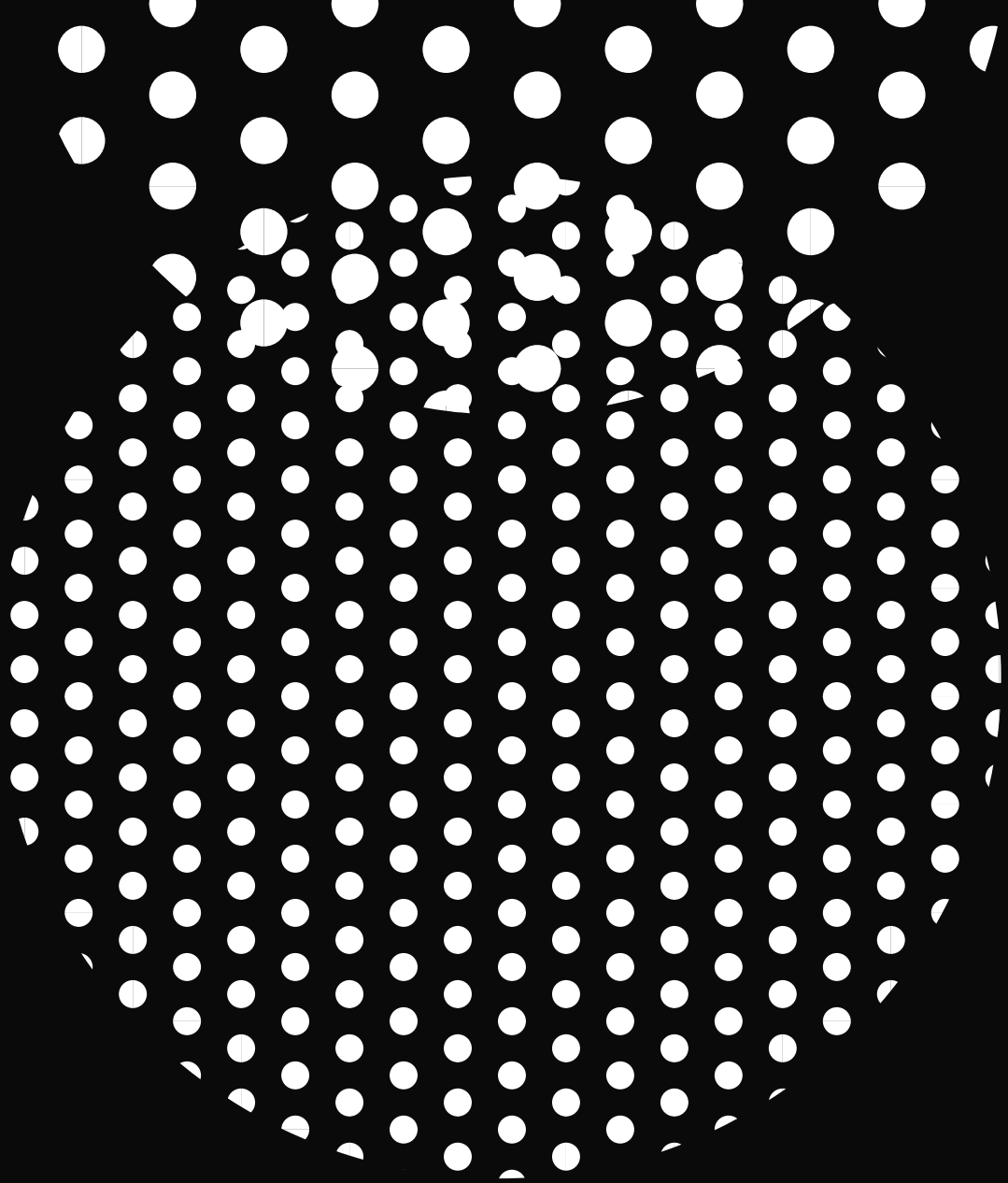




SDE $\rightarrow$ DIFFUSION MODEL

$P(x) \rightarrow$ TARGET

$P_\theta(x)$ Our approximation of $P(x)$

$$P_\theta(x) = \frac{e^{-f_\theta(x)}}{Z_\theta} \Rightarrow \text{good way to parameterize}$$

because $-f_\theta(x)$ force everything to be non negative. Z_θ normalizer such that $\int P_\theta(x) = 1$

Example:

$$N(x/\mu, \sigma) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$

$f_\theta(x)$ = how dense the probability is around x .

Calculating Z_θ requires us to integrate over the entire space of $x \Rightarrow$ TOO EXPENSIVE
so we need the score

$$\nabla_x \log P_\theta(x) = \nabla_x \log \frac{e^{-f_\theta(x)}}{Z_\theta} =$$

$$= \nabla_x \log e^{-f_\theta(x)} - \cancel{\nabla_x \log z_\theta} = -\nabla_x f_\theta(x)$$

$$= s_\theta(x)$$

SCORE MATCHING:

$$\frac{1}{2} \mathbb{E}_{p(x)} [\| \nabla_x \log p(x) - s_\theta(x) \|_2^2]$$

Score function = where to move to increase the probability of a data point

But WHO IS $\nabla_x \log(p(x))$?

$$\frac{1}{2} \mathbb{E}_{p(x)} [\| \nabla_x \log(p(x)) - s_\theta(x) \|_2^2] =$$

$$= \frac{1}{2} \int p(x) (\nabla_x \log(p(x)) - s_\theta(x))^2 dx =$$

$$= \frac{1}{2} \int p(x) \nabla_x^2 \log p(x) + p(x) s_\theta^2(x) - 2 \nabla_x \log p(x) \cdot s_\theta(x) dx$$

$$= \frac{1}{2} \int p(x) \nabla_x^2 \log p(x) dx + \frac{1}{2} \int p(x) s_\theta^2(x) dx - \int p(x) \nabla_x \log p(x) s_\theta(x) dx$$

We know that: $\nabla_x \log p(x) = \frac{\nabla_x p(x)}{p(x)}$

So we have that:

$$\int p_x \nabla_x \log p(x) s_\theta(x) dx = \int \cancel{p_x} \frac{\nabla_x p(x)}{\cancel{p(x)}} s_\theta(x) dx =$$

$$= \int \nabla_x p(x) s_\theta(x) dx = \text{Integration by parts}$$

$$\int_a^b dv u = uv \Big|_a^b - \int v du \quad \text{So we have:}$$

$$= \cancel{p(x) s_\theta(x)} \Big|_{-\infty}^{+\infty} - \int p_x \nabla_x s_\theta(x) dx =$$

$\downarrow$
 because $x \rightarrow +\infty \quad p(x) \rightarrow 0$

$\Rightarrow$ Going back we have that:

$$\frac{1}{2} \mathbb{E}_{p(x)} [\| \nabla_x \log p(x) - s_\theta(x) \|_2^2] =$$

$$= \frac{1}{2} \int \cancel{p(x) \nabla_x^2 \log(p(x))} dx + \frac{1}{2} \int p(x) s_\theta^2(x) dx -$$

$$+ \int p(x) \nabla_x s_\theta(x) dx$$

$\cancel{+}$ = there is no θ so it is a constant and for optimization problem it is 0

$$= \frac{1}{2} p(x) s_\theta^2(x) dx + \int p(x) \nabla_x s_\theta(x) dx$$

and $\mathbb{E}_{p(x)} = \int p(x) dx$

$$= \frac{1}{2} \mathbb{E}_{p(x)} [S_0^2(x)] + \mathbb{E}_{p(x)} [\nabla_x S_0(x)]$$

We see that even if we don't have $p(x)$

$$\frac{1}{2} \mathbb{E}_{p(x)} [\|\nabla_x \log p(x) - S_0(x)\|_2^2] =$$

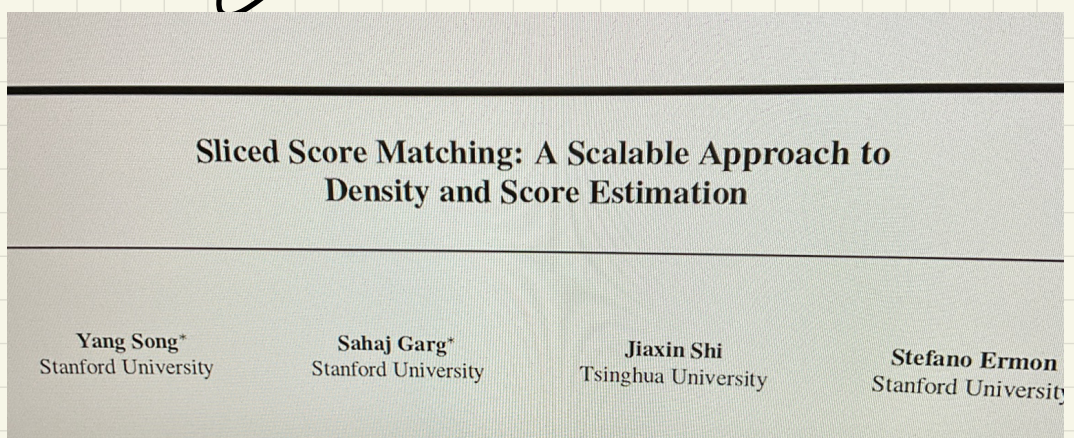
It doesn't $\leftarrow \frac{1}{2} \mathbb{E}_{p(x)} [S_0^2(x)] + \mathbb{E}_{p(x)} [\nabla_x S_0(x)] + C$
depend on $p(x)$

We want to minimize it so it means we want $\nabla_x S_0(x) \rightarrow 0 \rightarrow$ local maximum

while the fitter term $\mathbb{E}_{p(x)} [S_0^2(x)]$ is 0 at data points

PROBLEM: $\nabla_x S_0(x)$ Very expensive

It involves computing the Jacobian especially for images. This has been resolved in:



2nd Problem: Data space coverage $\Rightarrow$

$\Rightarrow$ model is good for data in the data space near the crowded ones but not in others.

To solve this problem we **NOISE** the images (Adding Gaussian noise)

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

$p(x) = p_\sigma(\tilde{x})$ & we solve problem 2.

How much noise depends on σ so:

$$\frac{1}{2} \mathbb{E}_{p_\sigma(x)} [\|\nabla_{\tilde{x}} \log p_\sigma(\tilde{x}) - \nabla_{\tilde{x}} s_\theta(\tilde{x})\|_2^2]$$

If σ is very large the space is too different

If σ is very small we have 2nd problem

$$\frac{1}{2} \mathbb{E}_{p_\sigma(\tilde{x})} [s_\theta^2(\tilde{x})] + \mathbb{E}_{p_\sigma(x)} [\nabla_x s_\theta(x)]$$

How do we compute $\nabla_x s_\theta(\tilde{x})$?

Fischer Vincent 2010 saw a connection

Score Matching $\Leftrightarrow$ Denoising Autoencoders

Train a regular AE but add noise

In the bottleneck $\Rightarrow$ being able to

separate noise from data it was possible

that the model could learn important features. So we start from

$$\frac{1}{2} \mathbb{E}_{p_{\sigma}(x)} [\| \nabla_{\tilde{x}} \log p_{\sigma}(x) - s_{\theta}(x) \|^2_2]$$

and our goal is:

$$\frac{1}{2} \mathbb{E}_{x \sim p(x), \tilde{x} \sim p_{\sigma}(\tilde{x}|x)} [\| s_{\theta}(\tilde{x}) - \nabla_{\tilde{x}} \log p_{\sigma}(\tilde{x}|x) \|^2_2]$$

So for

$$\frac{1}{2} \mathbb{E}_{p_{\sigma}(x)} [\| \nabla_{\tilde{x}} \log p_{\sigma}(x) - s_{\theta}(x) \|^2_2] =$$

$$= \frac{1}{2} \int p_{\sigma}(\tilde{x}) \left| \nabla_{\tilde{x}} \log p_{\sigma}(\tilde{x}) - s_{\theta}(\tilde{x}) \right|^2 d\tilde{x},$$

$$= \frac{1}{2} \int p_{\sigma}(\tilde{x}) \nabla_{\tilde{x}}^2 \log p_{\sigma}(\tilde{x}) d\tilde{x} +$$

$$+ \frac{1}{2} \int p_{\sigma}(\tilde{x}) s_{\theta}^2(\tilde{x}) d\tilde{x} +$$

$$= \int p_{\theta}(\tilde{x}) \frac{\nabla_{\tilde{x}} \log p_{\theta}(\tilde{x})}{\frac{\nabla_{\tilde{x}} p_{\theta}(\tilde{x})}{p_{\theta}(\tilde{x})}} S_{\theta}(\tilde{x}) d\tilde{x}$$

$$\downarrow$$

$$\int \cancel{p_{\theta}(\tilde{x})} \frac{\nabla_{\tilde{x}} p_{\theta}(\tilde{x})}{\cancel{p_{\theta}(\tilde{x})}} S_{\theta}(\tilde{x}) d\tilde{x}$$

Use marginalization and so:

$$p_{\theta}(\tilde{x}) = \int p(x) p_{\theta}(\tilde{x}|x) dx$$

$$= \int \nabla_{\tilde{x}} \left(\int p(x) p_{\theta}(\tilde{x}|x) dx \right) S_{\theta}(\tilde{x}) d\tilde{x}$$

$\downarrow$

Leibniz integration rule:

$$= \int \left(\int p(x) \nabla_{\tilde{x}} p_{\theta}(\tilde{x}|x) dx \right) S_{\theta}(\tilde{x}) d\tilde{x}$$

$$= \int \left(\int p(x) p_{\theta}(\tilde{x}|x) \nabla_{\tilde{x}} \log p_{\theta}(\tilde{x}|x) dx \right) S_{\theta}(\tilde{x}) d\tilde{x}$$

= For beauty of integrals we move $s_\theta(\tilde{x})$

$$= \iint p(x) p_\theta(\tilde{x}|x) \nabla \log p_\theta(\tilde{x}|x) s_\theta(\tilde{x}) dx d\tilde{x}$$

Now we bring everything back:

$$\begin{aligned} & \frac{1}{2} \mathbb{E}_{p_\theta(\tilde{x})} [\|\nabla_{\tilde{x}} \log p_\theta(\tilde{x}) - s_\theta(\tilde{x})\|_2^2] = \\ & = \frac{1}{2} \int p_\theta(\tilde{x}) \nabla_{\tilde{x}}^2 \log p_\theta(\tilde{x}) d\tilde{x} + \frac{1}{2} \int p_\theta(\tilde{x}) s_\theta^2(\tilde{x}) d\tilde{x} \\ & + \iint p(x) p_\theta(\tilde{x}|x) \nabla \log p_\theta(\tilde{x}|x) s_\theta(\tilde{x}) dx d\tilde{x} \\ & = \frac{1}{2} \mathbb{E}_{p_\theta(\tilde{x})} [\|\nabla_{\tilde{x}} \log p_\theta(\tilde{x})\|_2^2] + \\ & + \frac{1}{2} \mathbb{E}_{p_\theta(\tilde{x})} [\|s_\theta(\tilde{x})\|_2^2] + \\ & - \mathbb{E}_{x \sim p(x), \tilde{x} \sim p_\theta(\tilde{x}|x)} [\nabla_{\tilde{x}} \log(p_\theta(\tilde{x}|x)) \cdot s_\theta(\tilde{x})] \end{aligned}$$

We open up the expectation

$$- = \frac{1}{2} \mathbb{E}_{x \sim p(x), \tilde{x} \sim p_\theta(x|x)} [\|s_\theta^2(\tilde{x})\|_2^2]$$

and so everything becomes:

$$= \frac{1}{2} \mathbb{E}_{p_\theta(\tilde{x})} [\|\nabla_{\tilde{x}} \log p_\theta(\tilde{x})\|_2^2] +$$

$$+ \frac{1}{2} \mathbb{E}_{x \sim p(x), \tilde{x} \sim p_\theta(x|x)} [\|s_\theta^2(\tilde{x})\|_2^2 - 2 \nabla_{\tilde{x}} \log p_\theta(\tilde{x}|x) \cdot s_\theta(\tilde{x})]$$

We sum and remove $\|\nabla_{\tilde{x}} \log p_\theta(\tilde{x}|x)\|_2^2$

$$= \frac{1}{2} \mathbb{E}_{p_\theta(\tilde{x})} [\|\nabla_{\tilde{x}} \log p_\theta(\tilde{x})\|_2^2] +$$

$$+ \frac{1}{2} \mathbb{E}_{x \sim p(x), \tilde{x} \sim p_\theta(x|x)} [\|s_\theta^2(\tilde{x})\|_2^2 - 2 \nabla_{\tilde{x}} \log p_\theta(\tilde{x}|x) \cdot s_\theta(\tilde{x})]$$

$$+ \|\nabla_{\tilde{x}} \log p_\theta(\tilde{x}|x)\|_2^2 - \|\nabla_{\tilde{x}} \log p_\theta(\tilde{x}|x)\|_2^2$$

$$\frac{1}{2} \mathbb{E}_{x \sim p(x), \tilde{x} \sim p_\theta(\tilde{x}|x)} [\|s_\theta(\tilde{x}) - \nabla_{\tilde{x}} \log p_\theta(\tilde{x}|x)\|_2^2]$$

$$\begin{aligned}
 &= \frac{1}{2} \mathbb{E}_{p_\theta(\tilde{x})} [\|\nabla_{\tilde{x}} \log p_\theta(\tilde{x})\|_2^2] + \\
 &+ \frac{1}{2} \mathbb{E}_{x \sim p(x), \tilde{x} \sim p_\theta(\tilde{x}|x)} [\|s_\theta(\tilde{x}) - \nabla_{\tilde{x}} \log p_\theta(\tilde{x}|x)\|_2^2] + \\
 &- \mathbb{E}_{x \sim p(x), \tilde{x} \sim p_\theta(\tilde{x}|x)} [\|\nabla_{\tilde{x}} \log p_\theta(\tilde{x}|x)\|_2^2]
 \end{aligned}$$

$$\downarrow \\
 - \frac{1}{2} \mathbb{E}_{x \sim p(x), \tilde{x} \sim p_\theta(\tilde{x}|x)} [\|\nabla_{\tilde{x}} \log p_\theta(\tilde{x}|x)\|_2^2]$$

But we see that $=$ and $-$ do not depend on $\theta \Rightarrow$ if we want to minimise they are constant and so our objective becomes:

$$= \mathbb{E}_{x \sim p(x), \tilde{x} \sim p_\theta(\tilde{x}|x)} [\|s_\theta(\tilde{x}) - \nabla_{\tilde{x}} \log p_\theta(\tilde{x}|x)\|_2^2]$$

Remember $\tilde{x} = x + \epsilon$

$$p_\theta(\tilde{x}|x) = \frac{1}{(2\pi)^{d/2} \sigma^2} e^{-\frac{1}{2\sigma^2} \|\tilde{x} - x\|^2}$$

So we have that $\nabla_{\tilde{x}} \log p_{\sigma}(\tilde{x}|x) =$

$$= \nabla_{\tilde{x}} \log \frac{1}{2\pi^{d/2} \sigma^d} \exp\left(-\frac{1}{2\sigma^2} \|\tilde{x} - x\|^2\right) =$$

$$= \cancel{\nabla_{\tilde{x}} \log \frac{1}{2\pi^{d/2} \sigma^d}} + \nabla_{\tilde{x}} \log e^{-\frac{1}{2\sigma^2} \|\tilde{x} - x\|^2} =$$

$$= -\frac{1}{\sigma^2} 2\|\tilde{x} - x\| = -\frac{\tilde{x} - x}{\sigma^2} =$$

$$= \frac{x - \tilde{x}}{\sigma^2}$$

So we have that:

$$\frac{1}{2} \mathbb{E}_{x \sim p(x), \tilde{x} \sim p_{\sigma}(\tilde{x}|x)} \left[\left\| \nabla_{\tilde{x}} \log p_{\sigma}(\tilde{x}|x) - \frac{1}{\sigma^2} (x - \tilde{x}) \right\|_2^2 \right]$$

Remember that $\tilde{x} = x + \epsilon$

$$= \frac{1}{2} \mathbb{E}_{x \sim p(x), \tilde{x} \sim p_{\sigma}(\tilde{x}|x)} \left[\left\| \nabla_{\tilde{x}} \log p_{\sigma}(\tilde{x}|x) + \frac{\epsilon}{\sigma^2} \right\|_2^2 \right]$$

So we can see the our model S_0 predicts the negative of the noise that we added. The negative of the noise points in the direction of the data manifold where the data likelihood grows

HOW DO WE GENERATE NEW DATA!

If we do just see there is a high chance we overshoot the data and fail so we do it step by step with small step

$$\tilde{x}_{i+1} \leftarrow \tilde{x}_i + \alpha \cdot S_0(\tilde{x}_i)$$

the size of the step is determined by α

The noise points to the case where the likelihood is high because

that's where we have higher density
of data

There is an easy fix, it is called
LANGEVIAN DYNAMICS

$$\tilde{x}_{i+1} \leftarrow \tilde{x}_i + \alpha \nabla_{\theta} \ell(\tilde{x}_i) + \sqrt{2\alpha} \cdot \epsilon$$

we add this

$\epsilon \sim \text{Random Gaussian noise}$

$\sqrt{2\alpha}$ very small

In this way we have nice outputs
that do not collapse

Considering σ changing slowly,
so the model can adapt
on all over the process

So we have $\tilde{x} = x + \epsilon$ $\epsilon \sim \mathcal{N}(0, \sigma_{\epsilon}^2 I)$

so σ is not fixed anymore

If you increase $t \rightarrow \infty$ the noise perturbation becomes stochastic

So we can use SDEs to model stochastic processes.

$$dx = \underbrace{f(x, t)}_{\text{drift}} \underbrace{dt}_{\text{infinitesimal changes in time}} + \underbrace{g(t)}_{\text{diffusion}} \underbrace{dw}_{\text{Wiener process / infinitesimal changes in noise}}$$

drift is fixed and describes if and how we change x in a deterministic way

diffusion describes the influence of the stochasticity over time

the more volatility we have

$$\tilde{x} = x + G \quad \epsilon \sim \mathcal{N}(0, \sigma_t^2 I)$$

So change of x is only influenced by a stochastic term so $f(x, t) = 0$

$$dx = g(t) dw \quad \text{so as } t \rightarrow +\infty$$

$$\sigma_t \rightarrow \sigma(t)$$

$$\downarrow$$

$$G \in [0, 1]$$

N.B.

IN DDPM $f(x, t) \neq 0$

This SDE corresponds to the process of noising our data **FORWARD SDE**

$$dx = g(t) dw$$

REVERSE SDE (with LANGEVIN DYNAMICS)

↓
sampling

If FORWARD $dx = f(x, t) dt + g(t) dw$

the REVERSE SDE is:

$$dx = [f(x, t) - g^2(t) (\nabla_x \log p_\theta(x))] dt + g(t) dw$$

Score function

So to reverse we can use the
SCORE FUNCTION

CONNECTION TO DDPM

Score matching $\tilde{x} = x + \epsilon \quad \epsilon \sim \mathcal{N}(0, \sigma_t^2 I)$

DDPM $x_t = \sqrt{1 - \beta_t} x_{t-1} + \sqrt{\beta_t} \epsilon \quad \epsilon \sim \mathcal{N}(0, I)$

So we see DDPM is non zero

If we compute the reverse SDE for
DDPM we have that

$$dx = [f(x, t) - g^2(t) \nabla_x \log p_\theta(x)] dt + g(t) dw$$

$$dx = x_t - x_{t-1}$$

$$= \sqrt{1 - \beta_t} x_{t-1} + \sqrt{\beta_t} \epsilon - x_{t-1}$$

$\downarrow$ Taylor expansion $= 1 - \frac{1}{2} \beta_t$

$$= \left(1 - \frac{1}{2} \beta_t\right) x_{t-1} + \sqrt{\beta_t} \epsilon - x_{t-1}$$

$$= \cancel{x_{t-1}} - \frac{1}{2} \beta_t x_{t-1} + \sqrt{\beta_t} \epsilon - \cancel{x_{t-1}}$$

$$= -\frac{1}{2} \beta_t x_{t-1} + \sqrt{\beta_t} \epsilon$$

So e, $x_t - x_{t-1} \rightarrow 0$, $t \rightarrow \infty$

$$dx = -\frac{1}{2} \beta(t) x dt + \sqrt{\beta(t)} dw$$

FORWARD SDE

REVERSE SDE

$$f(x, t) = -\frac{1}{2} \beta(t) x$$

$$g(t) = \sqrt{\beta(t)}$$

We apply the formula

$$dx = \left[-\frac{1}{2} \beta(t) x - \beta(t) \nabla_x \log p_\theta(x) \right] dt + \sqrt{\beta(t)} dw$$

If we discretize it:

↓

DDPM SAMPLER

$$x_{t-1} = \frac{1}{\sqrt{1-\beta_t}} \left(x_t - \frac{\beta_t}{\sqrt{1-\bar{\alpha}_t}} \nabla \log p_\theta(x_t, t) \right) + \sqrt{\beta_t} z$$

We discretize t with the

Euler-Maruyama Method