

On the Accuracy of Diffusion Models in Bayesian Image Inverse Problems: A Gaussian Case Study

Émile Pierret^a, supervised by Bruno Galerne^{a,b}

French Biennial of Applied and Industrial Mathematics (SMAI), June 6, 2025

^a Institut Denis Poisson – Université d'Orléans, Université de Tours, CNRS

^b Institut universitaire de France (IUF)

Introduction

$$v = Ax + \sigma n, \quad n \sim \mathcal{N}_0 \quad (1)$$

where A is an inpainting, super-resolution, or blur operator.

Clean image



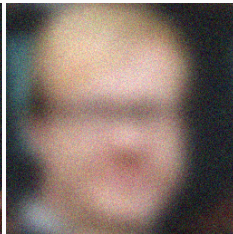
Super-resolution ($\times 4$)



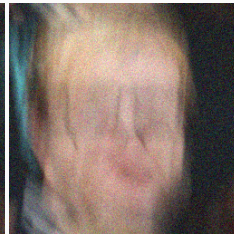
Inpainting



Gaussian blur



Motion blur



$$v = Ax + \sigma n, \quad n \sim \mathcal{N}_0 \quad (1)$$

where A is an inpainting, super-resolution, or blur operator.

Clean image



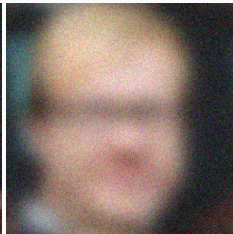
Super-resolution ($\times 4$)



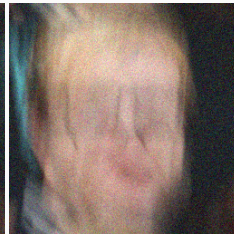
Inpainting



Gaussian blur



Motion blur



$\Rightarrow$ One popular solution: diffusion models.

Ground truth

Degraded image

DDRM

[Kawar et al., 2022]

DPS

[Chung et al., 2023]

IIGDM

[Song et al., 2023]



Ground truth

Degraded image

DDRM

[Kawar et al., 2022]

DPS

[Chung et al., 2023]

IIGDM

[Song et al., 2023]



Ground truth

Degraded image

DDRM

[Kawar et al., 2022]

DPS

[Chung et al., 2023]

IIGDM

[Song et al., 2023]



Ground truth

Degraded image

DDRM

[Kawar et al., 2022]

DPS

[Chung et al., 2023]

IIGDM

[Song et al., 2023]



One pending question

$$\boldsymbol{v} = \boldsymbol{A}\boldsymbol{x} + \sigma\boldsymbol{n}, \quad \boldsymbol{n} \sim \mathcal{N}_0 \quad (2)$$

where $\boldsymbol{A}$ inpainting, SR or blur operator.

¹E.P & Galerne, B. (2025). Diffusion models for gaussian distributions: Exact solutions and wasserstein errors. *Forty-second International Conference on Machine Learning*

One pending question

$$\boldsymbol{v} = \boldsymbol{A}\boldsymbol{x} + \sigma\boldsymbol{n}, \quad \boldsymbol{n} \sim \mathcal{N}_0 \quad (2)$$

where $\boldsymbol{A}$ inpainting, SR or blur operator.

Aim: Sample from $p(\boldsymbol{x} \mid \boldsymbol{v})$

¹E.P & Galerne, B. (2025). Diffusion models for gaussian distributions: Exact solutions and wasserstein errors. *Forty-second International Conference on Machine Learning*

One pending question

$$\boldsymbol{v} = \boldsymbol{A}\boldsymbol{x} + \sigma\boldsymbol{n}, \quad \boldsymbol{n} \sim \mathcal{N}_0 \quad (2)$$

where $\boldsymbol{A}$ inpainting, SR or blur operator.

Aim: Sample from $p(\boldsymbol{x} \mid \boldsymbol{v})$

Question: To what extent do diffusion models allow sampling from the target posterior distribution?

¹E.P & Galerne, B. (2025). Diffusion models for gaussian distributions: Exact solutions and wasserstein errors. *Forty-second International Conference on Machine Learning*

$$\boldsymbol{v} = \boldsymbol{A}\boldsymbol{x} + \sigma\boldsymbol{n}, \quad \boldsymbol{n} \sim \mathcal{N}_0 \quad (2)$$

where $\boldsymbol{A}$ inpainting, SR or blur operator.

Aim: Sample from $p(\boldsymbol{x} \mid \boldsymbol{v})$

Question: To what extent do diffusion models allow sampling from the target posterior distribution?

Idea: Observe what happens for Gaussian image distributions, that allow for feasible calculations.
(follows the previous work [E.P and Galerne, 2025]¹)

¹E.P & Galerne, B. (2025). Diffusion models for gaussian distributions: Exact solutions and wasserstein errors. *Forty-second International Conference on Machine Learning*

Diffusion models for image generation

Forward process

$$\mathbf{x}_t = \sqrt{1 - \beta_t} \mathbf{x}_{t-1} + \sqrt{\beta_t} \mathbf{z}_t, \quad 1 \leq t \leq T, \quad \mathbf{z}_t \sim \mathcal{N}_0, \quad \mathbf{x}_0 \sim p_{\text{data}}, \quad (3)$$

One can write

$$\mathbf{x}_t = \sqrt{\bar{\alpha}_t} \mathbf{x}_0 + \sqrt{1 - \bar{\alpha}_t} \varepsilon_t, \quad \varepsilon_t \sim \mathcal{N}_0 \quad (4)$$

²Ho, J., Jain, A., & Abbeel, P. (2020). Denoising diffusion probabilistic models. *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020*

Forward process

$$\mathbf{x}_t = \sqrt{1 - \beta_t} \mathbf{x}_{t-1} + \sqrt{\beta_t} \mathbf{z}_t, \quad 1 \leq t \leq T, \quad \mathbf{z}_t \sim \mathcal{N}_0, \quad \mathbf{x}_0 \sim p_{\text{data}}, \quad (3)$$

One can write

$$\mathbf{x}_t = \sqrt{\bar{\alpha}_t} \mathbf{x}_0 + \sqrt{1 - \bar{\alpha}_t} \varepsilon_t, \quad \varepsilon_t \sim \mathcal{N}_0 \quad (4)$$

Backward process

By learning ε_θ such that $\varepsilon_\theta(\mathbf{x}_t, t) \approx \varepsilon_t$,

$$\mathbf{y}_{t-1} = \frac{1}{\sqrt{\alpha_t}} \left(\mathbf{y}_t - \frac{\beta_t}{\sqrt{1 - \bar{\alpha}_t}} \varepsilon_\theta(\mathbf{y}_t, t) \right) + \sqrt{\tilde{\beta}_t} \mathbf{z}_t, \quad \mathbf{z}_t \sim \mathcal{N}_0. \quad (5)$$

²Ho, J., Jain, A., & Abbeel, P. (2020). Denoising diffusion probabilistic models. *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020*

Forward process

$$\mathbf{x}_t = \sqrt{1 - \beta_t} \mathbf{x}_{t-1} + \sqrt{\beta_t} \mathbf{z}_t, \quad 1 \leq t \leq T, \quad \mathbf{z}_t \sim \mathcal{N}_0, \quad \mathbf{x}_0 \sim p_{\text{data}}, \quad (3)$$

One can write

$$\mathbf{x}_t = \sqrt{\bar{\alpha}_t} \mathbf{x}_0 + \sqrt{1 - \bar{\alpha}_t} \varepsilon_t, \quad \varepsilon_t \sim \mathcal{N}_0 \quad (4)$$

Backward process

By learning ε_θ such that $\varepsilon_\theta(\mathbf{x}_t, t) \approx \varepsilon_t$,

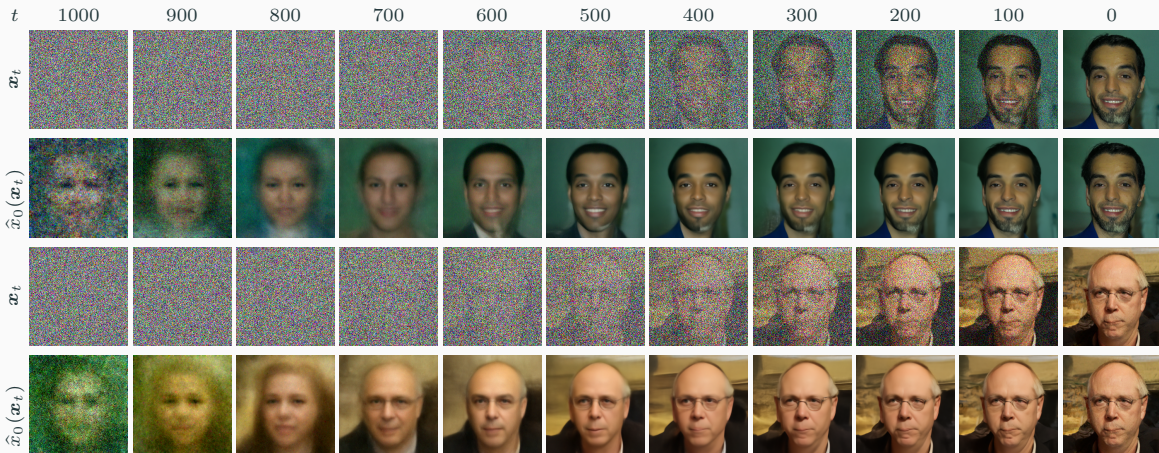
$$\mathbf{y}_{t-1} = \frac{1}{\sqrt{\alpha_t}} \left(\mathbf{y}_t - \frac{\beta_t}{\sqrt{1 - \bar{\alpha}_t}} \varepsilon_\theta(\mathbf{y}_t, t) \right) + \sqrt{\tilde{\beta}_t} \mathbf{z}_t, \quad \mathbf{z}_t \sim \mathcal{N}_0. \quad (5)$$

By denoting p_t the marginals of the forward process and learning $\mathbf{s}_\theta(\mathbf{x}, t) \approx \nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t)$,

$$\mathbf{y}_{t-1} = \frac{1}{\sqrt{\alpha_t}} (\mathbf{y}_t + \beta_t \mathbf{s}_{\theta^*}(\mathbf{y}_t, t)) + \sigma_t \mathbf{z}_t, \quad \mathbf{z}_t \sim \mathcal{N}_0, \quad 1 \leq t \leq T, \quad \mathbf{z}_t \sim \mathcal{N}_0, \quad \mathbf{y}_T \sim \mathcal{N}_0. \quad (6)$$

²Ho, J., Jain, A., & Abbeel, P. (2020). Denoising diffusion probabilistic models. *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020*

Generation examples



Diffusion models for inverse problems

$$v = Ax_0 + \sigma n, \quad x_0 \sim p_0, n \sim \mathcal{N}_0 \quad (7)$$

Aim: Sampling $p_0(\cdot \mid v)$

$$v = Ax_0 + \sigma n, \quad x_0 \sim p_0, n \sim \mathcal{N}_0 \quad (7)$$

Aim: Sampling $p_0(\cdot \mid v)$

$\implies$ conditional forward process

$$\tilde{x}_t = \sqrt{1 - \beta_t} \tilde{x}_{t-1} + \sqrt{\beta_t} z_t, \quad 1 \leq t \leq T, \quad z_t \sim \mathcal{N}_0, \quad \tilde{x}_0 \sim p_{\text{data}}(\cdot \mid v), \quad (8)$$

And for solving image problems ?

$$\boldsymbol{v} = \boldsymbol{A}\boldsymbol{x}_0 + \sigma\boldsymbol{n}, \quad \boldsymbol{x}_0 \sim p_0, \boldsymbol{n} \sim \mathcal{N}_0 \quad (7)$$

Aim: Sampling $p_0(\cdot \mid \boldsymbol{v})$

$\Rightarrow$ conditional forward process

$$\tilde{\boldsymbol{x}}_t = \sqrt{1 - \beta_t}\tilde{\boldsymbol{x}}_{t-1} + \sqrt{\beta_t}\boldsymbol{z}_t, \quad 1 \leq t \leq T, \quad \boldsymbol{z}_t \sim \mathcal{N}_0, \quad \tilde{\boldsymbol{x}}_0 \sim p_{\text{data}}(\cdot \mid \boldsymbol{v}), \quad (8)$$

$\Rightarrow$ conditional backward process

$$\tilde{\boldsymbol{y}}_{t-1} = \frac{1}{\sqrt{\alpha_t}} (\tilde{\boldsymbol{y}}_t + \beta_t \nabla \log \tilde{p}_t(\tilde{\boldsymbol{y}}_t)) + \sigma_t \boldsymbol{z}_t, \quad \boldsymbol{z}_t \sim \mathcal{N}_0, 1 \leq t \leq T, \boldsymbol{z}_t \sim \mathcal{N}_0, \tilde{\boldsymbol{y}}_T \sim \mathcal{N}_0 \quad (9)$$

Our point of interest is $\nabla \log \tilde{p}_t(\mathbf{x}_t) = \nabla \log p_t(\mathbf{x}_t \mid \mathbf{v})$.

Our point of interest is $\nabla \log \tilde{p}_t(\mathbf{x}_t) = \nabla \log p_t(\mathbf{x}_t \mid \mathbf{v})$. By Bayes' formula,

$$\nabla_{\mathbf{x}} \log \tilde{p}_t(\mathbf{x}_t) = \nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t) + \nabla_{\mathbf{x}} \log p_t(\mathbf{v} \mid \mathbf{x}_t), \quad (10)$$

Our point of interest is $\nabla \log \tilde{p}_t(\mathbf{x}_t) = \nabla \log p_t(\mathbf{x}_t \mid \mathbf{v})$. By Bayes' formula,

$$\nabla_{\mathbf{x}} \log \tilde{p}_t(\mathbf{x}_t) = \nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t) + \nabla_{\mathbf{x}} \log p_t(\mathbf{v} \mid \mathbf{x}_t), \quad (10)$$

Algorithm 5

Unconditional DDPM backward process

```
1:  $\mathbf{y}_T \sim \mathcal{N}_0$ 
2: for  $t=T$  to 1 do
3:    $\mathbf{z}_t \sim \mathcal{N}_0$ 
4:    $\mathbf{y}_{t-1} = \frac{1}{\sqrt{\alpha_t}} (\mathbf{y}_t + \beta_t s_{\theta^*}(\mathbf{y}_t, t)) + \sigma_t \mathbf{z}_t$ 
5: end for
```

Algorithm 6 Conditional DDPM backward process

```
1:  $\mathbf{y}_T \sim \mathcal{N}_0$ 
2: for  $t = T$  to 1 do
3:    $\hat{\mathbf{x}}_0(\mathbf{x}_t) = \frac{1}{\sqrt{\alpha_t}} (\mathbf{x}_t + (1 - \bar{\alpha}_t) s_{\theta^*}(\mathbf{y}_t, t))$ 
4:    $\tilde{s}_{\theta}(\mathbf{y}_t, t) = s_{\theta^*}(\mathbf{y}_t, t) + \nabla \log p_t(\mathbf{x} \mid \mathbf{v})$ 
5:    $\mathbf{z}_t \sim \mathcal{N}_0$ 
6:    $\mathbf{y}_{t-1} = \frac{1}{\sqrt{\alpha_t}} (\mathbf{y}_t + \beta_t \tilde{s}_{\theta}(\mathbf{y}_t, t)) + \sigma_t \mathbf{z}_t$ 
7: end for
```

Description of two algorithms from the literature

$$\mathbf{v} = \mathbf{A}\mathbf{x}_0 + \sigma\mathbf{n}, \quad \mathbf{x}_0 \sim p_0, \mathbf{n} \sim \mathcal{N}_0 \quad (11)$$

$$\nabla_{\mathbf{x}} \log \tilde{p}_t(\mathbf{x}_t) = \nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t) + \nabla_{\mathbf{x}} \log p_t(\mathbf{v} \mid \mathbf{x}_t), \quad (12)$$

Some assumptions lead to " $p_t(\mathbf{v} \mid \mathbf{x}_t)$ is Gaussian".

$$\mathbf{v} = \mathbf{A}\mathbf{x}_0 + \sigma\mathbf{n}, \quad \mathbf{x}_0 \sim p_0, \mathbf{n} \sim \mathcal{N}_0 \quad (11)$$

$$\nabla_{\mathbf{x}} \log \tilde{p}_t(\mathbf{x}_t) = \nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t) + \nabla_{\mathbf{x}} \log p_t(\mathbf{v} \mid \mathbf{x}_t), \quad (12)$$

Some assumptions lead to " $p_t(\mathbf{v} \mid \mathbf{x}_t)$ is Gaussian". Let note that an approximation of $\mathbb{E}(\mathbf{v} \mid \mathbf{x}_t)$ is known.

By Tweedie's formula, that is,

$$\hat{\mathbf{x}}_0(\mathbf{x}_t) := \mathbb{E}[\mathbf{x}_0 \mid \mathbf{x}_t] = \frac{1}{\sqrt{\bar{\alpha}_t}} (\mathbf{x}_t + (1 - \bar{\alpha}_t) \nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t)). \quad (13)$$

Mean of $p_t(\mathbf{v} \mid \mathbf{x}_t)$

$$\mathbf{v} = \mathbf{A}\mathbf{x}_0 + \sigma\mathbf{n}, \quad \mathbf{x}_0 \sim p_0, \mathbf{n} \sim \mathcal{N}_0 \quad (11)$$

$$\nabla_{\mathbf{x}} \log \tilde{p}_t(\mathbf{x}_t) = \nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t) + \nabla_{\mathbf{x}} \log p_t(\mathbf{v} \mid \mathbf{x}_t), \quad (12)$$

Some assumptions lead to " $p_t(\mathbf{v} \mid \mathbf{x}_t)$ is Gaussian". Let note that an approximation of $\mathbb{E}(\mathbf{v} \mid \mathbf{x}_t)$ is known.

By Tweedie's formula, that is,

$$\hat{\mathbf{x}}_0(\mathbf{x}_t) := \mathbb{E}[\mathbf{x}_0 \mid \mathbf{x}_t] = \frac{1}{\sqrt{\bar{\alpha}_t}} (\mathbf{x}_t + (1 - \bar{\alpha}_t) \nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t)). \quad (13)$$

As a consequence, $\mathbb{E}[\mathbf{v} \mid \mathbf{x}_t]$ is given by

$$\mathbb{E}[\mathbf{v} \mid \mathbf{x}_t] = \mathbf{A}\hat{\mathbf{x}}_0(\mathbf{x}_t) = \frac{1}{\sqrt{\bar{\alpha}_t}} \mathbf{A} (\mathbf{x}_t + (1 - \bar{\alpha}_t) \nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t)). \quad (14)$$

Mean of $p_t(\mathbf{v} \mid \mathbf{x}_t)$

$$\mathbf{v} = \mathbf{A}\mathbf{x}_0 + \sigma\mathbf{n}, \quad \mathbf{x}_0 \sim p_0, \mathbf{n} \sim \mathcal{N}_0 \quad (11)$$

$$\nabla_{\mathbf{x}} \log \tilde{p}_t(\mathbf{x}_t) = \nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t) + \nabla_{\mathbf{x}} \log p_t(\mathbf{v} \mid \mathbf{x}_t), \quad (12)$$

Some assumptions lead to " $p_t(\mathbf{v} \mid \mathbf{x}_t)$ is Gaussian". Let note that an approximation of $\mathbb{E}(\mathbf{v} \mid \mathbf{x}_t)$ is known.

By Tweedie's formula, that is,

$$\hat{\mathbf{x}}_0(\mathbf{x}_t) := \mathbb{E}[\mathbf{x}_0 \mid \mathbf{x}_t] = \frac{1}{\sqrt{\bar{\alpha}_t}} (\mathbf{x}_t + (1 - \bar{\alpha}_t) \nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t)). \quad (13)$$

As a consequence, $\mathbb{E}[\mathbf{v} \mid \mathbf{x}_t]$ is given by

$$\mathbb{E}[\mathbf{v} \mid \mathbf{x}_t] = \mathbf{A}\hat{\mathbf{x}}_0(\mathbf{x}_t) = \frac{1}{\sqrt{\bar{\alpha}_t}} \mathbf{A} (\mathbf{x}_t + (1 - \bar{\alpha}_t) \nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t)). \quad (14)$$

Finally,

$$p(\mathbf{x}_t \mid \mathbf{v}) = \mathcal{N}(\mathbf{A}\hat{\mathbf{x}}_0(\mathbf{x}_t), \mathbf{C}_{\mathbf{v}|\mathbf{t}}) \quad (15)$$

with $\mathbf{C}_{\mathbf{v}|\mathbf{t}}$ to fix.

Covariance of $p_t(\boldsymbol{v} \mid \boldsymbol{x})$

We denote it $\boldsymbol{C}_{\boldsymbol{v} \mid t}$.

- Denoising Posterior Sampling (DPS) algorithm [Chung et al., 2023]³

$$\log p(\boldsymbol{v} \mid \boldsymbol{x}_t) \approx \log p(\boldsymbol{v} \mid \boldsymbol{x}_0 = \hat{\boldsymbol{x}}_0(\boldsymbol{x}_t)) \quad (16)$$

³Chung, H., Kim, J., Mccann, M. T., Klasky, M. L., & Ye, J. C. (2023). Diffusion posterior sampling for general noisy inverse problems. *The Eleventh International Conference on Learning Representations*

We denote it $\boldsymbol{C}_{\boldsymbol{v} \mid t}$.

- Denoising Posterior Sampling (DPS) algorithm [Chung et al., 2023]³

$$\log p(\boldsymbol{v} \mid \boldsymbol{x}_t) \approx \log p(\boldsymbol{v} \mid \boldsymbol{x}_0 = \hat{\boldsymbol{x}}_0(\boldsymbol{x}_t)) \quad (16)$$

$$p(\boldsymbol{v} \mid \boldsymbol{x}_0) = \mathcal{N}(\boldsymbol{A}\boldsymbol{x}_0, \sigma^2 \boldsymbol{I}) \text{ (by } \boldsymbol{v} = \boldsymbol{A}\boldsymbol{x}_0 + \sigma \boldsymbol{n}\text{).} \quad (17)$$

³Chung, H., Kim, J., Mccann, M. T., Klasky, M. L., & Ye, J. C. (2023). Diffusion posterior sampling for general noisy inverse problems. *The Eleventh International Conference on Learning Representations*

We denote it $\mathbf{C}_{\mathbf{v}|t}$.

- Denoising Posterior Sampling (DPS) algorithm [Chung et al., 2023]³

$$\log p(\mathbf{v} \mid \mathbf{x}_t) \approx \log p(\mathbf{v} \mid \mathbf{x}_0 = \hat{\mathbf{x}}_0(\mathbf{x}_t)) \quad (16)$$

$$p(\mathbf{v} \mid \mathbf{x}_0) = \mathcal{N}(\mathbf{A}\mathbf{x}_0, \sigma^2 \mathbf{I}) \text{ (by } \mathbf{v} = \mathbf{A}\mathbf{x}_0 + \sigma \mathbf{n} \text{).} \quad (17)$$

$$\nabla_{\mathbf{x}_t} \log p(\mathbf{v} \mid \mathbf{x}_0 = \hat{\mathbf{x}}_0(\mathbf{x}_t)) = -\frac{1}{2\sigma^2} \nabla_{\mathbf{x}_t} \|\mathbf{v} - \mathbf{A}\hat{\mathbf{x}}_0(\mathbf{x}_t)\|^2. \quad (18)$$

³Chung, H., Kim, J., Mccann, M. T., Klasky, M. L., & Ye, J. C. (2023). Diffusion posterior sampling for general noisy inverse problems. *The Eleventh International Conference on Learning Representations*

Covariance of $p_t(\mathbf{v} \mid \mathbf{x})$

We denote it $\mathbf{C}_{\mathbf{v}|t}$.

- Denoising Posterior Sampling (DPS) algorithm [Chung et al., 2023]³

$$\log p(\mathbf{v} \mid \mathbf{x}_t) \approx \log p(\mathbf{v} \mid \mathbf{x}_0 = \hat{\mathbf{x}}_0(\mathbf{x}_t)) \quad (16)$$

$$p(\mathbf{v} \mid \mathbf{x}_0) = \mathcal{N}(\mathbf{A}\mathbf{x}_0, \sigma^2 \mathbf{I}) \text{ (by } \mathbf{v} = \mathbf{A}\mathbf{x}_0 + \sigma \mathbf{n} \text{)}. \quad (17)$$

$$\nabla_{\mathbf{x}_t} \log p(\mathbf{v} \mid \mathbf{x}_0 = \hat{\mathbf{x}}_0(\mathbf{x}_t)) = -\frac{1}{2\sigma^2} \nabla_{\mathbf{x}_t} \|\mathbf{v} - \mathbf{A}\hat{\mathbf{x}}_0(\mathbf{x}_t)\|^2. \quad (18)$$

In practice,

$$\nabla_{\mathbf{x}_t} \log p(\mathbf{v} \mid \mathbf{x}_0 = \hat{\mathbf{x}}_0(\mathbf{x}_t)) \approx -\frac{\alpha_{\text{DPS}}}{2\sigma^2} \nabla_{\mathbf{x}_t} \|\mathbf{v} - \mathbf{A}\hat{\mathbf{x}}_0(\mathbf{x}_t)\|^2. \quad (19)$$

$$\implies \mathbf{C}_{\mathbf{v}|t}^{\text{DPS}} = \frac{\sigma^2}{\alpha_{\text{DPS}}} \mathbf{I}$$

³Chung, H., Kim, J., Mccann, M. T., Klasky, M. L., & Ye, J. C. (2023). Diffusion posterior sampling for general noisy inverse problems. *The Eleventh International Conference on Learning Representations*

We denote it $C_{v|t}$.

- Pseudo-Guided Diffusion model (PIGDM) algorithm [Song et al., 2023]⁴

$$p(x_0 \mid x_t) \approx \mathcal{N}(\hat{x}_0(x_t), r_t^2 \mathbf{I}). \quad (20)$$

⁴Song, J., Vahdat, A., Mardani, M., & Kautz, J. (2023). Pseudoinverse-guided diffusion models for inverse problems. *International Conference on Learning Representations*

We denote it $C_{v|t}$.

- Pseudo-Guided Diffusion model (PIGDM) algorithm [Song et al., 2023]⁴

$$p(x_0 | x_t) \approx \mathcal{N}(\hat{x}_0(x_t), r_t^2 \mathbf{I}). \quad (20)$$

Consequently,

$$p(v | x_t) \approx \mathcal{N}\left(\mathbf{A}\hat{x}_0(x_t), r_t^2 \mathbf{A}\mathbf{A}^T + \sigma^2 \mathbf{I}\right) \quad (21)$$

⁴Song, J., Vahdat, A., Mardani, M., & Kautz, J. (2023). Pseudoinverse-guided diffusion models for inverse problems. *International Conference on Learning Representations*

We denote it $C_{v|t}$.

- Pseudo-Guided Diffusion model (PIGDM) algorithm [Song et al., 2023]⁴

$$p(x_0 \mid x_t) \approx \mathcal{N}(\hat{x}_0(x_t), r_t^2 \mathbf{I}). \quad (20)$$

Consequently,

$$p(v \mid x_t) \approx \mathcal{N}(\mathbf{A}\hat{x}_0(x_t), r_t^2 \mathbf{A}\mathbf{A}^T + \sigma^2 \mathbf{I}) \quad (21)$$

$$\implies C_{v|t}^{\text{PIGDM}} = r_t^2 \mathbf{A}\mathbf{A}^T + \sigma^2 \mathbf{I}$$

⁴Song, J., Vahdat, A., Mardani, M., & Kautz, J. (2023). Pseudoinverse-guided diffusion models for inverse problems. *International Conference on Learning Representations*

Under Gaussian assumption

By adding the assumption: x_0 follows a Gaussian distribution $\mathcal{N}(\mu, \Sigma)$.

By adding the assumption: $\mathbf{x}_0$ follows a Gaussian distribution $\mathcal{N}(\boldsymbol{\mu}, \boldsymbol{\Sigma})$.

$$p(\mathbf{v} \mid \mathbf{x}_t) = \mathcal{N}\left(\mathbf{A}\hat{\mathbf{x}}_0(\mathbf{x}_t), (1 - \bar{\alpha}_t)\mathbf{A}\boldsymbol{\Sigma}\boldsymbol{\Sigma}_t^{-1}\mathbf{A}^T + \sigma^2\mathbf{I}\right), \quad (22)$$

$$\text{with } \hat{\mathbf{x}}_0(\mathbf{x}_t) = \boldsymbol{\mu} + \sqrt{\bar{\alpha}_t}\boldsymbol{\Sigma}\boldsymbol{\Sigma}_t^{-1}(\mathbf{x}_t - \sqrt{\bar{\alpha}_t}\boldsymbol{\mu}). \quad (23)$$

We call this setting "Conditional Gaussian Diffusion Model" (CGDM) $\implies \mathbf{C}_{\mathbf{v}|t}^{\text{CGDM}} = (1 - \bar{\alpha}_t)\mathbf{A}\boldsymbol{\Sigma}\boldsymbol{\Sigma}_t^{-1}\mathbf{A}^T + \sigma^2\mathbf{I}$

Comparison of the choice made by different algorithms

	$C_{v t}$ (Covariance of $p(v x_t)$)
DPS Chung et al., 2023	$\frac{\sigma^2}{\alpha_{\text{DPS}}} \mathbf{I}$
Π GDM Song et al., 2023	$(1 - \bar{\alpha}_t) \mathbf{A} \mathbf{A}^T + \sigma^2 \mathbf{I}$
CGDM	$(1 - \bar{\alpha}_t) \mathbf{A} \Sigma \Sigma_t^{-1} \mathbf{A}^T + \sigma^2 \mathbf{I}, \quad \Sigma_t = \bar{\alpha}_t \Sigma + (1 - \bar{\alpha}_t) \mathbf{I}$

Comparison of the choice made by different algorithms

	$C_{v t}$ (Covariance of $p(v x_t)$)
DPS Chung et al., 2023	$\frac{\sigma^2}{\alpha_{\text{DPS}}} \mathbf{I}$
PIGDM Song et al., 2023	$(1 - \bar{\alpha}_t) \mathbf{A} \mathbf{A}^T + \sigma^2 \mathbf{I}$
CGDM	$(1 - \bar{\alpha}_t) \mathbf{A} \Sigma \Sigma_t^{-1} \mathbf{A}^T + \sigma^2 \mathbf{I}, \quad \Sigma_t = \bar{\alpha}_t \Sigma + (1 - \bar{\alpha}_t) \mathbf{I}$

Algorithm 7

Unconditional DDPM backward process

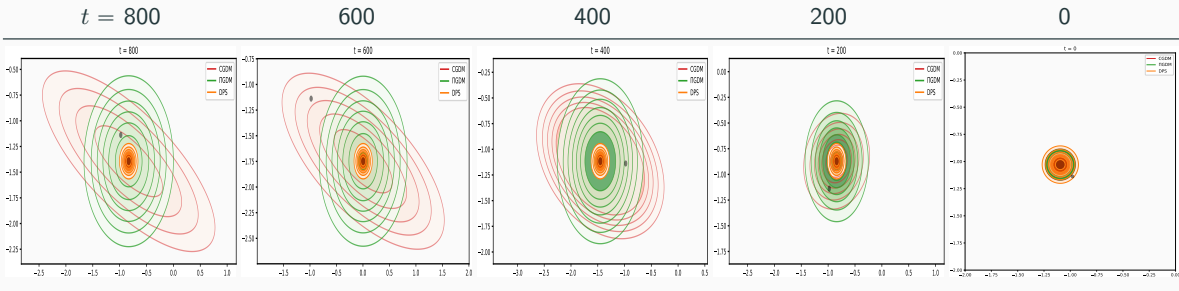
- 1: $\mathbf{y}_T \sim \mathcal{N}_0$
- 2: **for** $t=T$ **to** 1 **do**
- 3: $\mathbf{z}_t \sim \mathcal{N}_0$
- 4: $\mathbf{y}_{t-1} = \frac{1}{\sqrt{\alpha_t}} (\mathbf{y}_t + \beta_t s_{\theta^*}(\mathbf{y}_t, t)) + \sigma_t \mathbf{z}_t$
- 5: **end for**

Algorithm 8 Conditional DDPM backward process

- 1: $\mathbf{y}_T \sim \mathcal{N}_0$
- 2: **for** $t = T$ **to** 1 **do**
- 3: $\hat{\mathbf{x}}_0(\mathbf{x}_t) = \frac{1}{\sqrt{\bar{\alpha}_t}} (\mathbf{x}_t + (1 - \bar{\alpha}_t) s_{\theta^*}(\mathbf{y}_t, t))$
- 4: $\tilde{s}_{\theta}(\mathbf{y}_t, t) = s_{\theta^*}(\mathbf{y}_t, t) - \frac{1}{2} \nabla_{\mathbf{x}_t} \|\mathbf{v} - \mathbf{A} \hat{\mathbf{x}}_0(\mathbf{x}_t)\|_{C_{v|t}^{-1}}^2$
- 5: $\mathbf{z}_t \sim \mathcal{N}_0$
- 6: $\mathbf{y}_{t-1} = \frac{1}{\sqrt{\alpha_t}} (\mathbf{y}_t + \beta_t \tilde{s}_{\theta^*}(\mathbf{y}_t, t)) + \sigma_t \mathbf{z}_t$
- 7: **end for**

Comparison of the choice made by different algorithms

	$C_{v t}$ (Covariance of $p(v x_t)$)
DPS Chung et al., 2023	$\frac{\sigma^2}{\alpha_{\text{DPS}}} I$
Π GDM Song et al., 2023	$(1 - \bar{\alpha}_t) A A^T + \sigma^2 I$
CGDM	$(1 - \bar{\alpha}_t) A \Sigma \Sigma_t^{-1} A^T + \sigma^2 I, \quad \Sigma_t = \bar{\alpha}_t \Sigma + (1 - \bar{\alpha}_t) I$



At $t \approx 0$, $p(v | x_0) = \mathcal{N}(Ax_0, \sigma^2 I)$

Comparison of the algorithms via 2-Wasserstein distance

The Asymptotic Discrete Spot Noise (ADSN) model [Galerne et al., 2011b] ⁶

Let $\mathbf{u} \in \mathbb{R}^{\Omega_{M,N}}$ be a grayscale image, m its grayscale mean and $\mathbf{t} = \frac{1}{\sqrt{MN}}(\mathbf{u} - m)$ its associated texton. Let $\mathbf{w}$ be a white Gaussian noise,

$$\mathbf{X} = \mathbf{t} \star \mathbf{w} \sim \text{ADSN}(\mathbf{u}) = \mathcal{N}(\mathbf{0}, \Gamma) \quad \text{which is a stationary law}$$

$\mathbf{u}$

Samples of $\text{ADSN}(\mathbf{u})$



Image extracted from [Galerne et al., 2011a] ⁵

⁵Galerne, B., Gousseau, Y., & Morel, J.-M. (2011a). Micro-texture synthesis by phase randomization. *Image Processing On Line*

⁶Galerne, B., Gousseau, Y., & Morel, J.-M. (2011b). Random Phase Textures: Theory and Synthesis. *IEEE Transactions on Image Processing*, 20(1), 257–267

The Asymptotic Discrete Spot Noise (ADSN) model [Galerne et al., 2011b] ⁶

Let $\mathbf{u} \in \mathbb{R}^{\Omega_{M,N}}$ be a grayscale image, m its grayscale mean and $\mathbf{t} = \frac{1}{\sqrt{MN}}(\mathbf{u} - m)$ its associated texton. Let $\mathbf{w}$ be a white Gaussian noise,

$$\mathbf{X} = \mathbf{t} \star \mathbf{w} \sim \text{ADSN}(\mathbf{u}) = \mathcal{N}(\mathbf{0}, \mathbf{\Gamma}) \quad \text{which is a stationary law}$$

$\mathbf{\Gamma}$ represents the convolution by the kernel $\gamma = \mathbf{t} \star \check{\mathbf{t}}$.

$\mathbf{u}$

Samples of $\text{ADSN}(\mathbf{u})$



Image extracted from [Galerne et al., 2011a]⁵

⁵Galerne, B., Gousseau, Y., & Morel, J.-M. (2011a). Micro-texture synthesis by phase randomization. *Image Processing On Line*

⁶Galerne, B., Gousseau, Y., & Morel, J.-M. (2011b). Random Phase Textures: Theory and Synthesis. *IEEE Transactions on Image Processing*, 20(1), 257–267

The Asymptotic Discrete Spot Noise (ADSN) model [Galerne et al., 2011b] ⁶

Let $\mathbf{u} \in \mathbb{R}^{\Omega_{M,N}}$ be a grayscale image, m its grayscale mean and $\mathbf{t} = \frac{1}{\sqrt{MN}}(\mathbf{u} - m)$ its associated texton. Let $\mathbf{w}$ be a white Gaussian noise,

$$\mathbf{X} = \mathbf{t} \star \mathbf{w} \sim \text{ADSN}(\mathbf{u}) = \mathcal{N}(\mathbf{0}, \Gamma) \quad \text{which is a stationary law}$$

Γ represents the convolution by the kernel $\gamma = \mathbf{t} \star \check{\mathbf{t}}$.

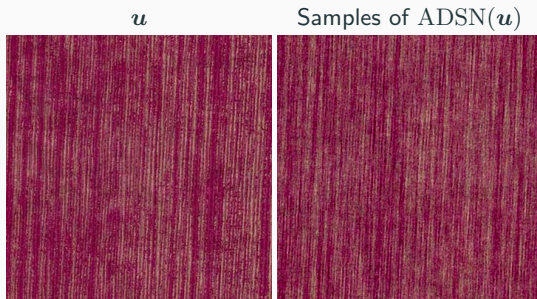


Image extracted from [Galerne et al., 2011a]⁵

⁵Galerne, B., Gousseau, Y., & Morel, J.-M. (2011a). Micro-texture synthesis by phase randomization. *Image Processing On Line*

⁶Galerne, B., Gousseau, Y., & Morel, J.-M. (2011b). Random Phase Textures: Theory and Synthesis. *IEEE Transactions on Image Processing*, 20(1), 257–267

The Asymptotic Discrete Spot Noise (ADSN) model [Galerne et al., 2011b] ⁶

Let $\mathbf{u} \in \mathbb{R}^{\Omega_{M,N}}$ be a grayscale image, m its grayscale mean and $\mathbf{t} = \frac{1}{\sqrt{MN}}(\mathbf{u} - m)$ its associated texton. Let $\mathbf{w}$ be a white Gaussian noise,

$$\mathbf{X} = \mathbf{t} \star \mathbf{w} \sim \text{ADSN}(\mathbf{u}) = \mathcal{N}(\mathbf{0}, \Gamma) \quad \text{which is a stationary law}$$

Γ represents the convolution by the kernel $\gamma = \mathbf{t} \star \check{\mathbf{t}}$.

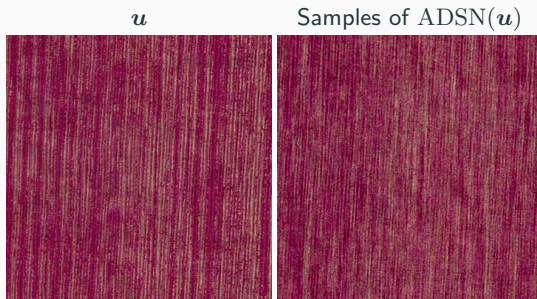


Image extracted from [Galerne et al., 2011a]⁵

⁵Galerne, B., Gousseau, Y., & Morel, J.-M. (2011a). Micro-texture synthesis by phase randomization. *Image Processing On Line*

⁶Galerne, B., Gousseau, Y., & Morel, J.-M. (2011b). Random Phase Textures: Theory and Synthesis. *IEEE Transactions on Image Processing*, 20(1), 257–267

The Asymptotic Discrete Spot Noise (ADSN) model [Galerne et al., 2011b] ⁶

Let $\mathbf{u} \in \mathbb{R}^{\Omega_{M,N}}$ be a grayscale image, m its grayscale mean and $\mathbf{t} = \frac{1}{\sqrt{MN}}(\mathbf{u} - m)$ its associated texton. Let $\mathbf{w}$ be a white Gaussian noise,

$$\mathbf{X} = \mathbf{t} \star \mathbf{w} \sim \text{ADSN}(\mathbf{u}) = \mathcal{N}(\mathbf{0}, \Gamma) \quad \text{which is a stationary law}$$

Γ represents the convolution by the kernel $\gamma = \mathbf{t} \star \check{\mathbf{t}}$.

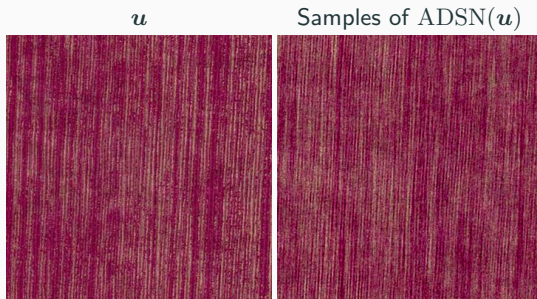


Image extracted from [Galerne et al., 2011a]⁵

⁵Galerne, B., Gousseau, Y., & Morel, J.-M. (2011a). Micro-texture synthesis by phase randomization. *Image Processing On Line*

⁶Galerne, B., Gousseau, Y., & Morel, J.-M. (2011b). Random Phase Textures: Theory and Synthesis. *IEEE Transactions on Image Processing*, 20(1), 257–267

The Asymptotic Discrete Spot Noise (ADSN) model [Galerne et al., 2011b] ⁶

Let $\mathbf{u} \in \mathbb{R}^{\Omega_{M,N}}$ be a grayscale image, m its grayscale mean and $\mathbf{t} = \frac{1}{\sqrt{MN}}(\mathbf{u} - m)$ its associated texton. Let $\mathbf{w}$ be a white Gaussian noise,

$$\mathbf{X} = \mathbf{t} \star \mathbf{w} \sim \text{ADSN}(\mathbf{u}) = \mathcal{N}(\mathbf{0}, \Gamma) \quad \text{which is a stationary law}$$

Γ represents the convolution by the kernel $\gamma = \mathbf{t} \star \check{\mathbf{t}}$.

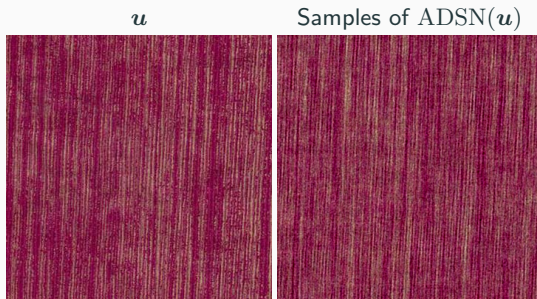


Image extracted from [Galerne et al., 2011a]⁵

⁵Galerne, B., Gousseau, Y., & Morel, J.-M. (2011a). Micro-texture synthesis by phase randomization. *Image Processing On Line*

⁶Galerne, B., Gousseau, Y., & Morel, J.-M. (2011b). Random Phase Textures: Theory and Synthesis. *IEEE Transactions on Image Processing*, 20(1), 257–267

The Asymptotic Discrete Spot Noise (ADSN) model [Galerne et al., 2011b] ⁶

Let $\mathbf{u} \in \mathbb{R}^{\Omega_{M,N}}$ be a grayscale image, m its grayscale mean and $\mathbf{t} = \frac{1}{\sqrt{MN}}(\mathbf{u} - m)$ its associated texton. Let $\mathbf{w}$ be a white Gaussian noise,

$$\mathbf{X} = \mathbf{t} \star \mathbf{w} \sim \text{ADSN}(\mathbf{u}) = \mathcal{N}(\mathbf{0}, \Gamma) \quad \text{which is a stationary law}$$

Γ represents the convolution by the kernel $\gamma = \mathbf{t} \star \check{\mathbf{t}}$.

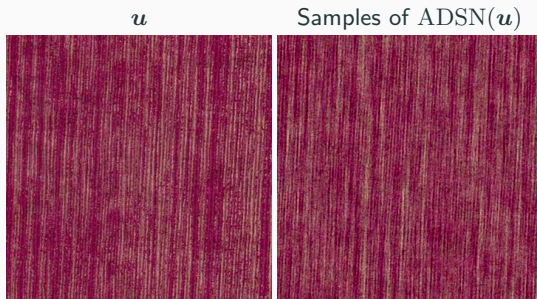
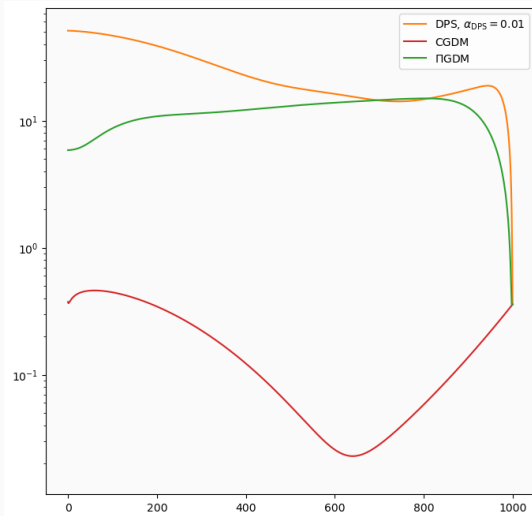
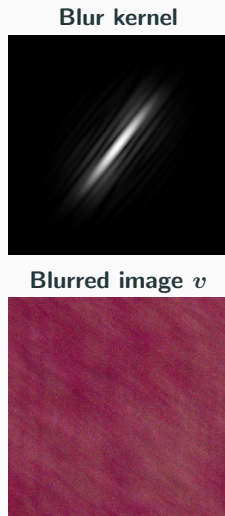


Image extracted from [Galerne et al., 2011a]⁵

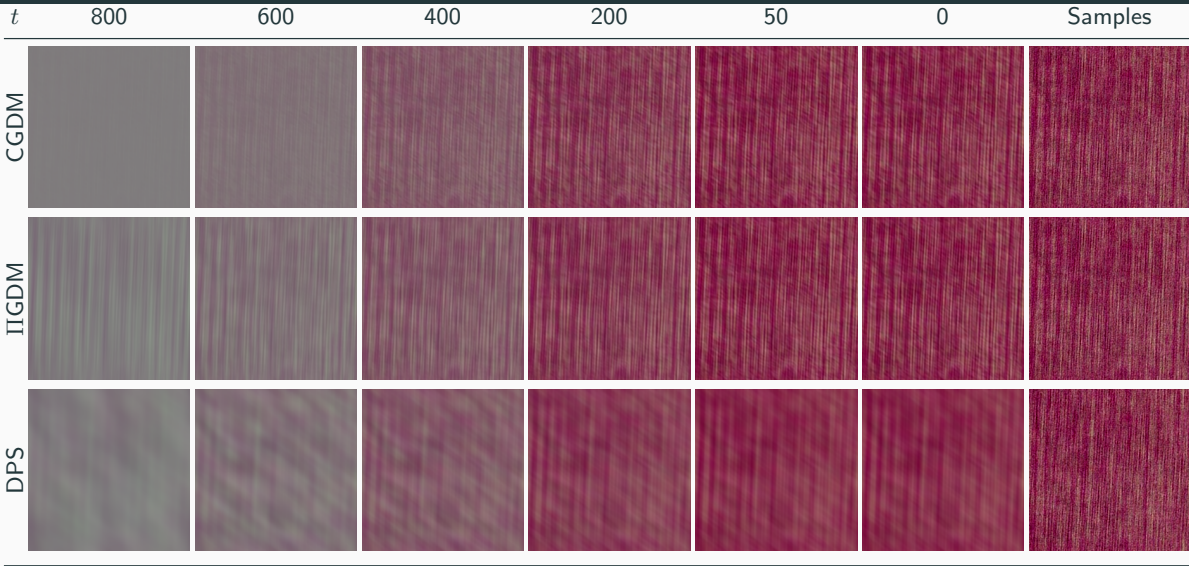
⁵Galerne, B., Gousseau, Y., & Morel, J.-M. (2011a). Micro-texture synthesis by phase randomization. *Image Processing On Line*

⁶Galerne, B., Gousseau, Y., & Morel, J.-M. (2011b). Random Phase Textures: Theory and Synthesis. *IEEE Transactions on Image Processing*, 20(1), 257–267

Exact Wasserstein error for deblurring



Study of the bias



We use

$$\nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t \mid \mathbf{v}) = \nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t) + \nabla_{\mathbf{x}} \log p_t(\mathbf{v} \mid \mathbf{x}_t), \quad (24)$$

where $\nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t)$ is known and $\nabla_{\mathbf{x}} \log p_t(\mathbf{v} \mid \mathbf{x}_t)$ is modeled (with a covariance matrix $\mathbf{C}_{\mathbf{v}|t}$).

We use

$$\nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t | \mathbf{v}) = \nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t) + \nabla_{\mathbf{x}} \log p_t(\mathbf{v} | \mathbf{x}_t), \quad (24)$$

where $\nabla_{\mathbf{x}} \log p_t(\mathbf{x}_t)$ is known and $\nabla_{\mathbf{x}} \log p_t(\mathbf{v} | \mathbf{x}_t)$ is modeled (with a covariance matrix $\mathbf{C}_{\mathbf{v}|\mathbf{t}}$). From these expressions, we can write $p_t(\mathbf{x}_t | \mathbf{v}) = \mathcal{N}(\boldsymbol{\mu}_{t|\mathbf{v}}, \mathbf{C}_{t|\mathbf{v}})$ with

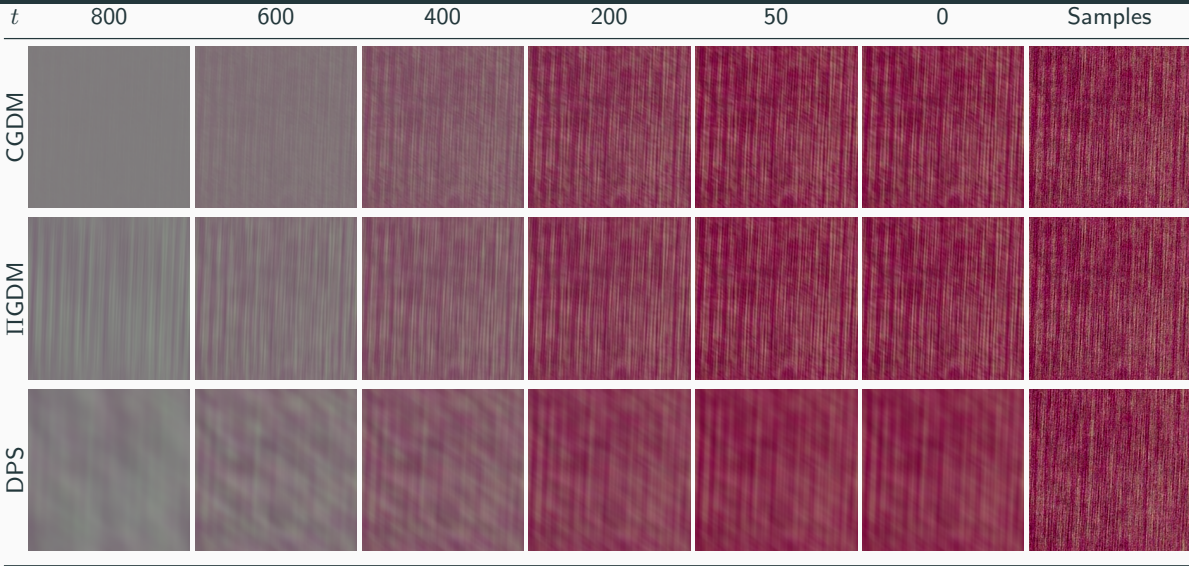
$$\mathbf{C}_{t|\mathbf{v}}^{\text{DPS}} = \bar{\alpha}_t \left[\boldsymbol{\Sigma} - \boldsymbol{\Sigma} \mathbf{A}^T \left(\frac{\sigma^2}{\alpha_{\text{DPS}}} \mathbf{I} + \bar{\alpha}_t \mathbf{A} \boldsymbol{\Sigma}^2 \boldsymbol{\Sigma}_t^{-1} \mathbf{A}^T \right)^{-1} \mathbf{A} \boldsymbol{\Sigma} \right] + (1 - \bar{\alpha}_t) \mathbf{I} \quad (25)$$

$$\mathbf{C}_{t|\mathbf{v}}^{\text{IIGDM}} = \bar{\alpha}_t \left[\boldsymbol{\Sigma} - \boldsymbol{\Sigma} \mathbf{A}^T \left(\sigma^2 \mathbf{I} + (1 - \bar{\alpha}_t) \mathbf{A} \mathbf{A}^T + \bar{\alpha}_t \mathbf{A} \boldsymbol{\Sigma}^2 \boldsymbol{\Sigma}_t^{-1} \mathbf{A}^T \right)^{-1} \mathbf{A} \boldsymbol{\Sigma} \right] + (1 - \bar{\alpha}_t) \mathbf{I} \quad (26)$$

$$\mathbf{C}_{t|\mathbf{v}}^{\text{CGDM}} = \bar{\alpha}_t \left[\boldsymbol{\Sigma} - \boldsymbol{\Sigma} \mathbf{A}^T (\sigma^2 \mathbf{I} + \mathbf{A} \boldsymbol{\Sigma} \mathbf{A}^T)^{-1} \mathbf{A} \boldsymbol{\Sigma} \right] + (1 - \bar{\alpha}_t) \mathbf{I} \quad (27)$$

At $t \approx 0$, $\mathbf{C}_{t|\mathbf{v}}^{\text{DPS}} \approx \mathbf{C}_{t|\mathbf{v}}^{\text{IIGDM}} \approx \mathbf{C}_{t|\mathbf{v}}^{\text{CGDM}}$.

Study of the bias



Conclusion

- Study of the DPS in practice,

$$\nabla_{\mathbf{x}_t} \log p(\mathbf{v} \mid \mathbf{x}_0 = \hat{\mathbf{x}}_0(\mathbf{x}_t)) \approx -\frac{\alpha_{\text{DPS}}}{2\sigma^2} \nabla_{\mathbf{x}_t} \|\mathbf{v} - \mathbf{A}\hat{\mathbf{x}}_0(\mathbf{x}_t)\|^{\color{red}1}. \quad (28)$$

- Generalization to other inverse problems.
- Generalization to multimodal distributions.
- An important direction of research [Rozet et al., 2024]⁷:

$$\text{Cov}(\mathbf{x} \mid \mathbf{x}_t) = \sigma_t^2 + \sigma_t^4 \nabla_{\mathbf{x}}^2 \log p_t(\mathbf{x}_t), \quad (29)$$

⁷Rozet, F., Andry, G., Lanusse, F., & Louppe, G. (2024). Learning diffusion priors from observations by expectation maximization. *The Thirty-eighth Annual Conference on Neural Information Processing Systems*

- Study of the DPS in practice,

$$\nabla_{\mathbf{x}_t} \log p(\mathbf{v} \mid \mathbf{x}_0 = \hat{\mathbf{x}}_0(\mathbf{x}_t)) \approx -\frac{\alpha_{\text{DPS}}}{2\sigma^2} \nabla_{\mathbf{x}_t} \|\mathbf{v} - \mathbf{A}\hat{\mathbf{x}}_0(\mathbf{x}_t)\|^{\color{red}1}. \quad (28)$$

- Generalization to other inverse problems.
- Generalization to multimodal distributions.
- An important direction of research [Rozet et al., 2024]⁷:

$$\text{Cov}(\mathbf{x} \mid \mathbf{x}_t) = \sigma_t^2 + \sigma_t^4 \nabla_{\mathbf{x}}^2 \log p_t(\mathbf{x}_t), \quad (29)$$

Thank you for your attention !

⁷Rozet, F., Andry, G., Lanusse, F., & Louppe, G. (2024). Learning diffusion priors from observations by expectation maximization. *The Thirty-eighth Annual Conference on Neural Information Processing Systems*

References

- Chung, H., Kim, J., Mccann, M. T., Klasky, M. L., & Ye, J. C. (2023). Diffusion posterior sampling for general noisy inverse problems. *The Eleventh International Conference on Learning Representations*.
- E.P & Galerne, B. (2025). Diffusion models for gaussian distributions: Exact solutions and wasserstein errors. *Forty-second International Conference on Machine Learning*.
- Galerne, B., Gousseau, Y., & Morel, J.-M. (2011a). Micro-texture synthesis by phase randomization. *Image Processing On Line*.
- Galerne, B., Gousseau, Y., & Morel, J.-M. (2011b). Random Phase Textures: Theory and Synthesis. *IEEE Transactions on Image Processing*, 20(1), 257–267.
- Ho, J., Jain, A., & Abbeel, P. (2020). Denoising diffusion probabilistic models. *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020*.
- Kawar, B., Elad, M., Ermon, S., & Song, J. (2022). Denoising diffusion restoration models. *ICLR Workshop on Deep Generative Models for Highly Structured Data*.

- Rozet, F., Andry, G., Lanusse, F., & Louppe, G. (2024). Learning diffusion priors from observations by expectation maximization. *The Thirty-eighth Annual Conference on Neural Information Processing Systems*.
- Song, J., Vahdat, A., Mardani, M., & Kautz, J. (2023). Pseudoinverse-guided diffusion models for inverse problems. *International Conference on Learning Representations*.