

today: exp. family
estimation for PGM

Exponential family

a (flat/canonical) exponential family on X

is a parametric family of dist. on X defined by two quantities

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

reference density
base measure

counting measure (discrete R.V.) $\rightarrow \sum_x$ pmf
Lebesgue " (cts. R.V.) $\rightarrow \int_x$ pdf

II) $T: X \rightarrow \mathbb{R}^p$ called "sufficient statistics" vector
aka. feature vector
members of the family will have pmf/pdf

$$p(x; \eta) d\mu(x) = \exp(\underbrace{\eta^T T(x)}_{\text{canonical parameter}} - \underbrace{A(\eta)}_{\text{log-normalization or cumulant generating function}}) \underbrace{h(x) d\mu(x)}_{\text{defining pieces } (+\Omega_X)}$$

if Ω_X is discrete; then $p(x; \eta)$ is a pmf

" " cts. ; " " " pdf

$$\textcircled{*} \text{ want } 1 = \int_X p(x; \eta) d\mu(x)$$

$$= \int_X \exp(\eta^T T(x)) e^{-A(\eta)} h(x) d\mu(x)$$

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

$$\text{or } \sum_x f(x) \mu(x) \text{ if } \mu \text{ is discrete}$$

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

$$\Rightarrow A(\eta) \triangleq \log \left(\underbrace{\int_X \exp(\eta^T T(x)) h(x) d\mu(x)}_{Z(\eta)} \right)$$

$$\text{domain } \Omega \triangleq \{ \eta \in \mathbb{R}^p \mid A(\eta) < \infty \}$$

set of valid canonical parameters

choice of notation:

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

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

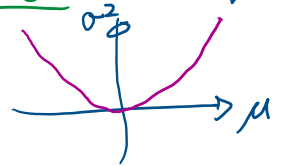
$\textcircled{*}$ more generally, consider a reparameterization of a subset of the flat exponential family
by defining $\eta: \underbrace{\Theta}_{\text{new set}} \rightarrow \Omega$

by defining $\eta: \Theta \rightarrow \Sigma$
 Θ
 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 Σ)

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



* note: only single dist $p(x)$ can be part of an exponential family

by using $h(x) = p(x)$

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

• mixture of Gaussians
(latent variable model)

Example 1: (multinomial)

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

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

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

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

$x \in X$ "think as '0'"

$$= \exp\left(\sum_j x_j \log \pi_j - 0\right)$$

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

$$T(x) = x$$

$$\Omega_X \in \mathbb{R}^k$$

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

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

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

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

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

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

$$\Omega \subseteq \mathbb{R}^k \rightarrow \dim k$$

we do not have a "minimal exp family"

we do not have a "minimal exp family"

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

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

affine linear dep between components of T

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

e.g. $(n + \underbrace{1}_{\text{vector } (1)})^T T(x) = n^T T(x) + \underbrace{1}_{\text{not a minimal exp. family}}$

⊗ for a multinomial; a min. exp family

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

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

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

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

$$\begin{aligned} \text{for multinomial; } \frac{\partial A}{\partial \eta_j} &= \frac{1}{Z(\eta)} e^{\eta_j} = p("x=j" | \eta) \\ &= \mathbb{E}_{p(x; \eta)} [T(x)] \\ &\quad \text{as required} \end{aligned}$$

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

e.g. gamma dist $T(\alpha, \beta)$ 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

15h35

example 2: 1d Gaussian

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

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

$$= \exp\left(-\frac{x^2}{2} \underbrace{\left[\frac{1}{\sigma^2}\right]}_{\eta_2} + x \underbrace{\left[\frac{\mu}{\sigma^2}\right]}_{\eta_1} - \left[\frac{\mu^2}{2\sigma^2} + \frac{1}{2} \log(2\pi\sigma^2)\right] \right) \quad A(\eta)$$

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

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

$$\eta_1 = \eta_2 \mu$$

$$\Omega = \{(\eta_1, \eta_2) : \eta_2 > 0, \eta_1 \in \mathbb{R}\} \quad h(x) = 1 \quad (\text{but some people use } h(x) = \frac{1}{\sqrt{2\pi}})$$

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

Example 3: discrete UGM

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

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

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

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

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

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

$$\eta(\Theta) = \begin{pmatrix} \vdots \\ \log \psi_c(y_c) \end{pmatrix} \leftarrow "$$

[not a minimal representation]

notes: a) mult(x) 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

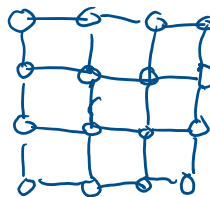
feature on x_i & x_{i+1} e.g. $\mathbb{I} \{ x_i \text{ is a verb} \}$
 $x_{i+1} \text{ "is a noun"}$

→ much smaller set of parameters

"parameter sharing"

c) binary Ising model

$$x_i \in \{-1, 1\} \quad |V| \leq 2$$



suppose use nodes & pairs (edges)
as cliques

$$\Rightarrow \text{dimension of } T(x) \quad 2|V| + 4|E|$$

"overparameterized exp. family"

$$\sum_{y \in \mathcal{Y}} T_{y, y}(x) = 1 \quad \text{for any } x$$

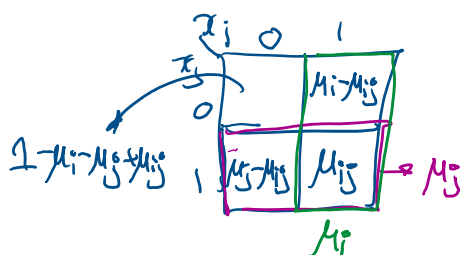
⇒ not a min. exp. fam

* a minimal representation

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

$$\downarrow \{x_i = 1, x_j = 1\}$$

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



properties of A (for generic flat exponential family)

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

$$\bullet \left(\frac{\partial^2 A(n)}{\partial \eta_i \partial \eta_j} \right)_{(i, j) \in (1:p)^2} = \mathbb{E}_{p(x; n)} [(T(x) - \mu(n))(T(x) - \mu(n))^T] = \text{cov}(T(x))$$

(proof as exercise)

"cumulant generating fct."

Estimation of parameters DBM/UGM

independent + . . .

Estimation of parameters DGM/UGM

DGM
(fully observed)

parametric family $\mathcal{P}_\Theta = \{p_\Theta(x) = \prod_i p(x_i | x_{\pi_i}, \theta_i)\} \ni$
 $(\theta_1, \dots, \theta_{|V|}) \in \Theta$ independent parametrization
 $\Theta = \Theta_1 \times \Theta_2 \times \dots \times \Theta_{|V|}$ se. no. of parameters

$\Rightarrow$ MLE decouples in $|V|$ independent MLE problems:

$$\{x^{(i)}\}_{i=1}^n \quad p(\text{data} | \Theta) = \prod_{i=1}^n p(x^{(i)} | \Theta) = \prod_{i=1}^n \prod_{j=1}^{|V|} p(x_j^{(i)} | x_{\pi_j}^{(i)}, \theta_j)$$

$$\log(\quad) = \sum_{j=1}^{|V|} \underbrace{\left(\sum_{i=1}^n \log p(x_j^{(i)} | x_{\pi_j}^{(i)}, \theta_j) \right)}_{f_j(\theta_j)}$$

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

$$= \frac{\#(x_j = k, x_{\pi_j} = \text{stuff})}{\#(x_{\pi_j} = \text{stuff})}$$

(fully observed DGM is relatively easy)
often in closed form?

⊕ if have latent variable (i.e. unobserved variables)

$\Rightarrow$ use EM. (like HMM)

UGM

example for exp. family

$$p(x; \eta) = \exp\left(\sum_c \langle \eta_c, T_c(x) \rangle - A(\eta)\right)$$

$\rightarrow$ unlike in a DGM, $\log p(x; \eta)$ does not separate $\sum_c f_c(\eta_c)$

gradient ascent on log-likelihood

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

$$\nabla_{\eta_c} [\quad] = \hat{\mu} - \mu_c(\eta)$$

$\downarrow$
 $\mathbb{E}_{p(x; \eta)}[T_c(x)]$

to compute this, need inference

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

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

here need approximate inference

← sampling [Gibbs sampling]
variation [mean field]