

Generative Models

W

GAN

Generative Adversarial Nets

W

Implicit Generative Model

$$p_{\theta}(x)$$

- **Goal:** a sampler $g(\cdot)$ to generate images
- A simple generator $g(z; \theta)$:
 - $z \sim N(0, I)$
 - $x = g(z; \theta)$ deterministic transformation
- Likelihood-free training:
 - Given a dataset from some distribution p_{data}
 - Goal: $g(z; \theta)$ defines a distribution, we want this distribution $\approx p_{data}$
 - Training: minimize $D(g(z; \theta), p_{data})$
 - D is some distance metric (not likelihood)
 - Key idea: **Learn a differentiable D**

$$\{x_1, \dots, x_N\}$$

metrics:

- ① KL-divergence
- ② Total variation
- ③ Wasserstein dist

$$\left(\frac{1}{\epsilon}\right)^d$$

GAN (Goodfellow et al., '14)

- Parameterize the discriminator $D(\cdot; \phi)$ with parameter ϕ
- **Goal:** learn ϕ such that $D(x; \phi)$ measures how likely x is from p_{data}
 - $D(x, \phi) = 1$ if $x \sim p_{data}$
 - $D(x, \phi) = 0$ if $x \not\sim p_{data}$
 - a.k.a., a binary classifier
- GAN: use a neural network for $D(\cdot; \phi)$

true $\{(x_1, 1), (x_2, 1), \dots, (x_N, 1)\}$
artificial $\{(x'_1, 0), (x'_2, 0), \dots, (x'_N, 0)\}$
- **Training:** need both negative and positive samples
 - Positive samples: just the training data
 - Negative samples: use our sampler $g(\cdot; z)$ (can provide infinite samples).
- **Overall objectives:**
 - Generator: $\theta^* = \max_{\theta} D(g(z; \theta); \phi)$
 - Discriminator uses MLE Training:
$$\phi^* = \max_{\phi} \mathbb{E}_{x \sim p_{data}} [\log D(x; \phi)] + \mathbb{E}_{\hat{x} \sim g(\cdot)} [\log(1 - D(\hat{x}; \phi))]$$

GAN (Goodfellow et al., '14)

- Generator $G(z; \theta)$ where $z \sim N(0, I)$
 - Generate realistic data

- Discriminator $D(x; \phi)$
 - Classify whether the data is real (from p_{data}) or fake (from G)

- Objective function:

$$L(\theta, \phi) = \min_{\theta} \max_{\phi} \mathbb{E}_{x \sim p_{data}} [\log D(x; \phi)] + \mathbb{E}_{\hat{x} \sim G} [\log(1 - D(\hat{x}; \phi))]$$

- Training procedure:

- Collect dataset $\{(x, 1) \mid x \sim p_{data}\} \cup \{(\hat{x}, 0) \sim g(z; \theta)\}$
- Train discriminator

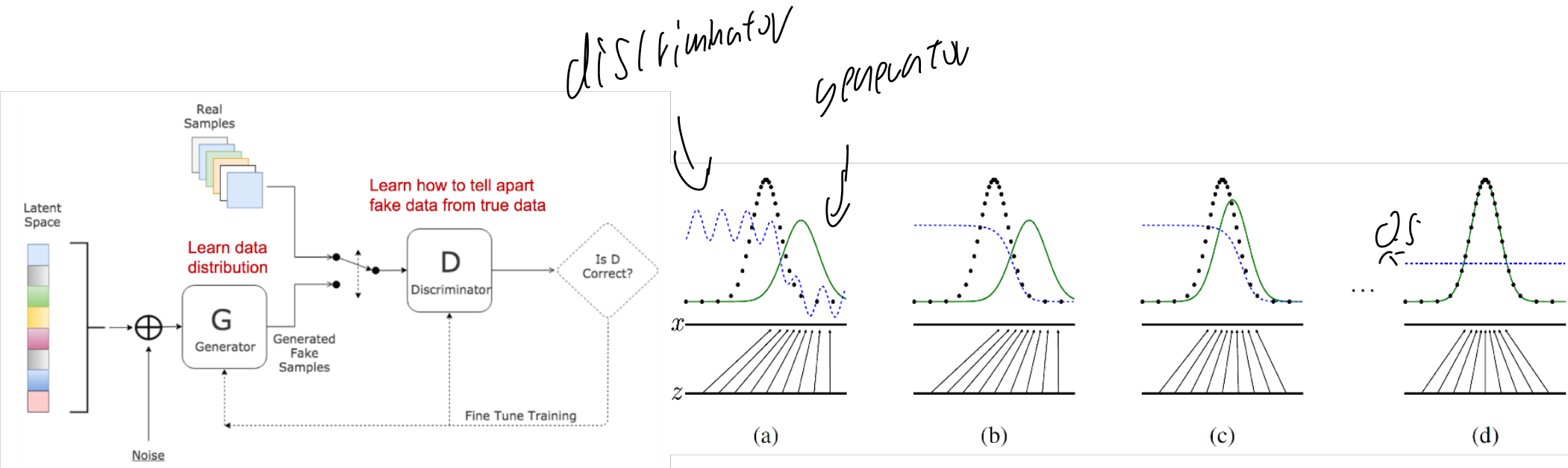
$$D : L(\phi) = \mathbb{E}_{x \sim p_{data}} [\log D(x; \phi)] + \mathbb{E}_{\hat{x} \sim G} [\log(1 - D(\hat{x}; \phi))]$$

- Train generator $G : L(\theta) = \mathbb{E}_{z \sim N(0, I)} [\log D(G(z; \theta), \phi)]$
- Repeat

GAN (Goodfellow et al., '14)

- Objective function:

$$L(\theta, \phi) = \min_{\theta} \max_{\phi} \mathbb{E}_{x \sim p_{data}} [\log D(x; \phi)] + \mathbb{E}_{\hat{x} \sim G} [\log(1 - D(\hat{x}; \phi))]$$



Math Behind GAN

$$\mathcal{L}(\theta, \phi) = \min_{\theta} \max_{\phi} \mathbb{E}_{x \sim p_{\text{data}}} [\log D(x; \phi)] + \mathbb{E}_{x \sim g(\cdot)} [\log (1 - D(x; \phi))]$$

Let D^* , g^* be the solution to

Optimal D^* , for a given X

$$\mathcal{L}_X(D) = p_{\text{data}}(X) \cdot \log(D(X)) + p_g(X) \cdot \log(1 - D(X))$$

$$\Rightarrow \frac{\partial \mathcal{L}}{\partial D} = 0, \quad \text{first-order condition}$$

$$\Rightarrow \frac{p_{\text{data}}(X)}{D^*(X)} - \frac{p_g(X)}{1 - D^*(X)} = 0$$

$$\Rightarrow D^*(X) = \frac{p_{\text{data}}(X)}{p_{\text{data}}(X) + p_g(X)}$$

$$\star \text{ if } (p_{\text{data}} = p_g) \Rightarrow D^*(X) = 0.5$$

Math Behind GAN

consider optimal g^* , given optimal D^*

$$\mathcal{L}(\theta, \phi) = \mathbb{E}_{x \sim p_{\text{data}}} \left[\log \frac{p_{\text{data}}(x)}{p_{\text{data}}(x) + p_g(x)} \right] + \mathbb{E}_{x \sim g} \left[\log \frac{p_g(x)}{p_{\text{data}}(x) + p_g(x)} \right]$$

$$= \mathbb{E}_{x \sim p_{\text{data}}} \left[\log \frac{p_{\text{data}}(x)}{\frac{p_{\text{data}}(x) + p_g(x)}{2}} \right] - \log 2$$

$$+ \mathbb{E}_{x \sim g} \left[\log \frac{p_g(x)}{\frac{p_{\text{data}}(x) + p_g(x)}{2}} \right] - \log 2$$

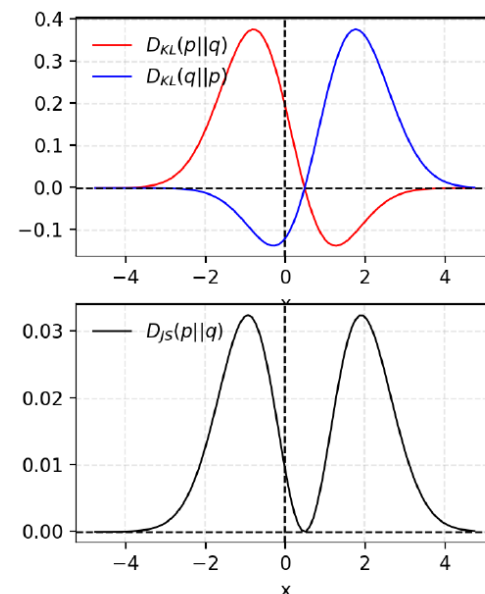
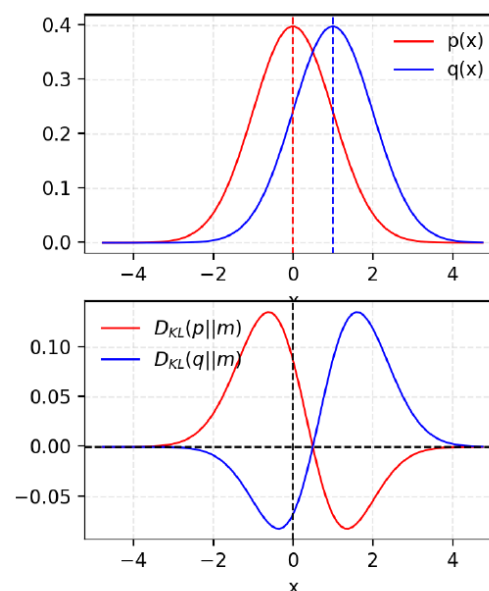
$$= \underbrace{K\left(p_{\text{data}} \middle| \frac{1}{2} (p_{\text{data}} + p_g)\right) + K\left(p_g \middle| \frac{1}{2} (p_{\text{data}} + p_g)\right)}_{\text{JSD}(p_{\text{data}}, p_g)} - \log 4$$

2. Jensen Shannon Divergence
JSD(p_{data}, p_g)

$K(p|q)$
 $= \mathbb{E}_{x \sim p} \left[\log \frac{p(x)}{q(x)} \right]$

KL-Divergence and JS-Divergence

$JSD \geq 0$
if $JSD(p||q) = 0$
 $\Leftrightarrow p = q$



Math Behind GAN

=> Given D^*

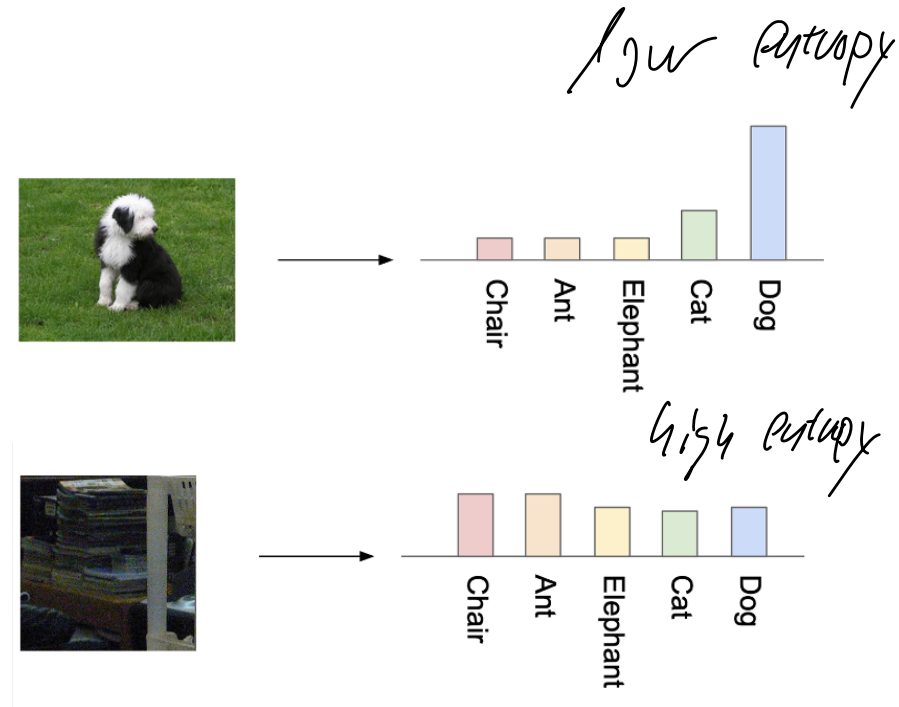
$$\min_g \mathcal{L}(g) = 2 \cdot \text{JS}(P_g \| P_{data}) - \log 4$$

if g^* is optimal

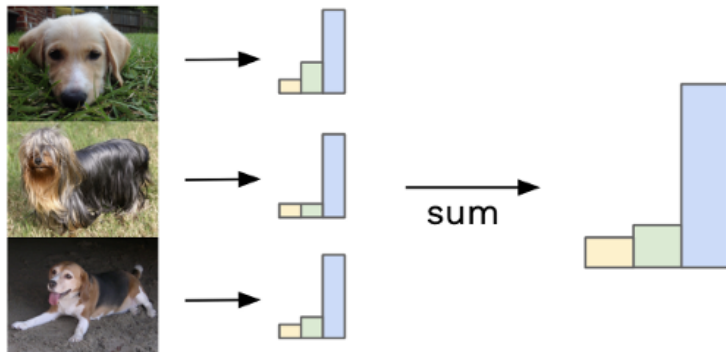
$$\Rightarrow P_{g^*} = P_{data}$$

Evaluation of GAN

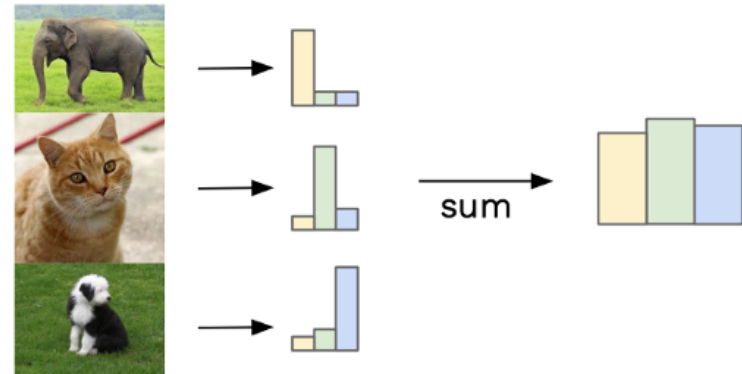
- No $p(x)$ in GAN.
- Idea: use a trained classifier $f(y | x)$:
- If $x \sim p_{data}$, $f(y | x)$ should have low entropy
 - Otherwise, $f(y | x)$ close to uniform.
- Samples from G should be diverse:
 - $p_f(y) = \mathbb{E}_{x \sim G}[f(y | x)]$ close to uniform.



Similar labels sum to give focussed distribution

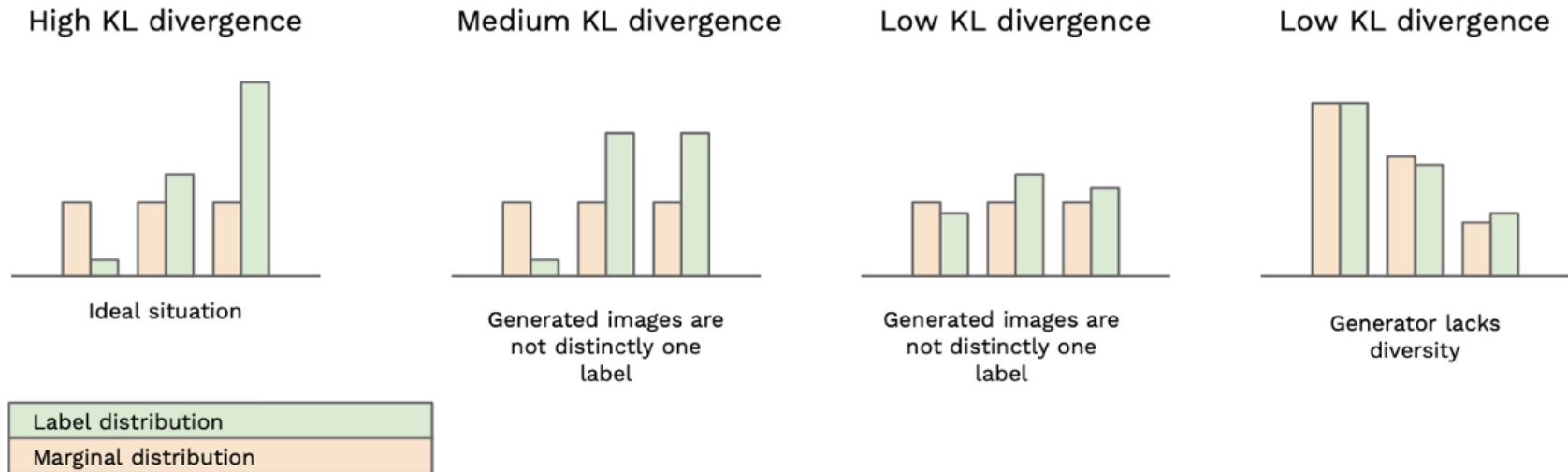


Different labels sum to give uniform distribution



Evaluation of GAN

- **Inception Score (IS, Salimans et al. '16)**
 - Use Inception V3 trained on ImageNet as $f(y|x)$
 - $IS = \exp \left(\mathbb{E}_{x \sim G} \left[\underbrace{KL(f(y|x) || p_f(y))}_{\text{individual}} \right] \right)$
 - Higher the better



Comments on GAN

- Other evaluation metrics:
 - Fréchet Inception Distance (FID): Wasserstein distance between Gaussians
- Mode collapse:
 - The generator only generate a few type of samples.
 - Or keep oscillating over a few modes.
- Training instability:
 - Discriminator and generator may keep oscillating
 - Example: $-xy$, generator x , discriminatory. NE: $x = y = 0$ but GD oscillates.
 - No stopping criteria.
 - Use Wasserstein GAN (Arjovsky et al. '17):
$$\min_G \max_{f: \text{Lip}(f) \leq 1} \mathbb{E}_{x \sim p_{data}} [f(x)] - \mathbb{E}_{\hat{x} \sim p_G} [f(\hat{x})]$$
 - And need many other tricks...

Energy-Based Models



Energy-based Models

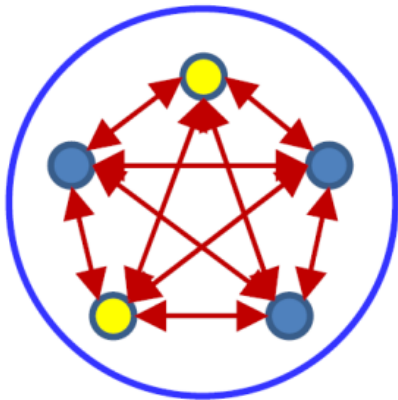
- Goal of generative models:
 - a probability distribution of data: $P(x)$
- Requirements
 - $P(x) \geq 0$ (non-negative)
 - $\int_x P(x)dx = 1$
- Energy-based model:
 - Energy function: $E(x; \theta)$, parameterized by θ
 - $P(x) = \frac{1}{z} \exp(-E(x; \theta))$ (why exp?)
 - $z = \int_x \exp(-E(x; \theta))dx$

Boltzmann Machine

$$y \in \mathbb{R}^d, \quad y_i \in \{-1, 1\}$$

- Generative model
 - $E(y) = -\frac{1}{2}y^\top W y$
 - $P(y) = \frac{1}{Z} \exp(-\frac{E(y)}{T})$, T : temperature hyper-parameter
 - W : parameter to learn
- When y_i is binary, patterns are affecting each other through W

$$y_1, \dots, y_i, \dots, y_d$$



$$z_i = \frac{1}{T} \sum_j w_{ji} s_j$$

$$P(s_i = 1 | s_{j \neq i}) = \frac{1}{1 + e^{-z_i}}$$

Boltzmann Machine: Training

$$y \in \mathbb{R}^d$$

$$y_i \in \{-1, 1\}$$

- Objective: maximum likelihood learning (assume $T=1$):
 - Probability of one sample:

$$P(y) = \frac{\exp(\frac{1}{2} y^T W y)}{\sum_{y'} \exp(y'^T W y')}$$

2^d values of y

- Maximum log-likelihood:

$$\max_W L(W) = \frac{1}{N} \sum_{y \in D} \frac{1}{2} y^T W y - \log \sum_{y'} \exp(\frac{1}{2} y'^T W y')$$

$$D = \{y^1, \dots, y^N\}$$

$$W \in d \times d$$

Boltzmann Machine: Training

$$L(W) = \frac{1}{N} \sum_{y \in D} \frac{1}{2} y^T W y - \log \sum_{y'} \exp \left(\frac{1}{2} y'^T W y' \right)$$

$$\nabla_{W_{ij}} L = \frac{1}{N} \sum_{y \in D} y_i y_j - \sum_{y'} \frac{\exp \left(\frac{1}{2} y'^T W y' \right) y_i' y_j'}{\sum_{y''} \exp \left(\frac{1}{2} y''^T W y'' \right)}$$

$P(y')$

$$\Rightarrow \approx \frac{1}{N} \sum_{y \in D} y_i y_j - \frac{1}{N} \sum_{y' \in S} y_i' y_j'$$

$S = \text{Sample } \{y^1, y^2, \dots, y^N\}$
 $\mathbb{E}_{y'} [y_i' \cdot y_j']$

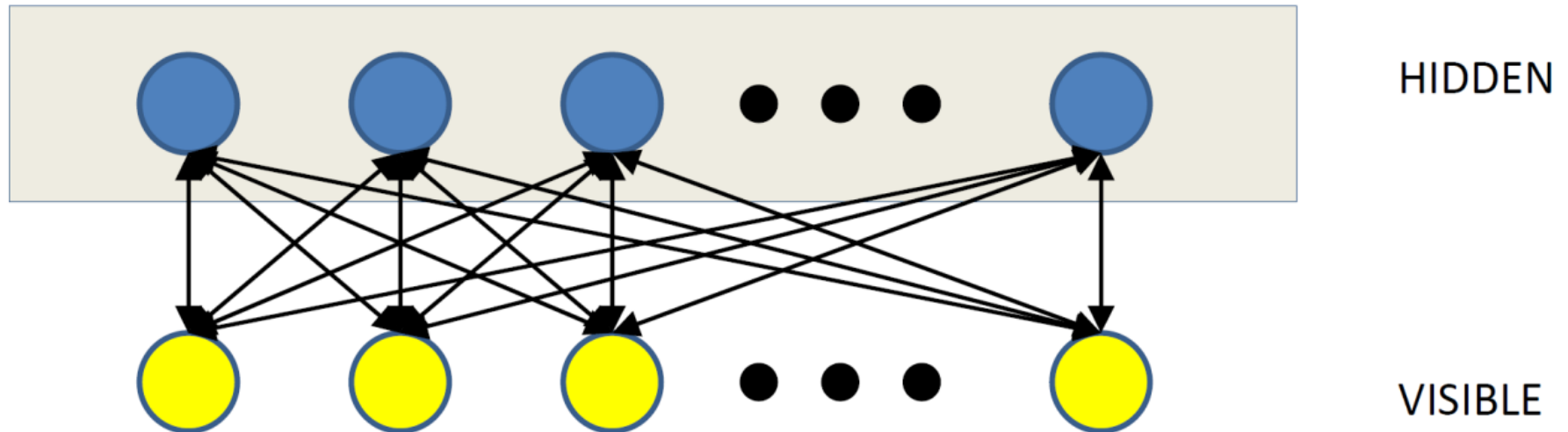
Sampling: $M(MC$
initial $y(u) \in \mathbb{R}^d$

for $t = 1, \dots, T$
 repeat $\{1, \dots, d\}$ conditional sampling
 $y_i(t) \sim p(\cdot | y_{i \neq j}(t-1))$

Theorem $T \rightarrow \infty$ $\{y(1), \dots, y(T)\}$
 $\rightarrow$ generator

Restricted Boltzmann Machine

- A structured Boltzmann Machine
 - Hidden neurons are only connected to visible neurons
 - No intra-layer connections
 - Invented by Paul Smolensky in '89
 - Became more practical after Hinton invested fast learning algorithms in mid 2000



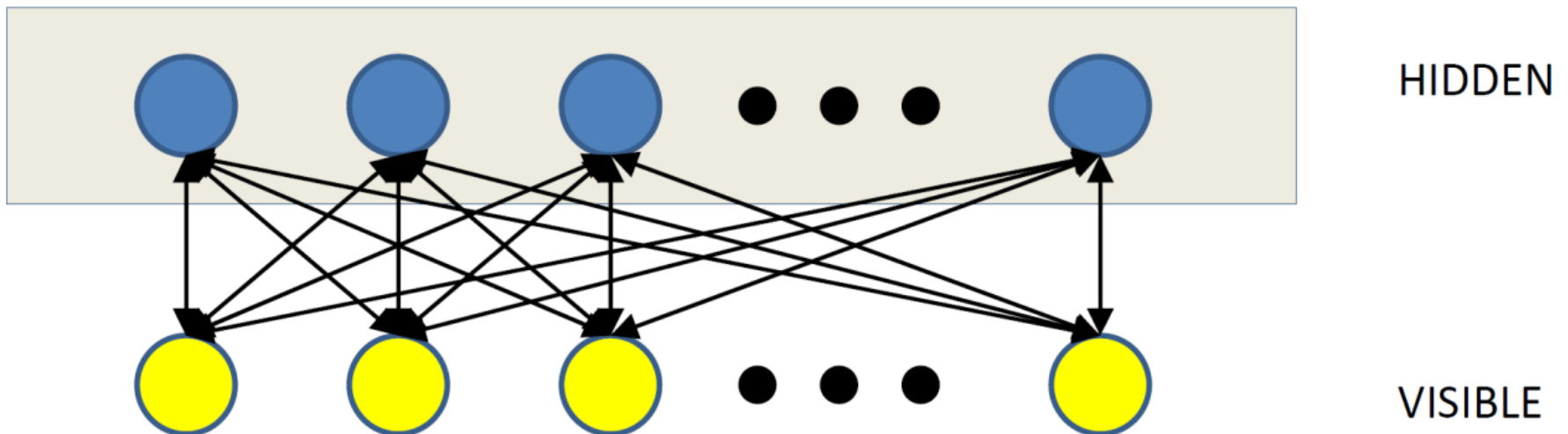
Restricted Boltzmann Machine

- Computation Rules

- Iterative sampling

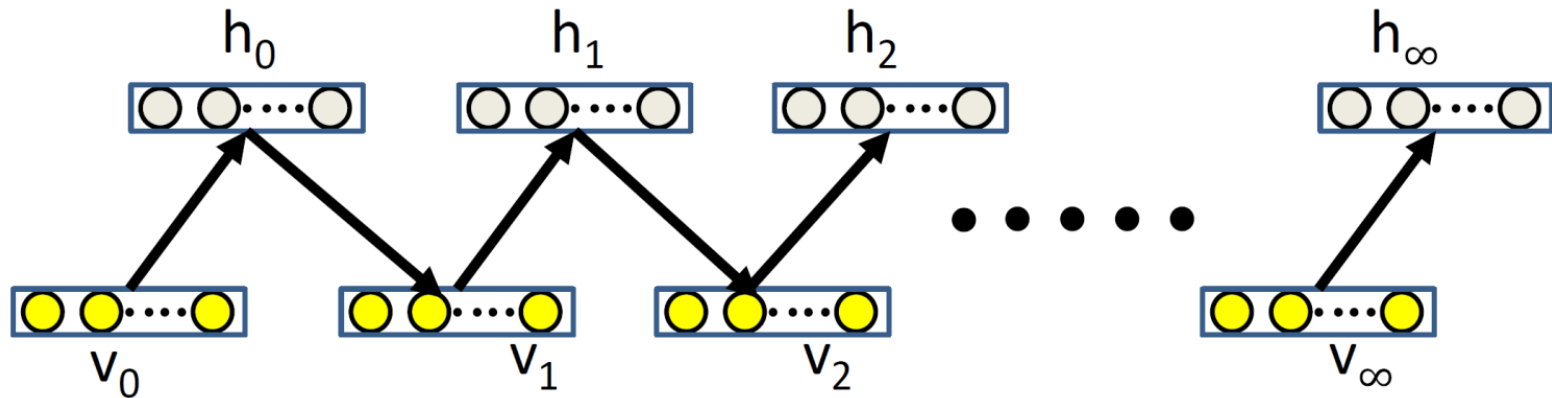
- Hidden neurons h_i : $z_i = \sum_j w_{ij} v_j$, $P(h_i | v) = \frac{1}{1 + \exp(-z_i)}$

- Visible neurons v_j : $z_j = \sum_i w_{ij} h_i$, $P(v_j | h) = \frac{1}{1 + \exp(-z_j)}$



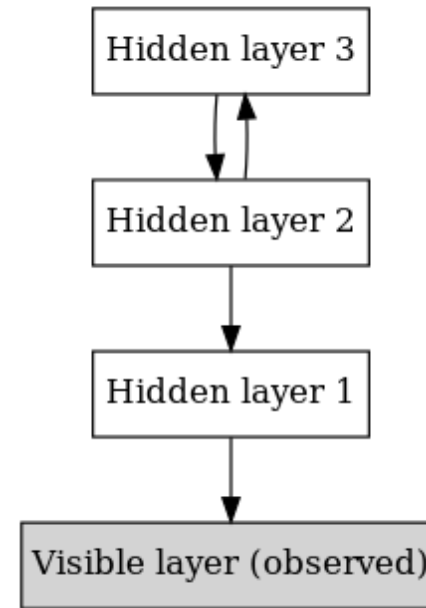
Restricted Boltzmann Machine

- Sampling:
 - Randomly initialize visible neurons v_0
 - Iterative sampling between hidden neurons and visible neurons
 - Get final sample (v_∞, h_∞)

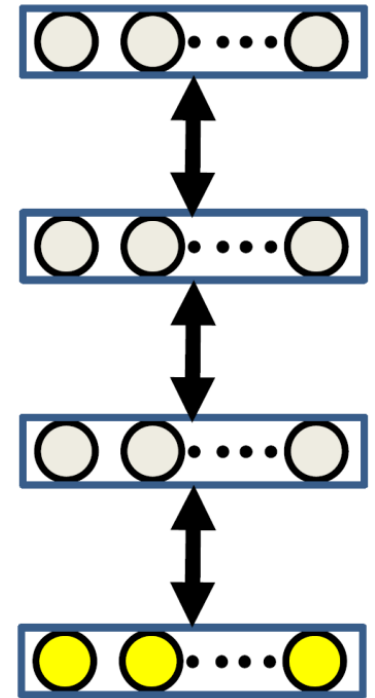


Deep Boltzmann Machine

- Can we have a **deep** version of RBM?
 - Deep Belief Net ('06)
 - Deep Boltzmann Machine ('09)
- Sampling?
 - Forward pass: bottom-up
 - Backward pass: top-down
- Deep Boltzmann Machine
 - The very first deep generative model
 - Salakhudinov & Hinton



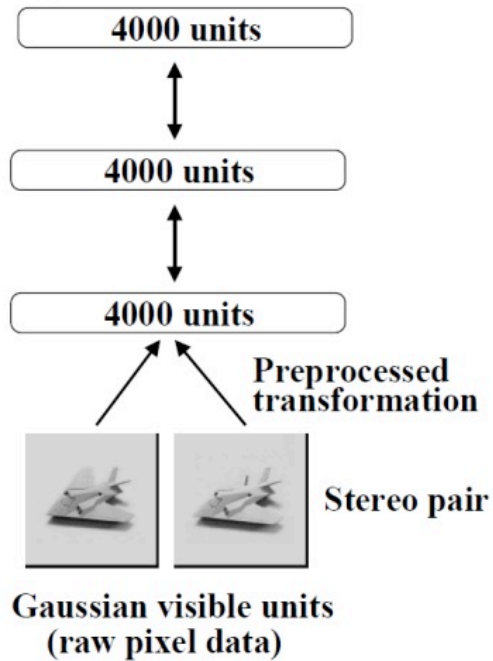
deep belief net



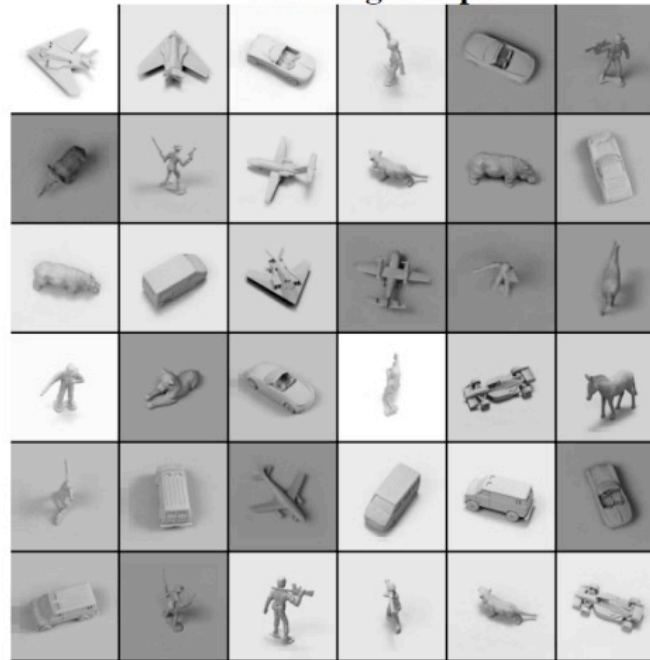
Deep Boltzmann Machine

Deep Boltzmann Machine

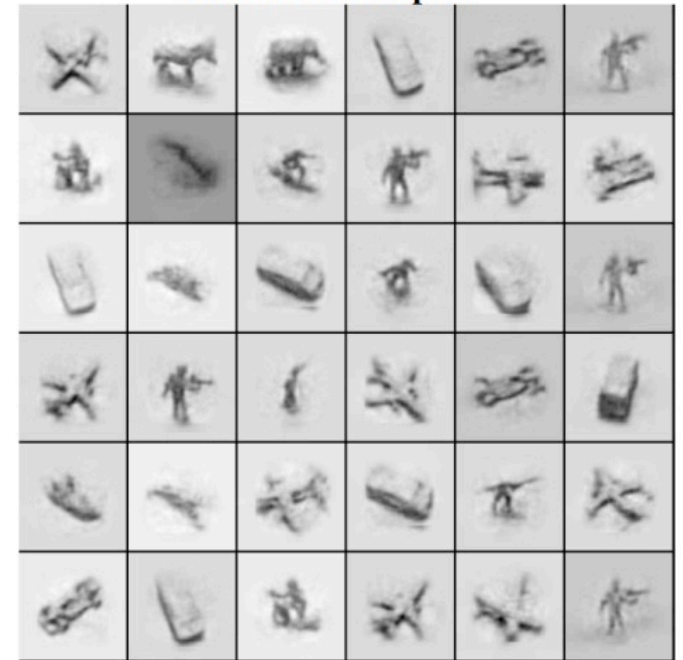
Deep Boltzmann Machine



Training Samples



Generated Samples



Summary

- Pros: powerful and flexible
 - An arbitrarily complex density function $p(x) = \frac{1}{z} \exp(-E(x))$
- Cons: hard to sample / train
 - Hard to sample:
 - MCMC sampling
 - Partition function
 - No closed-form calculation for likelihood
 - Cannot optimize MLE loss exactly
 - MCMC sampling