

Bayesian Non-parametrics

Plan:

- Bayesian Lin. Regression
- Gaussian Process
- Dirichlet Process / DP mixtures
- de Finetti's Representation Theorem (dFRT).

Linear Regression:

$$y(x, \omega) = \omega_0 + \omega_1 x_1 + \dots + \omega_D x_D$$

$$y(x, \omega) = \omega_0 + \sum_{j=1}^{M-1} \omega_j \underbrace{\phi_j(x)}_{\substack{\text{basis} \\ \text{fn.}}}$$

$$= \sum_{j=0}^{M-1} \omega_j \phi_j(x)$$

where $\phi_0(x) = 1$

$$= \omega^T \phi(x)$$

Gaussian basis fn:

$$\phi_j(x) = \exp \left\{ -\frac{(x - \mu_j)^2}{2\sigma^2} \right\}$$

MLE: $t = y(x, \omega) + \epsilon$ Zero mean
 β Precision

$$P(t | x, \omega, \beta) = \mathcal{N}(t | y(x, \omega), \beta^{-1})$$

$$= \prod_{n=1}^N \mathcal{N}(t_n | \omega^T \phi(x_n), \beta^{-1})$$

$$\omega_{ML} = (\Phi^T \Phi)^{-1} \Phi^T t.$$

Bayesian Linear Regression:

β is known.

ω - parameter to estimate.

$$p(\omega) = \mathcal{N}(\omega | m_0, S_0)$$

Posterior = Prior $\times$ likelihood.

$$p(\omega | t) = \mathcal{N}(\omega | m_N, S_N)$$

) ex.

$$m_N = \frac{S_N (S_0^{-1} m_0 + \beta \phi^T t)}{S_N^{-1}}$$

$$S_N^{-1} = S_0^{-1} + \beta \phi^T \phi.$$

$$\omega_{MAP} = m_N$$

$$S_0 = \alpha^{-1} I \quad \text{with } \alpha \rightarrow 0 \quad m_N \rightarrow \omega_{ML}$$

$$\text{assume: } p(\omega | \alpha) = N(\omega | 0, \alpha^{-1} I)$$

$$m_N = \beta S_N \phi^T t$$

$$S_N^{-1} = \alpha I + \beta \phi^T \phi.$$

$$\ln p(\omega | t) = -\frac{\beta}{2} \sum_{n=1}^N \{t_n - \omega^T \phi(x_n)\}^2$$

$$- \frac{\alpha}{2} \omega^T \omega + \text{const.}$$

$$P(\bar{t} | \underline{t}, \alpha, \beta) = \int P(t | \bar{\omega}, \beta) P(\omega | \bar{t}, \alpha, \beta) d\omega$$

$$P(\bar{t} | x, \underline{t}, \alpha, \beta) = \mathcal{N}(\bar{t} | m_N^T \phi(x), \sigma_N^2(x))$$

$$\sigma_N^2(x) = \frac{1}{\beta} + \phi(x)^T S_N \phi(x)$$

kernel
view

$$y(x, m_N) = m_N^T \phi(x)$$

$$= \beta \phi(x)^T S_N \Phi^T \underline{t}$$

$$= \sum_{n=1}^N \beta \phi(x)^T S_N \phi(x_n) t_n$$

$$y(x, m_N) = \sum_{n=1}^N k(x, x_n) t_n$$

$$\text{where } k(x, x') = \beta \phi(x)^T S_N \phi(x')$$

↪ equivalent kernel /
Smoothing matrix.

$$\begin{aligned}
\text{Cov}(y(x), y(x')) &= \text{Cov}[\phi(x)^T \omega, \omega^T \phi(x')] \\
&= \phi(x)^T \Sigma_N \phi(x') \\
&= \beta^{-1} k(x, x')
\end{aligned}$$

Gaussian Processes:

functional view:

$$y(x) = \omega^T \phi(x)$$

$$p(\omega) = \mathcal{N}(\omega | 0, \alpha^{-1} \Sigma)$$

every $\omega \sim p(\omega) \rightarrow$ fn. $y(x)$

distribution over fn $y(x)$

$y(x)$

$y(x_1) \dots y(x_n)$

$$x_1 \dots x_n \rightarrow (y(x_1), \dots, y(x_n))$$

$$y = \Phi w$$

$$E[y] = \Phi E[w] = 0$$

$$\text{Cov}[y] = E[yy^T]$$

$$= \Phi E[ww^T] \Phi^T$$

$$= \frac{1}{\alpha} \Phi \Phi^T = K$$

→ Gram matrix.

$$K_{nm} = K(x_n, x_m) = \frac{1}{\alpha} \phi(x_n)^T \phi(x_m)$$

$$K(x, x') = \exp(-\theta |x - x'|)$$

How to use this GP for regression?

$$t_n = y_n + \epsilon_n$$

$$P(t_n | y_n) = \mathcal{N}(t_n | y_n, \beta^{-1})$$

$$P(\mathbf{t} | \mathbf{y}) = \mathcal{N}(\mathbf{t} | \mathbf{y}, \beta^{-1} \mathbf{I}_n)$$

$$P(\mathbf{y}) = \mathcal{N}(\mathbf{y} | \mathbf{0}, \mathbf{K})$$

$$P(\mathbf{t}) = \int P(\mathbf{t} | \mathbf{y}) P(\mathbf{y}) d\mathbf{y}$$

$$= \mathcal{N}(\mathbf{t} | \mathbf{0}, \mathbf{c}) \quad \text{ex.}$$

⊗ $C(x_n, x_m) = k(x_n, x_m) + \beta^{-1} \delta_{nm}$

$$P(\mathbf{t}_{n+1}) = \mathcal{N}(\mathbf{t}_{n+1} | \mathbf{0}, \mathbf{c}_{n+1})$$

$$\mathbf{c}_{n+1} = \begin{pmatrix} \mathbf{c}_n & \mathbf{k} \\ \mathbf{k}^T & c \end{pmatrix} \quad \text{ex.}$$

where $\mathbf{k} \rightarrow k(x_n, x_{n+1})$
 $c \rightarrow k(x_{n+1}, x_{n+1})$

$$P(t_{n+1} | t) = \mathcal{N}(m, \sigma^2)$$

$$m = K^T C_N^{-1} t$$

$$\sigma^2 = C - K^T C_N^{-1} K.$$

2nd.

one restriction!

Covariance matrix, $(*)$ should be

Positive definite.

C.P represent

base for learning

invert $N \times N$ matrix.

invert $M \times M$ matrix.

$$O(N^3)$$

$$O(M^3)$$

one step $\rightarrow$

$$M \ll N$$

every step $O(N^2)$

$$O(M^2)$$

Main advantage of GP:

you can use inf. dimensional
basis fn. $\rightarrow$ kernel.

Parametric Model

finite set of Param. θ

Given θ , Prediction
is indep. of D .

$$P(x|\theta, D) = P(x|\theta)$$

- even if data is unbounded,
complexity of model is
bounded.

- not flexible.

Non-parametric
Model.

- infinite dimensional
 θ .

$\theta \Rightarrow$ function.

$\rightarrow$ amt of info. θ
can capture grows
as D grows.

Example:

1) fn approximation

Poly. regression Vs. G.P.

2) Classification

log. regression Vs. GP classif.

3) Clustering.

mixture model, kMeans Vs.

DP mixtures .

4) time series .

HMM vs. infinite HMM.

de Finetti's Representation Theorem :

df-RT

Coin tossing:

I. I. D.

Independence:

$$P(x_{n+1} \mid x_1 \dots x_n) = P(x_{n+1})$$

only with I.D, we can learn.

Exchangeability:

$\{0, 1, 0, 0, 1, 0\}$

order does not matter.

$$P(F, F, F, G, G, G) >$$

$$P(G, G, G, F, F, F)$$

order matters.

exchangeable $\rightarrow$ weaker than
indep

Defn: A set of R.V. $\{X_n\}$ is
said to be exchangeable if
given the joint density $P(x_1, \dots, x_n)$
we have

$$P(x_1, \dots, x_n) = P(x_{\sigma(1)}, \dots, x_{\sigma(n)})$$

σ - perm. of $(1, \dots, n)$

$\text{IID} \rightarrow \text{exchangeable.}$

exchangeable $\not\rightarrow$ IID.

exchange $\rightarrow$ ID.

~~def~~ def dFRT for Bernoulli R.V

Let $\{x_i\}_{i=1}^{\infty}$ be a seq. of
finitely exchag. random variable.

i.e. $\forall n > 0$, each finite

Subsequence of $x_i\}_{i=1}^n$ is exchag.

Then $\exists$ a random Variable θ

and a distrib. fn $F(\theta)$ such that

$$P\left(\lim_{n \rightarrow \infty} \frac{\sum x_i}{n} = \theta\right) = 1$$

with $\theta \sim F(\theta)$

and

$$P(x_1, \dots, x_n) = \int_0^1 \prod_i \theta^{x_i} (1-\theta)^{1-x_i} dF(\theta)$$