

CS-E4890: Deep Learning

Deep autoencoders

Alexander Ilin

- Supervised learning problems: datasets consist of input-output pairs

$$(\mathbf{x}^{(1)}, \mathbf{y}^{(1)}), \dots, (\mathbf{x}^{(n)}, \mathbf{y}^{(n)})$$

- Deep learning: supervised learning solved.
- Unsupervised learning: Make computers learn from unlabeled data

$$\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(n)}$$

- Unsupervised learning seems important for building intelligent systems that can learn quickly. We humans learn a lot from unlabeled data.
- Unsupervised learning can be useful in practical applications:
 - detect samples that look different from training population (novelty/anomaly detection)
 - visualize data, discover patterns (information visualization)
 - generate new samples which look similar to the training data (generative models)

- We can use unlabeled data to do representation learning.
- Representation learning: extract features that may be useful for future (downstream) tasks

$$\mathbf{x} \xrightarrow{f} \mathbf{z}$$

- Extracted features might work better than raw data in supervised learning tasks (especially with little labeled data):

$$\mathbf{x} \xrightarrow{f} \mathbf{z} \rightarrow \mathbf{y}$$

- Problem: we do not know for which downstream tasks we need to prepare.
- We can extract patterns (features) that appear frequently in the data and hope that those features will be useful later.

Bottleneck autoencoders

Principal component analysis (PCA)

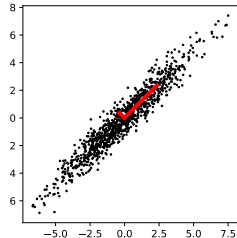
- One of the simplest methods of unsupervised learning is PCA.
- Assume that the data have been centered by subtracting its mean: $\mathbf{x} \leftarrow \mathbf{x} - \mathbb{E}\{\mathbf{x}\}$.
- The first principal component is found by maximizing the variance of the data multiplied by a unit-length vector $\mathbf{w}_1$:

$$y_1 = \mathbf{w}_1^\top \mathbf{x}, \quad \|\mathbf{w}_1\| = 1$$
$$\mathbb{E}\{y_1^2\} = \mathbb{E}\{\mathbf{x}\mathbf{x}^\top\} = \mathbf{w}_1^\top \mathbb{E}\{\mathbf{x}\mathbf{x}^\top\} \mathbf{w}_1 = \mathbf{w}_1^\top \mathbf{C}_x \mathbf{w}_1$$

$$\mathbf{w}_1^* = \arg \max_{\mathbf{w}_1} \mathbf{w}_1^\top \mathbf{C}_x \mathbf{w}_1, \quad \text{s.t. } \|\mathbf{w}_1\| = 1$$

The solution is given by the first dominant eigenvector of the covariance matrix $\mathbf{C}_x$.

- The second principal component is found by maximizing the variance in the subspace orthogonal to the first eigenvector of $\mathbf{C}_x$ (and so on).



PCA as minimum mean-square error compression

- We find an m -dimensional subspace spanned by orthonormal basis

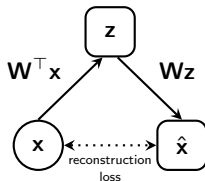
$$\mathbf{W}_{n \times m} = [\mathbf{w}_1 \quad \dots \quad \mathbf{w}_m] \quad \mathbf{W}^\top \mathbf{W} = \mathbf{I}$$

- We project n -dimensional data vectors $\mathbf{x}$ onto the subspace: $\mathbf{z} = \mathbf{W}^\top \mathbf{x}$.
- The reconstruction of $\mathbf{x}$ that stays within the m -dimensional subspace defined by $\mathbf{W}$:

$$\hat{\mathbf{x}} = \mathbf{W}\mathbf{z} = \sum_{i=1}^m (\mathbf{w}_i^\top \mathbf{x}) \mathbf{w}_i = \mathbf{W}\mathbf{W}^\top \mathbf{x}$$

- We find $\mathbf{W}$ such that the mean-square error between original data and reconstruction is minimized:

$$\mathbf{W}_{\text{PCA}} = \arg \min_{\mathbf{W}} \mathbb{E} \left\{ \left\| \mathbf{x} - \underbrace{\mathbf{W}\mathbf{W}^\top \mathbf{x}}_{\hat{\mathbf{x}}} \right\|^2 \right\}, \quad \text{s.t. } \mathbf{W}^\top \mathbf{W} = \mathbf{I}$$



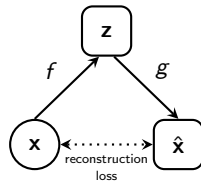
- PCA as an autoencoder: We learn a mapping from $\mathbf{x}$ to $\mathbf{x}$:

$$\hat{\mathbf{x}} = g(f(\mathbf{x}))$$

encoder: $f(\mathbf{x}) = \mathbf{W}_f \mathbf{x} + \mathbf{b}_f$

decoder: $\hat{\mathbf{x}} = g(\mathbf{z}) = \mathbf{W}_g \mathbf{z} + \mathbf{b}_g$

$$\mathcal{L} = \mathbb{E}\{\|\mathbf{x} - g(f(\mathbf{x}))\|^2\}$$



- If we do not restrict f and g , we can learn a trivial identity mapping:

$$\hat{\mathbf{x}} = g(f(\mathbf{x})) = (\mathbf{W}_g \mathbf{W}_f) \mathbf{x} + (\mathbf{W}_g \mathbf{b}_f + \mathbf{b}_g) = \mathbf{x}, \quad \text{if } \mathbf{W}_g = \mathbf{W}_f^{-1} \text{ and } \mathbf{b}_g = -\mathbf{W}_g \mathbf{b}_f$$

- If the dimensionality of $\mathbf{z}$ is smaller than the dimensionality of $\mathbf{x}$, autoencoding is useful: we compress the data.
 - $\mathbf{z}$ is often called a bottleneck.
 - Thus PCA can be implemented with a bottleneck autoencoder.

How can we improve compression?

- We have a linear autoencoder:

$$\hat{\mathbf{x}} = g(f(\mathbf{x}))$$

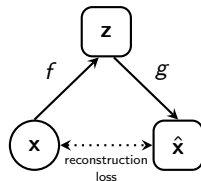
$$f(\mathbf{x}) = \mathbf{W}_f \mathbf{x} + \mathbf{b}_f$$

$$g(\mathbf{z}) = \mathbf{W}_g \mathbf{z} + \mathbf{b}_g$$

- How can we improve compression so that we get a smaller reconstruction error

$$\mathbb{E}\{\|\mathbf{x} - g(f(\mathbf{x}))\|^2\}$$

with a bottleneck layer of the same size?



How can we improve compression?

- We have a linear autoencoder:

$$\hat{\mathbf{x}} = g(f(\mathbf{x}))$$

$$f(\mathbf{x}) = \mathbf{W}_f \mathbf{x} + \mathbf{b}_f$$

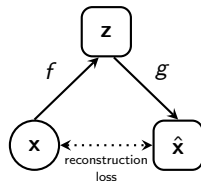
$$g(\mathbf{z}) = \mathbf{W}_g \mathbf{z} + \mathbf{b}_g$$

- How can we improve compression so that we get a smaller reconstruction error

$$\mathbb{E}\{\|\mathbf{x} - g(f(\mathbf{x}))\|^2\}$$

with a bottleneck layer of the same size?

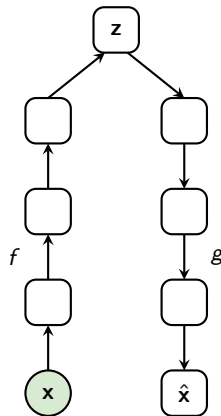
- We can use nonlinear encoder f and decoder g .



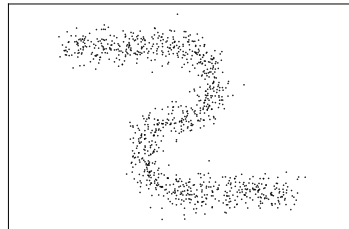
- Deep autoencoder: both the encoder and the decoder are deep neural networks.
- The optimization criterion is the mean-squared reconstruction error:

$$\theta_f, \theta_g = \arg \min_{\theta_f, \theta_g} \mathbb{E} \{ \| \mathbf{x} - g(\mathbf{z}, \theta_g) \|^2 \}, \quad \mathbf{z} = f(\mathbf{x}, \theta_f)$$

- Bottleneck autoencoder: To prevent learning a trivial (identity) function, we use $\mathbf{z}$ with fewer dimensions (a bottleneck layer).
- Bottleneck autoencoders were proposed by Bourlard and Kamp (1988), Oja (1991).



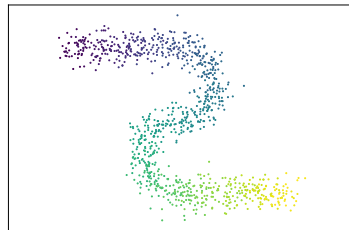
- In this hypothetical example, the data lie on one-dimensional manifold.
- Principal component analysis is not be able to learn the one-dimensional manifold because it is a linear model.



A one-dimensional data manifold in the two-dimensional space.

Deep autoencoders can learn complex data manifolds

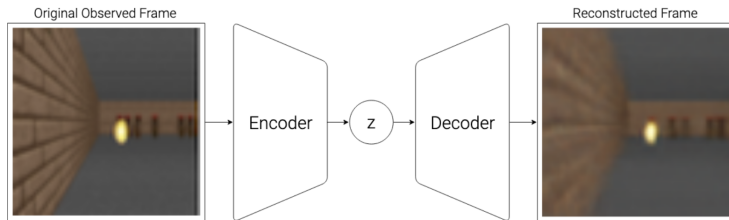
- In this hypothetical example, the data lie on one-dimensional manifold.
- Principal component analysis is not be able to learn the one-dimensional manifold because it is a linear model.
- With a nonlinear autoencoder, we can learn a curved data manifold.
- In our example, colors represents the values of the latent code z that may be found by an autoencoder.



A one-dimensional data manifold in the two-dimensional space.

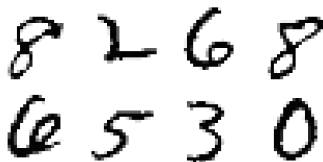
Deep autoencoder as feature extractor

- The most popular use case for deep autoencoders is data compression.
- Consider, for example, reinforcement learning (RL) tasks such as playing Doom ([Ha and Schmidhuber, 2018](#)).
- Learning from raw images (pixels) is likely to require a huge number of training episodes because the amount of input data is large.
- The authors first compress the data using an autoencoder and then train the agent using as observations compressed representations z .

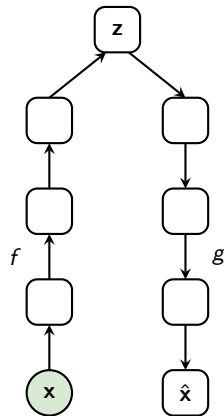
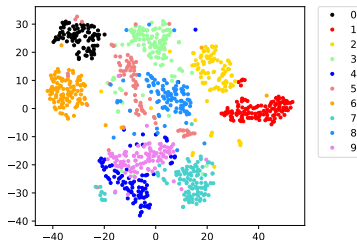


Deep bottleneck autoencoder: MNIST example

- In the home assignment, you will train a bottleneck autoencoder for the MNIST dataset.



- Visualization of the z -space using t-SNE:



Denoising autoencoders

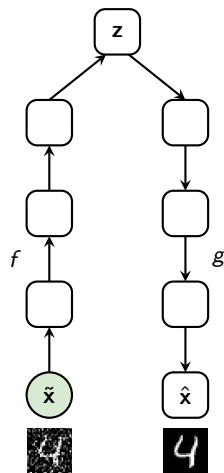
- Feed inputs corrupted with noise (for example, Gaussian):

$$\tilde{\mathbf{x}} = \mathbf{x} + \epsilon \quad \text{with} \quad \epsilon_i \sim \mathcal{N}(0, \sigma^2)$$

- Train $\hat{\mathbf{x}} = g(f(\tilde{\mathbf{x}}))$ to minimize the reconstruction error:

$$c = \mathbb{E}\{\|\hat{\mathbf{x}} - \mathbf{x}\|^2\}$$

- One can view adding noise to inputs as a way to regularize the autoencoder (regularization by noise injection) but there is more theory behind denoising autoencoders.



What does denoising autoencoder learn?

- For Gaussian corruption $\epsilon_i \sim \mathcal{N}(0, \sigma^2)$, the optimal denoising is

$$d(\tilde{\mathbf{x}}) = \tilde{\mathbf{x}} + \sigma^2 \nabla_{\tilde{\mathbf{x}}} \log p(\tilde{\mathbf{x}})$$

(see [Alain and Bengio, 2014](#), [Raphan and Simoncelli, 2011](#))

- $d(\cdot)$ learns to point towards higher probability density.
- Thus, by learning the optimal denoising function $d(\mathbf{x})$, we implicitly model the data distribution $p(\mathbf{x})$.

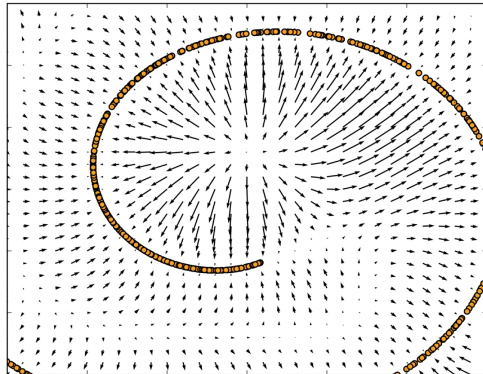
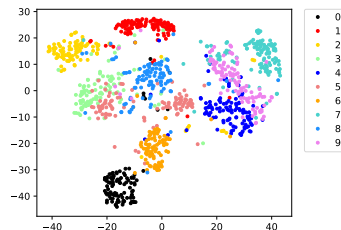
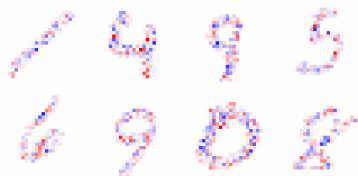


Image from ([Alain and Bengio, 2014](#))

Denoising autoencoder: variance MNIST example

- In the home assignment, we create a synthetic dataset (which we call variance MNIST).
- For this dataset, a vanilla bottleneck autoencoder with mean-squared error reconstruction loss cannot extract high-level features $\mathbf{z}$ that would capture the shapes of the digits.
- A denoising autoencoder can extract meaningful features. Visualization of the $\mathbf{z}$ -space using t-SNE:



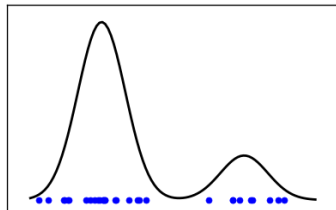
Converting autoencoders into
generative models with latent variables

- Generative models:
 - learn to represent the data distribution $p(\mathbf{x})$
 - can be used to generate new examples from $p(\mathbf{x})$.
- An example: a mixture-of-Gaussians model

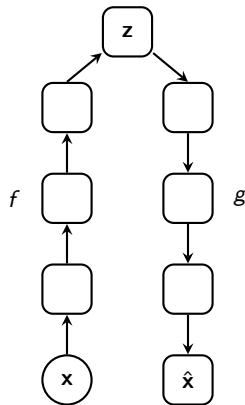
$$p(x | \theta) = w_1 \mathcal{N}(x | \mu_1, \sigma_1^2) + w_2 \mathcal{N}(x | \mu_2, \sigma_2^2)$$

Parameters $\theta = \{w_1, \mu_1, \sigma_1, w_2, \mu_2, \sigma_2\}$ can be estimated by maximum likelihood.

- This model is an example of an explicit density model:
 $p(\mathbf{x} | \theta)$ has an explicit parametric form.



- Vanilla autoencoders are not generative models.
 - We cannot generate new samples from $p(\mathbf{x})$.
 - We cannot compute the probability that a new sample $\mathbf{x}$ comes from the same distribution (e.g., for novelty detection).



- Vanilla autoencoders are not generative models.
 - We cannot generate new samples from $p(\mathbf{x})$.
 - We cannot compute the probability that a new sample $\mathbf{x}$ comes from the same distribution (e.g., for novelty detection).
- We can build a generative model, for example, in this way:

- Assume that variables $\mathbf{z}$ are normally distributed:

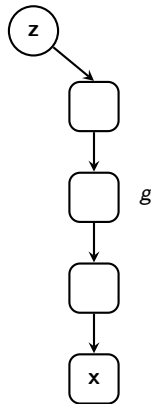
$$\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$$

- Data samples $\mathbf{x}$ are nonlinear transformations of latent variables $\mathbf{z}$:

$$\mathbf{x} = g(\mathbf{z}, \boldsymbol{\theta}) + \boldsymbol{\varepsilon}$$

with possibly noise added: $\boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$

- Function $g(\mathbf{z}, \boldsymbol{\theta})$ can be modeled as a neural network.
- Now we can draw samples from the model.



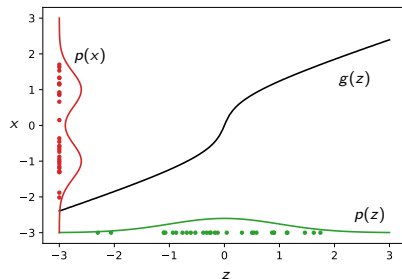
- Our model contains latent (unobserved) variables $\mathbf{z}$:

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

$$\mathbf{x} = g(\mathbf{z}, \boldsymbol{\theta}) + \boldsymbol{\epsilon}$$

$$\boldsymbol{\epsilon} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$$

- A simple example to illustrate the idea: We model one-dimensional data x as a Gaussian variable z transformed with nonlinearity g with some noise added.
- We need to learn the latent variable model from training data $\{\mathbf{x}_i\}$. We should tune parameters $\boldsymbol{\theta}, \sigma^2$ so that the training examples are likely to be produced by the model.



Learning the parameters of the latent variable model

- We can tune parameters θ, σ^2 by maximizing the probability of the training data (maximum likelihood estimate):

$$\theta_{\text{ML}} = \arg \max_{\theta} \log p(\mathbf{x}_1, \dots, \mathbf{x}_N \mid \theta)$$

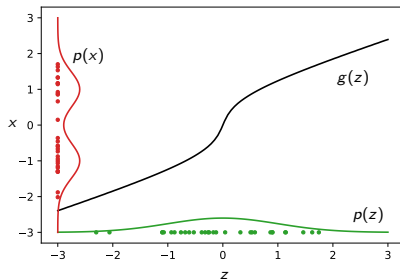
$$\log p(\mathbf{x}_1, \dots, \mathbf{x}_N \mid \theta) = \sum_{i=1}^N \log p(\mathbf{x}_i \mid \theta) = \sum_{i=1}^N \log \int p(\mathbf{x}_i \mid \mathbf{z}_i, \theta) p(\mathbf{z}_i) d\mathbf{z}$$

- The probability density functions are defined by our model:

$$p(\mathbf{x}_i \mid \mathbf{z}_i, \theta) = \mathcal{N}(\mathbf{x}_i \mid g(\mathbf{z}_i, \theta), \sigma^2 \mathbf{I})$$

$$p(\mathbf{z}_i) = \mathcal{N}(\mathbf{z}_i \mid 0, \mathbf{I})$$

- Direct optimization of $\log p(\mathbf{x}_1, \dots, \mathbf{x}_N \mid \theta)$ is difficult because the above integrals are intractable.

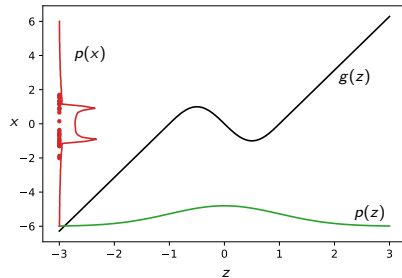


- The classical way to estimate parameters θ of a latent variable model

$$p(\mathbf{x}_1, \dots, \mathbf{x}_N, \mathbf{z}_1, \dots, \mathbf{z}_N \mid \theta) = \prod_{i=1}^N p(\mathbf{x}_i \mid \mathbf{z}_i, \theta) p(\mathbf{z}_i)$$

is the expectation-maximization (EM) algorithm.

- The EM-algorithm iterates between two steps: E-step and M-step.
 - E-step: Compute posterior probabilities $p(\mathbf{z}_i \mid \mathbf{x}_i, \theta)$ given current values of θ .
 - M-step: Update the values of θ using computed $p(\mathbf{z}_i \mid \mathbf{x}_i, \theta)$.



Consider our simple example. We initialize θ with values that give us $\mathbf{g}$ of the form shown in the figure.

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

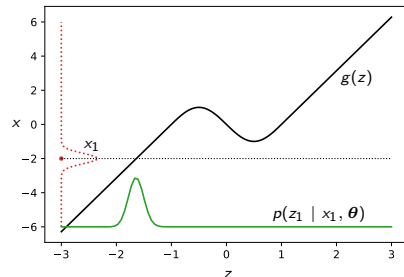
$$\mathbf{x} = g(\mathbf{z}, \boldsymbol{\theta}) + \boldsymbol{\varepsilon}$$

$$\boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$$

- The E-step: Compute the posterior probabilities of the unobserved latent variables $\mathbf{z}_i$ given the data and the current estimates of the model parameters $\boldsymbol{\theta}$:

$$q(\mathbf{z}_1, \dots, \mathbf{z}_N) = q(\mathbf{z}_1) \dots q(\mathbf{z}_N)$$

$$q(\mathbf{z}_i) = p(\mathbf{z}_i \mid \mathbf{x}_i, \boldsymbol{\theta})$$



E-step: For each training data point, find the distribution over the latent variables that could have produced that data point according to the model.

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

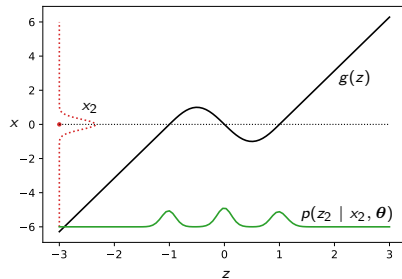
$$\mathbf{x} = g(\mathbf{z}, \boldsymbol{\theta}) + \boldsymbol{\varepsilon}$$

$$\boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$$

- The E-step: Compute the posterior probabilities of the unobserved latent variables $\mathbf{z}_i$ given the data and the current estimates of the model parameters $\boldsymbol{\theta}$:

$$q(\mathbf{z}_1, \dots, \mathbf{z}_N) = q(\mathbf{z}_1) \dots q(\mathbf{z}_N)$$

$$q(\mathbf{z}_i) = p(\mathbf{z}_i \mid \mathbf{x}_i, \boldsymbol{\theta})$$



E-step: For each training data point, find the distribution over the latent variables that could have produced that data point according to the model.

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

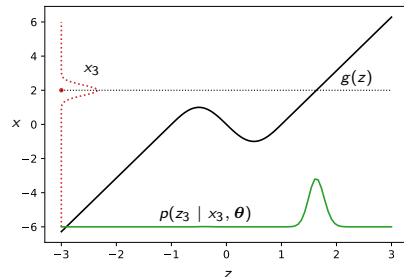
$$\mathbf{x} = g(\mathbf{z}, \boldsymbol{\theta}) + \boldsymbol{\varepsilon}$$

$$\boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$$

- The E-step: Compute the posterior probabilities of the unobserved latent variables $\mathbf{z}_i$ given the data and the current estimates of the model parameters $\boldsymbol{\theta}$:

$$q(\mathbf{z}_1, \dots, \mathbf{z}_N) = q(\mathbf{z}_1) \dots q(\mathbf{z}_N)$$

$$q(\mathbf{z}_i) = p(\mathbf{z}_i \mid \mathbf{x}_i, \boldsymbol{\theta})$$



E-step: For each training data point, find the distribution over the latent variables that could have produced that data point according to the model.

- In the M-step, we use the computed distributions $q(\mathbf{z}_i)$ to form the following objective function:

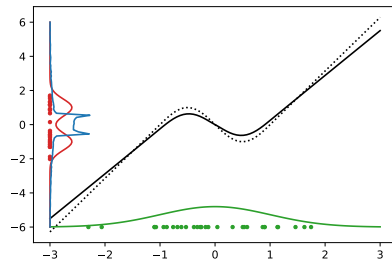
$$\begin{aligned}\mathcal{F}(\theta) &= \langle \log p(\mathbf{x}_1, \dots, \mathbf{x}_N, \mathbf{z}_1, \dots, \mathbf{z}_N \mid \theta) \rangle_{q(\mathbf{z}_1, \dots, \mathbf{z}_N)} \\ &= \sum_{i=1}^N \langle \log p(\mathbf{x}_i, \mathbf{z}_i \mid \theta) \rangle_{q(\mathbf{z}_i)} \\ &= \sum_{i=1}^N \int q(\mathbf{z}_i) \log p(\mathbf{x}_i, \mathbf{z}_i \mid \theta) d\mathbf{z}_i\end{aligned}$$

and maximize it wrt model parameters θ .

- We are guaranteed to improve the likelihood

$$\log p(\mathbf{x}_1, \dots, \mathbf{x}_N \mid \theta)$$

for each iteration of the EM-algorithm.



Iteration 1

- In the M-step, we use the computed distributions $q(\mathbf{z}_i)$ to form the following objective function:

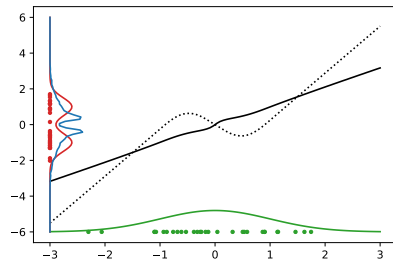
$$\begin{aligned}\mathcal{F}(\theta) &= \langle \log p(\mathbf{x}_1, \dots, \mathbf{x}_N, \mathbf{z}_1, \dots, \mathbf{z}_N \mid \theta) \rangle_{q(\mathbf{z}_1, \dots, \mathbf{z}_N)} \\ &= \sum_{i=1}^N \langle \log p(\mathbf{x}_i, \mathbf{z}_i \mid \theta) \rangle_{q(\mathbf{z}_i)} \\ &= \sum_{i=1}^N \int q(\mathbf{z}_i) \log p(\mathbf{x}_i, \mathbf{z}_i \mid \theta) d\mathbf{z}_i\end{aligned}$$

and maximize it wrt model parameters θ .

- We are guaranteed to improve the likelihood

$$\log p(\mathbf{x}_1, \dots, \mathbf{x}_N \mid \theta)$$

for each iteration of the EM-algorithm.



Iteration 2

- In the M-step, we use the computed distributions $q(\mathbf{z}_i)$ to form the following objective function:

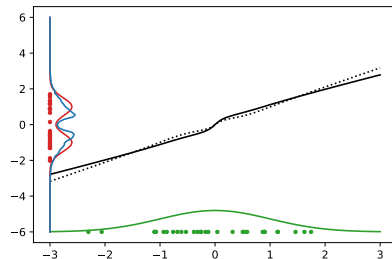
$$\begin{aligned}\mathcal{F}(\theta) &= \langle \log p(\mathbf{x}_1, \dots, \mathbf{x}_N, \mathbf{z}_1, \dots, \mathbf{z}_N \mid \theta) \rangle_{q(\mathbf{z}_1, \dots, \mathbf{z}_N)} \\ &= \sum_{i=1}^N \langle \log p(\mathbf{x}_i, \mathbf{z}_i \mid \theta) \rangle_{q(\mathbf{z}_i)} \\ &= \sum_{i=1}^N \int q(\mathbf{z}_i) \log p(\mathbf{x}_i, \mathbf{z}_i \mid \theta) d\mathbf{z}_i\end{aligned}$$

and maximize it wrt model parameters θ .

- We are guaranteed to improve the likelihood

$$\log p(\mathbf{x}_1, \dots, \mathbf{x}_N \mid \theta)$$

for each iteration of the EM-algorithm.



Iteration 3

- In the M-step, we use the computed distributions $q(\mathbf{z}_i)$ to form the following objective function:

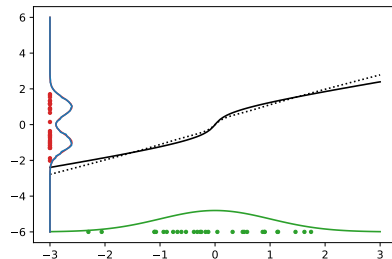
$$\begin{aligned}\mathcal{F}(\theta) &= \langle \log p(\mathbf{x}_1, \dots, \mathbf{x}_N, \mathbf{z}_1, \dots, \mathbf{z}_N \mid \theta) \rangle_{q(\mathbf{z}_1, \dots, \mathbf{z}_N)} \\ &= \sum_{i=1}^N \langle \log p(\mathbf{x}_i, \mathbf{z}_i \mid \theta) \rangle_{q(\mathbf{z}_i)} \\ &= \sum_{i=1}^N \int q(\mathbf{z}_i) \log p(\mathbf{x}_i, \mathbf{z}_i \mid \theta) d\mathbf{z}_i\end{aligned}$$

and maximize it wrt model parameters θ .

- We are guaranteed to improve the likelihood

$$\log p(\mathbf{x}_1, \dots, \mathbf{x}_N \mid \theta)$$

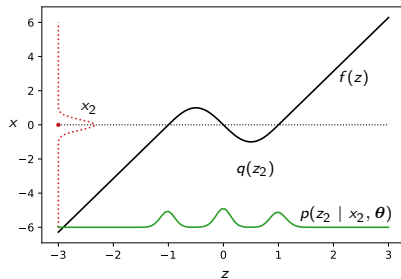
for each iteration of the EM-algorithm.



Iteration 4

Learning latent variable models
with variational approximations

- There are a few problems with the direct application of the EM-algorithm in nonlinear latent variable models.
- One problem is the intractability of the true conditional distributions $q(\mathbf{z}_i) = p(\mathbf{z}_i | \mathbf{x}_i, \boldsymbol{\theta})$ that we need to compute on the E-step.
- The true distributions can be very complex (for example, a multi-modal distribution in our simple example).



Example of multi-modal $p(\mathbf{z}_i | \mathbf{x}_i, \boldsymbol{\theta})$

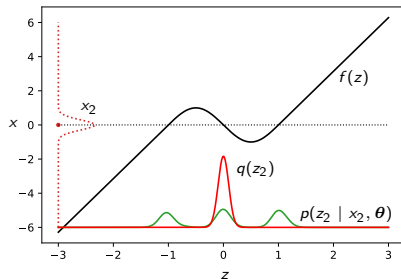
E-step: Variational approximations

- Solution: Instead of using true conditional distributions, use their approximations $q(\mathbf{z}_i) \approx p(\mathbf{z}_i | \mathbf{x}_i, \boldsymbol{\theta})$.
- $q(\mathbf{z}_i)$ is selected to have a simple form, most often a Gaussian:

$$q(\mathbf{z}_i) = \mathcal{N}(\mu_{\mathbf{z}_i}, \sigma_{\mathbf{z}_i}^2)$$

Note: we have two parameters $\mu_{\mathbf{z}_i}$ and $\sigma_{\mathbf{z}_i}^2$ describing $q(\mathbf{z}_i)$ for *each training sample*.

- Parameters describing the posterior distributions of the latent variables $\boldsymbol{\theta}_q = \{\mu_{\mathbf{z}_i}, \sigma_{\mathbf{z}_i}^2\}_{i=1}^N$ are called variational parameters.
- A popular way to find the approximation is by minimizing the Kullback-Leibler divergence between $q(\mathbf{z}_i)$ and $p(\mathbf{z}_i | \mathbf{x}_i, \boldsymbol{\theta})$.



- We can minimize the KL divergence between $q(\mathbf{z}_i)$ and $p(\mathbf{z}_i | \mathbf{x}_i, \boldsymbol{\theta})$ using the following trick:
 - Add to the objective function used in the M-step the entropies of the approximate distributions:

$$\begin{aligned}\mathcal{F}(\boldsymbol{\theta}, \boldsymbol{\theta}_q) &= \sum_{i=1}^N \underbrace{\int q(\mathbf{z}_i) \log p(\mathbf{x}_i, \mathbf{z}_i | \boldsymbol{\theta}) d\mathbf{z}_i}_{\text{what we had in the M-step}} - \underbrace{\int q(\mathbf{z}_i) \log q(\mathbf{z}_i) d\mathbf{z}_i}_{\text{entropy}} \\ &= \sum_{i=1}^N \int q(\mathbf{z}_i) \log \frac{p(\mathbf{x}_i, \mathbf{z}_i | \boldsymbol{\theta})}{q(\mathbf{z}_i)} d\mathbf{z}_i = \sum_{i=1}^N \int q(\mathbf{z}_i) \log \frac{p(\mathbf{z}_i | \mathbf{x}_i, \boldsymbol{\theta}) p(\mathbf{x}_i | \boldsymbol{\theta})}{q(\mathbf{z}_i)} d\mathbf{z}_i \\ &= \sum_{i=1}^N -D_{\text{KL}}(q(\mathbf{z}_i) \parallel p(\mathbf{z}_i | \mathbf{x}_i, \boldsymbol{\theta})) + \log p(\mathbf{x}_i | \boldsymbol{\theta})\end{aligned}$$

- One can see that maximizing $\mathcal{F}(\boldsymbol{\theta}, \boldsymbol{\theta}_q)$ wrt variational parameters $\boldsymbol{\theta}_q$ is equivalent to minimizing the KL divergence between $q(\mathbf{z}_i)$ and $p(\mathbf{z}_i | \mathbf{x}_i, \boldsymbol{\theta})$.

- We can now maximize a single function $\mathcal{F}$ wrt θ and θ_q jointly without the need to alternate between the E- and M-steps:

$$\begin{aligned}\mathcal{F}(\theta, \theta_q) &= \sum_{i=1}^N \int q(\mathbf{z}_i) \log p(\mathbf{x}_i, \mathbf{z}_i \mid \theta) d\mathbf{z}_i - \int q(\mathbf{z}_i) \log q(\mathbf{z}_i) d\mathbf{z}_i \\ &= \sum_{i=1}^N -D_{\text{KL}}(q(\mathbf{z}_i) \parallel p(\mathbf{z}_i \mid \mathbf{x}_i, \theta)) + \log p(\mathbf{x}_i \mid \theta)\end{aligned}$$

- Maximizing $\mathcal{F}(\theta, \theta_q)$ wrt θ is equivalent to the M-step.
- Maximizing $\mathcal{F}(\theta, \theta_q)$ wrt θ_q is done in the E-step with variational approximations.
- We can solve this optimization problem using any optimizer of our choice.

- The objective function

$$\mathcal{F}(\boldsymbol{\theta}, \boldsymbol{\theta}_q) = \sum_{i=1}^N -D_{\text{KL}}(q(\mathbf{z}_i) \parallel p(\mathbf{z}_i \mid \mathbf{x}_i, \boldsymbol{\theta})) + \log p(\mathbf{x}_i \mid \boldsymbol{\theta})$$

is the *lower bound* of the true likelihood that we want to optimize. Since $D_{\text{KL}}(q \parallel p) \geq 0$:

$$\mathcal{F}(\boldsymbol{\theta}, \boldsymbol{\theta}_q) \leq \sum_{i=1}^N \log p(\mathbf{x}_i \mid \boldsymbol{\theta}) = \log p(\mathbf{x}_1, \dots, \mathbf{x}_N \mid \boldsymbol{\theta})$$

- This function is often called *evidence lower bound* or ELBO.
- The closer our approximation $q(\mathbf{z}_i)$ to the true posterior $p(\mathbf{z}_i \mid \mathbf{x}_i, \boldsymbol{\theta})$, the tighter the bound.

- ELBO can be re-written in the following form:

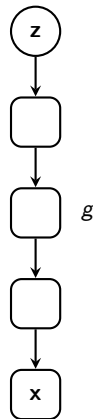
$$\mathcal{F}(\theta, \theta_q) = \sum_{i=1}^N \int q(\mathbf{z}_i) \log p(\mathbf{x}_i | \mathbf{z}_i, \theta) d\mathbf{z}_i - \int q(\mathbf{z}_i) \log \frac{q(\mathbf{z}_i)}{p(\mathbf{z}_i)} d\mathbf{z}_i \quad (1)$$

- Recall our deep generative model: $p(\mathbf{x}_i | \mathbf{z}_i, \theta) = \mathcal{N}(\mathbf{x}_i | g(\mathbf{z}_i, \theta), \sigma^2 \mathbf{I})$,
- The first term in equation (1) can be written as

$$\left\langle -\frac{D}{2} \log 2\pi\sigma^2 - \frac{1}{2\sigma^2} \sum_{d=1}^D (\mathbf{x}_i(d) - g_d(\mathbf{z}_i, \theta))^2 \right\rangle_{q(\mathbf{z}_i)}$$

where D is the number of dimensions in $\mathbf{x}$, $\mathbf{x}_i(d)$ is the d -th element of $\mathbf{x}_i$ and g_d is the d -th element of the output of function g .

- The first term contains the expected mean-squared error between data sample $\mathbf{x}_i$ and its reconstruction $g_d(\mathbf{z}_i, \theta)$ from the latent code $\mathbf{z}_i$.



$$\mathcal{F}(\boldsymbol{\theta}, \boldsymbol{\theta}_q) = \sum_{i=1}^N \underbrace{\int q(\mathbf{z}_i) \log p(\mathbf{x}_i | \mathbf{z}_i, \boldsymbol{\theta}) d\mathbf{z}_i}_{\text{minus mean-square reconstruction error}} - \underbrace{\int q(\mathbf{z}_i) \log \frac{q(\mathbf{z}_i)}{p(\mathbf{z}_i)} d\mathbf{z}_i}_{\text{regularization term}}$$

- The second term is minus KL-divergence between $q(\mathbf{z}_i)$ and the prior $p(\mathbf{z}_i) = \mathcal{N}(0, \mathbf{I})$:

$$- \int q(\mathbf{z}_i) \log \frac{q(\mathbf{z}_i)}{p(\mathbf{z}_i)} d\mathbf{z}_i = -D_{\text{KL}}(q(\mathbf{z}_i) \parallel p(\mathbf{z}_i))$$

- It is a kind of a regularization term: We want the conditional distributions $q(\mathbf{z}_i)$ to be close to the prior $p(\mathbf{z}_i) = \mathcal{N}(0, \mathbf{I})$.

Variational autoencoders

- The first algorithm for learning latent variable model

$$\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}) \quad \mathbf{x} = g(\mathbf{z}, \boldsymbol{\theta}) + \boldsymbol{\varepsilon} \quad \boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$$

using variational approximations was proposed in this university ([Lappalainen and Honkela, 2001](#)).

- The objective function was ELBO:

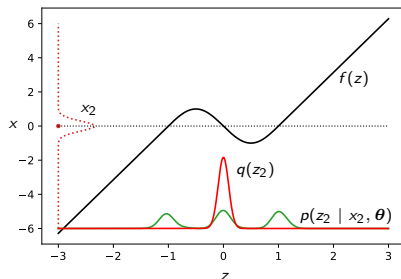
$$\mathcal{F}(\boldsymbol{\theta}, \boldsymbol{\theta}_q) = \sum_{i=1}^N \underbrace{\int q(\mathbf{z}_i) \log p(\mathbf{x}_i | \mathbf{z}_i, \boldsymbol{\theta}) d\mathbf{z}_i}_{\text{needs approximations}} - \underbrace{\int q(\mathbf{z}_i) \log \frac{q(\mathbf{z}_i)}{p(\mathbf{z}_i)} d\mathbf{z}_i}_{\text{can be computed analytically}}$$

- The posterior approximations were Gaussian $q(\mathbf{z}_i) = \mathcal{N}(\mu_{\mathbf{z}_i}, \sigma_{\mathbf{z}_i}^2)$. The number of variational parameters $\boldsymbol{\theta}_q = \{\mu_{\mathbf{z}_i}, \sigma_{\mathbf{z}_i}^2\}_{i=1}^N$ was proportional to the number of training samples.

- We want to get rid of the large number of variational parameters $\theta_q = \{\mu_{z_i}, \sigma_{z_i}^2\}_{i=1}^N$.
- For fixed model parameters θ , the optimal $q(\mathbf{z})$ only depends on $\mathbf{x}$. The inference procedure does the following mapping:

$$\mathbf{x} \rightarrow q(\mathbf{z})$$

For Gaussian approximation: $\mathbf{x} \rightarrow \mu_z, \sigma_z^2$.



- In variational autoencoders (VAE) (Kingma and Welling, 2014), mapping $\mathbf{x} \rightarrow q(\mathbf{z})$ is done using a neural network (encoder).
- The encoder performs so called *amortized inference*: When doing inference for a particular sample $\mathbf{x}_i$, we leverage the knowledge of the inference results for other samples. If two samples $\mathbf{x}_i$ and $\mathbf{x}_j$ are close to each other, the corresponding $q(\mathbf{z}_i)$, $q(\mathbf{z}_j)$ should be close as well.

Variational autoencoder (VAE): Encoder and decoder

- Our generative model is defined by the decoder.

$$\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}) \quad \mathbf{x} = g(\mathbf{z}, \boldsymbol{\theta}) + \boldsymbol{\varepsilon} \quad \boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$$

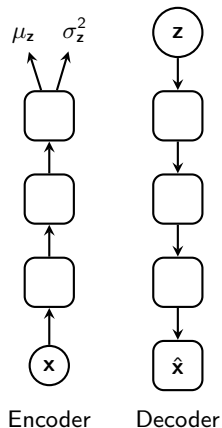
- Encoder is a neural network that is trained to perform variational inference:

$$\mathbf{x} \rightarrow q(\mathbf{z})$$

- For Gaussian approximation $q(\mathbf{z})$, the neural network needs to produce:

$$\mathbf{x} \rightarrow \mu_{\mathbf{z}}, \sigma_{\mathbf{z}}^2$$

- In practice, this is done using one neural network with two heads.
- The encoder is similar to the encoder in a bottleneck autoencoder but produces the mean and variance of the code $\mathbf{z}$.
- The encoder and decoder are two components of the variational autoencoder.



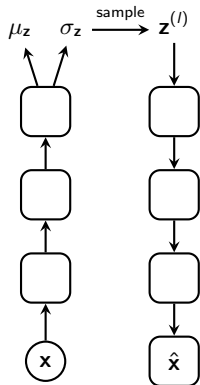
- The first term of the objective function cannot be computed analytically

$$\mathcal{F}(\theta, \theta_q) = \sum_{i=1}^N \underbrace{\int q(\mathbf{z}_i) \log p(\mathbf{x}_i | \mathbf{z}_i, \theta) d\mathbf{z}_i}_{\text{needs approximations}} - \underbrace{\int q(\mathbf{z}_i) \log \frac{q(\mathbf{z}_i)}{p(\mathbf{z}_i)} d\mathbf{z}_i}_{\text{can be computed analytically}}$$

- Kingma and Welling (2014) proposed to use Monte Carlo estimates:

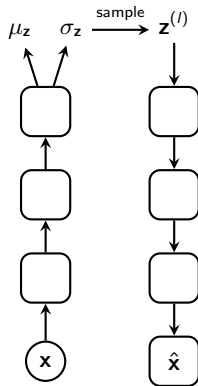
$$\int q(\mathbf{z}_i) \log \mathcal{N}(\mathbf{x}_i | g(\mathbf{z}_i, \theta), \sigma^2 \mathbf{I}) d\mathbf{z}_i \approx \frac{1}{L} \sum_{l=1}^L \log \mathcal{N}(\mathbf{x}_i | g(\mathbf{z}_i^{(l)}, \theta), \sigma^2 \mathbf{I})$$

where $\mathbf{z}_i^{(l)}$ are drawn from $q(\mathbf{z}_i)$. Using $L = 1$ works well in practice.

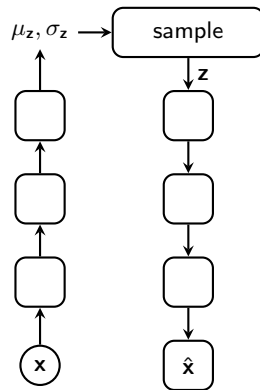


$$\mathcal{F}(\theta, \theta_q) = \sum_{i=1}^N \underbrace{\log \mathcal{N}(\mathbf{x}_i \mid g(\mathbf{z}_i^{(l)}, \theta), \sigma^2 \mathbf{I})}_{\text{Monte Carlo estimate}} - \underbrace{\int q(\mathbf{z}_i) \log \frac{q(\mathbf{z}_i)}{p(\mathbf{z}_i)} d\mathbf{z}_i}_{\text{can be computed analytically}}$$

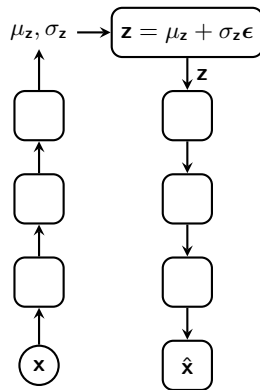
- For each training example $\mathbf{x}_i$:
 - compute means $\mu_{\mathbf{z}_i}$ and $\sigma_{\mathbf{z}_i}$ using the encoder
 - compute the second term analytically
 - draw $L = 1$ samples $\mathbf{z}_i^{(l)}$ from $q(\mathbf{z}_i) = \mathcal{N}(\mu_{\mathbf{z}_i}, \sigma_{\mathbf{z}_i}^2)$
 - propagate $\mathbf{z}_i^{(l)}$ through the decoder and compute the first term
- Problem: We can use backpropagation to compute the derivatives wrt the parameters of the decoder but we need an extra trick to propagate derivatives through the encoder.



- We need a computational block that would
 - take as inputs μ_z and σ_z
 - produce a sample from distribution $\epsilon \sim \mathcal{N}(\mu_{z_i}, \sigma_{z_i})$
 - would be differentiable wrt μ_z and σ_z



- We need a computational block that would
 - take as inputs μ_z and σ_z
 - produce a sample from distribution $\epsilon \sim \mathcal{N}(\mu_{z_i}, \sigma_{z_i})$
 - would be differentiable wrt μ_z and σ_z
- We can obtain this with the reparameterization trick:
 - Sample $\epsilon \sim \mathcal{N}(0, \mathbf{I})$
 - Compute $\mathbf{z}_i = \mu_{z_i} + \sigma_{z_i} \epsilon_i$
- Now we can also backpropagate through the sampling block and then further through the encoder.

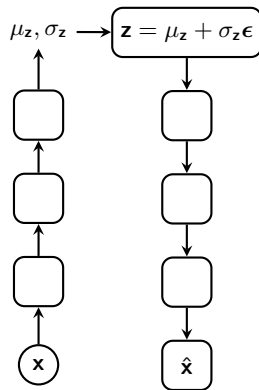


- VAE training algorithm:

- Take a mini-batch $\{\mathbf{x}_i\}$ of training samples.
- Use the encoder to compute means $\mu_{\mathbf{z}_i}$ and standard deviations $\sigma_{\mathbf{z}_i}$ for each sample $\mathbf{x}_i$ in the mini-batch.
- Draw $\epsilon_i \sim \mathcal{N}(0, \mathbf{I})$ and compute samples $\mathbf{z}_i = \mu_{\mathbf{z}_i} + \sigma_{\mathbf{z}_i} \epsilon_i$
- Propagate samples $\mathbf{z}_i$ through the decoder to compute reconstructions $\hat{\mathbf{x}}_i$.
- Compute the loss which is the negative of

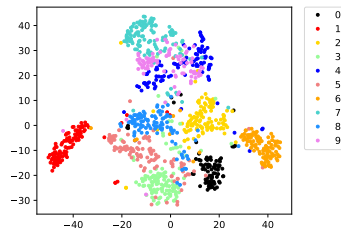
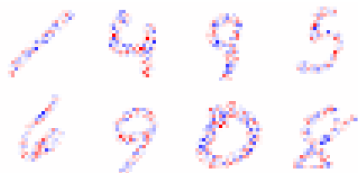
$$\mathcal{F}(\theta, \theta_q) = \frac{1}{n} \sum_{i=1}^n \underbrace{\log \mathcal{N}(\mathbf{x}_i \mid g(\mathbf{z}_i^{(l)}, \theta), \sigma^2 \mathbf{I})}_{\text{Monte Carlo estimate}} - \underbrace{\int q(\mathbf{z}_i) \log \frac{q(\mathbf{z}_i)}{p(\mathbf{z}_i)} d\mathbf{z}_i}_{\text{can be computed analytically}}$$

- Perform backpropagation and update the parameters of the encoder and the decoder.



- In the home assignment, we train a variational autoencoder on a synthetic (variance MNIST) dataset.
- In order to extract meaningful features for this dataset, we need to use a generator (decoder) that models the variances of pixel intensities:

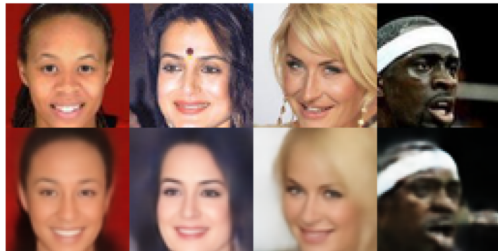
$$\begin{aligned} \mathbf{z} &\sim \mathcal{N}(\mathbf{0}, \mathbf{I}) & \mathbf{x} &\sim \mathcal{N}(\boldsymbol{\mu}(\mathbf{z}), \text{diag}(\boldsymbol{\sigma}(\mathbf{z}))) \\ \boldsymbol{\mu}(\mathbf{z}) &= g_{\mu}(\mathbf{z}, \boldsymbol{\theta}) & \boldsymbol{\sigma}(\mathbf{z}) &= \exp(g_{\sigma}(\mathbf{z}, \boldsymbol{\theta})) \end{aligned}$$



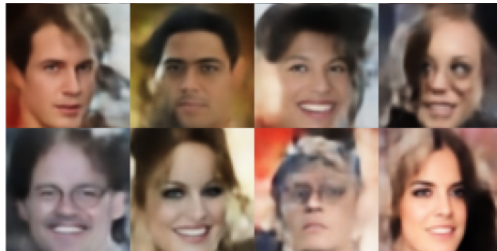
- VAE is more complex than a simple bottleneck autoencoder. Do we need these complications?
- As we will see in the home assignment, VAEs are more powerful. In some problems when vanilla autoencoders fail, VAEs can develop useful representations.
- The problem of the vanilla autoencoder is the mean-squared error loss, which makes too simplistic assumptions about the data distribution.
- One advantage of VAE is in greater flexibility in defining the generative model.
- Note that denoising autoencoders are more powerful than standard autoencoders even though they also use the mean-squared error loss.

- The main benefit of VAEs is that we can encode data into a lower-dimensional representation.
- But VAEs are generative models and we can draw samples using VAEs.
- So far, the quality of the VAE-generated samples have not been very impressive (the samples as well as reconstructions usually look blurry).

Reconstructions



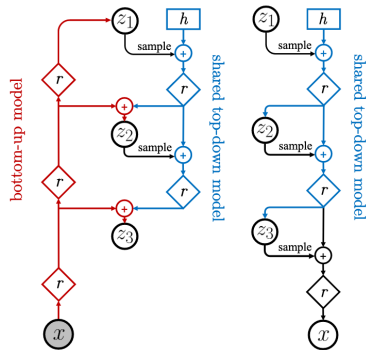
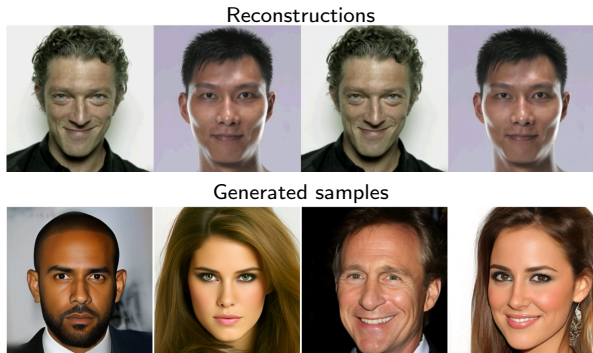
Generated samples



Images from (Tolstikhin et al., 2017)

Nouveau VAE (NVAE; Vahdat and Kautz, 2020)

- Vahdat and Kautz (2020) presented a VAE model that is able to generate high-quality images.
- It is a hierarchical latent variable model, that is there are multiple levels of latent variables.



Home assignment

- In the home assignment, you will have to implement three types of autoencoders:
 1. Vanilla bottleneck autoencoder
 2. Denoising autoencoder
 3. Variational autoencoder

- Chapter 14 of the Deep Learning book
- Papers cited in the lecture slides