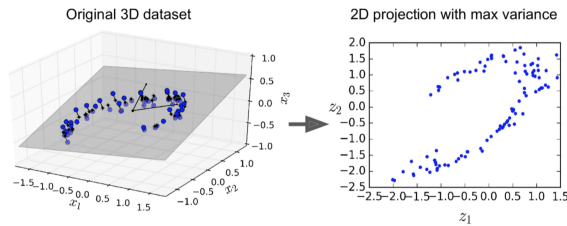


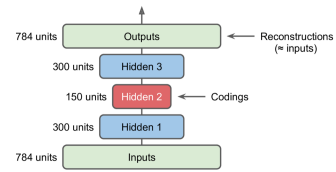
Basic Idea

Example: Principal Component Analysis



- A 3-2-3 autoencoder with linear units and square loss performs **principal component analysis**: Find linear transformation of data to maximize variance

Stacked Autoencoders



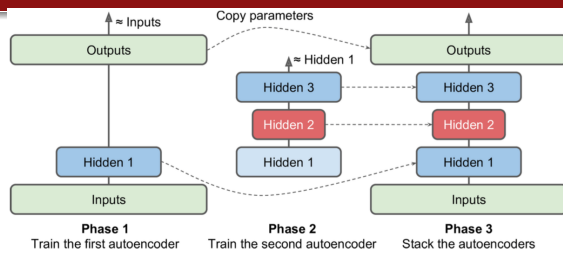
- A **stacked autoencoder** has multiple hidden layers

- Can share parameters to reduce their number by exploiting symmetry: $W_4 = W_1^T$ and $W_3 = W_2^T$

```
weights1 = tf.Variable(weights1_init, dtype=tf.float32, name="weights1")
weights2 = tf.Variable(weights2_init, dtype=tf.float32, name="weights2")
weights3 = tf.transpose(weights2, name="weights3") # shared weights
weights4 = tf.transpose(weights1, name="weights4") # shared weights
```

Stacked Autoencoders

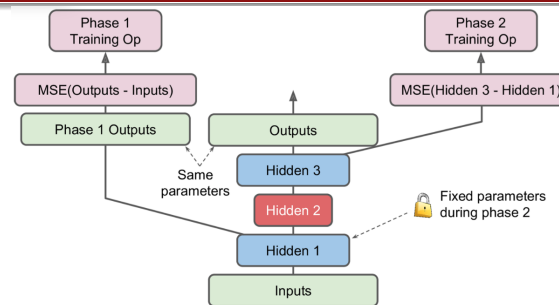
Incremental Training



- Can simplify training by starting with single hidden layer H_1
- Then, train a second AE to mimic the output of H_1
- Insert this into first network
- Can build by using H_1 's output as training set for Phase 2

Stacked Autoencoders

Incremental Training (Single TF Graph)



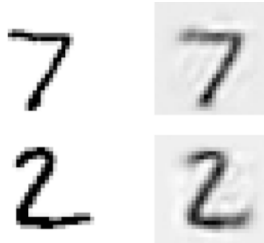
- Previous approach requires multiple TensorFlow graphs
- Can instead train both phases in a single graph: First left side, then right

Stacked Autoencoders

Visualization

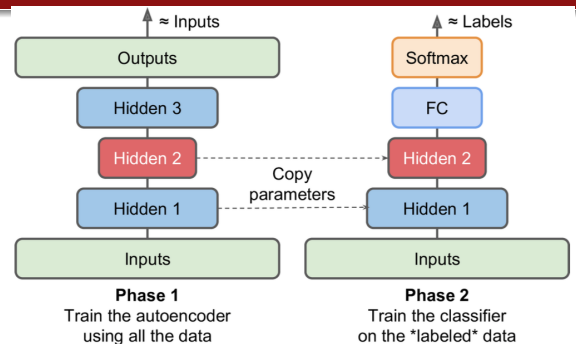
Input MNIST Digit

Network Output

Weights (features selected) for five nodes from H_1 :

Stacked Autoencoders

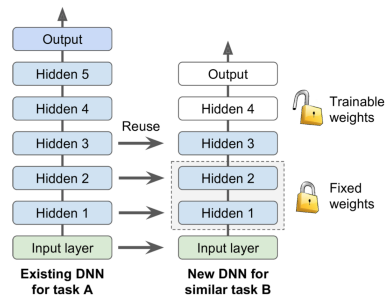
Semi-Supervised Learning



- Can **pre-train** network with unlabeled data
⇒ learn useful features and then train "logic" of dense layer with labeled data

Transfer Learning from Trained Classifier

- Can also transfer from a classifier trained on different task, e.g., transfer a GoogleNet architecture to ultrasound classification
- Often choose existing one from a **model zoo**

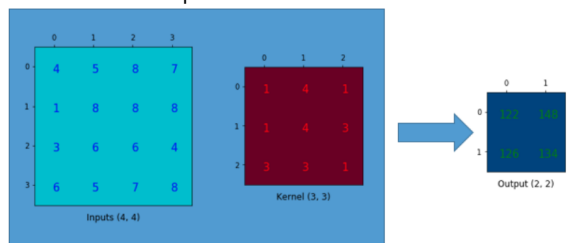


Transposed Convolutions

- What if some encoder layers are convolutional? How to upsample to original resolution?
- Can use, e.g., **linear interpolation**, **bilinear interpolation**, etc.
- Or, **transposed convolution**, e.g., `tf.layers.conv2d_transpose`

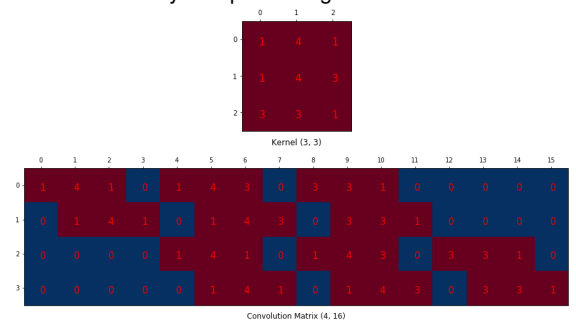
Transposed Convolutions (2)

Consider this example convolution



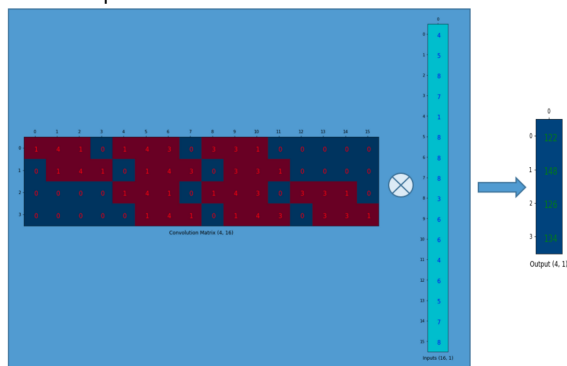
Transposed Convolutions (3)

An alternative way of representing the kernel



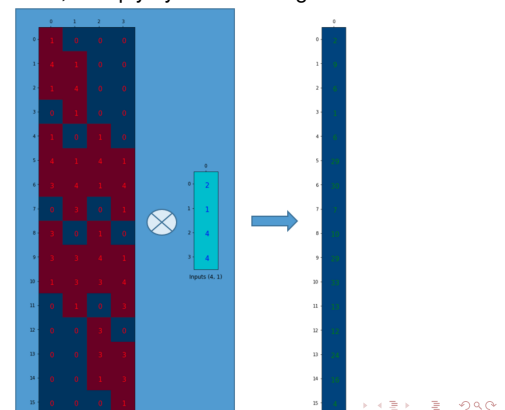
Transposed Convolutions (4)

This representation works with matrix multiplication on flattened input:



Transposed Convolutions (5)

Transpose kernel, multiply by flat 2×2 to get flat 4×4

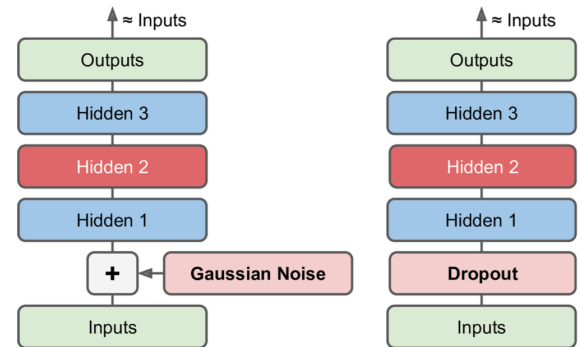


Denoising Autoencoders

Vincent et al. (2010)

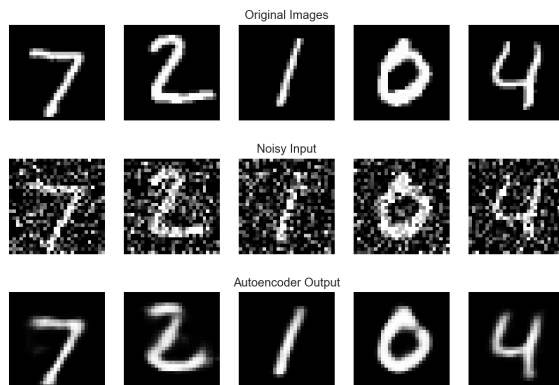
- Can train an autoencoder to learn to **denoise** input by giving input **corrupted** instance $\tilde{x}$ and targeting **uncorrupted** instance x
- Example noise models:
 - Gaussian noise:** $\tilde{x} = x + z$, where $z \sim \mathcal{N}(0, \sigma^2 I)$
 - Masking noise:** zero out some fraction ν of components of x
 - Salt-and-pepper noise:** choose some fraction ν of components of x and set each to its min or max value (equally likely)

Denoising Autoencoders



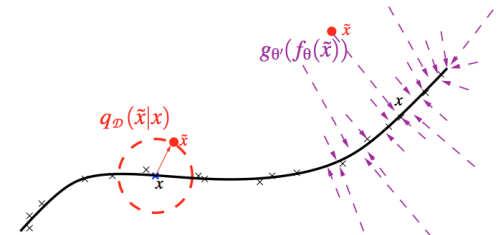
Denoising Autoencoders

Example



Denoising Autoencoders

- How does it work?
- Even though, e.g., MNIST data are in a 784-dimensional space, they lie on a low-dimensional **manifold** that captures their most important features
- Corruption process** moves instance x off of manifold
- Encoder f_θ and decoder $g_{\theta'}$ are trained to project $\tilde{x}$ back onto manifold



Sparse Autoencoders

- An overcomplete architecture
- Regularize outputs of hidden layer to enforce **sparsity**:

$$\tilde{\mathcal{J}}(x) = \mathcal{J}(x, g(f(x))) + \alpha \Omega(h),$$

where $\mathcal{J}$ is loss function, f is encoder, g is decoder, $h = f(x)$, and Ω penalizes non-sparsity of h

- E.g., can use $\Omega(h) = \sum_i |h_i|$ and ReLU activation to force many zero outputs in hidden layer
- Can also measure average activation of h_i across mini-batch and compare it to user-specified **target sparsity** value p (e.g., 0.1) via square error or **Kullback-Leibler divergence**:

$$p \log \frac{p}{q} + (1-p) \log \frac{1-p}{1-q},$$

where q is average activation of h_i over mini-batch

Contractive Autoencoders

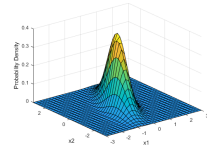
- Similar to sparse autoencoder, but use

$$\Omega(h) = \sum_{j=1}^m \sum_{i=1}^n \left(\frac{\partial h_i}{\partial x_j} \right)^2$$

- I.e., penalize large partial derivatives of encoder outputs wrt input values
- This **contracts** the output space by mapping input points in a neighborhood near x to a smaller output neighborhood near $f(x)$
 - $\Rightarrow$ Resists perturbations of input x
- If h has sigmoid activation, encoding near binary and a CE pushes embeddings to corners of a hypercube

Variational Autoencoders

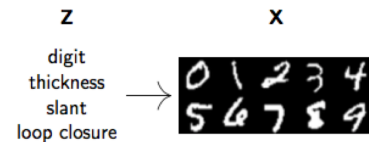
- VAE is an autoencoder that is also **generative model**
 - ⇒ Can generate new instances according to a probability distribution
 - E.g., hidden Markov models, Bayesian networks
 - Contrast with **discriminative models**, which predict classifications
- Encoder f outputs $[\mu, \sigma]^T$
 - Pair (μ_i, σ_i) parameterizes Gaussian distribution for dimension $i = 1, \dots, n$
 - Draw $z_i \sim \mathcal{N}(\mu_i, \sigma_i)$
 - Decode this **latent variable** z to get $g(z)$



Variational Autoencoders

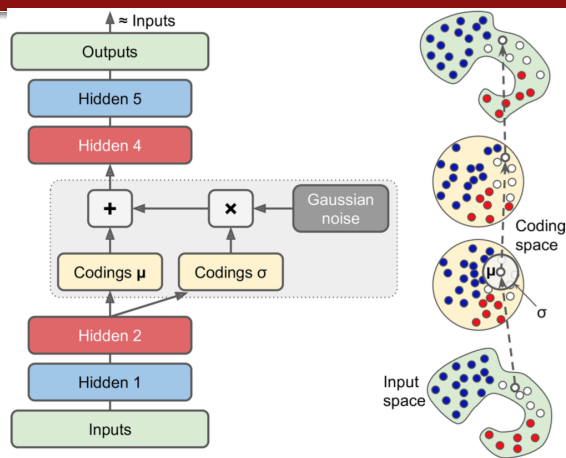
Latent Variables

- Independence of z dimensions makes it easy to generate instances wrt complex distributions via decoder g
- Latent variables can be thought of as values of attributes describing inputs
 - E.g., for MNIST, latent variables might represent "thickness", "slant", "loop closure"



Variational Autoencoders

Architecture



Variational Autoencoders

Optimization

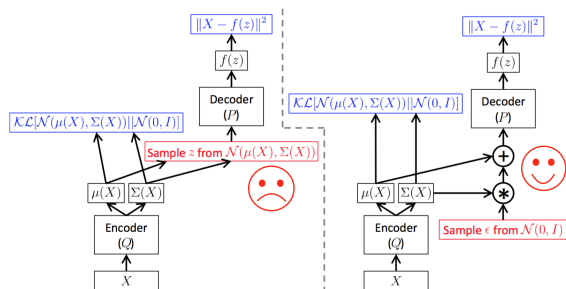
- **Maximum likelihood (ML)** approach for training generative models: find a model (θ) with maximum probability of generating the training set $\mathcal{X}$
- Achieve this by minimizing the sum of:
 - End-to-end AE loss (e.g., square, cross-entropy)
 - **Regularizer** measuring distance (K-L divergence) from latent distribution $q(z|x)$ and $\mathcal{N}(\mathbf{0}, I)$ (= standard multivariate Gaussian)
- $\mathcal{N}(\mathbf{0}, I)$ also considered the **prior distribution** over z (= distribution when no x is known)

```
eps = 1e-10
latent_loss = 0.5 * tf.reduce_sum(
    tf.square(hidden3_sigma) + tf.square(hidden3_mean)
    - 1 - tf.log(eps + tf.square(hidden3_sigma)))
```

Variational Autoencoders

Reparameterization Trick

- Cannot backprop error signal through random samples
- **Reparameterization trick** emulates $z \sim \mathcal{N}(\mu, \sigma)$ with $\epsilon \sim \mathcal{N}(0, 1)$, $z = \epsilon\sigma + \mu$



Variational Autoencoders

Example Generated Images: Random

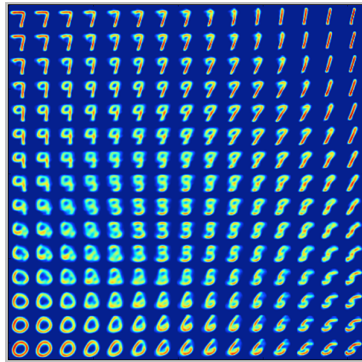
- Draw $z \sim \mathcal{N}(\mathbf{0}, I)$ and display $g(z)$



Variational Autoencoders

Example Generated Images: Manifold

- Uniformly sample points in (2-dimensional) z space and decode

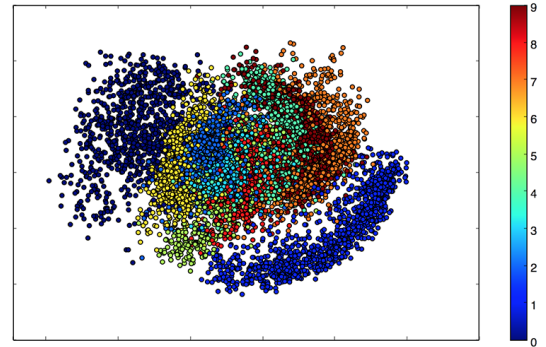


Navigation icons: back, forward, search, etc.

Variational Autoencoders

2D Cluster Analysis

- Cluster analysis by digit (2D latent space)



Navigation icons: back, forward, search, etc.

Aside: Visualizing with t-SNE

van der Maaten and Hinton (2008)

- Visualize high-dimensional data, e.g., embedded representations
- Want low-dimensional representation to have similar neighborhoods as high-dimensional one
- Map each high-dimensional $x_1, \dots, x_N$ to low-dimensional $y_1, \dots, y_N$ via matching **pairwise distributions** based on distance
 - $\Rightarrow$ Probability p_{ij} pair (x_i, x_j) chosen similar to probability q_{ij} pair (y_i, y_j) chosen
- Set $p_{ij} = (p_{ji} + p_{ij}) / (2N)$ where

$$p_{ji} = \frac{\exp(-\|x_i - x_j\|^2 / (2\sigma_i^2))}{\sum_{k \neq i} \exp(-\|x_i - x_k\|^2 / (2\sigma_i^2))}$$

and σ_i chosen to control density of the distribution

- I.e., p_{ji} is probability of x_i choosing x_j as its neighbor if chosen in proportion of Gaussian density centered at x_i

Navigation icons: back, forward, search, etc.

Aside: Visualizing with t-SNE (2)

van der Maaten and Hinton (2008)

- Also, define q via student's t distribution:

$$q_{ij} = \frac{(1 + \|y_i - y_j\|^2)^{-1}}{\sum_{k \neq \ell} (1 + \|y_k - y_\ell\|^2)^{-1}}$$

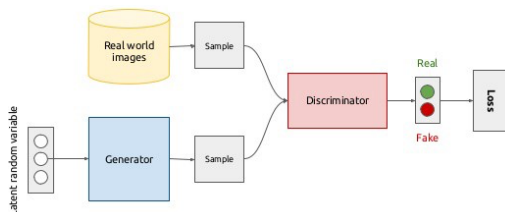
- Using student's t instead of Gaussian helps address **crowding problem** where distant clusters in x space squeeze together in y space
- Now choose y values to match distributions p and q via **Kullback-Leibler divergence**:

$$\sum_{i \neq j} p_{ij} \log \frac{p_{ij}}{q_{ij}}$$

Navigation icons: back, forward, search, etc.

Generative Adversarial Network

- GANs are also generative models, like VAEs
- Models a game between two players
 - Generator** creates samples intended to come from training distribution
 - Discriminator** attempts to discern the "real" (original training) samples from the "fake" (generated) ones
- Discriminator trains as a binary classifier, generator trains to fool the discriminator



Navigation icons: back, forward, search, etc.

Generative Adversarial Network

How the Game Works

- Let $D(x)$ be discriminator parameterized by $\theta^{(D)}$
 - Goal: Find $\theta^{(D)}$ minimizing $J^{(D)}(\theta^{(D)}, \theta^{(G)})$
- Let $G(z)$ be generator parameterized by $\theta^{(G)}$
 - Goal: Find $\theta^{(G)}$ minimizing $J^{(G)}(\theta^{(D)}, \theta^{(G)})$
- A **Nash equilibrium** of this game is $(\theta^{(D)}, \theta^{(G)})$ such that each $\theta^{(i)}$, $i \in \{D, G\}$ yields a local minimum of its corresponding J

Navigation icons: back, forward, search, etc.

Generative Adversarial Network

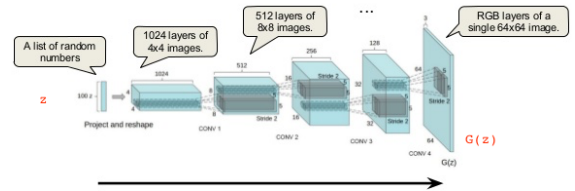
Training

- Each training step:
 - Draw a minibatch of x values from dataset
 - Draw a minibatch of z values from prior (e.g., $\mathcal{N}(\mathbf{0}, I)$)
 - Simultaneously update $\theta^{(G)}$ to reduce $J^{(G)}$ and $\theta^{(D)}$ to reduce $J^{(D)}$, via, e.g., Adam
- For $J^{(D)}$, common to use cross-entropy where label is 1 for real and 0 for fake
- Since generator wants to trick discriminator, can use $J^{(G)} = -J^{(D)}$
 - Others exist that are generally better in practice, e.g., based on ML

Generative Adversarial Network

DCGAN: Radford et al. (2015)

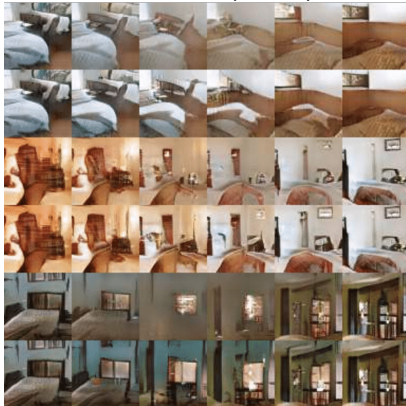
- “Deep, convolution GAN”
- Generator uses **transposed convolutions** (e.g., `tf.layers.conv2d_transpose`) without pooling to upsample images for input to discriminator



Generative Adversarial Network

DCGAN Generated Images: Bedrooms

Trained from LSUN dataset, sampled z space



Generative Adversarial Network

DCGAN Generated Images: Adele Facial Expressions

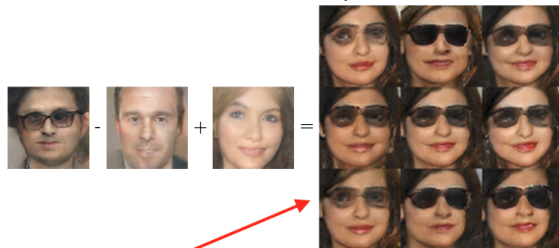
Trained from frame grabs of interview, sampled z space



Generative Adversarial Network

DCGAN Generated Images: Latent Space Arithmetic

Performed semantic arithmetic in z space!



(Non-center images have noise added in z space; center is noise-free)