

Machine Learning (part II)

Deep Generative Models

Angelo Ciaramella

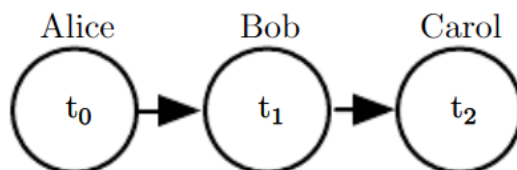
Structured probabilistic models

- Structured probabilistic models
 - the **structure of the model** is defined by a **graph**, these models are often also referred to as **graphical models**
- Challenge of Unstructured Modeling
 - Density estimation
 - Data generating distribution
 - Denoising
 - Missing value imputation
 - Distribution over unobserved elements
 - Sampling
 - Speech synthesis



Direct models

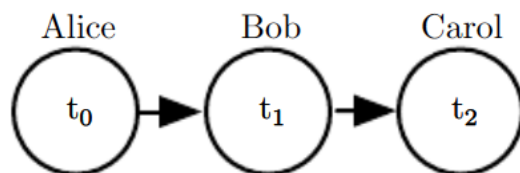
- Directed graphical model
 - belief network or Bayesian network
- Example - relay race
 - suppose we name Alice's finishing time t_0 , Bob's finishing time t_1 , and Carol's finishing time t_2
 - Our estimate of t_2 depends directly on t_1 but only indirectly on t_0



Direct models

- Directed graphical model
 - General local conditional probability distributions

$$p(\mathbf{x}) = \prod_i p(x_i \mid \text{Pa}_{\mathcal{G}}(x_i))$$



$$p(t_0, t_1, t_2) = p(t_0)p(t_1 \mid t_0)p(t_2 \mid t_1)$$



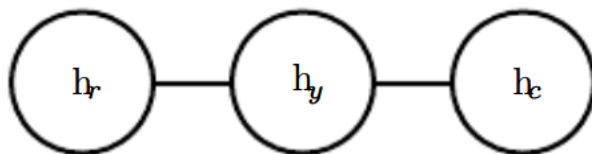
Undirect models

■ Undirect models

- Markov random fields (MRFs) or Markov networks
- when the interactions seem to have no intrinsic direction, or to operate in both directions, it may be more appropriate to use an undirected model

■ Example

- whether or not you are sick, whether or not your coworker is sick, and whether or not your roommate is sick (coworker and roommate do not know each other)



Undirect models

■ Undirected graphical model

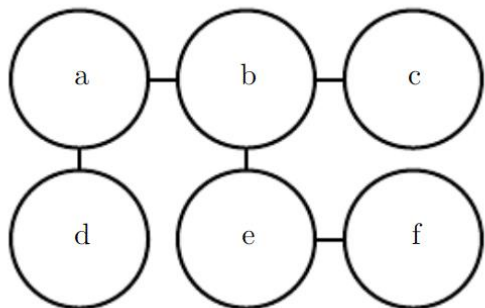
- structured probabilistic model defined on an undirected graph

- Clique potential

 - Factor $\phi(\mathcal{C})$

 - the affinity of the variables in that clique for being in each of their possible joint states

$$\tilde{p}(\mathbf{x}) = \prod_{\mathcal{C} \in \mathcal{G}} \phi(\mathcal{C}) \quad \text{Unnormalized probability distribution}$$



$$\frac{1}{Z} \phi_{a,b}(a, b) \phi_{b,c}(b, c) \phi_{a,d}(a, d) \phi_{b,e}(b, e) \phi_{e,f}(e, f)$$



Partition function

- Unnormalized probability distribution
 - guaranteed to be non-negative everywhere, it is not guaranteed to sum or integrate to 1
 - valid probability distribution, we must use the corresponding normalized probability distribution

$$p(\mathbf{x}) = \frac{1}{Z} \tilde{p}(\mathbf{x})$$

$$Z = \int \tilde{p}(\mathbf{x}) d\mathbf{x} \quad \text{intractable to compute}$$

Z as a constant when the ϕ functions are held constant. Note that if the ϕ functions have parameters, then Z is a function of those parameters



Energy-based models

- Many interesting theoretical results about undirected models depend on the assumption that

$$\forall \mathbf{x}, \tilde{p}(\mathbf{x}) > 0$$

- A convenient way to enforce this condition is to use an **Energy-Based Model (EBM)** where

$$\tilde{p}(\mathbf{x}) = \exp(-E(\mathbf{x}))$$

Boltzmann
distribution

Energy function

Because $\exp(z)$ is positive for all z , this guarantees that no energy function will result in a probability of zero for any state x



Energy-based models

- Energy-based model is just a special kind of Markov network
 - the exponentiation makes each term in the energy function correspond to a factor for a different clique

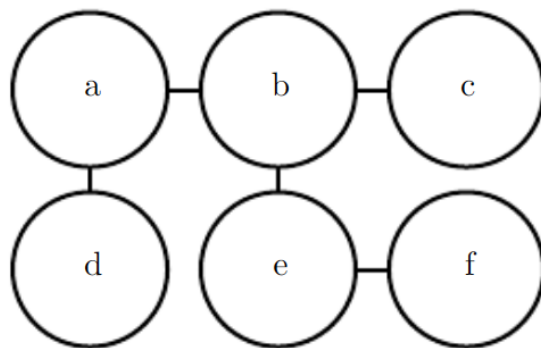


Figure 16.5: This graph implies that $E(a, b, c, d, e, f)$ can be written as $E_{a,b}(a, b) + E_{b,c}(b, c) + E_{a,d}(a, d) + E_{b,e}(b, e) + E_{e,f}(e, f)$ for an appropriate choice of the per-clique energy functions. Note that we can obtain the ϕ functions in Fig. 16.4 by setting each ϕ to the exponential of the corresponding negative energy, e.g., $\phi_{a,b}(a, b) = \exp(-E(a, b))$.



Deep Belief Networks

- Deep belief networks (DBN)
 - first non-convolutional models to successfully admit training of deep architectures
 - Deep belief networks demonstrated that deep architectures can be successful
 - Generative models with several layers of latent variables
 - Latent variables are typically binary
 - Visible units may be binary or real
 - no intra-layer connections
 - Usually, every unit in each layer is connected to every unit in each neighboring layer
 - construct more sparsely connected DBNs
 - The connections between the top two layers are undirected
 - A DBN with only one hidden layer is a Restricted Boltzmann Machine (RBM)



Boltzmann Machines

■ BM

- connectionist approach to learning arbitrary probability distributions over binary vectors

■ Energy function

$$P(\mathbf{x}) = \frac{\exp(-E(\mathbf{x}))}{Z}$$

$$\mathbf{x} \in \{0, 1\}^d$$

joint probability
distribution over the
observed variables

$$E(\mathbf{x}) = -\mathbf{x}^\top \mathbf{U} \mathbf{x} - \mathbf{b}^\top \mathbf{x}$$



Boltzmann Machines

■ Powerful when

- not all the variables are observed
- latent variables act similarly to **hidden units** in MLP
- Universal approximator of probability mass functions over discrete variables

■ Energy function

$$E(\mathbf{v}, \mathbf{h}) = -\mathbf{v}^\top \mathbf{R} \mathbf{v} - \mathbf{v}^\top \mathbf{W} \mathbf{h} - \mathbf{h}^\top \mathbf{S} \mathbf{h} - \mathbf{b}^\top \mathbf{v} - \mathbf{c}^\top \mathbf{h}$$

■ Learning

- intractable partition function



Intractable partition functions

- Normalized probability distribution

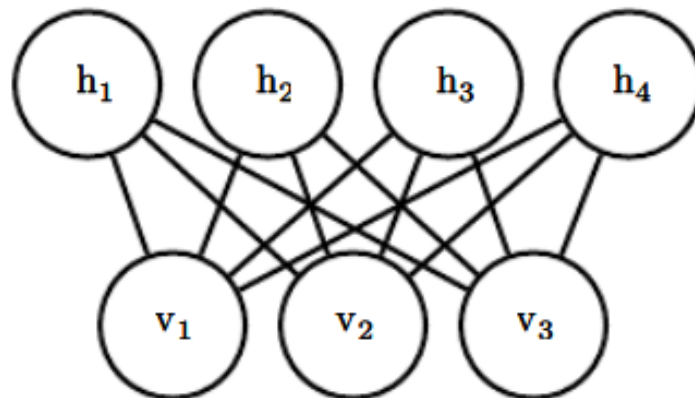
$$p(\mathbf{x}; \boldsymbol{\theta}) = \frac{1}{Z(\boldsymbol{\theta})} \tilde{p}(\mathbf{x}; \boldsymbol{\theta})$$

- Techniques used for training and evaluating models
 - Log-Likelihood Gradient
 - Stochastic Maximum Likelihood
 - Markov Chain Monte-Carlo sampling
 - Contrastive Divergence (CD)
 - Pseudolikelihood
 - Score Matching and Ratio Matching
 - Noise-Contrastive Estimation
 - Annealed Importance Sampling
 - Bridge Sampling



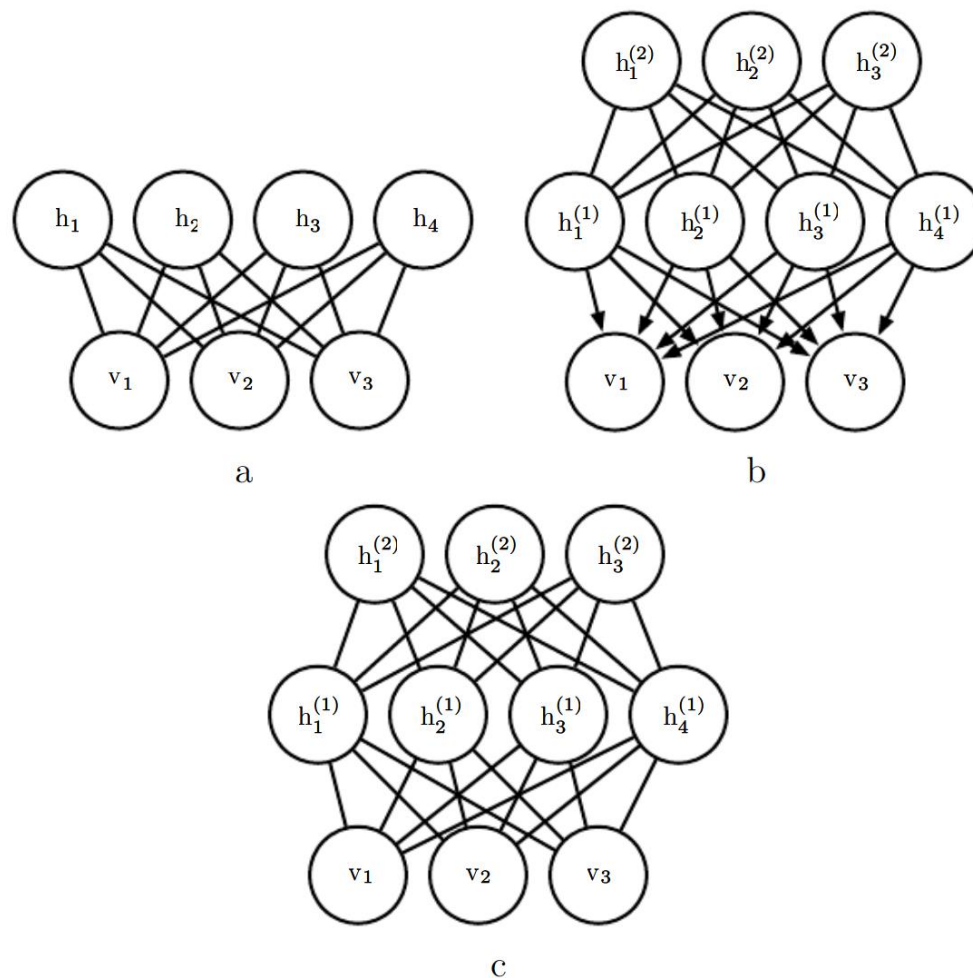
Restricted Boltzmann Machines

- Restricted Boltzmann Machines (RBM)
 - Energy based model (named Harmonium, 1986)
 - undirected probabilistic graphical model
 - a layer of observable variables and a single layer of latent variables
 - **Restricted**: No direct interactions between any two visible units or between any two hidden units



Restricted Boltzmann Machines

■ Stacked Boltzmann Machines (deeper models)



RBM

- Joint probability distribution

$$P(\mathbf{v} = \mathbf{v}, \mathbf{h} = \mathbf{h}) = \frac{1}{Z} \exp(-E(\mathbf{v}, \mathbf{h}))$$

- Energy function

$$Z = \sum_{\mathbf{v}} \sum_{\mathbf{h}} \exp\{-E(\mathbf{v}, \mathbf{h})\}$$

$$E(\mathbf{v}, \mathbf{h}) = -\mathbf{b}^\top \mathbf{v} - \mathbf{c}^\top \mathbf{h} - \mathbf{v}^\top \mathbf{W} \mathbf{h}$$



RBM properties

■ Structure

$$p(\mathbf{h} \mid \mathbf{v}) = \prod_i p(h_i \mid \mathbf{v})$$

$$p(\mathbf{v} \mid \mathbf{h}) = \prod_i p(v_i \mid \mathbf{h}).$$

$$P(h_i = 1 \mid \mathbf{v}) = \sigma \left(\mathbf{v}^\top \mathbf{W}_{:,i} + b_i \right),$$

$$P(h_i = 0 \mid \mathbf{v}) = 1 - \sigma \left(\mathbf{v}^\top \mathbf{W}_{:,i} + b_i \right)$$

Block Gibbs sampling

■ Easy to take its derivatives

$$\frac{\partial}{\partial W_{i,j}} E(\mathbf{v}, \mathbf{h}) = -v_i h_j$$



■ Deriving the conditional distributions

$$\begin{aligned}P(\mathbf{h} \mid \mathbf{v}) &= \frac{P(\mathbf{h}, \mathbf{v})}{P(\mathbf{v})} \\&= \frac{1}{P(\mathbf{v})} \frac{1}{Z} \exp \left\{ \mathbf{b}^\top \mathbf{v} + \mathbf{c}^\top \mathbf{h} + \mathbf{v}^\top \mathbf{W} \mathbf{h} \right\} \\&= \frac{1}{Z'} \exp \left\{ \mathbf{c}^\top \mathbf{h} + \mathbf{v}^\top \mathbf{W} \mathbf{h} \right\} \\&= \frac{1}{Z'} \exp \left\{ \sum_{j=1}^{n_h} c_j h_j + \sum_{j=1}^{n_h} \mathbf{v}^\top \mathbf{W}_{:,j} h_j \right\} \\&= \frac{1}{Z'} \prod_{j=1}^{n_h} \exp \left\{ c_j h_j + \mathbf{v}^\top \mathbf{W}_{:,j} h_j \right\}\end{aligned}$$



Training RBM

- RBM training

- training models that have **intractable partition functions**

- CD

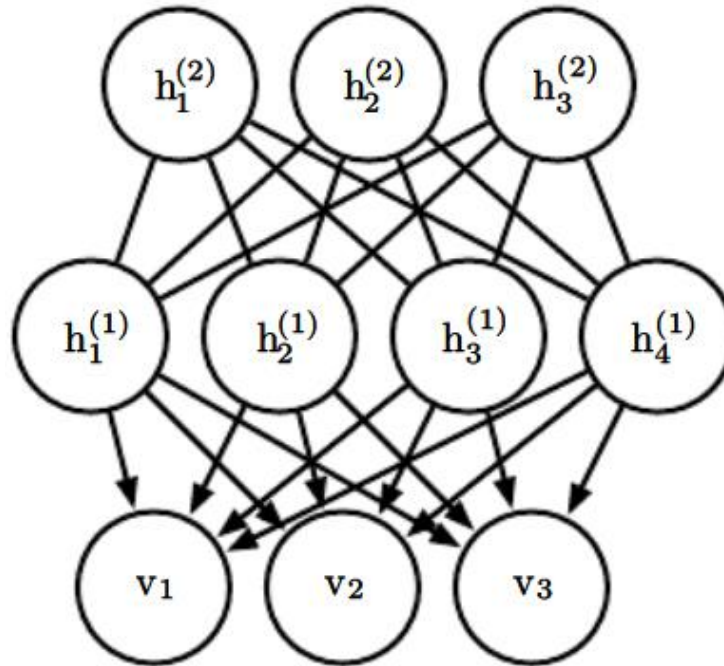
- SML

- Ratio matching

- ...



DBNs



The connections between the top two layers are undirected.

The connections between all other layers are directed

A deep belief network is a hybrid graphical model involving both directed and undirected connections

Deep belief networks demonstrated that deep architectures can be successful



■ Probability distribution

$$P(\mathbf{h}^{(l)}, \mathbf{h}^{(l-1)}) \propto \exp \left(\mathbf{b}^{(l)\top} \mathbf{h}^{(l)} + \mathbf{b}^{(l-1)\top} \mathbf{h}^{(l-1)} + \mathbf{h}^{(l-1)\top} \mathbf{W}^{(l)} \mathbf{h}^{(l)} \right),$$

$$P(h_i^{(k)} = 1 \mid \mathbf{h}^{(k+1)}) = \sigma \left(b_i^{(k)} + \mathbf{W}_{:,i}^{(k+1)\top} \mathbf{h}^{(k+1)} \right) \forall i, \forall k \in 1, \dots, l-2,$$

$$P(v_i = 1 \mid \mathbf{h}^{(1)}) = \sigma \left(b_i^{(0)} + \mathbf{W}_{:,i}^{(1)\top} \mathbf{h}^{(1)} \right) \forall i.$$

$$\mathbf{v} \sim \mathcal{N} \left(\mathbf{v}; \mathbf{b}^{(0)} + \mathbf{W}^{(1)\top} \mathbf{h}^{(1)}, \beta^{-1} \right)$$

Real-valued visible units



Training DBN

- Learning

- Training an RBM to maximize by CD or SML

$$\mathbb{E}_{\mathbf{v} \sim p_{\text{data}}} \log p(\mathbf{v})$$

- Second RBM maximize

$$\mathbb{E}_{\mathbf{v} \sim p_{\text{data}}} \mathbb{E}_{\mathbf{h}^{(1)} \sim p^{(1)}(\mathbf{h}^{(1)}|\mathbf{v})} \log p^{(2)}(\mathbf{h}^{(1)})$$

- procedure can be repeated indefinitely on many layers



DBF as generative model

- Improve classification models

- weights from the DBN and use them to define an MLP

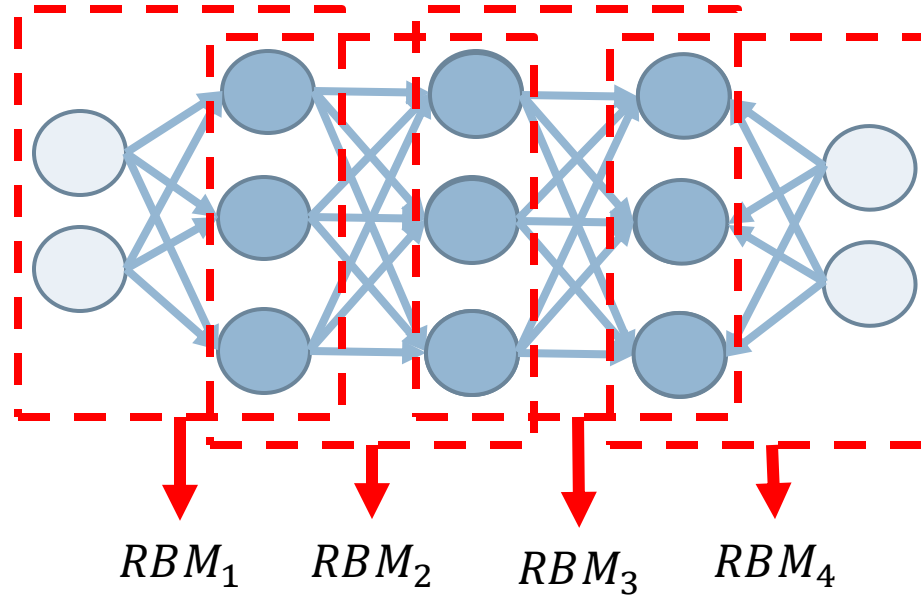
$$\mathbf{h}^{(1)} = \sigma \left(\mathbf{b}^{(1)} + \mathbf{v}^\top \mathbf{W}^{(1)} \right).$$

$$\mathbf{h}^{(l)} = \sigma \left(\mathbf{b}_i^{(l)} + \mathbf{h}^{(l-1)\top} \mathbf{W}^{(l)} \right) \forall l \in 2, \dots, m,$$

- Training MLP for classification tasks (**discriminative fine-tuning**)



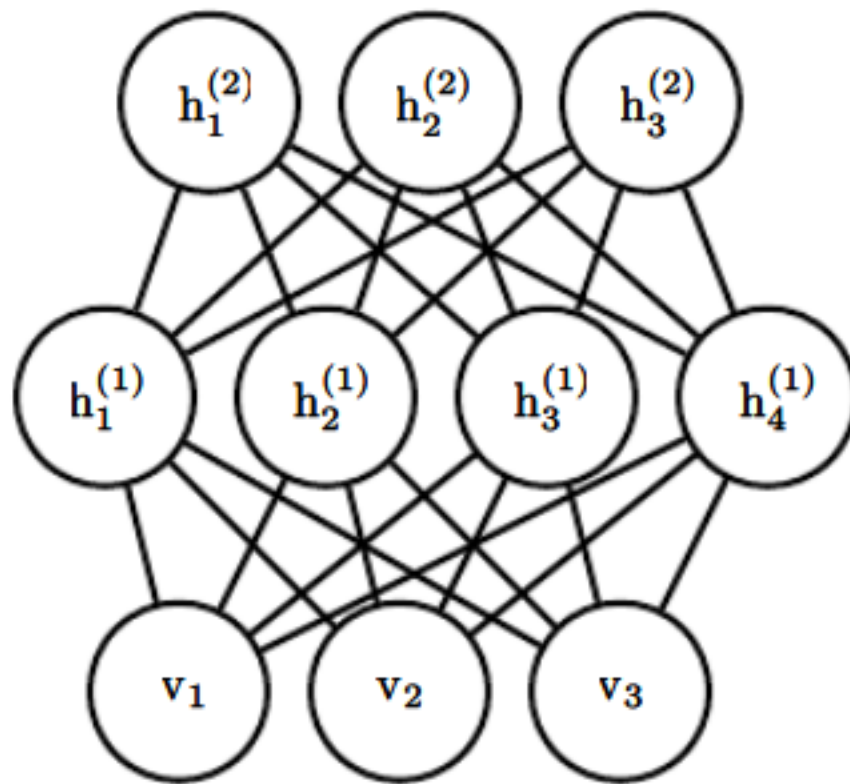
Deep Belief Networks



unsupervised, layer-wise, greedy pre-training



Deep Boltzmann Machines



entirely undirected model



Deep Boltzmann Machines

- joint probability

$$P(v, h^{(1)}, h^{(2)}, h^{(3)}) = \frac{1}{Z(\theta)} \exp(-E(v, h^{(1)}, h^{(2)}, h^{(3)}; \theta))$$

- Energy function

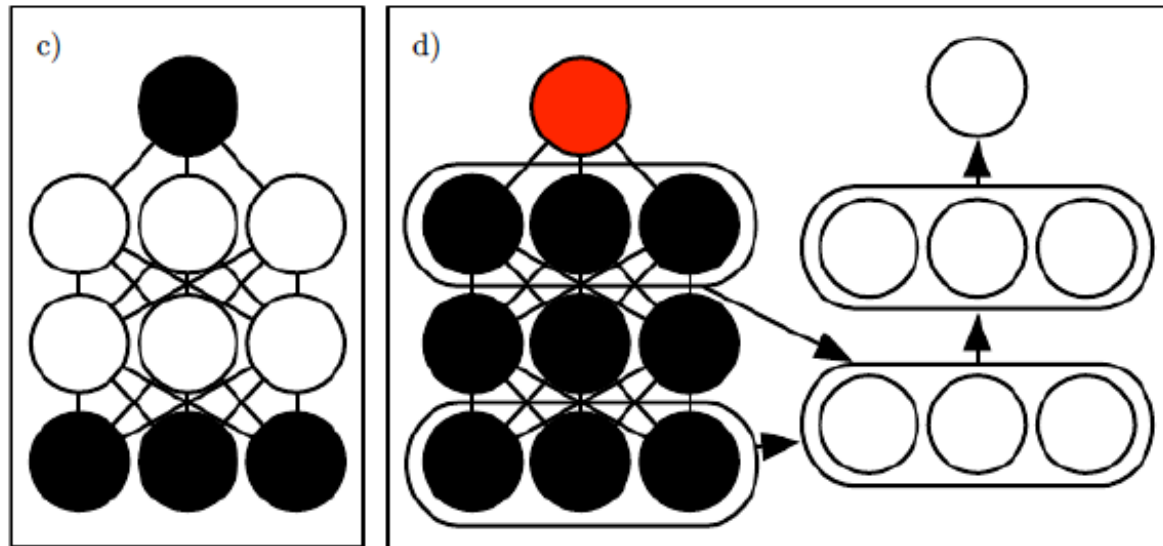
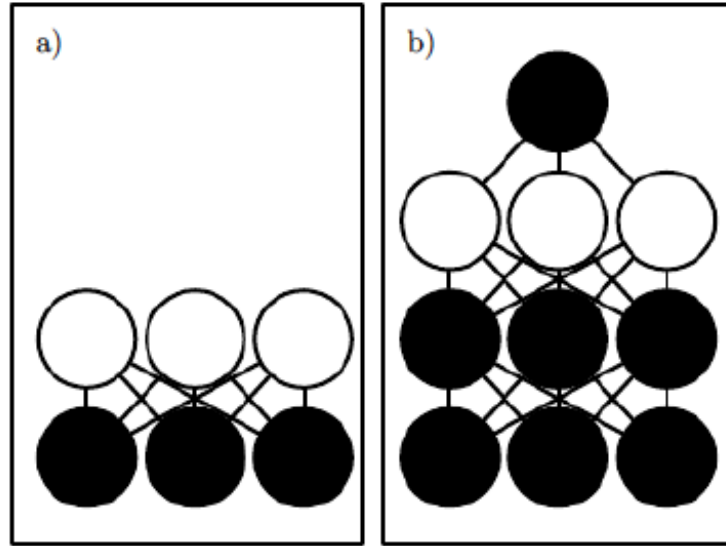
$$E(v, h^{(1)}, h^{(2)}, h^{(3)}; \theta) = -v^\top W^{(1)} h^{(1)} - h^{(1)\top} W^{(2)} h^{(2)} - h^{(2)\top} W^{(3)} h^{(3)}$$

- Learning

- challenge of an intractable partition function
- challenge of an intractable posterior distribution



Jointly Training DBM



Generative Adversarial Networks

■ GANs

■ game theoretic scenario

- generator network must compete against an adversary

$$x = \hat{g}(z; \theta^{(\tilde{g})})$$

Generator network

$$d(x; \theta^{(d)})$$

Discriminator network

probability that x is a real training example
rather than a fake sample drawn from the
model



Generative Adversarial Networks

- Zero-sum game

$$v(\bar{\theta}^{(g)}, \bar{\theta}^{(d)})$$

payoff of the
discriminator

- to maximize its own payoff

$$g^* = \arg \min_g \max_d v(g, d)$$

$$v(\theta^{(g)}, \theta^{(d)}) = \mathbb{E}_{\mathbf{x} \sim p_{\text{data}}} \log d(\mathbf{x}) + \mathbb{E}_{\mathbf{x} \sim p_{\text{model}}} \log (1 - d(\mathbf{x}))$$

This drives the discriminator to attempt to learn to correctly classify samples as real or fake. Simultaneously, the generator attempts to fool the classifier into believing its samples are real.



GANs



Images generated by GANs trained on the LSUN dataset

