



Given an autoregressive model $P(x) = \prod_{i=1}^n P(x_i | p(x_{1:i-1}); \theta_i)$

$p(x_i) = x_1 \dots x_{i-1}$ in an autoregressive model

dataset $D = \{x^{(1)} \dots x^{(m)}\}$

$$L(\theta) = \prod_{j=1}^m \prod_{i=1}^n P(x_j^{(i)} | p(x_j^{(i)}); \theta_j)$$

$$-L(\theta) = \sum_{j=1}^m \sum_{i=1}^n \log P(x_j^{(i)} | p(x_j^{(i)}); \theta_j)$$

↓

$$L_j(\theta_j) = \sum_{i=1}^n \log P(x_j^{(i)} | p(x_j^{(i)}); \theta_j)$$

$$\max_{\theta_j} \sum_{i=1}^n \log P(x_j^{(i)} | p(x_j^{(i)}); \theta_j)$$

no close-form solution

decomposition of Bayesian Network
train each model $P(x_i | p(x_i); \theta_i)$ separately

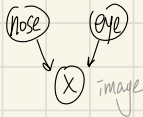
$$\begin{aligned} \nabla_{\theta_j} L_j(\theta_j) &= \sum_{i=1}^n \nabla_{\theta_j} \log P(x_j^{(i)} | p(x_j^{(i)}); \theta_j) \\ &= m \cdot \frac{1}{m} \sum_{i=1}^n \nabla_{\theta_j} \log P(x_j^{(i)} | p(x_j^{(i)}); \theta_j) \\ &= m \cdot E_{x_j \sim p} [\nabla_{\theta_j} \log P(x_j | p(x_j); \theta_j)] \end{aligned}$$

} Monte-Carlo view of stochastic gradient descent

Latent Variable Models

eg. human face images

the variability in images due to some high-level features (eye color, facial expression, etc.)
these factors are not explicitly given (latent)



latent variables correspond to high-level features

if z chose properly $p(x|z)$ is much simpler than $p(x)$
having trained these model, can identify features via $p(z|x)$

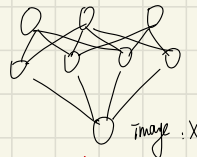
} ∴ hard to specify latent vars

Deep Latent Variable Models

$$z \sim \mathcal{N}(0, I)$$

$$x|z \sim \mathcal{N}(U_0(z), \Sigma_0(z)) \text{ where } (U_0, \Sigma_0) \text{ are neural networks}$$

z will hopefully corresponds to meaningful latent variations (features)
features can be computed via $p(z|x)$



unsupervised representation learning

eg. Mixture of Gaussian is a shallow Latent Variable Model

Variational AutoEncoder

$$z \sim \mathcal{N}(0, I)$$

$$x|z \sim \mathcal{N}(\mu_\theta(z), \Sigma_\theta(z)) \quad \text{where } \mu_\theta(\cdot), \Sigma_\theta(\cdot) \text{ are neural networks}$$

$$\mu_\theta(z) = \phi(Az + c)$$

$$\Sigma_\theta(z) = \text{diag}(\exp(Bz + d))$$

} mixture of infinite Gaussians

training is hard

$$\max_{\theta} p(x) = \max_{\theta} \int_{\mathbb{R}} p(x) p(x|z; \theta) dz \quad \text{infeasible}$$