

today: exp. families
estimation for PGM

Exponential Family

a (flat/canonical) exponential family on $\mathcal{X}$

is a parametric family of dist. on $\mathcal{X}$ defined by two quantities

I) $h(x) d\mu(x) \rightarrow$ reference measure

reference base
density measure

$\rightarrow$ counting measure (discrete R.V.) $\rightarrow \sum_x$ pmf
 $\rightarrow$ Lebesgue " (cts R.V.) $\rightarrow \int_{\mathcal{X}}$ pdf

II) $T: \mathcal{X} \rightarrow \mathbb{R}^p$ called "sufficient statistics" vector
aka. feature vector

members of the family will have pmf/pdf

$$p(x; \eta) d\mu(x) = \exp(\underbrace{\eta^T T(x)}_{\text{"canonical parameter"}} - A(\eta)) \underbrace{h(x) d\mu(x)}_{\text{defining pieces } (+ \Omega_{\mathcal{X}})}$$

log-normalization or cumulant generating functions
log partition fct.

$$\mathbb{E}[f] = \int f(x) d\mu(x) / \int f(x) d\mu(x)$$

• if $\Omega_{\mathcal{X}}$ is discrete, then $p(x; \eta)$ is a pmf

" " cts. ; " " " a pdf

$$\star \text{ want } 1 = \int_{\mathcal{X}} p(x; \eta) d\mu(x) = \int_{\mathcal{X}} \exp(\eta^T T(x)) e^{-A(\eta)} h(x) d\mu(x)$$

$$[\text{discrete: } \sum_x p(x; \eta)]$$

$$\Rightarrow A(\eta) \triangleq \log \left(\underbrace{\int_{\mathcal{X}} \exp(\eta^T T(x)) h(x) d\mu(x)}_{z(\eta)} \right)$$

domain $\Omega \triangleq \{\eta \in \mathbb{R}^p \mid A(\eta) < \infty\}$
set of valid canonical parameters

class of notation:

$$\Omega \neq \Omega_{\mathcal{X}} \subseteq \mathcal{X} \subseteq \mathbb{R}^p$$

note: $A(\eta)$ is convex in $\eta \Rightarrow \Omega$ is convex

⊗ more generally, consider a reparameterization of a subset of the flat exponential family

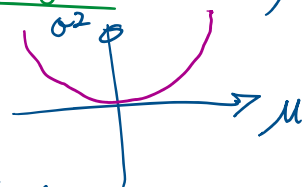
$$h. \dots \eta. \cap \mathbb{R} \rightarrow \cap$$

⊗ more generally, consider a reparameterization of a subset of the flat exponential family by defining $\eta: \Theta \rightarrow \Omega$
new set of parameters

$$p(x; \theta) \triangleq p(x; \eta(\theta)) \text{ for } \theta \in \Theta$$

(get a "curved exponential family" if $\eta(\Theta)$ is curved manifold in Ω)

↳ e.g. could consider Gaussians where $N(\mu, \mu^2)$



*: note: any single dist. $p(x)$ can be part of an exponential family by using $\eta(x) = p(x)$

two examples of family not an exp family: • $\text{unif}(0, \theta)$

• mixture of Gaussians (latent variable model)

Example 1: (multinomial)

$$X \sim \text{Mult}(\pi) \quad X = \{0, 1\}^k$$

$$\Omega_X = \Delta_K \cap X \quad (\text{one hot encodings})$$

parameter $\pi \in \Delta_K$; suppose $\pi_i > 0 \forall i$

$$p(x; \pi) = \prod_{j=1}^k \pi_j^{x_j} = \exp\left(\sum_{j=1}^k x_j \log \pi_j\right)$$

"think as 0"

$$= \exp\left(\sum_{j=1}^k x_j \log \pi_j - 0\right)$$

$x \in X$

$$\text{we have } \eta_j(\pi) = \log \pi_j$$

$$T(x) = x$$

$d\mu(x)$ = counting measure on X

$$\eta(x) = \mathbb{1}\{x \in \Omega_X\} = \mathbb{1}\{x \in \Delta_K \cap X\}$$

$$A(\eta) = 0??$$

$$\Theta = \text{int}(\Delta_K) \quad A(\eta(\pi)) = 0 \quad \forall \pi \in \Theta$$

$$\Theta \rightarrow \dim k-1$$

$$\eta(\Theta) \rightarrow \text{" " " "}$$

$$\Omega_X \rightarrow \mathbb{R}^K$$

we do not have a "minimal exp family"

note: here, for any x s.t. $h(x) \neq 0$

$$\sum_{j=1}^k T_j(x) = \sum_j x_j = 1$$

affine linear dep. between components of T

$\Rightarrow$ multiple η 's give rise to same distribution
 $\rightarrow$ "overparameterization"

$$\text{eg } (n+1)\alpha)^T T(x) = n^T T(x) + \alpha^T T(x)$$

not a minimal exp. family

⊛ for a multinomial; a min. exp family

$$T(x) = \begin{pmatrix} x_1 \\ \vdots \\ x_{k-1} \end{pmatrix} \quad \left["x_k = 1 - \sum_{j=1}^{k-1} x_j" \right]$$

$$Z(\eta) = \sum_{x \in \Omega_X} \exp(\eta^T T(x)) = \sum_{j=1}^{k-1} e^{\eta_j} + 1$$

$$p(x; \eta) = \exp \left(\sum_{j=1}^{k-1} \eta_j x_j - \log \left(1 + \sum_{j=1}^{k-1} e^{\eta_j} \right) \right)$$

recall: $\nabla_{\eta} A(\eta) = \mathbb{E}_{p(x; \eta)} [T(x)]$ (valid for $\eta \in \text{int}(\Omega)$)

for multinomial; $\frac{\partial A}{\partial \eta_j} = \frac{1}{Z(\eta)} e^{\eta_j} = p("x=j" | \eta)$

$= \mathbb{E}_{p(x; \eta)} [T(x)]$
as required //

moment matching can be different than MLE in exp. family (with wrong moments)

gamma dist $\Gamma(\alpha, \beta)$ has $T(x) = \begin{bmatrix} \log x \\ x \end{bmatrix}$

so moment matching with $T(x) = \begin{bmatrix} x \\ x^2 \end{bmatrix}$ will yield a different estimate than MLE

15h28

example 2: 1d Gaussian

$X \sim N(\mu, \sigma^2)$ $X = \mathbb{R}$ $\Theta = (\mu, \sigma^2)$ "moment parametrization"

$$p(x; (\mu, \sigma^2)) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(x-\mu)^2}{2\sigma^2}\right)$$

$$= \exp\left(-\frac{x^2}{2} \left[\frac{1}{\sigma^2}\right] + x \left[\frac{\mu}{\sigma^2}\right] - \left[\frac{\mu^2}{2\sigma^2} + \frac{1}{2} \log(2\pi\sigma^2)\right]\right)$$

$$T(x) = \begin{bmatrix} x \\ -\frac{x^2}{2} \end{bmatrix} \quad \eta(\theta) = \begin{bmatrix} \mu/\sigma^2 \\ 1/\sigma^2 \end{bmatrix} = \begin{bmatrix} \eta_1 \\ \eta_2 \end{bmatrix}$$

$$\eta_2 = \frac{1}{\sigma^2} = \text{precision} > 0$$

$$\eta_1 = \eta_2 \cdot \mu$$

$$\Omega = \{(\eta_1, \eta_2) : \eta_2 > 0, \eta_1 \in \mathbb{R}\}$$

$h(x) = 1$ (but some people use $h(x) = \frac{1}{\sqrt{2\pi}}$ for Gaussian)

[we'll see later: Multivariate Gaussian $T(x) = \begin{bmatrix} x \\ -\frac{xx^T}{2} \end{bmatrix}$ $\eta_1 = \mathbb{E}[\mu] = \mathbb{E}^T \mu$ $\eta_2 = \Sigma = \mathbb{E}^{-1}$]

Example 3: discrete UGM

let $p \in \mathcal{P}(\mathcal{G})$ $\mathcal{G}$ is undirected

with $\psi_c(x_c) > 0 \forall c, x_c$

$$p(x) = \frac{1}{Z} \prod_{c \in \mathcal{C}} \psi_c(x_c) = \exp\left(\sum_c \log \psi_c(x_c) - \log Z\right)$$

$$= \exp\left(\sum_{c \in \mathcal{C}} \sum_{y_c \in \mathcal{X}_c} \underbrace{\mathbb{1}\{y_c = x_c\}}_{T_{c,y_c}(x)} \underbrace{\log \psi_c(y_c)}_{\eta_{c,y_c}} - \log Z\right)$$

$$T(x) = \begin{pmatrix} \vdots \\ \mathbb{1}\{x_c = y_c\} \\ \vdots \end{pmatrix} \leftarrow \begin{matrix} y_c \in \mathcal{X}_c \\ c \in \mathcal{C} \end{matrix}$$

$$\eta(\theta) = \begin{pmatrix} \vdots \\ \log \psi_c(y_c) \\ \vdots \end{pmatrix} \leftarrow \begin{matrix} y_c \in \mathcal{X}_c \\ c \in \mathcal{C} \end{matrix}$$

$$\mathcal{X}_c = \{(y_i)_{i \in \mathcal{C}} : y_i \in \mathcal{X}_i\}$$

[not a minimal representation]

notes: a) Mult(π) is a special case where have complete graph (1 big clique)

notes: a) $\text{Mult}(\pi)$ is a special case where have complete graph (1 big clique)

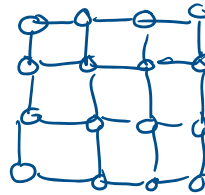
b) feature perspective: instead of using all possible indicators $\mathbb{1}\{y_c = x_c\}$ you could use a subset or a function of them for a task

for example: suppose x is a sentence
 x_i is a word

feature on x_i & x_{i+1} e.g. $\mathbb{1}\{x_i \text{ is a verb, } x_{i+1} \text{ is "noun"}\}$
 $\leadsto$ much smaller set of parameters
 "parameter sharing"

c) binary Ising model

$$x_i \in \{0, 1\} \quad |C| \leq 2$$



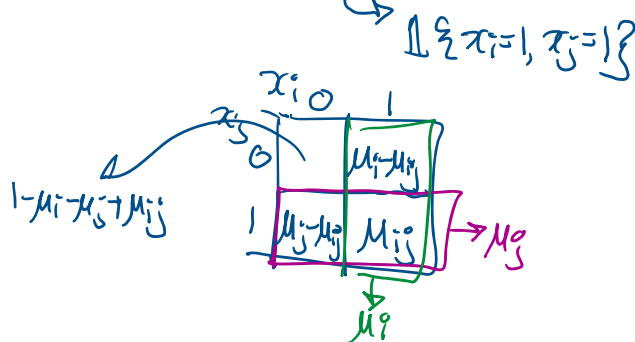
suppose use nodes & pairs (edges)
 as cliques

$\Rightarrow$ dimension of $T(x)$ $2|V| + 4|E| \leadsto$ "overparameterized exp. family"
 $\sum_{y_c} T_{c, y_c}(x) = 1$ for any c
 $\Rightarrow$ not a min. exp. fam.

* a minimal representation

$$T(x) = \begin{pmatrix} (x_i)_{i \in V} \\ (x_i x_j)_{i, j \in E} \end{pmatrix} \begin{matrix} \eta_i \\ \eta_{ij} \end{matrix} \rightarrow \dim |V| + |E|$$

$$\mathbb{E} T(x) = \begin{pmatrix} (\mu_i)_{i \in V} \\ (\mu_{ij})_{i, j \in E} \end{pmatrix} \quad \downarrow \quad P(x_i=1, x_j=1)$$



properties of A (for general Lat exponential family)

properties of A (for generalised exponential family)

• $\nabla_n A(n) = \mathbb{E}_{p(x;n)} [T(x)] \triangleq \mu(n)$ "moment vector" (for $n \in \text{int}(\Omega)$)

• $\left(\frac{\partial^2 A}{\partial n_i \partial n_j} \right)_{(i,j) \in \mathbb{I}^2} = \mathbb{E}_{p(x;n)} [(T(x) - \mu(n)) (T(x) - \mu(n))^T] = \text{cov}(T(x))$
(proof as exercise)

"cumulant generating fct."

Estimation of parameters DOM / UGM

DOM:

(fully observed)

parametric family $\mathcal{P}_{\Theta} = \{ p_{\theta}(x) = \prod_i p(x_i | x_{\pi_i}, \theta_i) : (\theta_1, \dots, \theta_{|V|}) \in \Theta \}$ i.e. varying parameters
 $\Theta = \Theta_1 \times \Theta_2 \times \dots \times \Theta_{|V|}$

$\Rightarrow$ MLE decouple in $|V|$ independent MLE problems

$\{x^{(i)}\}_{i=1}^n$ $p(\text{data} | \theta) = \prod_{i=1}^n p(x^{(i)} | \theta) = \prod_{i=1}^n \prod_{j=1}^{|V|} p(x_j^{(i)} | x_{\pi_j}^{(i)}; \theta_j)$
 $\log(\quad) = \sum_{j=1}^{|V|} \underbrace{\left(\sum_{i=1}^n \log p(x_j^{(i)} | x_{\pi_j}^{(i)}; \theta_j) \right)}_{\sum_i \log p_j}$

example: for discrete P.V. $\Rightarrow \hat{\theta}_j^{\text{MLE}} = \text{proportion of observations}$
 (multinomial conditionals) $\hat{p}(x_j = k | x_{\pi_j} = \text{stuff})$
 $= \frac{\#(x_j = k, x_{\pi_j} = \text{stuff})}{\#(x_{\pi_j} = \text{stuff})}$

(fully observed DOM is relatively easy)
often in closed form!

⊛ if have latent variables (ie unobserved variables)

$\Rightarrow$ use E.M. (like EM)

UGM:

example for exp. family

$p(x; n) = \exp \left(\sum_c \langle n_c, T_c(x_c) \rangle - A(n) \right)$

$$p(x; n) = \exp\left(\sum_c \langle n_c, T_c(x_c) \rangle - A(n)\right)$$

→ unlike in a DGM, $\log p(x; n)$ does not separate $\sum_c f_c(n_c)$

gradient ascent on log likelihood:

$$\frac{1}{n} \sum_{i=1}^n \log p(x^{(i)} | n) = \sum_c n_c^T \left[\underbrace{\frac{1}{n} \sum_{i=1}^n T_c(x_c^{(i)})}_{\hat{\mu}_c} \right] - \frac{n A(n)}{n}$$

$$\nabla_{n_c} [\quad] = \hat{\mu}_c - \mu_c(n)$$

$\hookrightarrow \mathbb{E}_{p(x; n)} [T_c(x_c)]$
 to compute this, need inference

e.g. Ising model $T_{ij}(x_i, x_j) = x_i x_j$

$$\mathbb{E}[T_{ij}] = \mu_{ij} = p(x_i=1, x_j=1 | n)$$

here need approximate inference $\begin{cases} \text{sampling [Gibbs sampling]} \\ \text{variational [mean field]} \end{cases}$