


· Likelihood - Based Training

$$\begin{cases} p(x) \geq 0 \\ \int_{\mathcal{X}} p(x) dx = 1 \end{cases} \quad \left(\int_{\mathcal{X}} p(x) dx = 1 \right)$$

the "volume" is fixed, when the MLE training increase $P(x_{\text{train}})$,
it decrease $P(x_{\text{test}})$ normally $\therefore$

coming up with a non-negative function $p(x)$ is easy (eg. $g(x) = f(x)^2$ or $\exp(f(x))$ where f is a neural net.)
but $g(x)$ may not sum/integrate to 1 ($\int g(x) = z(x) \neq 1$)

· Energy - Based Model

$$p(x) = \frac{1}{z(x)} \exp(f(x)) \quad z(x) = \int_{\mathcal{X}} \exp(f(x)) dx$$

why exponential:

1. want to capture large variations in probability (otherwise we need highly non-smooth f)
2. log-prob is the natural that we want to work with
3. exponential family: many common distributions can be written in this form
4. these distributions arise under fairly general assumptions in statistical physics $-f(x)$ is called the energy

Pros: can we specify much any $f(x)$

Cons: sampling is hard

evaluating and optimizing likelihood $p(x)$ is hard (need to evaluate $\int \exp(f(x)) dx$)

no feature learning (but can add latent variables)

Given x and x' , evaluating $P_0(x)$ and $P_0(x')$ is hard.

However, their ratio $P_0(x)/P_0(x') = \exp(f_0(x) - f_0(x'))$ does not involve $z(0)$

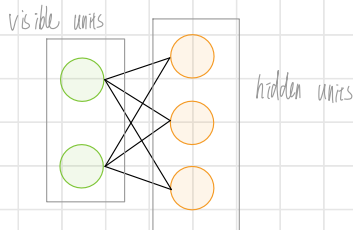
can do anomaly detection, denoising

Restricted Boltzmann Machine

RBM: Energy-based model with latent variable

$$x \in \{0,1\}^n \quad z \in \{0,1\}^m$$

the joint distribution is $P(x, z; w, b, c) = \frac{1}{Z} \exp(x^T w z + b^T x + \tilde{c}^T z)$



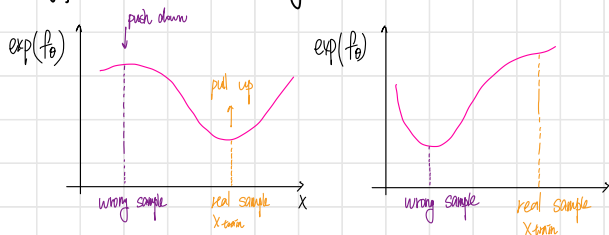
called restrict because there's no x - x connection or z - z connection

(can stack multiple RBM together)

i. the normalizing constant is hard to compute: $Z(w, b, c) = \sum_{x \in \{0,1\}^n} \sum_{z \in \{0,1\}^m} \exp(x^T w z + b^T x + \tilde{c}^T z)$

optimizing the unnormalized prob $\exp(\cdot)$ wrt w, b, c is easy, but optimizing $P(x, z; w, b, c)$, which depends on Z , is hard

Energy-based Model Training



Goal: maximize $\frac{1}{Z(0)} \exp(f_0(x_{\min}))$ increase the numerator, decrease the denominator

Because the model is unnormalized, increasing $\exp(f_0(x_{\min}))$ wrt θ does not guarantee that $x_{\min}$ become more likely compared to other x 's

contrastive divergence:

instead of evaluating $z(\theta)$ exactly, use a Monte-Carlo estimation

contrastive divergence: sample $x_{\text{sample}} \sim p_\theta$

take step on $\nabla_\theta (f_\theta(x_{\text{train}}) - f_\theta(x_{\text{sample}}))$

'i' sampling is hard

Sampling from energy-based models

no direct way to sample from p_θ since we cannot evaluate how likely each possible sample is

However, we can easily compare two samples

Markov chain Monte Carlo:

initialize x^0 randomly, too
repeat {

$$x' = x^0 + \text{noise}$$

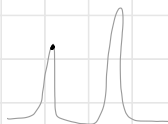
if $f_\theta(x') > f_\theta(x^0)$ let $x^{t+1} = x'$

else let $x^{t+1} = x^0$ with prob $\exp(f_\theta(x') - f_\theta(x^0))$

}

or flip some pixels
gaussian noise

run this procedure for sufficiently large number of iterations
this will give a sample from the distribution



stochastic local sample, may be stuck at local maximum
and only get one with very low prob