0$$

Others: linear regression

→ Semidefinite program.