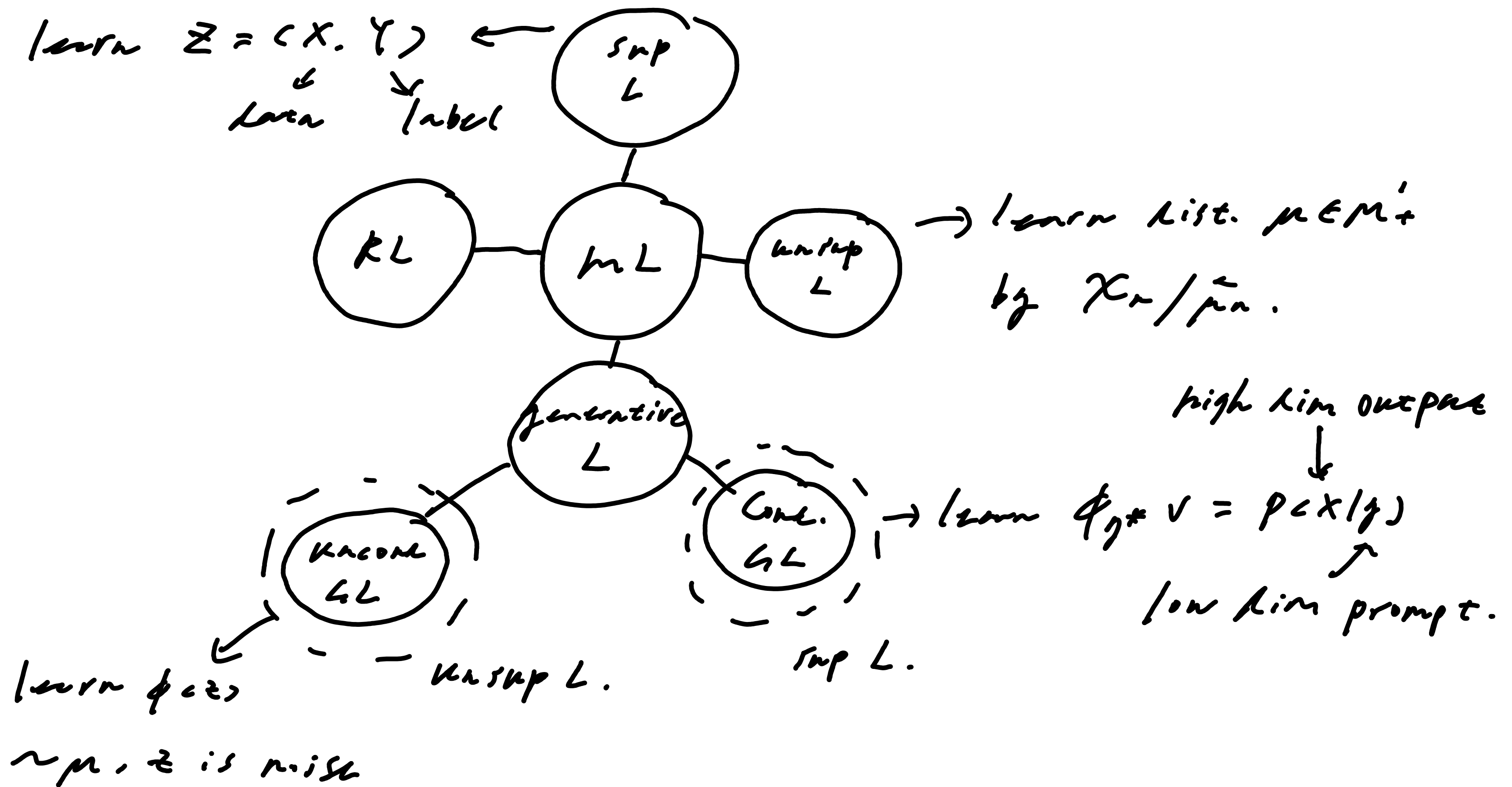


Supervised Learning



Remark: Supervised learning is easier than unsupervised one since the dim. of label is low in common but data of unsupervised has high dim.

In supervised learning, we observe $Z_j = (Y_j, X_j)$
 $=(\mathcal{L}, \mathcal{X}, \mathcal{Y}) \rightarrow (\mathbb{R}^s \times \mathbb{R}^d, B_{\mathbb{R}^s \times \mathbb{R}^d})$ in sample $X_n = (Z_1, \dots, Z_n)$, Z_k i.i.d. where Y_j is label/output and X_j is covariate/input.

Our goal is not to learn distribution of Z_j

rather to learn the dist. of $Y|X$, i.e.
 propose a dist. $\hat{\mu}_{Y|X, \mathcal{X}_n}$ from $\mathcal{X}_n = (Z_1, \dots, Z_n)$
 to get close to $\mu_{Y|X} = P(Y|X)$.

e.g. Y is r.v. of products customer will buy
 X is shopping history ... info. of customer
 $\Rightarrow$ interested in $P(Y = \text{"soap"} | X) = ?$

(1) Regular conditional probn.:

Lemma. For $\tilde{A} \subset A$. $\Rightarrow E(\cdot | \tilde{A}) = \text{Proj.} : L^2(\mathcal{A}, \mathcal{A}, P; \mathbb{R}^d) \rightarrow L^2(\mathcal{A}, \tilde{A}, P; \mathbb{R}^d)$. ortho. proj.

This fixes $E(\cdot | \tilde{A})$. P -a.s.

RMK: Recall i) $\text{Proj } X = \arg \min_{Y \in L^2(\tilde{A})} E\|X - Y\|^2$.

ii) $\text{Proj} \circ \text{Proj} = \text{Proj}$. iii) $\text{Proj} = i_{L^2(\tilde{A})}$

iv) $E\langle X, \text{Proj } Y \rangle = E\langle \text{Proj } X, Y \rangle$

(By ortho. decomposition of X, Y)

Lemma. If $Y \in \sigma(X)$, $Y: \Omega \rightarrow \mathbb{R}^d$ r.v. $X: \Omega \rightarrow \mathbb{R}^d$ r.v. Then: $\exists g: (\mathbb{R}^d, \mathcal{B}_{\mathbb{R}^d}) \rightarrow (\mathbb{R}^d, \mathcal{B}_{\mathbb{R}^d})$
 measurable. st. $Y = g(X)$. g is unique

$\mathbb{P}_X$ -a.s. i.e. $Y' = Y$. $\mathbb{P}_X$ -a.s. $\Rightarrow g(x) = Y$. $\mathbb{P}$ -a.s.

Def: i) $\{P_{Y|X}(A)\}_{A \in \mathcal{A}}$ is regular conditional probability (r.c.p.) if it's induced by a

Markov kernel $K(x, B) = P_{Y|X=x}(B)$. $\mathbb{P}_X$ -a.s. x &

a) $B \mapsto K(x, B)$ is p.m. for $\forall x \in \mathbb{R}^d$.

b) $x \mapsto K(x, B)$ is $\mathcal{B}_{\mathbb{R}^d}$ -measurable. $\forall B \in \mathcal{A}$.

ii) $P_{Y|X=x}(B)$ admits a regular version if

$\bar{P}_{Y|X}(\cdot)$ is r.c.p. and $P_{Y|X} = \bar{P}_{Y|X}$ a.s.

Remark: Note if let $P_{Y|X}(A) = \mathbb{E}[\mathbb{I}_{\{Y \in A\}}$

$|\mathcal{G}(X)\rangle$. then: $P_{Y|X}(\cdot)$ also only

satisfies property of p.m. $\mathbb{P}$ -a.s.

But the null set N will depend

on $\{A_n\}$ we choose if it's not r.c.p.

$\Rightarrow$ it may not be p.m. for $\mathbb{P}_X$ -a.s. x .

Supervised learning is to learn Markov kernel.

Ex: $f(y|x) = \frac{1}{(\sqrt{2\pi}\sigma)^2} e^{-\frac{1}{2}(\frac{y-x}{\sigma})^2} k_y \Rightarrow$

$K(x, B) = \int_B f(y|x) k_y$ is a Markov kernel.

for linear regression

Thm. $\mathcal{Z} = (Y, X) : (\mathcal{X}, \mathcal{A}, \mathbb{P}) \rightarrow (\mathcal{Y} \times \mathcal{X}^s, \mathcal{B}_{\mathcal{Y} \times \mathcal{X}^s})$ s.t.
 $Y \in L^2$ -r.v. $\Rightarrow$ There exists a regular
 version $P_{Y|X}(A)$.

Lemma. The list. of $\mathcal{Z} = (Y, X) : P_{\mathcal{Z}}$ is uniquely
 determined by its marginal $\mathbb{P}_X$ and r.c.p.
 $\mathbb{P}_{Y|X=X}(\cdot)$. i.e. we have: $\forall B \in \mathcal{B}_{\mathcal{Y} \times \mathcal{X}^s}$.

$$\mathbb{P}_{\mathcal{Z}}(B) = \int_{\mathcal{X}^s} \left(\int_{\mathcal{Y}}, I_B(\eta, x) \mathbb{P}_{Y|X=X}(\eta) \right) \mathbb{P}_X(x).$$

Lemma. If $\mathcal{Z} = (X, Y)$ has density $f_{\mathcal{Z}}(x, \eta)$. let

$$f_X(x) = \int_{\mathcal{Y}} f_{\mathcal{Z}}(x, \eta) \mathbb{P}_{\mathcal{Y}}(\eta) \text{ and let}$$

$$f(\eta|x) = \begin{cases} f_{\mathcal{Z}}(x, \eta) / f_X(x) & \text{if } f_X(x) > 0 \\ \mathbb{I}_{\{0,1\}^S}(\eta) & \text{otherwise.} \end{cases}$$

Then: $P_{Y|X=X}(B) = \int_{\mathcal{Y}} I_B(\eta) f(\eta|x) \mathbb{P}_{\mathcal{Y}}(\eta)$ is
 a r.c.p. of Y over $X=x$.

(2) Divergence:

Next we want to introduce a divergence to
 measure success of learning from $\hat{P}_{Y|X, x_n}$
 to true Markov kernel $P_{Y|X}$.

Def: For two Markov kernel $\mu|x=x$, $V|x=x$ and divergence $K(\cdot||\cdot)$ on $\mathcal{M}_1^+(\mathcal{R}')$, let $x \in \mathcal{R}^d$
 $\mapsto K(\mu|x=x || V|x=x) \in (\mathcal{R}', B_{\mathcal{R}'})$ is measurable
 for $\forall \mu|x, V|x \in \mathcal{K}_1^+(\mathcal{R}^d, B_{\mathcal{R}'})$ space of
 Markov kernels. Set:

$$D_p(\mu|x || V|x) = D_{p,x,K}(\mu|x || V|x)$$

$$= \mathbb{E}_x(K(\mu|x || V|x)^p)^{\frac{1}{p}}$$

$$= \left(\int_{\mathcal{R}^d} K(\mu|x=x || V|x=x)^p K(p_x(x)) \right)^{\frac{1}{p}}, 1 \leq p < \infty.$$

$$D_\infty(\mu|x || V|x) = \operatorname{ess\,sup}_{x \in \mathcal{R}^d} K(\mu|x=x || V|x=x)$$

Remarks: i) $p \uparrow$, then sensitivity of $K(\mu|x=x || V|x=x)$ at each pt. $x \uparrow$.

ii) $D_{1,TV}$ is weaker than $D_{2,KL}$.

$\mathcal{D}$ will mostly inherit the strength of K . But it's not from Pinsker's)

iii) If $x \sim p_x$ but we let $x \sim \bar{p}_x$ knowing inference. Suppose that

$$D_{1,x \sim p_x,K}(\mu|x || \hat{\mu}_{1x,n}) \rightarrow 0 \text{ a.s.}$$

$$\text{Note } D_{1,x \sim \bar{p}_x,K}(\mu|x || \hat{\mu}_{1x,n}) =$$

$$\mathbb{E}_{x \sim P_x} \left(L(\mu_{1x} \| \tilde{\mu}_{1x,n}, \frac{L\bar{P}_x}{L P_x}) \right) \leq$$

$$\left\| \frac{L\bar{P}_x}{L P_x} \right\|_n D(x \sim P_x, L(\mu_{1x} \| \tilde{\mu}_{1x,n})) \xrightarrow{n.s.} 0 \quad \text{if } \left\| \frac{L\bar{P}_x}{L P_x} \right\|_n < \infty$$

So the real event can still be covered in training

Def. Denote $V_{1x} P_x(B) \stackrel{\Delta}{=} \int \int I_B(y, x) dV_{1x=x}(y) dP_x(x)$

For $V_{1x}, \mu_{1x} \in \mathcal{K}_1^+(\mathbb{R}^L, \mathcal{B}_{\mathbb{R}^L})$. We have:

$$D(x, K_L(\mu_{1x} \| V_{1x})) = K_L(\mu \| V), \text{ where } \mu_{1x} = \mu_{1x} P_x, \quad V = V_{1x} P_x.$$

rmk: In supervised learning we always know μ and V .

Pf: i) We first prove: $\frac{L V}{L \mu}(y, x) \stackrel{n.s.}{=} \frac{L V_{1x=x}}{L \mu_{1x=x}}(y)$.

$$a) \mu \ll V \Rightarrow \mu_{1x=x} \ll V_{1x=x}, \quad P_x \text{ a.s. } x.$$

otherwise, $\exists B_1$ s.t. $P_x(B_1) > 0$ satisfy:

$$\exists B_2(x), \text{ s.t. } \mu_{1x=x}(B_2(x)) > 0 = V_{1x=x}(B_2(x))$$

$$\text{Set } B = \{(y, x) : x \in B_1, y \in B_2(x)\}$$

$$\Rightarrow \mu(B) = \int \mu_{1x=x}(B_2(x)) I_{B_1}(x) dP_x > 0$$

$$\text{But } V(B) = \int V_{1x=x}(B_2(x)) I_{B_1}(x) dP_x = 0.$$

$\Rightarrow$ contradiction.

$$\begin{aligned}
b) \text{ And we see } & \int I_B(y, x) \frac{\mu_{V|x=x}}{\mu_{\mu|x=x}}(y) \mu(y, x) \\
&= \int \left(\int I_B(y, x) \frac{\mu_{V|x=x}}{\mu_{\mu|x=x}}(y) \mu_{\mu|x=x}(y) \right) \mu_P(x) \\
&= \int \left(\int I_B \mu_{V|x=x}(y) \right) \mu_P(x) = V(B)
\end{aligned}$$

$$\int \mu_{\mu|x=x} \leq V_{\mu|x=x}, \mu\text{-a.s. } x \Rightarrow \mu \leq V.$$

$$\text{and } \frac{\mu_V}{\mu_P}(y, x) = \frac{\mu_{V|x=x}}{\mu_{\mu|x=x}}(y), \mu\text{-a.s.}$$

$$\begin{aligned}
2) \text{ } KL(\mu || V) &\stackrel{i)}{=} - \int \log \left(\frac{\mu_{V|x=x}}{\mu_{\mu|x=x}}(y) \right) \mu(y, x) \\
&= \int \left(- \int \log \left(\frac{\mu_{V|x=x}}{\mu_{\mu|x=x}}(y) \right) \mu_{\mu|x=x}(y) \right) \mu_P(x) \\
&= D_{KL}(\mu_{\mu|x} || \mu_{V|x}).
\end{aligned}$$

(3) Framework:

Denote $\mathcal{J} \subseteq \mathcal{K}_1^+(\mathcal{K}^t, \mathcal{B}_{\mathcal{K}^t})$ is set of target Markov kernels. $\mathcal{X}_n = (z_1, \dots, z_n)$ is samples and $\hat{\mu}_{n|x} \stackrel{a)}{=} \hat{\mu}_{n|x} \circ \mathcal{X}_n$. For D div. We'll require: $\mathcal{X} \in \mathcal{K}^{(s+t) \times n} \mapsto D(\mu_{\mu|x} || \hat{\mu}_{n|x}(x)) \in \mathcal{R}^+$ is $\mathcal{B}_{\mathcal{K}^{(s+t) \times n}}$ -measurable.

Def: $\mathcal{K}_n \subseteq \mathcal{K}_1^+(\mathcal{K}^t, \mathcal{B}_{\mathcal{K}^t})$ hypothesis space of Markov kernel, $z_j = (Y_j, X_j) \stackrel{i.i.d}{\sim} \mu = \mu_{\mu} \mu_P$

Let $D = D_{p,x,d}$ div. on $\mathcal{K}_i^+$.

i) Empirical risk function $\hat{\mathcal{I}}_n = \hat{\mathcal{I}}_n(v_{|X}, \mathcal{X}_n)$

satisfies:

a) $\mathbb{R}^{(1+1)n} \ni \mathcal{X}_n \mapsto \hat{\mathcal{I}}_n(v_{|X}, \mathcal{X}_n)$ is measurable
for $\forall v_{|X} \in \mathcal{H}_n$

b) $\exists h(\cdot)$ on $\mathcal{M}_i^+(\mathcal{C}_n) \subset \mathbb{R}^+$ s.t.

$$c_n \hat{\mathcal{I}}_{n_k}(v_{|X}, \mathcal{X}_{n_k}) + h_{n_k}(\mu) \xrightarrow{pr} D_{p,x,d}(\mu_{|X} \| v_{|X})$$

$\forall (n_k)$. s.t. $v_{|X} \in \mathcal{H}_{n_k}$.

ii) $\hat{\mathcal{I}}_n$ is unbiased if $\mathbb{E}_{\mathbb{Z}}(c_n \hat{\mathcal{I}}_{n_k}(v_{|X}, \mathcal{X}_{n_k}) + h_{n_k}(\mu))$
 $= D_{p,x,d}(\mu_{|X} \| v_{|X})$.

iii) $\hat{\mu}_{n|X}$ is supervised ERM-learning for $\hat{\mathcal{I}}_n$
if $\hat{\mu}_{n|X} \in \operatorname{argmin}_{v_{|X} \in \mathcal{H}_n} \hat{\mathcal{I}}_n(v_{|X}, \mathcal{X}_n)$.

Note we still have error decomposition:

$$0 \leq D_{p,x,d}(\mu_{|X} \| \hat{\mu}_{n|X}) \leq \varepsilon_{n,\text{mod}}(\mu_{|X}) + \varepsilon_{n,\text{learn}} + 2\varepsilon_{n,\text{samp}}$$

$$\varepsilon_{n,\text{mod}}(\mu) = \inf_{v_{|X} \in \mathcal{H}_n} D_{p,x,d}(\mu_{|X} \| v_{|X});$$

$$\varepsilon_{n,\text{learn}} = c_n \left(\hat{\mathcal{L}}_n(\hat{\mu}_{n|X}, \mathcal{X}_n) - \inf_{v_{|X} \in \mathcal{H}_n} \hat{\mathcal{L}}_n(v_{|X}, \mathcal{X}_n) \right);$$

$$\varepsilon_{n,\text{samp}} = \sup_{v_{|X} \in \mathcal{H}_n} \left| D(\mu_{|X} \| v_{|X}) - (c_n \hat{\mathcal{L}}_n(v_{|X}, \mathcal{X}_n) + h_n(\mu)) \right|.$$

Cor. For $\mathcal{J} \subset \mathcal{K}$, $|\mathcal{K}| < \infty$. Then: $\mathcal{J}$ is PAC-learnable.

Cor. $\mathcal{K}$ is opt. $\mathcal{J} \subset \mathcal{K}$. For $\hat{\mathcal{I}}_n(V_{1:n}, \mathcal{X}_n) = \sum_{j=1}^n \ell(z_j | V_{1:n})$ with $\ell(z | V_{1:n}) \leq k(z) \in L'(\mathbb{R}^{s+k}, \mu)$, $\forall \mu = \mu_{1:n} p_x$, $\mu_{1:n} \in \mathcal{K}$. If $V_{1:n} \in \mathcal{K} \mapsto \ell(z | V_{1:n})$ is $D_{p,x}(\cdot, \|\cdot\|)$ -conti.

Then $\mathcal{J}$ is learned by ERM learner $\hat{\mu}_n$.

Thm. $\mathcal{J}, \mathcal{K} = \mathcal{K}_1^T \subset \mathbb{R}^{\hat{d}}, \mathcal{B}_{\mathcal{K}_1} \cap \{ \ell(V_{1:n}(y) = f_{V_{1:n}}(y|x)) \}$

If $\ell(z | V_{1:n}) \stackrel{\Delta}{=} -\log f_{V_{1:n}}(y|x) \in L'(\mathbb{R}^{s+k}, \mu)$

$\forall \mu = \mu_{1:n} p_x$, $\forall \mu_{1:n} \in \mathcal{J}$, $\forall V_{1:n} \in \mathcal{K}$. And let:

$\hat{\mathcal{I}}_n(V_{1:n}, \mathcal{X}_n) = \sum_{j=1}^n \ell(z_j | V_{1:n})$. Then:

i) $\hat{\mathcal{I}}_n$ is unbiased ERMF with $\mathcal{L}_n = n^{-1}$ and

$$h_n(\mu) = h(\mu) = \int_{\mathbb{R}^{s+k}} \log f_{\mu_{1:n}=x}(y|x) d\mu_{1:n}=x(y) d p_x(x)$$

ii) If $X \sim f_x(x) dx$, $f_v(z) \stackrel{\Delta}{=} f_{V_{1:n}}(y|x) f_x(x)$

$\bar{\mathcal{J}} \stackrel{\Delta}{=} \{ \mu = \mu_{1:n} p_x \mid \mu_{1:n} \in \mathcal{K} \}$.

$\bar{\mathcal{K}} \stackrel{\Delta}{=} \{ V = V_{1:n} p_x \mid V_{1:n} \in \mathcal{J} \}$. Then:

$\hat{\mu}_n$ is ERM learner w.r.t. $\mathcal{K} \subset \mathcal{L}$ i.e. $\hat{\mu}_n \in$

$\arg \min_{v \in \bar{\mathcal{K}}} - \sum_{j=1}^n \log f_v(z_j) \Leftrightarrow \hat{\mu}_{n|X}$ r.c.p. of

$\hat{\mu}_n$ w.r.t $\mathcal{G}(X)$ is ERM learner for $\hat{I}_n$

Pf: i) $\hat{E}_n(\hat{I}_n(V_{1X}, \mathcal{X}_n) + h_n(\mu)) \stackrel{\text{Len.}}{=}$

$$= \int \log(f_{V_{1X}}(y|x) / f_{\mu_{1X}}(y|x)) d\mu_{X=X}(y) dP_X(x) \\ = \int A_{KL}(\mu_{1X=X} \| V_{1X=X}) dP_X(x) = D_{1X, KL}(\mu_{1X} \| V_{1X})$$

ii) Note $f_v(z) = f_{V_{1X}}(y|x) f_X(x)$

$$= \sum_j \log f_v(z_j) = \hat{I}_n(V_{1X}, \mathcal{X}_n) - \sum_j \log f_X(x_j)$$

$\Rightarrow$ The 2nd term doesn't depend on V_{1X} with last Len. of (2). and def of $\hat{\mu}_n$. We obtain the corresponding.

Remark: We can see supervised learning is as unsupervised learning with part of dist.

(4) Common Principles for models:

For $(P(y|x))$ Markov kernel. Next, we want to know what kind type of y :

a) continuous b) discrete: $\mathbb{N}$, $\{1, 2, \dots, q\}$, $\{0, 1\}, \dots$

$P(y|x, \theta)$ should be constructed to give all cond. proba. and label space of y .

e.g. i) (Binary labels)

$B(p) = \text{Bernoulli dist. on } \{0,1\}$. i.e.

$$B(p)(1) = p = 1 - B(p)(0).$$

Note all dist. on $\{0,1\}$ can be repr
in $\{B(p(x)) \mid p(x) \text{ is measurable}\}$ all Markov
kernels over $\mathbb{R}^k$ and label space $= \{0,1\}$.

ii) (Conti. dist.: Classical regression)

$$y \in \mathbb{R}. \quad p(y|x) = (\sqrt{2\pi}\sigma)^{-1} \exp\left(-\frac{1}{2}\left(\frac{y - m(x)}{\sigma}\right)^2\right)$$

Rmk. If we use it to model all conti.

dist. $\Rightarrow$ error will be from:

- a) Model error (may not be Gaussian)
- b) Deviation of $m(x)$

Steps:

1) identify the data structure of labels:
i.e. conti. or discrete?

2) choose a parametric family of dist. $\{M_\xi$:
 $\xi \in \Xi$. $\xi \mapsto M_\xi(B)$ is measurable. $\forall B \in \mathcal{A}\}$.

3) choose a parametric class of functions $\{m_\theta$:

$\theta \in \mathbb{R}^L \mapsto \exists m_\theta \text{ is measurable}$. Derive the hypothesis space $\mathcal{H} = \{ \mu_{m_\theta} : \theta \in \Theta, \mu_{x=x}(B) := \mu_{m_\theta}(B) \text{ is Markov kernel} \}$.

4) Choose a divergence $\rho(\cdot, \cdot)$ (may KL) and corresponding ERF $\tilde{L}_n$ (may neg. log likelihood)

5) Train with ERM.

① Binary classification: (logistic regression)

label space = $\{0, 1\}$. $\sigma(z) = e^z / (1 + e^z)$.

$m_\theta(x) = \bar{\theta}^T x + \theta_0$. $x \in \mathbb{R}^L$. $\theta = (\theta_0, \bar{\theta}) \in \mathbb{R}^{L+1}$.

(θ_0 is bias and $\bar{\theta}$ is weight.)

Set $B(p(x)) \equiv 1 = p(x) = \sigma(m_\theta(x))$ and use

Neg. log likelihood:

$\tilde{L}_n((y_j, x_j)_{j=1}^n) = - \sum_{j=1}^n y_j \log p(x_j) + (1 - y_j) \log (1 - p(x_j))$

$\Rightarrow$ Choose $y_j = \arg \max_{y \in \{0, 1\}} B(p(x_j)) \equiv y_j$

Or we can use other functions:

$m_\theta(x) = \phi_{\theta_L}^{(L)} \circ \dots \circ \phi_{\theta_1}^{(1)}$. $\{\phi_{\theta_k}^{(k)}\}$ is layers. $\theta =$

$(\theta_1, \dots, \theta_L)$. L is depth. $\theta_j = (\bar{\theta}_j, \theta_j^0)$. s.t.

$\bar{\theta}_j \in \mathbb{R}^{L+1} \times \mathbb{R}^{L_j}$. $\theta_j^0 \in \mathbb{R}^{L_j}$.

We set $\phi_{\theta_j}^{(j)}(x^{(j-1)}) = \chi(\theta_j^T x^{(j-1)} + \theta_j^0) \in \mathbb{R}^{k_j}$

where $\chi: \mathbb{R}^{k_j} \rightarrow \mathbb{R}^1$, $\chi(z) = (\chi(z_1) \dots \chi(z_{k_j}))^A$

for $z \in \mathbb{R}^{k_j}$, $\forall j$. activation function

e.g. $\chi(z_j) = e^{z_j} / (1 + e^{z_j})$ sigmoid function

and $\text{ReLU}(z) = z \vee 0$.

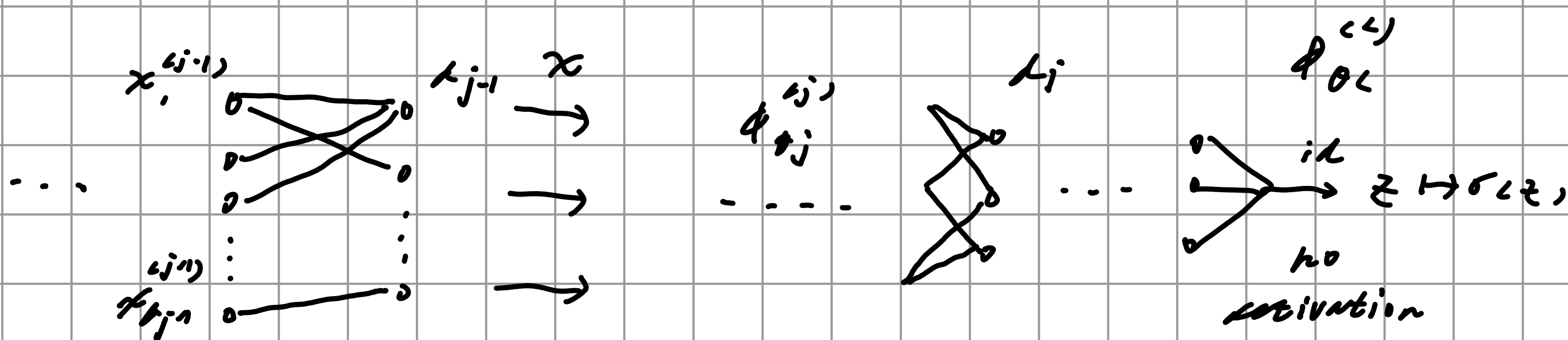
Remark: i) We don't use polynomials generally

because of comp. ineff. and hard to optimize. ($|Z(p(x))|$ is large)

ii) $\max \{, \}$ isn't differentiable. So:

optimal can't be trained to get

Fully connected Neural Network (FCNN):



Multi-class Classification:

$q \in \mathcal{L} = \{1, 2, \dots, \ell\}$. Set $\text{softmax}(z)_q =$

$$e^{z_q} / \sum_{j \in \mathcal{L}} e^{z_j}. \quad z \in \mathbb{R}^{\ell}.$$

Remark: Note $\beta + c \mathbb{I}$ still corresponds the same $\langle P(\cdot | \beta) \rangle$. So it's not identifiable

Hence we generally force $\beta_1 = 0$.

Next, we do a logistic regression: $\theta = (\bar{\theta}, \theta_0)$

$\bar{\theta} \in \text{Mat}_{1 \times d-1}(\mathbb{R})$, $\theta_0 \in \mathbb{R}^{1 \times 1}$. $m_{\theta}(x) = \bar{\theta}^T x + \theta_0$.

$$\mu_{1|x=x}^{\theta}(\{y\}) = P_{\theta}(\{y\} | x) := \text{softmax}(0, m_{\theta}(x))_y$$

So: $\sum_y P_{\theta}(\{y\} | x) = 1 \Rightarrow$ it's dist. on $\mathcal{L}$.

And we get Markov kernel $\mu_x(\cdot)$.

Remark: All dist. on $\mathcal{L}$ can be represented by one β .

$$\begin{aligned} \tilde{J}_n(\langle y_j, x_j \rangle) &= - \sum_{j=1}^n \log(\text{softmax}(0, \bar{\theta}^T x_j + \theta_0)_{y_j}) \\ &= - \sum_{j=1}^n \sum_{i=1}^2 \delta_{i, y_j} \log \square. \end{aligned}$$

③ Count regression:

$P_{\theta}(y | \{x\}) := e^{-s} s^y / y!$. $y \in \mathcal{L} = \mathbb{N}$. Set:

$$m_{\theta}(x) = \begin{cases} \bar{\theta}^T x + \theta_0 & \text{or} \\ \varphi_{\theta}(x) \cdot (FCNN) \end{cases}$$

mult Markov kernel $\mu_{1|x=x}^{\theta}(y) = p_{\theta}(y|x) =$
 $e^{-m_{\theta}(x)} m_{\theta}(x)^y / y!$

$$\Rightarrow \hat{\mathcal{L}}((y_j, x_j), \theta) = \sum_j -m_{\theta}(x_j) + y_j \log m_{\theta}(x_j) - \log y_j.$$

④ Regression:

Markov kernel is $(\sqrt{2\pi}\sigma)^{-1} e^{-\frac{1}{2}(\frac{y-m_{\theta}(x)}{\sigma})^2} \frac{1}{\sigma}$

$$\begin{aligned} \hat{\mathcal{L}}_n((y_j, x_j), \theta, \sigma) &= -\sum -\frac{1}{2} \left(\frac{y_j - m_{\theta}(x_j)}{\sigma} \right)^2 - \log \sqrt{2\pi}\sigma \\ &= (2\sigma^2)^{-1} \sum (y_j - m_{\theta}(x_j))^2 + n \log \sqrt{2\pi}\sigma \\ &= \frac{n}{2\sigma^2} \mathcal{L}_n^{MSE}(\theta) + n \log \sqrt{2\pi}\sigma. \end{aligned}$$

Def: For $m_{\theta}(x) = \theta x + \theta_0$. We call it ordinary

least square (OLS) linear regression:

if $\theta_0 = 0$, then $\theta^* = (X^T X)^{-1} X^T Y$ is the

MLE if $|X^T X| \neq 0$. Where $X = (x_1, \dots, x_n)$, $Y =$

$(y_1, \dots, y_n)$. Also: $\sigma_x^2 = \mathcal{L}_n^{MSE}(\theta^*)$ is optimal.