

$x \in \mathbb{R}^D$: observation

$z \in \mathbb{R}^K$: latent variable

$$p(x, z) = \frac{1}{\sum_z} \exp(-E(x, z))$$

正規化定数

energy function

parameterized by DNN

training dataset of N example ($x^{(i)}$ with $i \in [1, N]$)

gradient descent の学習過程を記述する

(log probability of $p(x)$
 を最大にするように)

Answer this problem can be formed as follows

$$\hat{\theta} = \underset{\theta}{\operatorname{argmin}} p_{\theta}(x, z)$$

$$= \underset{\theta}{\operatorname{argmin}} \log p(x, z)$$

by using this grad. run GD. $\rightarrow$

$$\frac{\partial \log p_{\theta}(x, z)}{\partial \theta}$$

$$\mathbb{E}_{p_{\text{data}}} \left[\frac{\partial E}{\partial \theta} \right] - \mathbb{E}_{p_{\theta}(x)} \left[\frac{\partial E}{\partial \theta} \right]$$

x を観測 $\rightarrow$ 変数の集合 y を出力

EBM は x と相性の良い y を見つける問題を考える

$$\hat{y} = \underset{y}{\operatorname{argmin}} F(x, y) \quad / \quad F: x \times y \rightarrow \mathbb{R}$$

$F(x, y)$ が低い y を見つける問題

e.g. テキスト y は テキスト x の良い翻訳であるか?

勾配法による推論, 互換性のあつ y を見つける

潜在変数を持つ EBM z = (latent variable)

$$y, z = \underset{y, z}{\operatorname{argmin}} F(x, y, z)$$

エネルギーの確率密度関数 $p_\theta(x)$

$$p_\theta(x) = \frac{\exp(-E_\theta(x))}{Z(\theta)}$$

where $Z(\theta) = \int \exp(-E_\theta(x)) dx$

↑ 分母に関数

E : エネルギー関数

密度比推定

エネルギー関数は密度の対数の値に比例

ただし $Z(\theta)$ があつたため密度自体は計算できない

↳ MCE の推定はできる

どう学習するか?

→ 最大推定 train data に対する

好適性度 $\log p_\theta(x)$ を最大化する

↳ $Z(\theta)$ の計算を避けたい

SGDの学習の場合、収束速度を上げるために

1000-10000回程度学習させる

$$p_{\theta} = \frac{\exp(-E_{\theta}(x))}{\int \exp(-E_{\theta}(x)) dx}$$

$$\hookrightarrow \frac{\partial \log p_{\theta}(x)}{\partial \theta}$$

$$\boxed{\begin{aligned} \log(x t^{-1}) \\ = \log x - \log t \end{aligned}}$$

$$= \mathbb{E}_{p_{\theta}(x)} \left[\frac{\partial E_{\theta}(x)}{\partial \theta} \right] + \mathbb{E}_{p_{data}(x)} \left[\frac{\partial E_{\theta}(x)}{\partial \theta} \right]$$

$$\begin{aligned} \log p_{\theta}(x) &= \log e^{-E_{\theta}(x)} - \log \int e^{-E_{\theta}(x)} dx \\ &= -E_{\theta}(x) + \int E_{\theta}(x) dx \end{aligned}$$

$$\begin{aligned} \frac{\partial \log p_{\theta}(x)}{\partial \theta} &= \mathbb{E}_{p_{data}(x)} \left[\frac{\partial E_{\theta}(x)}{\partial \theta} \right] \quad \text{データ分布上の勾配の期待値} \\ &\quad - \mathbb{E}_{p_{\theta}(x)} \left[\frac{\partial E_{\theta}(x)}{\partial \theta} \right] \quad \text{モデル分布上の勾配の期待値} \end{aligned}$$

$p_{\theta}(x)$ のサンプリングはMCMCなどでできる?

$\hookrightarrow$ 1000-10000回程度学習させる $p_{\theta}(x)$ のサンプリングはできない

分配関数 $z(\theta)$ 自体を別のパラメータ c で推定

$$\log p_{\theta}(x) = -\bar{L}_{\theta}(x) - c \text{ を最大化する}$$

$$J(\theta) = \mathbb{E}_{p_{\text{data}}(x)} \left[\log \frac{p_{\theta}(x)}{p_{\theta}(x) + q(x)} \right] \\ + \mathbb{E}_{q(x)} \left[\log \frac{q(x)}{p_{\theta}(x) + q(x)} \right]$$

$q(x)$ はなんでもいい

ノイズ分布, e.g.

Gaussian,

$J(\theta)$ の最大化が θ は対数尤度が最大化され

$$c = z(\theta) = -\text{損失を減らす}$$

直感的には テンソルなどのパラメータノイズを

見合わせるように学習

実は GAN とよく関係

EBMは 計算神経科学での 脳のモデル化に
いまだに よく使われる

最近の EBM : エネルギー関数

その中で NN で定義

→ 全体に 100% 決定的な関数

エネルギー関数の出力を 2 つの層に二分