

Review of Denoising Diffusion Probabilistic Models

Minji Kim Hyeon Lee
University of North Carolina, Chapel Hill

1 Introduction

What do generative models do? In the machine learning literature, “generation” refers to sampling X from the ground truth probability distribution of the target data $\mathbb{P}_{\text{data}}$,

$$X \sim \mathbb{P}_{\text{data}},$$

obtaining a data point that was not observed or used during the learning process. In practice, X usually represents complicated data such as images, texts, audios, or even movies. An important class of generative models is *probabilistic generative models*; they aim to learn the data distribution itself, finding a parameterized function

$$f_{\theta}(Z) \sim \mathbb{P}_{\text{data}},$$

where Z is a random variable from a noise space. In other words, it aims to learn a mapping from Z to a random variable of the data distribution.

There exist other types of generative models depending on the tasks. For example, *non-probabilistic generative models* such as Generative Adversarial Networks (GANs) (Goodfellow et al. (2014)) can be used when one requires samples from the population distribution but the distribution itself is not of interest. On the other hand, one may wish to request the model to generate images based on certain descriptions; these procedure is called *conditional generation*. The last example is *sequential generation*, which is useful for generating long text data or generating musics. In this report, we restrict our attention to the ground problem to generate samples just from the data distribution $\mathbb{P}_{\text{data}}$.

A Diffusion model is a type of probabilistic generative models which has garnered considerable attention from researchers due to its capacity for generating high-quality samples comparable to those produced by Generative Adversarial Networks (GANs), while also generating a diversity of samples akin to Variation Autoencoders (VAE) (Kingma and Welling (2014)) and Normalizing Flows (Rezende and Mohamed (2015)). Denoising diffusion models were initially introduced by Sohl-Dickstein et al. (2015), and early related work focusing on score-matching was developed by Song and Ermon (2019). Ho et al. (2020) significantly advanced this field by successfully producing high-quality training results in generative imaging.

In this report, we aim to deliver what are the diffusion models and how variational inference is employed in its learning procedure. To understand diffusion models’ strengths and weaknesses, we also review core ideas of the other three popular generative models mentioned above. Figure 1 summarizes their relative strengths and weaknesses, which will become clear as we learn their model structures in the later sections. Most of the contents in this report is derived from these papers and Prince (2023).

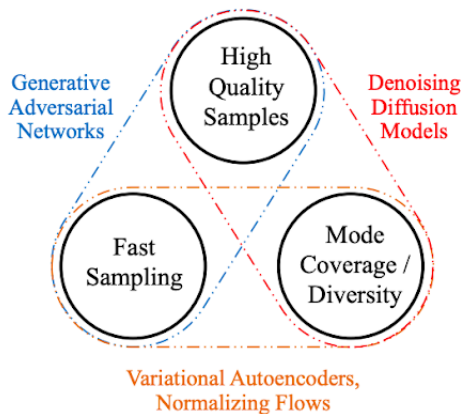


Figure 1: Strengths and Weaknesses of Generative Models, (Image from NVIDIA, link)

This report is organized as follows. We first discuss performance measures illustrated in Figure 1, and briefly review the core ideas of GAN, VAE, and normalizing flows. Section 3 provides a comprehensive overview and introduce the structure of the diffusion probabilistic model. Section 4 summarizes statistical techniques used for efficient optimization, while Section 5 delves into the reparameterization strategy that facilitates effective training and sampling algorithms. Additionally, we include special notes on two key aspects: First, we examine what information about the data distribution the model learns through the denoising process, revealing that it primarily learns the score function of the input data distribution. Second, we discuss how diffusion models approximate the solution to a specific stochastic differential equation.

2 Related Works

2.1 Performance measure for generative models

Goodness of generative models is subjective and there are different criteria depending on the specific tasks. However, four performance measures listed below are most fundamental and applicable to most of the generative models.

- **High Quality Samples:** The generated samples should be indistinguishable from the data that the model learned.

- **Fast Sampling** Once the model was trained, generating new samples from the model should be computationally efficient.
- **Mode Coverage/Diversity** The model should generate new samples from the entire data space, rather than drawing samples from a relatively small region of the data space.
- **Probabilistic Model** It is a strength for models to learn probability of data points, as the probability can be used for further analysis such as anomaly detection.

As we will see in the later sections, GAN is a non-probabilistic model and all other three models are probabilistic models. As the opposite of the term mode coverage, the phenomenon where a trained model generates samples from a small part of the data space is called *mode collapse*; GAN is well-known for its susceptibility to model collapse, because it aims to generate samples that are not distinguishable from the training data set without accounting for how the training data points are likely to occur within the model. In contrast, VAE and normalizing flows are fast and evenly generate samples from the data space, but they are limited to generate high quality samples due to their simplicity of the model and the latent space, respectively. A diffusion model lies in between them in the sense that it is capable of generating high quality samples from the entire data space, at the cost of higher computational expense.

2.2 Generative Adversarial Network

A *Generative Adversarial Network (GAN)* comprises two machines, a *generator* and a *discriminator*. A generator $\mathbf{g}(\mathbf{z}, \boldsymbol{\theta})$ takes a random noise $\mathbf{z}$ from a base distribution and generates a fake data $\mathbf{x} = \mathbf{g}(\mathbf{z}, \boldsymbol{\theta})$. A discriminator $\mathbf{f}(\mathbf{x}, \boldsymbol{\phi})$ takes either a real or a fake data $\mathbf{x}$ and returns the logit of the probability that the data comes from the real observations. Two machines go into a 2-player game: a discriminator tries to best classify inputs into real and fake observations, and a generator tries to forge new observations so the discriminator cannot distinguish them from the real data. Suppose we adopt the Bernoulli log-likelihood as the loss, that is,

$$l(\hat{p}, y) = -(1 - y) \log(1 - \hat{p}) - y \log(\hat{p}) \quad (1)$$

where $y = 0$ if the data was real and $y = 1$ if the data was fake.

In each iteration of training, the discriminator learns the parameter $\boldsymbol{\phi}$ from the set of real observations $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$ and the set of fake observations $\{\mathbf{x}_1^*, \dots, \mathbf{x}_m^*\}$ by minimizing the average loss as follows.

$$\hat{\boldsymbol{\phi}} := \arg \min_{\boldsymbol{\phi}} \left\{ \sum_{j=1}^m -\log \left[1 - \text{sign}(\mathbf{f}(\mathbf{x}_j^*, \boldsymbol{\phi})) \right] - \sum_{i=1}^n \log \left[\text{sign}(\mathbf{f}(\mathbf{x}_i, \boldsymbol{\phi})) \right] \right\} \quad (2)$$

The generator learns the parameter $\boldsymbol{\theta}$ by maximizing the minimum loss that a discriminator can

yield:

$$\hat{\boldsymbol{\theta}} := \arg \max_{\boldsymbol{\theta}} \left\{ \min_{\boldsymbol{\phi}} \left\{ \sum_{j=1}^m -\log \left[1 - \text{sign}(\mathbf{f}(\mathbf{g}(\mathbf{z}_j, \boldsymbol{\theta}), \boldsymbol{\phi})) \right] - \sum_{i=1}^n \log \left[\text{sign}(\mathbf{f}(\mathbf{x}_i, \boldsymbol{\phi})) \right] \right\} \right\}. \quad (3)$$

2.3 Normalizing Flow

The purpose of *normalizing flow* is to learn a function $\mathbf{f}(\cdot, \boldsymbol{\phi})$ such that the distribution of $\mathbf{f}(\mathbf{Z}, \boldsymbol{\phi})$ best mimics the sample distribution, where $\mathbf{Z}$ follows a pre-specified distribution which is usually chosen as standard multivariate normal distribution. If $\mathbf{f}(\cdot, \boldsymbol{\phi})$ is one-to-one, the density of $\mathbf{X} = \mathbf{f}(\mathbf{Z}, \boldsymbol{\phi})$ is given as

$$\mathbf{f}_{\mathbf{X}}(\mathbf{x}) = \left| \frac{\partial \mathbf{f}(\mathbf{z}, \boldsymbol{\phi})}{\partial \mathbf{z}} \right|^{-1} \mathbf{f}_{\mathbf{Z}}(\mathbf{z}), \mathbf{z} = \mathbf{f}^{-1}(\mathbf{x}, \boldsymbol{\phi}). \quad (4)$$

We estimate $\boldsymbol{\phi}$ by maximizing likelihood of observed data $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$. For $\mathbf{z}_i = \mathbf{f}^{-1}(\mathbf{x}_i, \boldsymbol{\phi}), i = 1, \dots, n$,

$$\begin{aligned} \hat{\boldsymbol{\phi}} &:= \arg \max_{\boldsymbol{\phi}} \left\{ \sum_{i=1}^n \log \left(\left| \frac{\partial \mathbf{f}(\mathbf{z}_i, \boldsymbol{\phi})}{\partial \mathbf{z}_i} \right|^{-1} \mathbf{f}_{\mathbf{Z}}(\mathbf{z}_i) \right) \right\} \\ &= \arg \min_{\boldsymbol{\phi}} \left\{ \sum_{i=1}^n \log \left(\left| \frac{\partial \mathbf{f}(\mathbf{z}_i, \boldsymbol{\phi})}{\partial \mathbf{z}_i} \right| \right) \right\} \end{aligned} \quad (5)$$

An important assumption for this formulation is that $\mathbf{f}(\cdot, \boldsymbol{\phi})$ is invertible. Because it is difficult to maintain a function being invertible while allowing it to be flexible at the same time, we factor the function as the composition of sequence of simple, invertible functions.

$$\mathbf{x} = \mathbf{f}(\mathbf{z}, \boldsymbol{\phi}) = \mathbf{f}_K(\mathbf{f}_{K-1}(\dots \mathbf{f}_2(\mathbf{f}_1(\mathbf{z}, \phi_1), \phi_2), \dots \phi_{k-1}), \phi_K). \quad (6)$$

Writing the sequence in the reverse order gives a function that maps $\mathbf{X}$ to $\mathbf{Z}$. This formulation explains the name “normalizing flow”, because the sequence of functions flows the sample distribution towards a normal distribution.

$$\mathbf{z} = \mathbf{f}^{-1}(\mathbf{x}, \boldsymbol{\phi}) = \mathbf{f}_1^{-1}(\mathbf{f}_2^{-1}(\dots \mathbf{f}_{K-1}^{-1}(\mathbf{f}_K^{-1}(\mathbf{z}, \phi_K), \phi_{K-1}), \dots \phi_2), \phi_1). \quad (7)$$

Normalizing flows compute exact likelihood of thus are capable of generating high-quality samples

2.4 Variational Autoencoder

A *Variational Autoencoder (VAE)* models the distribution of data $\mathbf{X}$ through a lower dimensional latent variable $\mathbf{Z}$ with a simple and known distribution. Suppose $p_{\boldsymbol{\phi}}(\mathbf{x}|\mathbf{z})$ models the conditional

density function of $\mathbf{x}$ given $\mathbf{z}$. Then we have

$$p_\phi(\mathbf{x}) = \int p_\phi(\mathbf{x}|\mathbf{z})p(\mathbf{z})d\mathbf{z} \quad (8)$$

and the model learns the parameter ϕ by maximizing log-likelihood of the sample $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$

$$\sum_{i=1}^n \log \int p_\phi(\mathbf{x}_i|\mathbf{z}_i)p(\mathbf{z}_i)d\mathbf{z}_i. \quad (9)$$

The distribution of $\mathbf{Z}$ is usually chosen to be a standard multivariate normal distribution.

Nevertheless, direct optimization of Equation (9) over ϕ is not straightforward because it involves an integral with respect to a complicated function $p_\phi(\mathbf{x}|\mathbf{z})$. VAE overcomes this challenge by approximating the reverse relationship $p_\phi(\mathbf{z}|\mathbf{x})$ using a normal distribution $q_\theta(\mathbf{z}|\mathbf{x})$ which is easy to manipulate. Let us denote by $\mathcal{N}_x(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ the density function of normal distribution with the corresponding mean and variance:

$$\mathcal{N}_x(\boldsymbol{\mu}, \boldsymbol{\Sigma}) = |\mathbf{2}\pi\boldsymbol{\Sigma}|^{-1/2} \exp\left(-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu})^\top \boldsymbol{\Sigma}^{-1}(\mathbf{x} - \boldsymbol{\mu})\right).$$

Then,

$$q_\theta(\mathbf{z}|\mathbf{x}) \approx p_\phi(\mathbf{z}|\mathbf{x}), \quad q_\theta(\mathbf{z}|\mathbf{x}) = \mathcal{N}_z(\boldsymbol{\mu}(\mathbf{x}, \boldsymbol{\theta}), \boldsymbol{\Sigma}(\mathbf{x}, \boldsymbol{\theta})). \quad (10)$$

Lastly, equipped with the manageable function q_θ , we make one more approximation to handle the likelihood function by introducing *evidence lower bound (ELBO)* as follows.

$$\begin{aligned} \log p_\phi(\mathbf{x}) &= \log \int p_\phi(\mathbf{x}|\mathbf{z})p(\mathbf{z})d\mathbf{z} \\ &= \log \int p_\phi(\mathbf{x}|\mathbf{z})p(\mathbf{z}) \frac{q_\theta(\mathbf{z}|\mathbf{x})}{q_\theta(\mathbf{z}|\mathbf{x})} d\mathbf{z} \\ &\geq \int q_\theta(\mathbf{z}|\mathbf{x}) \log \left(\frac{p_\phi(\mathbf{x}|\mathbf{z})p(\mathbf{z})}{q_\theta(\mathbf{z}|\mathbf{x})} \right) d\mathbf{z} \\ &= \int q_\theta(\mathbf{z}|\mathbf{x}) \log (p_\phi(\mathbf{x}|\mathbf{z})) d\mathbf{z} - D_{KL}(q_\theta(\mathbf{z}|\mathbf{x})||p(\mathbf{z})) \\ &=: ELBO(\mathbf{x}; \phi, \boldsymbol{\theta}). \end{aligned} \quad (11)$$

Here the inequality in the third line is given by Jensen's inequality. ELBO can be better estimated than $\log p_\phi(\mathbf{x})$ itself. First, the second Kullback-Leibler divergence term is instantly computable because both $q_\theta(\mathbf{z}|\mathbf{x})$ and $p(\mathbf{z})$ are normal density functions and Kullback-Leibler divergence between two normal distributions has a closed form. The first integral term still involves an integral, but now we can easily sample a point $\mathbf{z}$ from the normal distribution $q_\theta(\cdot|\mathbf{x}_i)$ for some randomly chosen $\mathbf{x}_i$ and estimate the integral by

$$\int q_\theta(\mathbf{z}|\mathbf{x}) \log (p_\phi(\mathbf{x}|\mathbf{z})) d\mathbf{z} \approx \log (p_\phi(\mathbf{x}_i|\mathbf{z})). \quad (12)$$

Two parameters θ and ϕ play different roles in maximizing log-likelihood. Given a fixed value of ϕ , maximizing ELBO with respect to θ tightens the lower bound. On the other hand, maximizing ELBO with respect to ϕ for the given θ maximizes the lower bound of log-likelihood. As a result, the estimated parameters are

$$(\hat{\phi}, \hat{\theta}) = \arg \max_{\phi, \theta} \left\{ \sum_{i=1}^n \int q_{\theta}(\mathbf{z}_i | \mathbf{x}_i) \log(p_{\phi}(\mathbf{x}_i | \mathbf{z}_i)) d\mathbf{z}_i - D_{KL}(q_{\theta}(\mathbf{z}_i | \mathbf{x}_i) \| p(\mathbf{z}_i)) \right\} \quad (13)$$

Once we have learned a VAE, the model generates a new sample point by drawing $\mathbf{z}$ from $N(0, \mathbf{I})$ and then drawing $\mathbf{x}$ from $p_{\phi}(\mathbf{x} | \mathbf{z})$.

The framework is named an autoencoder because it has the *encoder* $q_{\theta}(\mathbf{z} | \mathbf{x})$ and the *decoder* $p_{\phi}(\mathbf{x} | \mathbf{z})$. The model is variational because it approximates the encoding probability with a family of simple and known distributions.

3 Diffusion model

The main framework of a diffusion model is to construct a parameterized Markov chain and to train the chain through variational inference to produce samples matching the data distribution after finite time. Transitions of this chain are learned to reverse the diffusion process, which is a Markov chain that gradually adds noise to the data in the opposite direction of sampling until signal is destroyed. This is illustrated in Figure 2, where the latent variables obtained by adding noise to data sample $\mathbf{x}$ are denoted as $\mathbf{z}_1, \dots, \mathbf{z}_T$. The essential idea behind the diffusion model is to (i) **systematically and slowly destroy structure in a data distribution through an iterative forward diffusion process**, and then (ii) **learn a reverse diffusion process that restores structure in data** (Sohl-Dickstein et al. (2015)).

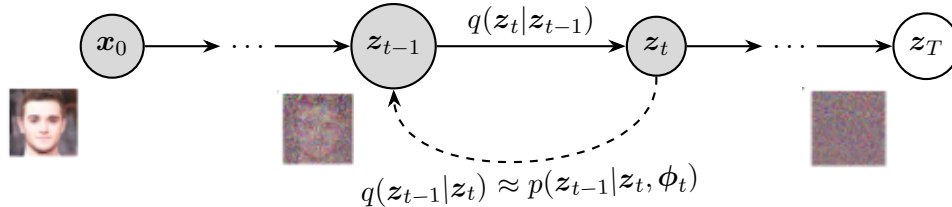


Figure 2: The directed graphical model illustrating diffusion model.

Technically, a diffusion model works as follows. Firstly, the diffusion model sequentially adds small Gaussian noises to the sample point $\mathbf{x}$ to obtain a sequence of noisy data $\mathbf{z}_1, \dots, \mathbf{z}_T$. We denote this process by $q(\mathbf{z}_t | \mathbf{z}_{t-1})$. After a large number of steps T , $\mathbf{z}_T$ becomes approximately a white noise. The goal of a diffusion model is to maximize the likelihood of the observed samples

$\mathbf{x}_1, \dots, \mathbf{x}_I$, assuming $\mathbf{z}_{iT}$ come from the standard normal distribution:

$$\begin{aligned}\hat{\phi}_{1,\dots,T} &:= \arg \max_{\phi_{1,\dots,T}} \sum_{i=1}^I q(\mathbf{x}_i) \\ &= \arg \max_{\phi_{1,\dots,T}} \sum_{i=1}^I \int q(\mathbf{x}_i | \mathbf{z}_{iT}) q(\mathbf{z}_{iT}) d\mathbf{z}_{iT}.\end{aligned}\tag{14}$$

Hypothetically, we can compute the density of the data from the reverse process $q(\mathbf{z}_{t-1} | \mathbf{z}_t)$ as follows:

$$\begin{aligned}q(\mathbf{x}) &= \int q(\mathbf{x} | \mathbf{z}_T) q(\mathbf{z}_T) d\mathbf{z}_T \\ &= \int q(\mathbf{x} | \mathbf{z}_1) q(\mathbf{z}_1 | \mathbf{z}_2) \cdots q(\mathbf{z}_{T-1} | \mathbf{z}_T) q(\mathbf{z}_T) d\mathbf{z}_1, \dots, \mathbf{z}_T.\end{aligned}\tag{15}$$

However, in practice, the exact form of $q(\mathbf{z}_{t-1} | \mathbf{z}_t)$ is neither known to us nor tractable. In addition, the high-dimensional integral is not straightforward to evaluate. A diffusion model overcomes these difficulties by first approximating each $q(\mathbf{z}_{t-1} | \mathbf{z}_t)$ by a normal density function

$$\begin{aligned}q(\mathbf{z}_{t-1} | \mathbf{z}_t) &\approx p(\mathbf{z}_{t-1} | \mathbf{z}_t, \phi_t) \\ &= \mathcal{N}_{\mathbf{z}_{t-1}} [\mathbf{f}_t(\mathbf{z}_t, \phi), \sigma_t^2 \mathbf{I}]\end{aligned}\tag{16}$$

where the mean $\mathbf{f}_t(\mathbf{z}_t, \phi)$ is trained as a neural network from the data and the variance σ_t^2 is pre-specified. The second challenge of handling the high-dimensional integral is resolved by substituting the integral by the Evidence Lower Bound (ELBO).

3.1 Encoder (Forward Pass, Model Setting)

A diffusion model is composed of an encoder and a decoder. Unlike the Variational Autoencoder (VAE) model, the encoder in a diffusion model is prespecified, representing the process of gradually adding noise to the data $\mathbf{x}$ and simulating its diffusion. To be specific, one can formulate the forward process as follows:

$$\begin{aligned}\mathbf{z}_1 &= \sqrt{1 - \beta_1} \mathbf{x} + \sqrt{\beta_1} \epsilon_1 \\ \mathbf{z}_t &= \sqrt{1 - \beta_t} \mathbf{z}_{t-1} + \sqrt{\beta_t} \epsilon_t, \quad \forall t \in 2, \dots, T,\end{aligned}\tag{17}$$

where ϵ_t follows a standard normal distribution. The hyperparameters $\beta_t \in [0, 1]$ determine the magnitude of the added noise. Here we note the following properties.

- We can derive the following distributions:

$$\begin{aligned}q(\mathbf{z}_1 | \mathbf{x}) &= \mathcal{N}(\sqrt{1 - \beta_1} \mathbf{x}, \beta_1 \mathbf{I}), \\ q(\mathbf{z}_t | \mathbf{z}_{t-1}) &= \mathcal{N}(\sqrt{1 - \beta_t} \mathbf{z}_{t-1}, \beta_t \mathbf{I}), \quad \forall t \in 2, \dots, T.\end{aligned}\tag{18}$$

This is a Markov chain.

- With enough steps, i.e. large T , all traces of the original data are removed, and the conditional distribution $q(\mathbf{z}_T|\mathbf{x})$ and marginal distribution $q(\mathbf{z}_T)$ both become the standard normal distribution. Indeed, the ultimate goal of this forward pass is to gradually convert the data distribution into a well behaved distribution by repeated application of Markov diffusion kernel.
- One can directly sample $\mathbf{z}_t$ given $\mathbf{x}$ without computing intermediate variables $\mathbf{z}_1, \dots, \mathbf{z}_{t-1}$. That is, by iteratively substituting variables, one can obtain

$$\begin{aligned}\mathbf{z}_t &= \sqrt{\alpha_t}\mathbf{x} + \sqrt{1 - \alpha_t}\boldsymbol{\epsilon}, \\ q(\mathbf{z}_t|\mathbf{x}) &= \mathcal{N}(\sqrt{\alpha_t}\mathbf{x}, (1 - \alpha_t)\mathbf{I}),\end{aligned}\tag{19}$$

where $\alpha_t = \prod_{s=1}^t (1 - \beta_s)$ and $\boldsymbol{\epsilon}$ is a sample from standard normal distribution. This $q(\mathbf{z}_t|\mathbf{x})$ is known as the diffusion kernel.

In summary, for any starting point $\mathbf{x}$, variable $\mathbf{z}_t$ is normally distributed with a known mean and variance.

3.2 Conditional Distributions

Before we proceed with learning the reverse process of the diffusion model, we examine here the true reverse probabilities that are of our interest. We first note that the marginal distribution of $\mathbf{z}_t$ can be computed using the diffusion kernel,

$$q(\mathbf{z}_t) = \int q(\mathbf{z}_t|\mathbf{x})p(\mathbf{x})d\mathbf{x}.\tag{20}$$

However, it is difficult to know $p(\mathbf{x})$ in practice, and consequently $q(\mathbf{z}_t)$ can not be obtained in a closed form. The conditional distribution of the reverse process can then be written as

$$q(\mathbf{z}_{t-1}|\mathbf{z}_t) = \frac{q(\mathbf{z}_t|\mathbf{z}_{t-1})q(\mathbf{z}_{t-1})}{q(\mathbf{z}_t)},\tag{21}$$

as a result of the Bayes' rule. Again, this is intractable since we cannot compute the marginal distributions in (20).

We then focus on the conditional diffusion distribution, defined as $q(\mathbf{z}_{t-1}|\mathbf{z}_t, \mathbf{x})$. As we do know $q(\mathbf{z}_{t-1}|\mathbf{x})$, which is just an diffusion kernel, we can utilize this to compute the conditional diffusion distribution in a closed form.

Lemma 1.

$$\begin{aligned}(a) \quad & \mathcal{N}(\mathbf{A}\mathbf{w}, \mathbf{B}) \propto \mathcal{N}\left((\mathbf{A}^\top \mathbf{B}^{-1} \mathbf{A})^{-1} \mathbf{A}^\top \mathbf{B}^{-1} \mathbf{v}, (\mathbf{A}^\top \mathbf{B}^{-1} \mathbf{A})^{-1}\right), \\ (b) \quad & \mathcal{N}(\mathbf{a}, \mathbf{A}) \cdot \mathcal{N}(\mathbf{b}, \mathbf{B}) \propto \mathcal{N}\left((\mathbf{A}^{-1} + \mathbf{B}^{-1})^{-1}(\mathbf{A}^{-1} \mathbf{a} + \mathbf{B}^{-1} \mathbf{b}), (\mathbf{A}^{-1} + \mathbf{B}^{-1})^{-1}\right)\end{aligned}\tag{22}$$

Theorem 2.

$$q(\mathbf{z}_{t-1}|\mathbf{z}_t, \mathbf{x}) = \mathcal{N}\left(\frac{1-\alpha_{t-1}}{1-\alpha_t}\sqrt{1-\beta_t}\mathbf{z}_t + \frac{\sqrt{\alpha_{t-1}}\beta_t}{1-\alpha_t}\mathbf{x}, \frac{\beta_t(1-\alpha_{t-1})}{1-\alpha_t}\mathbf{I}\right).$$

Proof.

$$\begin{aligned} q(\mathbf{z}_{t-1}|\mathbf{z}_t, \mathbf{x}) &= \frac{q(\mathbf{z}_t|\mathbf{z}_{t-1}, \mathbf{x})q(\mathbf{z}_{t-1}|\mathbf{x})}{q(\mathbf{z}_t|\mathbf{x})} \\ &\propto q(\mathbf{z}_t|\mathbf{z}_{t-1})q(\mathbf{z}_{t-1}|\mathbf{x}) \\ &= \mathcal{N}\left(\sqrt{1-\beta_t}\mathbf{z}_{t-1}, \beta_t\mathbf{I}\right)\mathcal{N}(\sqrt{\alpha_{t-1}}\mathbf{x}, (1-\alpha_{t-1})\mathbf{I}) \\ &\propto \mathcal{N}\left(\frac{1}{\sqrt{1-\beta_t}}\mathbf{z}_t, \frac{\beta_t}{1-\beta_t}\mathbf{I}\right)\mathcal{N}(\sqrt{\alpha_{t-1}}\mathbf{x}, (1-\alpha_{t-1})\mathbf{I}) \end{aligned} \tag{23}$$

Applying Lemma 1, we obtain the result. $\square$

We highlight here that this fact in Theorem 2 is then used to train the decoder.

3.3 Decoder (Backward pass, Model Learning)

When we train a diffusion model, it learns the reverse process, i.e. the backward map between $\mathbf{z}_t$ and $\mathbf{z}_{t-1}$. This network is trained to gradually remove noise, originating the term ‘denoising diffusion probabilistic model’. To generate a new data example $\mathbf{x}$, we draw a sample from $q(\mathbf{z}_T)$, which is simply a standard normal distribution, and then process it through the decoder models.

As seen in Section 3.2, it is difficult to know the true reverse distribution $q(\mathbf{z}_{t-1}|\mathbf{z}_t)$ of the diffusion process. We approximate these as normal distributions and set the following model:

$$\begin{aligned} p(\mathbf{z}_T) &= \mathcal{N}(\mathbf{0}, \mathbf{I}), \\ p(\mathbf{z}_{t-1}|\mathbf{z}_t, \phi_t) &= \mathcal{N}(\mathbf{f}_t(\mathbf{z}_t, \phi_t), \sigma_t^2\mathbf{I}) \\ p(\mathbf{x}|\mathbf{z}_1, \phi_1) &= \mathcal{N}(\mathbf{f}_1(\mathbf{z}_1, \phi_1), \sigma_1^2\mathbf{I}), \end{aligned} \tag{24}$$

where $\phi_t, t = 1, \dots, T$ are parameters to learn and $\mathbf{f}_t(\mathbf{z}_t, \phi_t)$ is a neural network model that predicts the mean of the normal distribution for the preceding latent variable $\mathbf{z}_{t-1}$ given $\mathbf{z}_t$. Here, $\{\sigma_t^2\}$ are predetermined parameters.

To train the model, we generate new examples from $p(\mathbf{x})$ as follows. First, sample $\mathbf{z}_T$ from $p(\mathbf{z}_T)$. Then, sample $\mathbf{z}_{T-1}$ from $p(\mathbf{z}_{T-1}|\mathbf{z}_T, \phi_T)$, and repeat this until we finally generate $\mathbf{x}$ from $p(\mathbf{x}|\mathbf{z}_1, \phi_1)$. Now we can approximate the probability of training dataset $\{\mathbf{x}_i\}$ as follows. First, the entire joint probability of $\mathbf{x}$ and latent variables $\mathbf{z}_{1,\dots,T}$ is

$$p(\mathbf{x}, \mathbf{z}_{1,\dots,T}|\phi_{1,\dots,T}) = p(\mathbf{x}|\mathbf{z}_1, \phi_1) \prod_{t=2}^T p(\mathbf{z}_{t-1}|\mathbf{z}_t, \phi_t) \cdot p(\mathbf{z}_T). \tag{25}$$

The marginal probability of $\mathbf{x}$ is obtained by integrating Equation (25) over all the latent variables:

$$p(\mathbf{x}|\phi_{1,\dots,T}) = \int p(\mathbf{x}, \mathbf{z}_{1,\dots,T}|\phi_{1,\dots,T}) d\mathbf{z}_{1,\dots,T}. \quad (26)$$

To train the model, we maximize the log-likelihood of the training data $\{\mathbf{x}_i\}$ with respect to the parameters ϕ :

$$\hat{\phi}_{1,\dots,T} = \arg \max_{\phi_{1,\dots,T}} \left[\sum_{i=1}^I \log [p(\mathbf{x}_i|\phi_{1,\dots,T})] \right]. \quad (27)$$

4 Evidence Lower Bound (ELBO)

The exact optimization problem in (27) is intractable because of the integral involved in Equation (26). We make a detour around this difficulty using *evidence lower bound (ELBO)*.

$$\begin{aligned} \log [p(\mathbf{x}|\phi_{1,\dots,T})] &= \log \left[\int p(\mathbf{x}, \mathbf{z}_1, \dots, \mathbf{z}_T|\phi_{1,\dots,T}) d\mathbf{z}_{1,\dots,T} \right] \\ &= \log \left[\int q(\mathbf{z}_{1,\dots,T}|\mathbf{x}) \frac{p(\mathbf{x}, \mathbf{z}_1, \dots, \mathbf{z}_T|\phi_{1,\dots,T})}{q(\mathbf{z}_{1,\dots,T}|\mathbf{x})} d\mathbf{z}_{1,\dots,T} \right] \\ &\geq \int q(\mathbf{z}_{1,\dots,T}|\mathbf{x}) \log \left[\frac{p(\mathbf{x}, \mathbf{z}_1, \dots, \mathbf{z}_T|\phi_{1,\dots,T})}{q(\mathbf{z}_{1,\dots,T}|\mathbf{x})} \right] d\mathbf{z}_{1,\dots,T} \\ &=: \text{ELBO}[\phi_{1,\dots,T}] \end{aligned} \quad (28)$$

Next, we simplify the whole joint probabilities of $\mathbf{x}$ and $\mathbf{z}_{1,\dots,T}$ in Equation (28) in terms of the conditional probabilities between $\mathbf{x}$, ϕ , and consecutive latent variables $\mathbf{z}_t$ and $\mathbf{z}_{t-1}$.

$$\begin{aligned} &\log \left[\frac{p(\mathbf{x}, \mathbf{z}_{1,\dots,T}|\phi_{1,\dots,T})}{q(\mathbf{z}_{1,\dots,T}|\mathbf{x})} \right] \\ &= \log \left[\frac{p(\mathbf{x}|\mathbf{z}_1, \phi_1) \prod_{t=2}^T p(\mathbf{z}_{t-1}|\mathbf{z}_t, \phi_t) \cdot p(\mathbf{z}_T)}{q(\mathbf{z}_1|\mathbf{x}) \prod_{t=2}^T q(\mathbf{z}_t|\mathbf{z}_{t-1})} \right] \\ &= \log \left[\frac{p(\mathbf{x}|\mathbf{z}_1, \phi_1)}{q(\mathbf{z}_1|\mathbf{x})} \right] + \log \left[\frac{\prod_{t=2}^T p(\mathbf{z}_{t-1}|\mathbf{z}_t, \phi_t)}{\prod_{t=2}^T q(\mathbf{z}_t|\mathbf{z}_{t-1})} \right] + \log [p(\mathbf{z}_T)]. \end{aligned} \quad (29)$$

In the middle term, the numerator $p(\mathbf{z}_{t-1}|\mathbf{z}_t, \phi_t)$ is given as the normal density with mean $\mathbf{f}[\mathbf{z}_t, \phi_t]$. To make the denominator $q(\mathbf{z}_t|\mathbf{z}_{t-1})$ more accessible, we rewrite the term as

$$q(\mathbf{z}_t|\mathbf{z}_{t-1}) = q(\mathbf{z}_t|\mathbf{z}_{t-1}, \mathbf{x}) = \frac{q(\mathbf{z}_{t-1}|\mathbf{z}_t, \mathbf{x})q(\mathbf{z}_t|\mathbf{x})}{q(\mathbf{z}_{t-1}|\mathbf{x})}.$$

As a result,

$$\begin{aligned}
& \log \left[\frac{p(\mathbf{x}, \mathbf{z}_{1,\dots,T} | \phi_{1,\dots,T})}{q(\mathbf{z}_{1,\dots,T} | \mathbf{x})} \right] \\
&= \log \left[\frac{p(\mathbf{x} | \mathbf{z}_1, \phi_1)}{q(\mathbf{z}_1 | \mathbf{x})} \right] + \log \left[\prod_{t=2}^T \frac{p(\mathbf{z}_{t-1} | \mathbf{z}_t, \phi_t)}{q(\mathbf{z}_{t-1} | \mathbf{z}_t, \mathbf{x})} \cdot \frac{q(\mathbf{z}_{t-1} | \mathbf{x})}{q(\mathbf{z}_t | \mathbf{x})} \right] + \log [p(\mathbf{z}_T)] \\
&= \log [p(\mathbf{x} | \mathbf{z}_1, \phi_1)] + \log \left[\prod_{t=2}^T \frac{p(\mathbf{z}_{t-1} | \mathbf{z}_t, \phi_t)}{q(\mathbf{z}_{t-1} | \mathbf{z}_t, \mathbf{x})} \right] + \log \left[\frac{p(\mathbf{z}_T)}{q(\mathbf{z}_T | \mathbf{x})} \right] \\
&= \log [p(\mathbf{x} | \mathbf{z}_1, \phi_1)] + \sum_{t=2}^T \log \left[\frac{p(\mathbf{z}_{t-1} | \mathbf{z}_t, \phi_t)}{q(\mathbf{z}_{t-1} | \mathbf{z}_t, \mathbf{x})} \right] + \log \left[\frac{p(\mathbf{z}_T)}{q(\mathbf{z}_T | \mathbf{x})} \right] \\
&\approx \log [p(\mathbf{x} | \mathbf{z}_1, \phi_1)] + \sum_{t=2}^T \log \left[\frac{p(\mathbf{z}_{t-1} | \mathbf{z}_t, \phi_t)}{q(\mathbf{z}_{t-1} | \mathbf{z}_t, \mathbf{x})} \right]
\end{aligned} \tag{30}$$

Here the last term $\log \left[\frac{p(\mathbf{z}_T)}{q(\mathbf{z}_T | \mathbf{x})} \right]$ is approximated to zero, because for sufficiently large T , both $p(\mathbf{z}_T)$ and $q(\mathbf{z}_T | \mathbf{x})$ are approximately standard normal densities. We obtain the following simplified version of ELBO by substituting (30) into the definition of ELBO:

$$\begin{aligned}
& \text{ELBO}[\phi_{1,\dots,T}] \\
&\approx \int q(\mathbf{z}_{1,\dots,T} | \mathbf{x}) \left(\log [p(\mathbf{x} | \mathbf{z}_1, \phi_1)] + \sum_{t=2}^T \log \left[\frac{p(\mathbf{z}_{t-1} | \mathbf{z}_t, \phi_t)}{q(\mathbf{z}_{t-1} | \mathbf{z}_t, \mathbf{x})} \right] \right) d\mathbf{z}_{1,\dots,T} \\
&= \int q(\mathbf{z}_1 | \mathbf{x}) \log [p(\mathbf{x} | \mathbf{z}_1, \phi_1)] d\mathbf{z}_1 \\
&\quad - \sum_{t=2}^T \int q(\mathbf{z}_t | \mathbf{x}) q(\mathbf{z}_{t-1} | \mathbf{z}_t, \mathbf{x}) \log \left[\frac{q(\mathbf{z}_{t-1} | \mathbf{z}_t, \mathbf{x})}{p(\mathbf{z}_{t-1} | \mathbf{z}_t, \phi_t)} \right] d\mathbf{z}_{t-1} d\mathbf{z}_t
\end{aligned} \tag{31}$$

The first integral term will be replaced by the empirical expectation of $\log [p(\mathbf{x} | \mathbf{z}_1, \phi_1)]$ with a single sample from the distribution $q(\mathbf{z}_1 | \mathbf{x})$, that is,

$$\int q(\mathbf{z}_1 | \mathbf{x}) \log [p(\mathbf{x} | \mathbf{z}_1, \phi_1)] d\mathbf{z}_1 \approx \log [p(\mathbf{x} | \mathbf{z}_1^*, \phi_1)] \tag{32}$$

where $\mathbf{z}_1^* \sim q(\mathbf{z}_1 | \mathbf{x})$. The second integral term can be further simplified as follows.

$$\begin{aligned}
& \int q(\mathbf{z}_t | \mathbf{x}) q(\mathbf{z}_{t-1} | \mathbf{z}_t, \mathbf{x}) \log \left[\frac{q(\mathbf{z}_{t-1} | \mathbf{z}_t, \mathbf{x})}{p(\mathbf{z}_{t-1} | \mathbf{z}_t, \phi_t)} \right] d\mathbf{z}_{t-1} d\mathbf{z}_t \\
&= \int q(\mathbf{z}_t | \mathbf{x}) D_{KL} [q(\mathbf{z}_{t-1} | \mathbf{z}_t, \mathbf{x}) \| p(\mathbf{z}_{t-1} | \mathbf{z}_t, \phi_t)] d\mathbf{z}_t \\
&= \mathbb{E}_{q(\mathbf{z}_t | \mathbf{x})} [D_{KL} [q(\mathbf{z}_{t-1} | \mathbf{z}_t, \mathbf{x}) \| p(\mathbf{z}_{t-1} | \mathbf{z}_t, \phi_t)]]
\end{aligned} \tag{33}$$

Recall that $q(\mathbf{z}_{t-1} | \mathbf{z}_t, \mathbf{x})$ and $p(\mathbf{z}_{t-1} | \mathbf{z}_t, \phi_t)$ are normal probability densities. Kullback-Leibler

divergence between two normal distributions can be calculated in a Rao-Blackwellized fashion with closed form expressions instead of high variance Monte Carlo estimates (Ho et al. (2020)).

Lemma 3. *Suppose $\mathbf{p}(\mathbf{x}) = \mathcal{N}_{\mathbf{x}}(\mathbf{a}, \mathbf{A})$ and $\mathbf{q}(\mathbf{x}) = \mathcal{N}_{\mathbf{x}}(\mathbf{b}, \mathbf{B})$ are d -dimensional normal distributions. Then the KL-divergence between $\mathbf{p}$ and $\mathbf{q}$ is given as follows.*

$$D_{KL}[\mathbf{p}||\mathbf{q}] = \frac{1}{2} \left(\text{tr}[\mathbf{B}^{-1}\mathbf{A}] - d + (\mathbf{a} - \mathbf{b})^\top \mathbf{B}^{-1}(\mathbf{a} - \mathbf{b}) + \log \left[\frac{|\mathbf{B}|}{|\mathbf{A}|} \right] \right). \quad (34)$$

If $\mathbf{A} = \sigma_a \mathbf{I}_d$ and $\mathbf{B} = \sigma_b \mathbf{I}_d$, KL-divergence is further simplified as

$$D_{KL}[\mathbf{p}||\mathbf{q}] = \frac{1}{2} \left(\frac{1}{\sigma_b} \|\mathbf{a} - \mathbf{b}\|^2 + d \left(\frac{\sigma_a}{\sigma_b} - 1 + \log \left(\frac{\sigma_a}{\sigma_b} \right) \right) \right).$$

Applying Lemma 3 to our problem yields:

$$\begin{aligned} L_{t-1} &:= D_{KL}[q(\mathbf{z}_{t-1}|\mathbf{z}_t, \mathbf{x})||p(\mathbf{z}_{t-1}|\mathbf{z}_t, \phi_t)] \\ &= \frac{1}{2\sigma_t^2} \left\| \frac{1 - \alpha_{t-1}}{1 - \alpha_t} \sqrt{1 - \beta_t} \mathbf{z}_t + \frac{\sqrt{\alpha_{t-1}}\beta_t}{1 - \alpha_t} \mathbf{x} - \mathbf{f}_t[\mathbf{z}_t, \phi_t] \right\|^2 + C \end{aligned} \quad (35)$$

for some constant C that does not depend on $\phi_{1,\dots,T}$. Finally, we again replace the expectation with respect to $\mathbf{z}_{1,\dots,T}$ in Equation (31) by the empirical expectation, by drawing $\mathbf{z}_{it}^*$ from $q(\mathbf{z}_{it}^*|\mathbf{x}_i)$, $i = 1, \dots, I, t = 1, \dots, T$.

$$\begin{aligned} L[\phi_{1,\dots,T}] &\approx \sum_{i=1}^I \left(-\log[\mathcal{N}_{\mathbf{x}_i}[\mathbf{f}_1[\mathbf{z}_{i1}^*, \phi_1], \sigma_1^2 \mathbf{I}]] \right. \\ &\quad \left. + \sum_{t=2}^T \frac{1}{2\sigma_t^2} \left\| \frac{1 - \alpha_{t-1}}{1 - \alpha_t} \sqrt{1 - \beta_t} \mathbf{z}_{it}^* + \frac{\sqrt{\alpha_{t-1}}\beta_t}{1 - \alpha_t} \mathbf{x}_i - \mathbf{f}_t[\mathbf{z}_{it}^*, \phi_t] \right\|^2 \right). \end{aligned} \quad (36)$$

5 Reparameterization and the resulting algorithms

The loss function in (36) can be decomposed in to the following parts:

$$\begin{aligned} &\text{Reconstruction term : } \log[\mathcal{N}_{\mathbf{x}_i}[\mathbf{f}_1[\mathbf{z}_{i1}, \phi_1], \sigma_1^2 \mathbf{I}]] \\ &(\text{Target}) \text{ mean of } q(\mathbf{z}_{t-1}|\mathbf{z}_t, \mathbf{x}) : \frac{1 - \alpha_{t-1}}{1 - \alpha_t} \sqrt{1 - \beta_t} \mathbf{z}_{it} + \frac{\sqrt{\alpha_{t-1}}\beta_t}{1 - \alpha_t} \mathbf{x}_i \\ &\text{Predicted } \mathbf{z}_{t-1} : \mathbf{f}_t[\mathbf{z}_{it}, \phi_t]. \end{aligned} \quad (37)$$

Here, each $\mathbf{x}_i$ is the i th data point and $\mathbf{z}_{it}$ is the associated latent variable at diffusion step t . This loss function can be used to train a network **for each diffusion time step**. It minimizes the difference between the estimated $\mathbf{f}_t[\mathbf{z}_{it}, \phi_t]$ of the hidden variable at the previous time step and the most likely value that it took given the ground truth de-noised data $\mathbf{x}$. However, the optimization

can be further simplified by reparameterizing the target and the network so that the model aims to predict the noise, instead of the hidden variable.

5.1 Reparameterization

The original diffusion update was given by (19). One can rewrite it in terms of the data $\mathbf{x}$ such that

$$\mathbf{x} = \frac{1}{\sqrt{\alpha_t}} \mathbf{z}_t + \frac{\sqrt{1-\alpha_t}}{\sqrt{\alpha_t}} \boldsymbol{\epsilon}. \quad (38)$$

Substituting this into the target terms in (37) yields

$$\begin{aligned} & \frac{1-\alpha_{t-1}}{1-\alpha_t} \sqrt{1-\beta_t} \mathbf{z}_t + \frac{\sqrt{\alpha_{t-1}} \beta_t}{1-\alpha_t} \mathbf{x} \\ &= \frac{1-\alpha_{t-1}}{1-\alpha_t} \sqrt{1-\beta_t} \mathbf{z}_t + \frac{\sqrt{\alpha_{t-1}} \beta_t}{1-\alpha_t} \left(\frac{1}{\sqrt{\alpha_t}} \mathbf{z}_t + \frac{\sqrt{1-\alpha_t}}{\sqrt{\alpha_t}} \boldsymbol{\epsilon} \right) \\ &= \left(\frac{(1-\alpha_{t-1})\sqrt{1-\beta_t}}{1-\alpha_t} + \frac{\beta_t}{(1-\alpha_t)\sqrt{1-\beta_t}} \right) \mathbf{z}_t - \frac{\beta_t}{\sqrt{1-\alpha_t}\sqrt{1-\beta_t}} \boldsymbol{\epsilon} \\ &= \frac{1}{\sqrt{1-\beta_t}} \mathbf{z}_t - \frac{\beta_t}{\sqrt{1-\alpha_t}\sqrt{1-\beta_t}} \boldsymbol{\epsilon}, \end{aligned} \quad (39)$$

where we have used the fact that $\frac{\sqrt{\alpha_t}}{\sqrt{\alpha_{t-1}}} = \sqrt{1-\beta_t}$.

Then, the network should also be reparameterized. We replace the model $\widehat{\mathbf{z}}_{t-1} = \mathbf{f}[\mathbf{z}_t, \boldsymbol{\phi}_t]$ with a new model $\boldsymbol{\epsilon} = \mathbf{g}_t[\mathbf{z}_t, \boldsymbol{\phi}_t]$, which predicts the noise $\boldsymbol{\epsilon}$ that was mixed with $\mathbf{x}$ to create $\mathbf{z}_t$:

$$\mathbf{f}[\mathbf{z}_t, \boldsymbol{\phi}_t] = \frac{1}{\sqrt{1-\beta_t}} \mathbf{z}_t - \frac{\beta_t}{\sqrt{1-\alpha_t}\sqrt{1-\beta_t}} \mathbf{g}_t[\mathbf{z}_t, \boldsymbol{\phi}_t]. \quad (40)$$

Combining the reparameterization of the target and the network, the loss function given in (36) can be rewritten as

$$\begin{aligned} & L[\boldsymbol{\phi}_{1,\dots,T}] \\ &= \sum_{i=1}^I \left(-\log [\mathcal{N}_{\mathbf{x}_i} [\mathbf{f}_1[\mathbf{z}_{i1}, \boldsymbol{\phi}_1], \sigma_1^2 \mathbf{I}]] \right) + \sum_{t=2}^T \frac{\beta_t^2}{(1-\alpha_t)(1-\beta_t)2\sigma_t^2} \|\mathbf{g}_t[\mathbf{z}_{it}, \boldsymbol{\phi}_t] - \boldsymbol{\epsilon}_{it}\|^2 \\ &= \sum_{i=1}^I \left(\frac{1}{2\sigma_1^2} \|\mathbf{x}_i - \mathbf{f}_1[\mathbf{z}_{i1}, \boldsymbol{\phi}_1]\|^2 + C_i \right) + \sum_{t=2}^T \frac{\beta_t^2}{(1-\alpha_t)(1-\beta_t)2\sigma_t^2} \|\mathbf{g}_t[\mathbf{z}_{it}, \boldsymbol{\phi}_t] - \boldsymbol{\epsilon}_{it}\|^2 \\ &= \sum_{i=1}^I \sum_{t=1}^T \frac{\beta_t^2}{(1-\alpha_t)(1-\beta_t)2\sigma_t^2} \|\mathbf{g}_t[\mathbf{z}_{it}, \boldsymbol{\phi}_t] - \boldsymbol{\epsilon}_{it}\|^2 + C_i, \end{aligned} \quad (41)$$

where we substituted in (38) and (40) for $t = 1$.

5.2 Algorithms

In practice, to solve the optimization problem, we can ignore the scaling factor and the constants. This leads to straightforward algorithms for both training the model and sampling as in Algorithms 1 and 2. These algorithms are simple to implement and naturally augments the dataset, as we can reuse every original data point $\mathbf{x}_i$ as many times as we want at each time step with different noise ϵ (Prince (2023)).

Algorithm 1 Training

```

1: repeat
2:    $\mathbf{x} \sim q(\mathbf{x})$ 
3:    $t \sim \text{Uniform}(\{1, \dots, T\})$  ▷ sample random timestep
4:    $\epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ 
5:   Take gradient descent step on
      $\nabla_{\theta} \|\epsilon - \mathbf{g}_t(\sqrt{\alpha_t}\mathbf{x} + \sqrt{1 - \alpha_t}\epsilon, \phi_t)\|^2$ 
6: until converged

```

Algorithm 2 Sampling

```

1:  $\mathbf{z}_T \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ 
2: for  $t = T, \dots, 2$  do
3:    $\epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$  if  $t > 1$ , else  $\epsilon = \mathbf{0}$ 
4:    $\hat{\mathbf{z}}_{t-1} = \frac{1}{\sqrt{1-\beta_t}} \left( \mathbf{z}_t - \frac{\beta_t}{\sqrt{1-\alpha_t}} \mathbf{g}_t[\mathbf{z}_t, \phi_t] \right)$ 
5:    $\mathbf{z}_{t-1} = \hat{\mathbf{z}}_{t-1} + \sigma_t \epsilon$  ▷ add noise to the next input
6: end for
7: return  $\mathbf{x} = \frac{1}{\sqrt{1-\beta_1}} \left( \mathbf{z}_1 - \frac{\beta_1}{\sqrt{1-\alpha_1}} \mathbf{g}_1[\mathbf{z}_1, \phi_1] \right)$ 

```

5.3 Special note on the connection to the Langevin dynamics

What information about the data distribution does the model learn through the denoising process? An interesting note can be made to say that “Diffusion models approximate the solution to stochastic differential equations.” To understand it, we note the following concepts.

- The score function is defined as the gradient vector of the log of the probability density function,

$$\nabla_{\mathbf{x}} \log p(\mathbf{x}). \tag{42}$$

The score matching (SM) technique aims to train the model to learn the Score function.

- Denoising autoencoder model (DAE) aims to learn the model $\hat{X} = f_{\theta}(Z)$, where Z is a noise random variable. This represents a one-step denoising process. In contrast, diffusion models employ a more gradual approach, iteratively adding noise and learning the reverse procedure.
- Interestingly, Vincent (2011) has shown that the denoising autoencoder model (DAE) shares the same objective function with the denoising score matching (DSM) technique, as well as

that of the SM technique.

- DSM aims to learn the Score function of the conditional probability distribution $q(\mathbf{z}|\mathbf{x})$. If the noise follows $\mathcal{N}(0, \sigma^2)$, then the score function of the conditional probability distribution is

$$s(\mathbf{z}|\mathbf{x}) = \frac{\mathbf{x} - \mathbf{z}}{\sigma^2}. \quad (43)$$

- When a machine learning model is optimized to learn the score function of the above conditional probability distribution, it becomes capable of estimating the relative magnitude of injected noise in comparison to the variance when it receives noise-added data as input. Consequently, the model learns the trajectory from the corrupted information Z towards the original data X , effectively learning the path of data restoration.
- DAE sharing the same objective function means that through the denoising process, the DAE model ultimately learns the score function of the input data distribution.
- Langevin dynamics can produce samples from a probability density $p(\mathbf{x})$ using the score function $\nabla_{\mathbf{x}} \log p(\mathbf{x})$. That is, given a fixed step size η , and an initial value $\mathbf{x} \sim \pi(\mathbf{x})$ with π being a prior distribution, the Langevin method recursively computes the following,

$$\mathbf{x}_t \leftarrow \mathbf{x}_{t-1} + \frac{\eta}{2} \nabla_{\mathbf{x}} \log p(\mathbf{x}_t) + \sqrt{\eta} \boldsymbol{\epsilon}_{t-1}, \quad (44)$$

where $\boldsymbol{\epsilon}_{t-1} \sim \mathcal{N}(0, I)$. Then, the distribution of $\mathbf{x}_T$ equals $p(\mathbf{x})$ when $\eta \rightarrow 0$ and $T \rightarrow \infty$.

- Let us revisit our sampling procedure in Algorithm 2. To sample $\mathbf{z}_{t-1} \sim p(\mathbf{z}_{t-1}|\mathbf{z}_t, \boldsymbol{\phi}_t)$, we compute

$$\mathbf{z}_{t-1} = \frac{1}{\sqrt{1 - \beta_t}} \left(\mathbf{z}_t - \frac{\beta_t}{\sqrt{1 - \alpha_t}} \mathbf{g}_t[\mathbf{z}_t, \boldsymbol{\phi}_t] \right) + \sigma_t \boldsymbol{\epsilon}. \quad (45)$$

As we have the fact that the denoising model ultimately learns the score function of the input data distribution, it can be viewed as the reverse Langevin procedure of (44).

- The squared norm objective loss given in (41) resembles denoising score matching (DSM) over multiple noise scales indexed by t . Also, the loss (41) can be viewed as the variational bound for the process defined in (40), which is a Langevin-like process. By combining these two observations, optimizing an objective resembling DSM is equivalent to using variational inference to fit the finite-time marginal of a sampling chain resembling Langevin dynamics.

In this part, we briefly introduced how this $\boldsymbol{\epsilon}$ -prediction parameterization resembles Langevin dynamics and simplifies the diffusion model’s variational bound to an objective that resembles denoising score matching (Ho et al. (2020)). Moreover, we note that the sampling algorithm of the Diffusion Model approximates the Langevin Monte Carlo (LMC) algorithm, which is a numerical solution to the stochastic differential equation known as the Langevin equation!

6 Conclusion

Diffusion models represent a significant breakthrough in the generative modeling domain, offering a novel approach to data generation that balances quality and diversity. Their method of mapping data through latent variables and the reverse denoising process demonstrate the advanced capabilities of these models. This report has followed a structured exploration, beginning with the initial introduction of denoising diffusion models by Sohl-Dickstein et al. (2015). The loss function could be simplified based on the evidence lower bound (ELBO), which leads to a least-squares formulation. Additionally, this report has elucidated the link to the score-matching method as introduced in Song and Ermon (2019) and its resemblance to Langevin dynamics (Ho et al. (2020)), further highlighting the nature of diffusion models.

As this field continues to mature, the potential applications of diffusion models in various domains like image and audio generation, and even in more complex data forms, are expansive. The ongoing development and refinement of these models will undoubtedly be a pivotal area of research, contributing substantially to the advancements in machine learning and artificial intelligence.

References

- Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A. and Bengio, Y. (2014), ‘Generative adversarial nets’, *Advances in neural information processing systems* **27**.
- Ho, J., Jain, A. and Abbeel, P. (2020), ‘Denoising diffusion probabilistic models’, *Advances in neural information processing systems* **33**, 6840–6851.
- Kingma, D. P. and Welling, M. (2014), Auto-encoding variational bayes, in ‘Proceedings of the International Conference on Learning Representations (ICLR)’.
- Prince, S. J. (2023), *Understanding Deep Learning.*, MIT PRESS.
- Rezende, D. and Mohamed, S. (2015), Variational inference with normalizing flows, in ‘International conference on machine learning’, PMLR, pp. 1530–1538.
- Sohl-Dickstein, J., Weiss, E., Maheswaranathan, N. and Ganguli, S. (2015), Deep unsupervised learning using nonequilibrium thermodynamics, in ‘International conference on machine learning’, PMLR, pp. 2256–2265.
- Song, Y. and Ermon, S. (2019), *Generative Modeling by Estimating Gradients of the Data Distribution*, Curran Associates Inc., Red Hook, NY, USA.
- Vincent, P. (2011), ‘A connection between score matching and denoising autoencoders’, *Neural Computation* **23**(7), 1661–1674.