

probabilistic model

$p(x|\theta)$ where θ is vector

N training examples $x_i, i \in [1, N]$

are drawn independently

from the

data distribution

6(a) scheme of estimating θ

to maximize the log likelihood

of N training
example.

ここで θ と θ' は互いに独立に θ と θ' と同様に

$p(C_k|x)$ は θ と θ' のどちらでも計算できる

→ θ と θ' は互いに独立に θ と θ' と同様に

$$p(C_k|x) = \frac{p(x|C_k) p(C_k)}{p(x)}$$

$$p(C_k(x)) = \frac{p(x|C_k) p(C_k)}{p(x)}$$

これは、クラスに属する

データの条件付き分布に等しい



クラスに属する 類似度の高いデータ への

重み付けをかける

識別モデル

$P(x|y)$

$$\hat{y} = \underset{y}{\operatorname{argmax}} P(y|x)$$

↑
ここを直接モデル化

識別モデルは直接確率を直接求めるモデル

生成モデルでは $P(y|x)$ を直接はモデル化しない

事後分布

$$P(y|x) =$$

$$\frac{\overset{\text{尤度}}{P(x|y)} \overset{\text{事前分布}}{P(y)}}{P(x)}$$

この値が最大になる

つまり y に分類する

分類時に

分母 $P(x)$ は分類には依存しないので

$$\hat{y} = \underset{y}{\operatorname{argmax}} P(x|y) P(y)$$

ケースの割合

全体のうちケース A の

ケースが 60% かつ

$$P(y) = 0.6$$

観測したケースは無作為

に生成されるのではなく

何らかの分布に基づいて生成

この分布が "y のケース" ごとに異なり、各ケースがデータを生成する確率

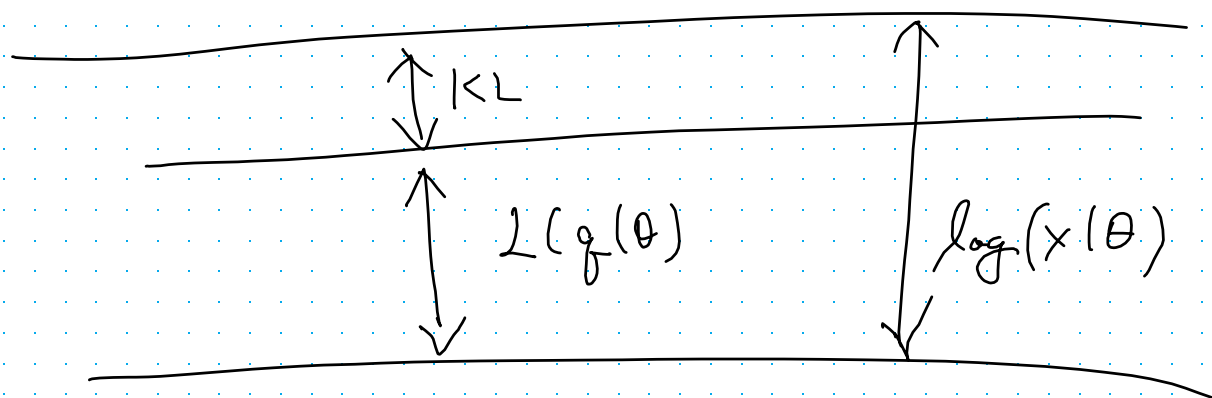
多峰性のある場合 点推定では

$$\log (X|\theta) = \underbrace{L(q|\theta)}_{\text{期待値}} + KL(q||p) \quad \begin{array}{l} \text{隠れた情報} \\ \text{を無視} \end{array}$$

expectation-maximization

期待値を最大にするアルゴリズム

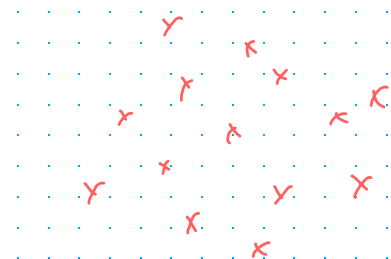
下界



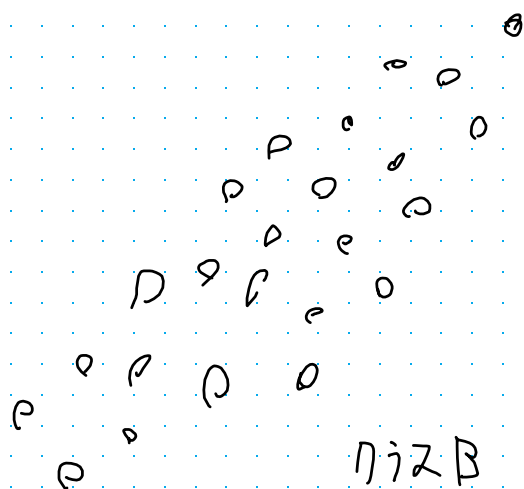
$$L(q, \theta) = \sum_z q(z) \log \left\{ \frac{p(X, z | \theta)}{q(z)} \right\}$$

$$KL(q||p) = \sum_z q(z) \log \left\{ \frac{q(z)}{p(z|X, \theta)} \right\}$$

E step θ を固定し、 KL 最小化する q を求めるM step 下界を (X, θ) に関して最大化

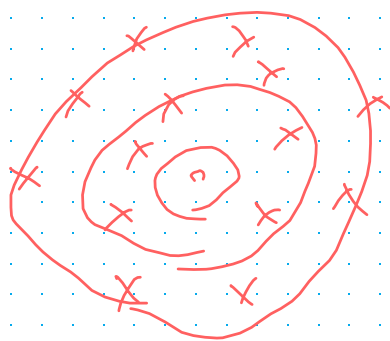


グループA

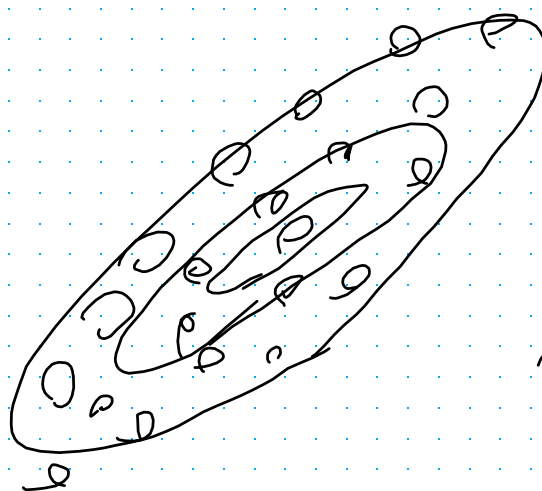


グループB

上記の2つのデータ群の観測値を1次元にまとめ



グループA



グループB

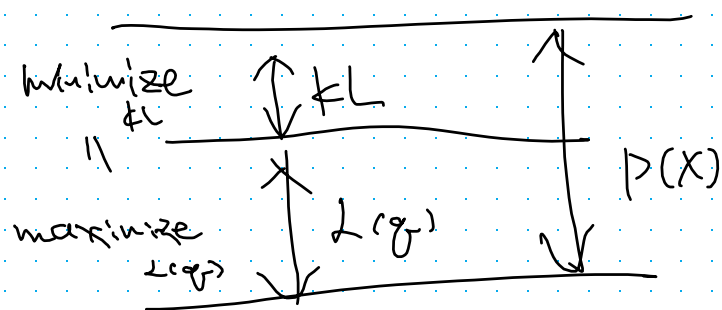
$P(y|x)$ は 与えられた x のときの y の確率, $P(x|y)$ $P(y)$ は y の確率
 グループAの観測値を1次元にまとめ、あるグループに属する
 観測値のデータ群を1次元にまとめよう

$$\log p(x) = \mathcal{L}(q) + KL(q||p)$$

変分モデル

エビデンス

0を更新して不変



$p(x)$ は q に依存しないので

q に関して KL を最小化すると

q に関して 下界 $\mathcal{L}(q)$ を

最大化するこれにあたる

$p(x|z)$ を求めたい

↓ 計算が難しい

$q(z)$ で近似 (下界)

↓

$KL(q||p)$ を最小化する q を
求める