

Exercises on Denoising diffusion probabilistic models

Gabriel V. Cardoso
STIM
Ecole des Mines Paris(PSL)
Fontainebleau, France
gabriel.victorino_cardoso@minesparis.psl.eu

September 15, 2026

Notation

Let $(\Omega, \mathcal{F}, \mathbb{P})$ be a probability space and $\mathbf{X} : (\Omega, \mathcal{F}) \rightarrow (\mathbb{R}^d, \mathcal{B}(\mathbb{R}^d))$ and $\mathbf{Y} : (\Omega, \mathcal{F}) \rightarrow (\mathbb{R}^m, \mathcal{B}(\mathbb{R}^m))$. We say that a Markov kernel $\mathbf{N} : \mathbb{R}^d \times \mathcal{B}(\mathbb{R}^m) \rightarrow \mathbb{R}_+$ is a conditional distribution of $\mathbf{Y}$ given $\mathbf{X}$ if

$$\mathbf{N}(\mathbf{X}, A) = \mathbb{E} [1_A(\mathbf{Y}) | \mathbf{X}] \quad \mathbb{P} - \text{almost surely.} \quad (1)$$

We use the notation $\mathbf{N} = \mathcal{L}(\mathbf{Y} | \mathbf{X})$.

Proposition 0.1. Suppose that for $B \in \mathcal{B}(\mathbb{R}^{2d})$, $\mathcal{L}(\mathbf{X}, \mathbf{Y})(B) = \int 1_B(\mathbf{x}, \mathbf{y}) p_{\mathbf{Y}|\mathbf{X}}(\mathbf{y}|\mathbf{x}) \mu(d\mathbf{x}) d\mathbf{y}^1$ where

$$p_{\mathbf{Y}|\mathbf{X}} : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}_+ \quad \text{such that} \quad \mu \left(\left\{ \mathbf{x} \in \mathbb{R}^d \mid \int p_{\mathbf{Y}|\mathbf{X}}(\mathbf{y}|\mathbf{x}) d\mathbf{y} = 1 \right\} \right) = 1. \quad (2)$$

Then

$$\mathbf{N} : \mathbb{R}^d \times \mathcal{B}(\mathbb{R}^d) \ni (\mathbf{x}, A) \rightarrow \int 1_A(\mathbf{y}) p_{\mathbf{Y}|\mathbf{X}}(\mathbf{y}|\mathbf{x}) d\mathbf{y}$$

is a conditional distribution (and the natural one) $\mathcal{L}(\mathbf{Y} | \mathbf{X})$.

Proof. It is straightforward to check that it is a Markov Kernel. We will check for the conditional expectation only.

We denote $\hat{f}(\mathbf{x}) = \int 1_A(\mathbf{y}) p_{\mathbf{Y}|\mathbf{X}}(\mathbf{y}|\mathbf{x}) d\mathbf{y}$. Note that it is enough (as we implicitly consider everywhere in the course that we are in L^2) that

$$\mathbb{E} \left[(1_A(\mathbf{Y}) - \hat{f}(\mathbf{X})) \hat{f}(\mathbf{X}) \right] = 0. \quad (3)$$

But

$$\mathbb{E} \left[(1_A(\mathbf{Y}) - \hat{f}(\mathbf{X})) \hat{f}(\mathbf{X}) \right] = \int (1_A(\mathbf{y}) - \hat{f}(\mathbf{x})) \hat{f}(\mathbf{x}) p_{\mathbf{Y}|\mathbf{X}}(\mathbf{y}|\mathbf{x}) \mu(d\mathbf{x}) d\mathbf{y} \quad (4)$$

$$= \int \underbrace{\left(\int 1_A(\mathbf{y}) p_{\mathbf{Y}|\mathbf{X}}(\mathbf{y}|\mathbf{x}) d\mathbf{y} \right)}_{\hat{f}(\mathbf{x})} \hat{f}(\mathbf{x}) \mu(d\mathbf{x}) - \int \hat{f}(\mathbf{x})^2 \mu(d\mathbf{x}) = 0. \quad (5)$$

Thus, $\hat{f}(\mathbf{X}) = \mathbf{N}(\mathbf{X}, A) = \mathbb{E} [1_A(\mathbf{Y}) | \mathbf{X}]$. □

Some observations are important:

- We can show that conditional distributions exist almost always (if the image space is nice enough) [1, Theorem B.3.11].

¹This simply means that the conditional admits a density with respect to the Lebesgue measure

1 Noising Markov Kernels

We define for $\mathbf{x} \in \mathbb{R}^d$ and $A \in \mathcal{B}(\mathbb{R}^d)$, $\alpha \in \mathbb{R}$ and $\beta \in \mathbb{R}_+$,

$$Q(\mathbf{x}, A) := \int 1_A(\mathbf{y}) \mathcal{N}(\mathbf{y}; \alpha \mathbf{x}, \beta^2 \mathbf{I}) d\mathbf{y}. \quad (6)$$

1. Give an upper bound on $KL(\mu Q \| \mathcal{N}(\mathbf{m}, \sigma^2 \mathbf{I}))$. Interpret what it means, at least for well chosen values of $\alpha, \beta, \mathbf{m}$ and σ .
2. Let $\eta = \mathcal{N}(0, \mathbf{I})$. Give a condition on α and β so that η is an invariant distribution for the Q , which means that

$$\eta Q = \eta. \quad (7)$$

Give an intuitive interpretation.

3. What is the relation between the random variables

$$\mathbf{X}_0 \sim \mu, \quad \mathbf{X}_1 = \alpha \mathbf{X}_0 + \beta \mathbf{Z}, \quad \mathbf{Z} \sim \mathcal{N}(0, \mathbf{I}), \quad \mathbf{X}_0 \perp\!\!\!\perp \mathbf{Z}, \quad (8)$$

and Q and μ ?

4. For $n \in \mathbb{N}$, give an expression for the distribution $Q^n(\mathbf{x})$.
5. Identify the cases where for every $\mathbf{x}$,

$$\lim_{n \rightarrow \infty} KL(Q^n(\mathbf{x}) \| \mathcal{N}(0, \mathbf{I})) = 0. \quad (9)$$

Does the speed of convergence depend on $\mathbf{x}$?

2 Bridge distribution on Gaussian Markov Chains

We consider now the following Markov Chain

$$\mathbf{X}_0 \sim \mu, \quad \mathbf{X}_1 = \mathbf{X}_0 + \sigma_1 \mathbf{Z}_1, \quad \mathbf{X}_2 = \mathbf{X}_1 + \sigma_2 \mathbf{Z}_2, \quad \mathbf{Z}_1 \perp\!\!\!\perp \mathbf{Z}_2 \perp\!\!\!\perp \mathbf{X}_0. \quad (10)$$

1. What is $\mathcal{L}(\mathbf{X}_i | \mathbf{X}_0)$ for $i = 1, 2$?

Solution: Let f be a continuous bounded function.

$$\begin{aligned} \mathbb{E}[f(\mathbf{X}_0, \mathbf{X}_1, \mathbf{X}_2)] &= \int f(\mathbf{x}_0, \mathbf{x}_0 + \sigma_1 \mathbf{z}_0, \mathbf{x}_0 + \sigma_2 \mathbf{z}_0 + \sigma \mathbf{z}_1) \mathcal{N}(\mathbf{z}_1; 0, \mathbf{I}) \mathcal{N}(\mathbf{z}_0; 0, \mathbf{I}) \mu(d\mathbf{x}_0) d\mathbf{z}_0 d\mathbf{z}_1 \\ &= \int f(\mathbf{x}_0, \mathbf{x}_1, \mathbf{x}_2) \mathcal{N}(\mathbf{x}_1; \mathbf{x}_0, \sigma_1^2 \mathbf{I}) \mathcal{N}(\mathbf{x}_2; \mathbf{x}_1, \sigma_2^2 \mathbf{I}) \mu(d\mathbf{x}_0) d\mathbf{x}_1 d\mathbf{x}_2. \end{aligned}$$

The joint vector $(\mathbf{X}_0, \mathbf{X}_1, \mathbf{X}_2)$ admits a density $(\mathbf{x}_0, \mathbf{x}_1, \mathbf{x}_2) \rightarrow \mathcal{N}(\mathbf{x}_1; \mathbf{x}_0, \sigma_1^2 \mathbf{I}) \mathcal{N}(\mathbf{x}_2; \mathbf{x}_1, \sigma_2^2 \mathbf{I})$ with respect to the product measure $\mu \otimes \text{Leb} \otimes \text{Leb}$, where Leb is the Lebesgue measure in $\mathbb{R}^d$.

Particularly, if we consider that μ admits a density f_μ with respect to the Lebesgue measure, then $(\mathbf{X}_0, \mathbf{X}_1, \mathbf{X}_2)$ admits a density $(\mathbf{x}_0, \mathbf{x}_1, \mathbf{x}_2) \rightarrow f_\mu(\mathbf{x}_0) \mathcal{N}(\mathbf{x}_1; \mathbf{x}_0, \sigma_1^2 \mathbf{I}) \mathcal{N}(\mathbf{x}_2; \mathbf{x}_1, \sigma_2^2 \mathbf{I})$ with respect to the product measure $\text{Leb} \otimes \text{Leb} \otimes \text{Leb}$ which is simply the Lebesgue measure in $\mathbb{R}^{3d}$.

We can now directly calculate the conditional distributions using proposition 0.1.

$$\mathcal{L}(\mathbf{X}_1 | \mathbf{X}_0)(\mathbf{x}_0, A) = \int 1_A(\mathbf{x}_1) \mathcal{N}(\mathbf{x}_1; \mathbf{x}_0, \sigma_1^2 \mathbf{I}) d\mathbf{x}_1, \quad (11)$$

which is equivalent to say that $\mathcal{L}(\mathbf{X}_1 | \mathbf{X}_0)(\mathbf{x}_0, \cdot) = \mathcal{N}(\mathbf{x}_0, \sigma_1^2 \mathbf{I})$, or even that $\mathcal{L}(\mathbf{X}_1 | \mathbf{X}_0) = \mathcal{N}(\mathbf{X}_0, \sigma_1^2 \mathbf{I})$.

For $\mathbf{X}_2$, simply note that $\int \mathcal{N}(\mathbf{x}_1; \mathbf{x}_0, \sigma_1^2 \mathbf{I}) \mathcal{N}(\mathbf{x}_2; \mathbf{x}_1, \sigma_2^2 \mathbf{I}) d\mathbf{x}_1 = \mathcal{N}(\mathbf{x}_2; \mathbf{x}_0, (\sigma_2^2 + \sigma_1^2) \mathbf{I})$, thus

$$\begin{aligned} \mathbb{E}[f(\mathbf{X}_0, \mathbf{X}_2)] &= \int f(\mathbf{x}_0, \mathbf{x}_2) \mathcal{N}(\mathbf{x}_1; \mathbf{x}_0, \sigma_1^2 \mathbf{I}) \mathcal{N}(\mathbf{x}_2; \mathbf{x}_1, \sigma_2^2 \mathbf{I}) \mu(d\mathbf{x}_0) d\mathbf{x}_0 d\mathbf{x}_1 \\ &= \int f(\mathbf{x}_0, \mathbf{x}_2) \mathcal{N}(\mathbf{x}_2; \mathbf{x}_0, (\sigma_2^2 + \sigma_1^2) \mathbf{I}) \mu(d\mathbf{x}_0) d\mathbf{x}_2. \end{aligned} \quad (12)$$

Thus, by proposition 0.1, $\mathcal{L}(\mathbf{X}_2 | \mathbf{X}_0) = \mathcal{N}(\mathbf{X}_0, (\sigma_1^2 + \sigma_2^2) \mathbf{I})$.

2. What is $\mathcal{L}(\mathbf{X}_1|\mathbf{X}_0, \mathbf{X}_2)$?

Solution: We start by noting that

$$\begin{aligned} & \int f(\mathbf{x}_0, \mathbf{x}_1, \mathbf{x}_2) \mathcal{N}(\mathbf{x}_1; \mathbf{x}_0, \sigma_1^2 \mathbf{I}) \mathcal{N}(\mathbf{x}_2; \mathbf{x}_1, \sigma_2^2 \mathbf{I}) \mu(d\mathbf{x}_0) d\mathbf{x}_1 d\mathbf{x}_2 \\ &= \int f(\mathbf{x}_0, \mathbf{x}_1, \mathbf{x}_2) \frac{\mathcal{N}(\mathbf{x}_1; \mathbf{x}_0, \sigma_1^2 \mathbf{I}) \mathcal{N}(\mathbf{x}_2; \mathbf{x}_1, \sigma_2^2 \mathbf{I})}{\mathcal{N}(\mathbf{x}_2; \mathbf{x}_0, (\sigma_1^2 + \sigma_2^2) \mathbf{I})} \mathcal{N}(\mathbf{x}_2; \mathbf{x}_0, (\sigma_2^2 + \sigma_1^2) \mathbf{I}) \mu(d\mathbf{x}_0) d\mathbf{x}_1 d\mathbf{x}_2 \\ &= \int f(\mathbf{x}_0, \mathbf{x}_1, \mathbf{x}_2) \mathcal{N}\left(\mathbf{x}_1; \mathbf{x}_0 \left(1 - \frac{\sigma_1^2}{\sigma_2^2 + \sigma_1^2}\right) + \mathbf{x}_2 \frac{\sigma_1^2}{\sigma_2^2 + \sigma_1^2}, \frac{\sigma_1^2 \sigma_2^2}{\sigma_2^2 + \sigma_1^2} \mathbf{I}\right) \mathcal{N}(\mathbf{x}_2; \mathbf{x}_0, (\sigma_2^2 + \sigma_1^2) \mathbf{I}) \mu(d\mathbf{x}_0) d\mathbf{x}_1 d\mathbf{x}_2. \end{aligned}$$

We apply proposition 0.1 with $\mathbf{X} = \mathbf{X}_1$ and $\mathbf{Y} = (\mathbf{X}_0, \mathbf{X}_2)$. Thus,

$$\mathcal{L}(\mathbf{X}_1|\mathbf{X}_0, \mathbf{X}_2)(\mathbf{x}_0, \mathbf{x}_1, A) = \int 1_A(\mathbf{x}_1) \mathcal{N}\left(\mathbf{x}_1; \mathbf{x}_0 \left(1 - \frac{\sigma_1^2}{\sigma_2^2 + \sigma_1^2}\right) + \mathbf{x}_2 \frac{\sigma_1^2}{\sigma_2^2 + \sigma_1^2}, \frac{\sigma_1^2 \sigma_2^2}{\sigma_2^2 + \sigma_1^2} \mathbf{I}\right) d\mathbf{x}_1. \quad (13)$$

3. Give an expression for $(\mathbf{m}, \Sigma) \in \arg \min \text{KL}(\mathcal{L}(\mathbf{X}_1|\mathbf{X}_2)(\mathbf{x}_2, \cdot) \|\mathcal{N}(\mathbf{m}, \Sigma))$ that depends only on the moments of $\mathcal{L}(\mathbf{X}_0|\mathbf{X}_2)$.

Solution: We know that $\mathbf{m} = \mathbb{E}[\mathbf{X}_1|\mathbf{X}_2]$ and that $\Sigma = \mathbb{V}[\mathbf{X}_1|\mathbf{X}_2]$. By the total expectation formula we have

$$\mathbf{m} = \mathbb{E}[\mathbf{X}_1|\mathbf{X}_2 = \mathbf{x}_2] = \mathbb{E}[\mathbb{E}[\mathbf{X}_1|\mathbf{X}_2, \mathbf{X}_0]|\mathbf{X}_2 = \mathbf{x}_2] = \mathbb{E}[\mathbf{X}_0|\mathbf{X}_2 = \mathbf{x}_2] \left(1 - \frac{\sigma_1^2}{\sigma_2^2 + \sigma_1^2}\right) + \mathbf{x}_2 \frac{\sigma_1^2}{\sigma_2^2 + \sigma_1^2}. \quad (14)$$

By the total variation formula, we have that

$$\Sigma = \mathbb{E}[\mathbb{V}[\mathbf{X}_1|\mathbf{X}_2, \mathbf{X}_0]|\mathbf{X}_2 = \mathbf{x}_2] + \mathbb{V}[\mathbb{E}[\mathbf{X}_1|\mathbf{X}_2, \mathbf{X}_0]|\mathbf{X}_2 = \mathbf{x}_2] \quad (15)$$

$$= \frac{\sigma_1^2 \sigma_2^2}{\sigma_2^2 + \sigma_1^2} \mathbf{I} + \left(1 - \frac{\sigma_1^2}{\sigma_2^2 + \sigma_1^2}\right)^2 \mathbb{V}[\mathbf{X}_0|\mathbf{X}_2 = \mathbf{x}_2] \quad (16)$$

$$= \frac{\sigma_1^2 \sigma_2^2}{\sigma_2^2 + \sigma_1^2} \mathbf{I} + \left(\frac{\sigma_1^2 \sigma_2^2}{\sigma_2^2 + \sigma_1^2}\right)^2 \frac{\mathbb{V}[\mathbf{X}_0|\mathbf{X}_2 = \mathbf{x}_2]}{\sigma_1^4} \dots \quad (17)$$

4. In which case can we neglect the terms that depend on $\mathbb{V}[\mathbf{X}_0|\mathbf{X}_2]$?

3 Generalizing and DDPM [2].

We now consider a generalized version of the previous exercise, where

$$\mathbf{X}_0 \sim \mu, \quad \mathbf{X}_k = \mathbf{X}_{k-1} + \sigma \mathbf{Z}_k, \quad \mathbf{Z}_k \perp \mathbf{Z}_\ell \perp \mathbf{X}_0 \quad \text{for } k \neq \ell, \quad (18)$$

and $k \in \mathbb{N}$. We suppose for simplicity that $\mathbb{E}[\mathbf{X}_0] = 0$ and $\mathbb{V}[\mathbf{X}_0] = \mathbf{I}$.

1. Calculate an upper bound for $\text{KL}(\mathcal{L}(\mathbf{X}_k) \|\mathcal{N}(0, (k\sigma^2 + 1) \mathbf{I}))$.

Solution: We consider the joint laws

$$\mu = \mathcal{L}(\mathbf{X}_0, \mathbf{X}_k) \quad \text{and} \quad \nu = p_{\text{data}} \otimes \mathcal{N}(0, (k\sigma^2 + 1) \mathbf{I}). \quad (19)$$

Then, by the data processing inequality,

$$\begin{aligned} \text{KL}(\mathcal{L}(\mathbf{X}_k) \|\mathcal{N}(0, (k\sigma^2 + 1) \mathbf{I})) &\leq \text{KL}(\mu \|\nu) = \mathbb{E} \left[\log \left(\frac{\mathcal{N}(\mathbf{X}_k; \mathbf{X}_0, k\sigma^2 \mathbf{I})}{\mathcal{N}(\mathbf{X}_k; 0, (k\sigma^2 + 1) \mathbf{I})} \right) \right] \\ &= \frac{1}{2} \left[d \left(\frac{k\sigma^2}{k\sigma^2 + 1} + \log \left(\frac{k\sigma^2}{k\sigma^2 + 1} \right) - 1 \right) + \frac{\mathbb{E}[\|\mathbf{X}_0\|^2]}{k\sigma^2 + 1} \right]. \end{aligned}$$

2. Suppose that you know for all k the functions $\mathbb{E}[\mathbf{X}_0|\mathbf{X}_k]$ and $\mathbb{V}[\mathbf{X}_0|\mathbf{X}_k]$. Can you propose a Markov chain $(\tilde{\mathbf{X}}_k)_{k=K}^0$ such that one might expect that

$$\text{KL}(\mathcal{L}(\mathbf{X}_k) \parallel \mathcal{L}(\tilde{\mathbf{X}}_k)) \approx 0 \quad (20)$$

for all k ?

Solution: Again, we use the data processing inequality to show that

$$\text{KL}(\mathcal{L}(\mathbf{X}_k) \parallel \mathcal{L}(\tilde{\mathbf{X}}_k)) \leq \text{KL}(\mathcal{L}(\mathbf{X}_k|\mathbf{X}_{k+1}) \parallel \mathcal{L}(\tilde{\mathbf{X}}_k|\tilde{\mathbf{X}}_{k+1})) .$$

Furthermore, we know by the previous exercise that

$$\begin{aligned} \mathbf{m}_{k+1}(\mathbf{x}_{k+1}) &= \left(1 - \frac{k}{k+1}\right) \mathbb{E}[\mathbf{X}_0|\mathbf{X}_{k+1} = \mathbf{x}_{k+1}] + \frac{k}{k+1} \mathbf{x}_{k+1} \\ \Sigma_{k+1}(\mathbf{x}_{k+1}) &= \left(1 - \frac{k}{k+1}\right)^2 \mathbb{V}[\mathbf{X}_0|\mathbf{X}_{k+1} = \mathbf{x}_{k+1}] + \frac{k}{k+1} \sigma^2 \mathbf{I} . \end{aligned}$$

is the minimizer of

$$\arg \min_{\mathbf{m}, \Sigma} \text{KL}(\mathcal{L}(\mathbf{X}_k|\mathbf{X}_{k+1} = \mathbf{x}_{k+1}) \parallel \mathcal{N}(\mathbf{m}_{k+1}, \Sigma_{k+1})) \approx 0$$

for every $\mathbf{x}_{k+1}$. Thus, if we take as the kernel for the Markov Chain

$$\mathbf{Q}(\mathbf{x}_{k+1}, \cdot) = \mathcal{N}(\cdot; \mathbf{m}_{k+1}(\mathbf{x}_{k+1}), \Sigma_{k+1}(\mathbf{x}_{k+1}))$$

and if $\text{KL}(\mathcal{L}(\mathbf{X}_{k+1}) \parallel \mathcal{L}(\tilde{\mathbf{X}}_{k+1})) \approx 0$, we will keep having small errors. This is of course a bit "hand wavy".

3. Provide an upper bound on $\text{KL}(\mathcal{L}(\mathbf{X}_0) \parallel \mathcal{L}(\tilde{\mathbf{X}}_0))$.

Solution: We start by the data processing inequality.

$$\begin{aligned} \text{KL}(\mathcal{L}(\mathbf{X}_0) \parallel \mathcal{L}(\tilde{\mathbf{X}}_0)) &\leq \text{KL}(\mathcal{L}(\mathbf{X}_0, \dots, \mathbf{X}_K) \parallel \mathcal{L}(\tilde{\mathbf{X}}_0, \dots, \tilde{\mathbf{X}}_K)) \\ &= \mathbb{E} \left[\log \left(\prod_{k=0}^{K-1} \frac{p_{k|k+1}(\mathbf{X}_k|\mathbf{X}_{k+1})}{\mathcal{N}(\mathbf{X}_k; \mathbf{m}_{k+1}(\mathbf{X}_{k+1}), \Sigma_{k+1}(\mathbf{X}_{k+1}))} \right) \frac{p_K(\mathbf{X}_K)}{\mathcal{N}(0, (K\sigma^2 + 1)\mathbf{I})} \right] \\ &\leq \sum_{k=0}^{K-1} \text{KL}(\mathcal{L}(\mathbf{X}_k|\mathbf{X}_{k+1}) \parallel \mathcal{N}(\mathbf{m}_{k+1}(\mathbf{X}_{k+1}), \Sigma_{k+1}(\mathbf{X}_{k+1}))) + \text{KL}(\mathcal{L}(\mathbf{X}_K) \parallel \mathcal{L}(\tilde{\mathbf{X}}_K)) \end{aligned}$$

It is now enough to upper bound $\text{KL}(\mathcal{L}(\mathbf{X}_k|\mathbf{X}_{k+1}) \parallel \mathcal{N}(\mathbf{m}_{k+1}(\mathbf{X}_{k+1}), \Sigma_{k+1}(\mathbf{X}_{k+1})))$ as the other term is already upper bounded by question 3.1.

To do so, we will use again the data processing inequality. Consider $\mu_{\mathbf{x}_{k+1}} = \mathcal{L}(\mathbf{X}_k, \mathbf{X}_0|\mathbf{X}_{k+1} = \mathbf{x}_{k+1})$ and $\nu_{\mathbf{x}_{k+1}} = \mathcal{N}(\mathbf{m}_{k+1}(\mathbf{x}_{k+1}), \Sigma_{k+1}(\mathbf{x}_{k+1})) \otimes \mathcal{L}(\mathbf{X}_0|\mathbf{X}_{k+1} = \mathbf{x}_{k+1})$.

Then,

$$\text{KL}(\mu_{\mathbf{x}_{k+1}} \parallel \nu_{\mathbf{x}_{k+1}}) = \int \log \left(\frac{p_{k|k+1,0}(\mathbf{x}_k|\mathbf{x}_k + 1, \mathbf{x}_0)p_{0|k+1}(\mathbf{x}_0|\mathbf{x}_{k+1})}{\mathcal{N}(\mathbf{m}_{k+1}(\mathbf{x}_{k+1}), \Sigma_{k+1}(\mathbf{x}_{k+1}))p_{0|k+1}(\mathbf{x}_0|\mathbf{x}_{k+1})} \right) p_{k|k+1,0}(\mathbf{x}_k|\mathbf{x}_k + 1, \mathbf{x}_0)p_{0|k+1}(\mathbf{x}_0|\mathbf{x}_{k+1}) d\mathbf{x}_0 d\mathbf{x}_k$$

where $p_{k|k+1,0}$ is the density of $\mathcal{L}(\mathbf{X}_k|\mathbf{X}_{k+1}, \mathbf{X}_0)$ and $p_{k|0}$ the density of $\mathcal{L}(\mathbf{X}_0|\mathbf{X}_k)$. Thus we have that

$$\begin{aligned} &\text{KL}(\mu_{\mathbf{x}_{k+1}} \parallel \nu_{\mathbf{x}_{k+1}}) \\ &= \mathbb{E} \left[\text{KL} \left(\mathcal{N} \left(\left(1 - \frac{k}{k+1}\right) \mathbf{X}_0 + \frac{k}{k+1} \mathbf{x}_{k+1}, \frac{k}{k+1} \sigma^2 \mathbf{I} \right) \parallel \mathcal{N}(\mathbf{m}_{k+1}(\mathbf{x}_{k+1}), \Sigma_{k+1}(\mathbf{x}_{k+1})) \right) \middle| \mathbf{X}_{k+1} = \mathbf{x}_{k+1} \right] . \end{aligned}$$

The rest follows from the KL formula for two Gaussians.

4. What if you know only $\mathbb{E}[\mathbf{X}_0|\mathbf{X}_k]$? Provide an upper bound for $\text{KL}\left(\mathcal{L}(\mathbf{X}_0)\left\|\mathcal{L}\left(\overleftarrow{\mathbf{X}}_0\right)\right.\right)$ in this case.

4 References

References

- [1] R. Douc et al. *Markov Chains*. Springer Series in Operations Research and Financial Engineering. Springer, Cham, 2018. xviii+757. ISBN: 978-3-319-97703-4 978-3-319-97704-1. DOI: 10.1007/978-3-319-97704-1. URL: <https://doi.org/10.1007/978-3-319-97704-1>.
- [2] Jonathan Ho, Ajay Jain, and Pieter Abbeel. “Denoising Diffusion Probabilistic Models”. In: *Advances in Neural Information Processing Systems*. Vol. 33. 2020, pp. 6840–6851.