

Boltzmann Machine

$x \in \{0, 1\}^N$ observation

$y \in \{0, 1\}^M$
 $z \in \{0, 1\}^K$ } latent variables

$$P(x, y, z) = \frac{1}{Z} \exp(-E(x, y, z))$$

where $Z = \sum_x \sum_y \sum_z \exp(-E(x, y, z))$

and $E(x, y, z) = \sum_{i,j,k} W_{ijk} x_i y_j z_k$

MCMC

→ 直接サンプリングが難しい確率分布の代わりに
そこから近似するサンプル列を生成する

↓
周辺分布、同時分布や期待値などの積分計算を近似するために
用いられる

e.g. Gibbs Sampling は統計的推定やベイズ推定の手法として
頻繁に用いられる → Metropolis-Hastings 法

Variational Bayes や EM アルゴリズムのような

統計的推定法のための理論的な方法の代替法

導出: 同時確率 $(x_1, \dots, x_n)$ から確率変数

$X = \{x_1, \dots, x_n\}$ を n サンプル得た...

n 次元の i 番目のサンプルを $X^{(i)} = \{x_1^{(i)}, \dots, x_n^{(i)}\}$ とする

step 0: 初期値 $X^{(0)}$ を設定

step 1: 条件付き確率 $p(x_j | x_1^{(i)}, \dots, x_{j-1}^{(i)}, x_{j+1}^{(i)}, \dots, x_n^{(i)})$
から i 番目のサンプル $x_j^{(i)}$ をサンプリング

step 2: 1 から以下各 step 1 を loop

$$\begin{aligned} P(x_j | x_1, \dots, x_{j-1}, x_{j+1}, \dots, x_n) \\ &= \frac{P(x_1, \dots, x_n)}{P(x_1, \dots, x_{j-1}, x_{j+1}, \dots, x_n)} = \frac{P(x_1, \dots, x_n)}{\int P(x_1, \dots, x_n) dx_j} \\ &= \frac{1}{Z} P(x_1, \dots, x_n) \\ &\propto P(x_1, \dots, x_n) \end{aligned}$$

April 17th

これらのサンプリング集合は全2の変数に
関する同時分布を近似する

サンプリングの初期化の手段に EM アルゴリズム or ジェネラル (一般化)
burn-in period を導入

また、この値に近づくための様々な方法がある

自己相関を減少させる方法として

- ① 焼きなまし法 (Simulated Annealing)
- ② Collapsed Gibbs Sampling
- ③ Block Gibbs Sampling

1. 状態の組合せ $\mathcal{X} = \{x_1, \dots, x_n\}$ における値が変化すると
エネルギー関数 E を定義

2. E が与えられた状態において生じやすいような
 $p(x|\theta)$ を定義

- 一般的なエネルギー関数

$$E(x, \theta) = - \sum_{i=1}^n b_i x_i - \sum_{(i,j) \in \mathcal{E}} w_{ij} x_i x_j$$

$\mathcal{E}$: 入力集合

b_i : i へのバイアス

w_{ij} : x_i, x_j の重み

$\mathcal{X}$: $\{x_1, \dots, x_n\}$

$\mathcal{X}$ における状態の
の確率分布 $p(x|\theta)$

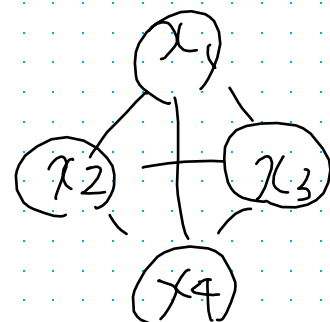
$$p(x|\theta) = \frac{1}{Z(\theta)} \exp(-E)$$

この $Z(\theta)$ は 正規化関数

$$Z(\theta) = \sum_x \exp(-E) = \sum_{x_1} \sum_{x_2} \dots \sum_{x_n} \exp(-E)$$

2^n 通りの計算

1. 対数尤度を最大化
2. $\therefore$ 1000-7 程度で grad 計算
3. GD で θ を update



$$\ln L(\theta) = \sum_{n=1}^N (-E(x_n, \theta) - \ln Z(\theta))$$

$$p(x|\theta) = \frac{1}{Z(\theta)} \exp(-E(x, \theta))$$

$$E(x, \theta) = -\sum_{i=1}^N b_i x_i - \sum_{(i,j) \in \{x, x_j\}} w_{ij}$$

$$\frac{\partial \ln L(\theta)}{\partial w_{ij}} = \sum_{n=1}^N \underbrace{x_n(i) x_n(j)} - N \mathbb{E}_\theta [x_i x_j]$$

x_n の添字の要否

$$x_i = \begin{pmatrix} x_i(0) \\ x_i(1) \\ \vdots \\ x_i(N) \end{pmatrix}$$

分布関数

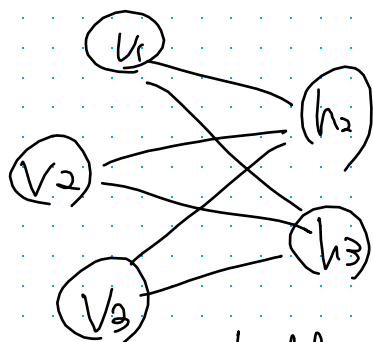
$$\mathbb{E}_\theta [x_i x_j] = \sum_x x_i x_j p(x|\theta)$$

の厳密な計算は難しい

$$\begin{aligned} & \left\langle x_i x_j \right\rangle_{\text{データ平均}} & \left\langle x_i x_j \right\rangle_{\text{モデル平均}} \\ & = \sum_x x_i x_j q(x) & - N \mathbb{E}_\theta [x_i x_j] \end{aligned}$$

visible variable

$$z = \{v_1, \dots, v_J, h_1, \dots, h_K\}$$



hidden variable

$$p(z|\theta) = \frac{1}{Z(\theta)} \exp(-E(z, \theta))$$

hidden layer/2L2

300-1500

$$p(z|\theta) = p(h, v|\theta) \quad h \text{ is hidden variable}$$

$$p(v|\theta) = \sum_h p(v, h|\theta)$$

$$\frac{\partial \ln L(\theta)}{\partial w_{ij}}$$

$$\propto \langle z_i z_j \rangle_{\text{data}} - \langle z_i z_j \rangle_{\text{model}}$$

Assumption each v_i, h_j doesn't have edgesbetween v_i, v_j ,
 h_i, h_j .

$$F(v, h|\theta) = -\sum_{i=1}^J a_i v_i - \sum_{i=1}^K b_i h_i - \sum_{(i,j) \in E} w_{ij} z_i z_j$$

$$p(h|v, \theta) = \prod_j p(h_j | v, \theta) \quad \text{conditional distribution}$$

$$p(h_j = 1 | v, \theta) = \sigma(b_j + \sum_i w_{ij} v_j)$$

$$\frac{1}{N} \frac{\partial \ln L(\theta)}{\partial w_{ij}} = \sum_{n=1}^N w_{ij} p(h_i = 1 | v)$$

$$- \sum_{n,h} w_{ij} h_i p(n, h | \theta)$$

に44,52変

Gibbs Sampling 2: 近似的

2: 1の条件を確率

$$p(x_i | x_{-i}, \theta) = \frac{\exp \{ (b_i + \sum_{j \in N} w_{ij} x_j) x_i \}}{1 + \exp \{ (b_i + \sum_{j \in N} w_{ij} x_j) x_i \}}$$

$$1 + \exp \{ (b_i + \sum_{j \in N} w_{ij} x_j) x_i \}$$

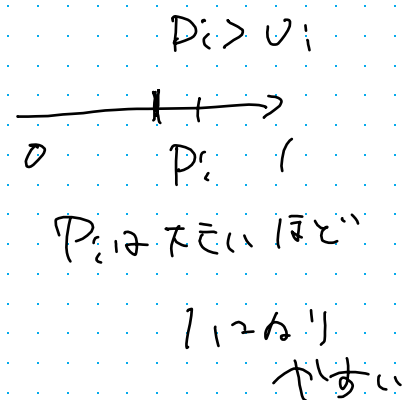
あ22: + x_i x_j x_j
 の x_i は 1 か 0 2: 1
 か 0 か 1 (x_i = 1)
 1 x_i x_j x_j x_j

$p(x_i | x_{-i}, \theta)$ の計算は

① $p_i \sim p(x_i = 1 | x_{-i}, \theta)$ を生成

② $u_i \sim U(0, 1)$
 if $p_i > u_i$
 $p_i = 0$?
 else
 $p_i = 1$?

p_i を利用して
 $\sigma = \sigma(1 - p_i)$ を



この $S_i = \sigma$ を生成した後は

$p(x_i | x_{-i}, \theta)$ に従う

実際の RBM では Gibbs Sampling の計算コストが高い
 ので

Contrastive Divergence を用いる

v と h の独立性を 利用

v と h を 交互に $\sigma(\cdot)$ を 用いる

1. $v^0 \leftarrow \text{random init}$

2. $v^0 \rightarrow h^0 \rightarrow v^1 \rightarrow h^1$ と 交互に $\sigma(\cdot)$ を 用いる

Contrastive Divergence

↳ Gibbs sampling の init を random として
互いの分布 p と q ($v^0 = v_u$) を 用いる

↓

T 回 交互に 用いる ($T = 1, 2, 3, \dots$)

この方法の収束速度を $\approx$ contrastive divergence

を用いる

これ (CD) に

・ weight decay

最適化を行う。

・ momentum

・ sparse regularization など 追加する

Persistent CD

通常のCDでは $v_0 = v_n$ と v

5) 10倍
高速

$v_0 \rightarrow h_0 \rightarrow v_1$ と $\gamma = \gamma' = \gamma''$ と v

PCDでは v_0 に 前回の γ の更新時に

$\gamma = \gamma' = \gamma''$ と v を用いる

深層学習モデルの構造

$$X = \{V, H\}$$

$$P(X=x|\theta) = P(V=v, H=h|\theta) = \frac{1}{Z(\theta)} \exp(-E(v, h, \theta))$$

$$p(x|\theta) = p(v|\theta) = \sum_h p(h, v|\theta)$$

h は周辺化した分布の最大値をとる

h は観測 v に対する

Restricted Boltzmann Machine

from p 3 f / p 66

Aizu Univ.

$$\frac{\partial \mathcal{L}_D(\theta)}{\partial w_{ij}} = N \sum_{u,h} x_i x_j p(h|u, \theta) g_D(u)$$

$$= N \mathbb{E}_{p(v, h | \theta)} [x_i x_j]$$

$$p(h|u, \theta) = \frac{p(u, h | \theta)}{p(u | \theta)}$$

石田 孝 - 3月2日

$$N \sum_{u,h} x_i x_j p(h|u, \theta) g_D(u)$$

$$= N \cdot \frac{1}{N} \sum_{u=1}^N \sum_h f(u^{(h)}, h) p(h|u^{(u)}, \theta)$$

$$= \frac{1}{N} \sum_{u=1}^N \mathbb{E}_{p(h|u^{(u)}, \theta)} [f(u^{(h)}, h)]$$

問題点: computational explosion.

⇒ Gibbs Sampling

BM の学習 $\rightarrow$ 最大推定 $\rightarrow$ 最適化

$$L_D(\theta) = \prod_{n=1}^N p(V = u^{(n)} | \theta)$$

$$L_D(\theta) = \ln L_D(\theta) = \sum_{n=1}^N \ln p(u^{(n)} | \theta)$$

$$\frac{\partial L_D(\theta)}{\partial w_{ij}} = \sum_{n=1}^N u_i^{(n)} v_j^{(n)} - N \mathbb{E}_{p(u|\theta)}[V_i V_j]$$

expectation data distribution

expectation of model distribution,

where

$$\mathbb{E}_{p(u|\theta)}[f(u)] = \sum_u f(u) p(u|\theta)$$

(learning equation of Boltzmann machine)

$$\frac{\partial L_D(\theta)}{\partial w_{ij}} = 0 \text{ as } \hat{=}$$

moment matching
 $\downarrow$

$$\frac{1}{N} \sum_{n=1}^N u_i^{(n)} v_j^{(n)} = \mathbb{E}_{p(u|\theta)}[u_i v_j]$$

Gradient Ascent

↳ need to calc

$$\frac{\partial \mathcal{L}_D(\theta)}{\partial \omega_{ij}} = \sum_{n=1}^N u_i^{(n)} u_j^{(n)} - \underbrace{N \mathbb{E}_{p(u|\theta)} [v_i v_j]}_{\downarrow}$$

combination

explosion.

$$\sum_n u_i u_j p(u|\theta)$$

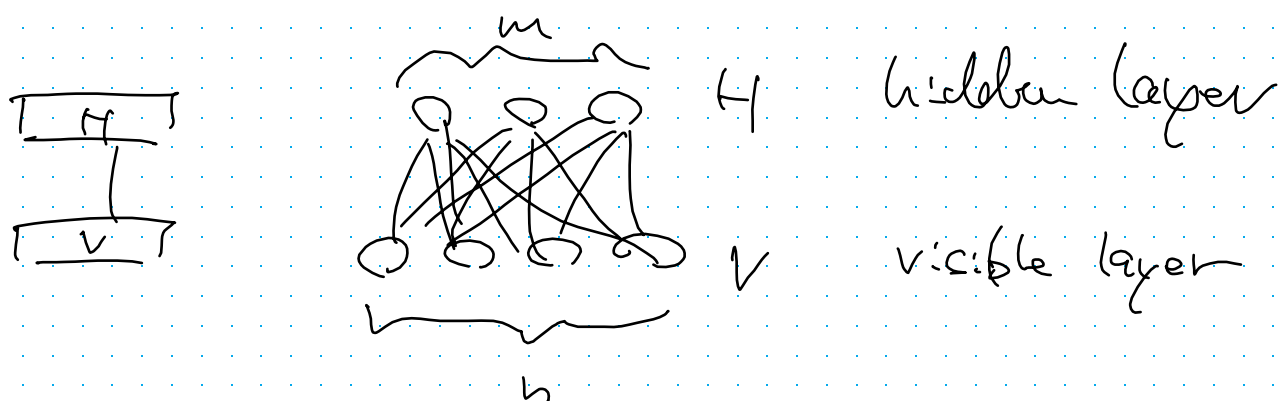
ホダ-が12個の項の和で計算する

(d は v の次元数
 $u_1 \dots u_d$ の d .)

↳ which is defined on complete bipartite graph

(完全二部グラフ)

as BM w/ hidden latent variable



as limitation,

there is no connection

between same layer.

n 個の要素を持つ観測

データ点の集合を $x_1, x_2, \dots, x_n$ とする

$$p(h|v, \theta) = \frac{p(v, h|\theta)}{\sum_h p(v, h|\theta)} = \frac{p(v, h|\theta)}{p(v|\theta)}$$

同様に

$$p(v|h, \theta)$$

$$= \prod_{j \in H} p_{sb}(h_j | \lambda_j)$$

$$= \prod_{j \in H} \underbrace{p_{st}(h_j | \lambda_j)}$$

$$\rightarrow p_{sb}(h_j | \lambda_j) = \frac{\exp(\lambda_j x_j)}{1 + \exp(\lambda_j)}$$

sigmoid belief

Training of RBM

$$p(x|\theta) = p(v|\theta) = \sum_h p(h, v|\theta)$$

$$= \frac{1}{Z(\theta)} \exp\left(\sum_{i \in V} b_i x_i + \sum_{j \in H} \lambda_j (1 + \exp(\lambda_j))\right)$$

↑

可変変数のみの周知

RBMと同様に重み合わせを調整する

↳ contrastive divergence 有るのと同じくらい必要

gibbs sampling は $t=1$ の 反復回数 T くらい必要

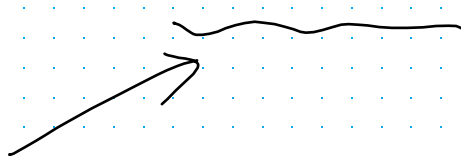
↓

1. $t=1$ で u を $u \sim U(0,1)$ の変数として与える

step 1

u を $u \sim U(0,1)$ と初期化して u を $u \sim U(0,1)$ と初期化する

↓
hidden layer u の $prob$ の $probabilities$



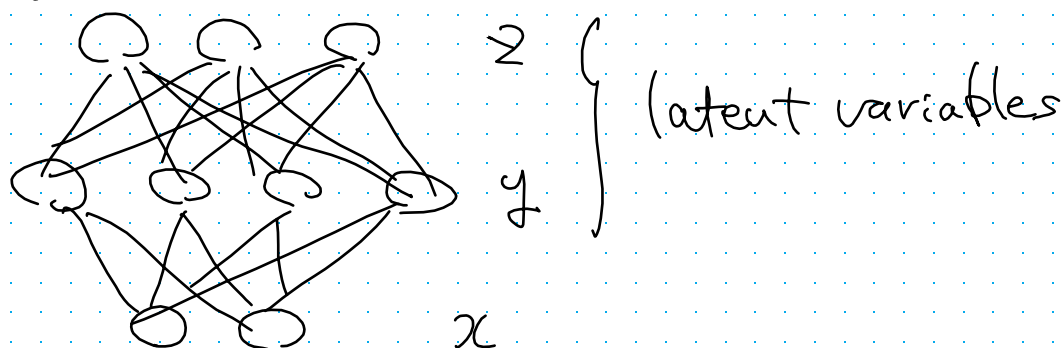
ここで $u \sim U(0,1)$ と初期化

$u > prob \rightarrow 1$

$u \leq prob \rightarrow 0$

Q1. Describe Block-Gibbs Sampling Strategy for RBM

Graph G



At first, we assume that there are no connections between the same layers.

So, now, we can divide each node into two groups as follows

$$\begin{aligned} g_1 &= \{x, z\} & x &= \{x_1, x_2, \dots\} \\ g_2 &= \{y\} & y &= \{y_1, y_2, \dots\} \\ & & z &= \{z_1, z_2, \dots\} \end{aligned}$$

by splitting G into two parts as bipartite graph.

Given g_1 , it is easy to compute the conditional probability of g_2

$$P(g_2 | g_1, \theta) = \prod_{i=1} P(g_{2i} | g_{1i}, \theta) = \frac{P(x, y, z)}{P(x, z)} \dots *1$$

and vice-versa

$$P(g_1 | g_2, \theta) = \prod_{j=1} P(g_{1j} | g_{2j}, \theta) = \frac{P(x, y, z)}{P(y)} \dots *2$$

Where

$$*1 = \frac{P(x, y, z)}{\sum_x \sum_z P(x, y, z)} = \frac{\exp(-\sum_{ijk} w_{xyz})}{\sum_x \sum_z (\exp(-\sum_{ijk} w_{xyz}))}$$

$O(2^{n+k})$

$$*2 = \frac{P(x, y, z)}{\sum_y P(x, y, z)} = \frac{\exp(-\sum_{ijk} w_{xyz})}{\sum_y (\exp(-\sum_{ijk} w_{xyz}))}$$

$O(2^m)$

April 18th

Block Gibbs Sampling Algorithm

step 1 initialize $x^{(0)}, y^{(0)}, z^{(0)}$ at random

$$t = 1$$

step t if $(1 < t < k)$

$$x^{(t+1)}, z^{(t+1)} \sim p(x, z | y^{(t)}, \theta) \text{ by using } \star 1$$

$$y^{(t+1)} \sim p(y | x^{(t+1)}, \theta) \text{ by using } \star 2$$

$$t = t + 1$$

Step k

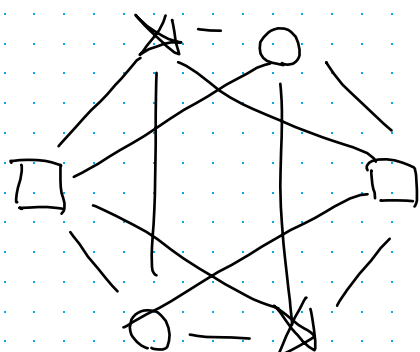
return

$$\left. \begin{array}{c} \textcircled{x^{(1)} y^{(1)} z^{(1)}} \\ \vdots \\ \textcircled{x^{(k)} y^{(k)} z^{(k)}} \end{array} \right\} k \text{ sample.}$$

April 18th

cost of block gibbs sampling

Multipartite Graph (2, 2, 2)



$$x = x$$

$$o = y$$

$$\square = z$$

$$P(x_i | y, z) = P(x_i | x_{-i}, y, z)$$

$$= \frac{P(x, y, z)}{P(x_{-i}, y, z)} \propto p(x, y, z)$$

$$\text{So } P(x_i = 0 | y, z) = \frac{\tilde{p}(x_i = 0 | y, z)}{\tilde{p}(x_i = 0 | y, z) + \tilde{p}(x_i = 1 | y, z)}$$

$$\tilde{p}(x_i = 1 | y, z) = \tilde{p}(x_i = 1, x_{-i} = 0, y = y, z = z)$$

$$\text{where } p = \frac{1}{2} \tilde{p}$$

for calculation

$$\hat{p}(x=x_1, x_{-1}, y=y_1, z=z_1)$$

but
we have to
wait mixing
time

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

$$\tilde{p}(x, y, z) = \exp(-E(x, y, z)) \leftarrow$$

$$E(x, y, z) = \sum_{i,j,k} w_{ijk} x_i y_j z_k$$

$$x \in \{0,1\}^N$$

$$y \in \{0,1\}^M$$

$$z \in \{0,1\}^K$$

it costs

$$O(N \times M \times K)$$

$$Z = \sum_x \sum_y \sum_z \exp(-E(x, y, z))$$

it costs $O(2^{M+N+K})$

we can save this cost
by block gibbs sampling