

Geometry-Aware Discretization Error of Diffusion Models

Samuel Hurault, Thomas Moreau, Gabriel Peyré

June, 19 2026



Generative models

given samples of p_{data} , generate new samples of p_{data}

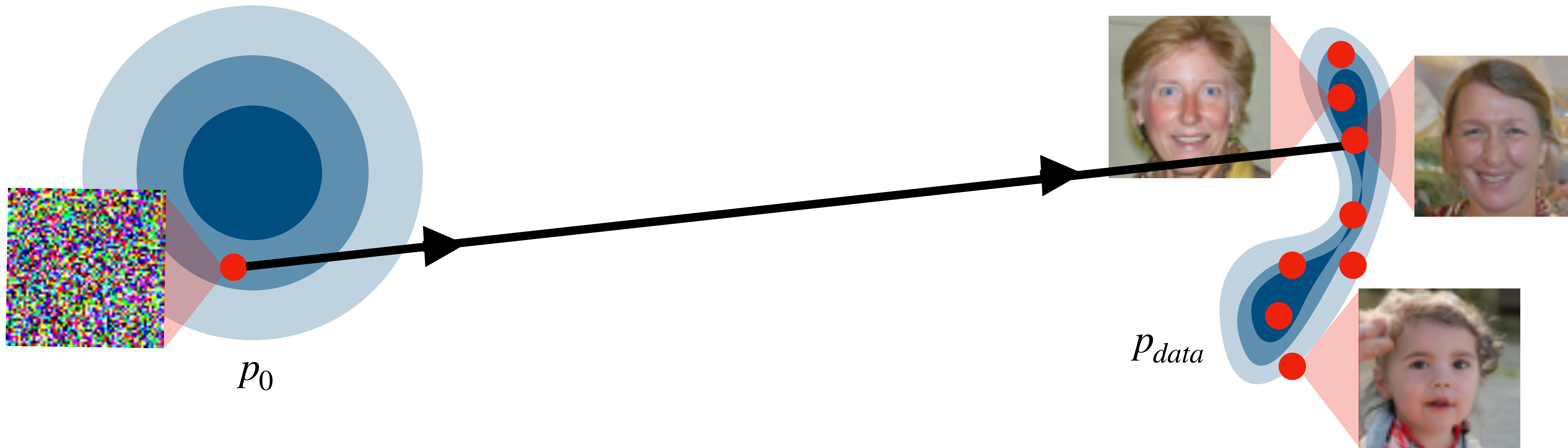
- Choose p_0 easy to sample
- Train T_θ for $T_\theta \# p_0 \approx p_{data}$

$$Y_0 \sim p_0, \quad Y = T_\theta(Y_0)$$

e.g. **Normalizing Flows**

[Rezende, & Mohamed, '15)]

$$\min_{\theta} KL(T_\theta \# p_0, p_{data})$$



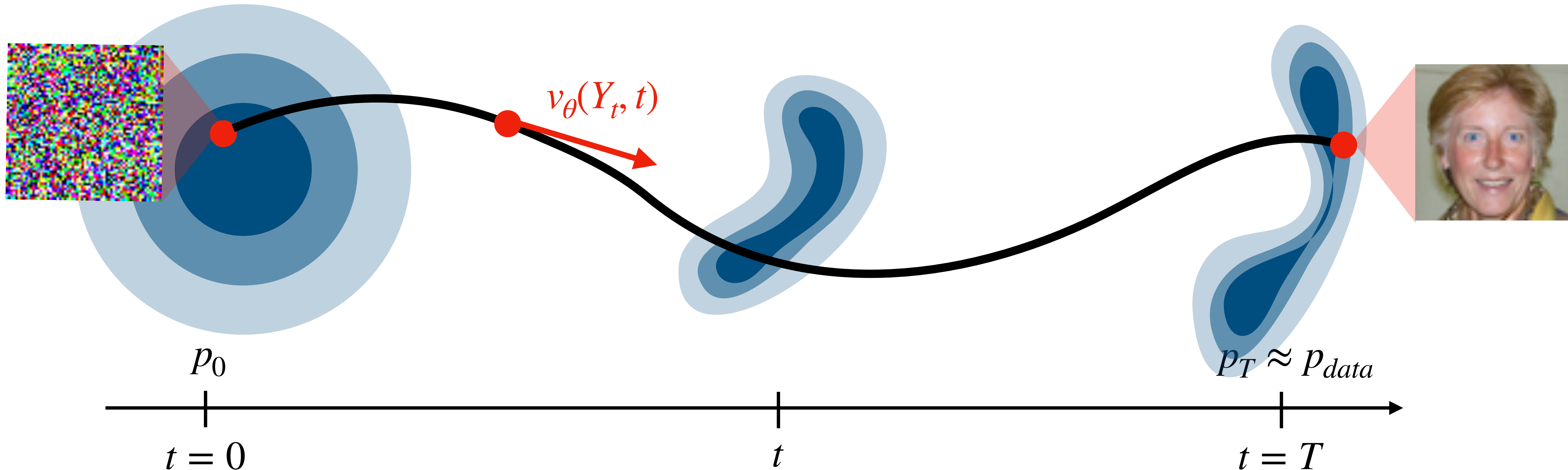
✗ Difficult to train a « single step » model : unstable training, mode collapse ...

Generative models

Diffusion models:

[Song et al., '21)]

$$Y_0 \sim p_0, \quad dY_t = v_\theta(Y_t, t) dt$$



✓ Continuous transformation
noise $\rightarrow$ data

✓ Easier sub-training at each time

$$\min_{\theta} \int_0^T \mathbb{E}_{x \sim p_{data}} \|v_\theta(x_t, t) - v^*(x_t, t)\|^2 dt$$

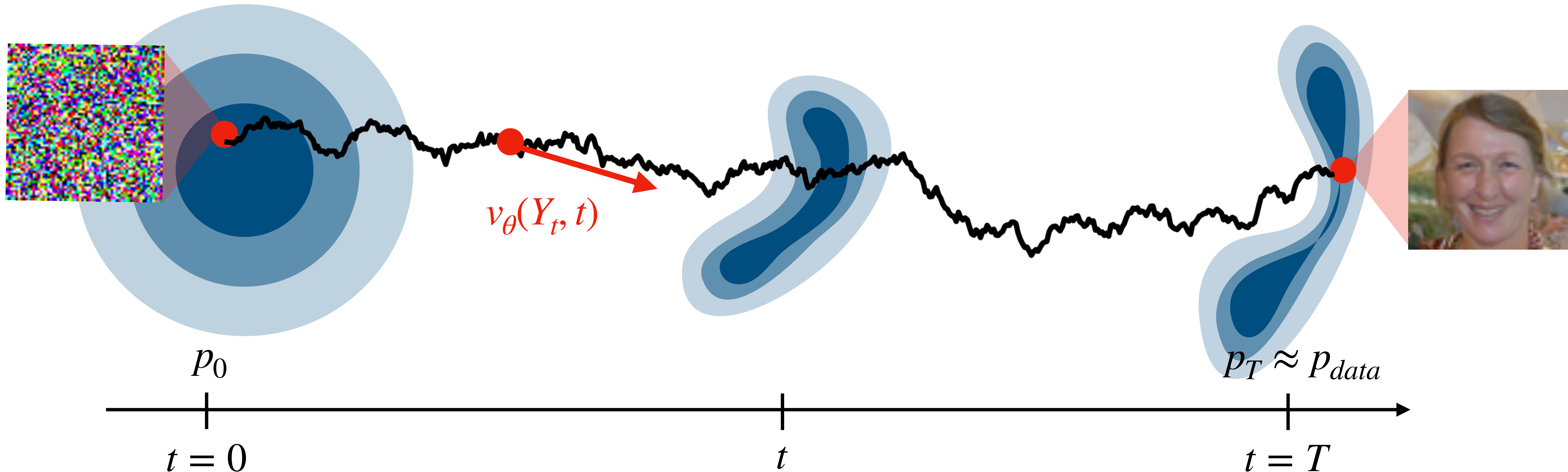
✗ Discretization error at
inference

Generative models

Diffusion models:

[Song et al., '21)]

$$Y_0 \sim p_0, \quad dY_t = v_\theta(Y_t, t) dt + \sqrt{2a_t} dW_t$$



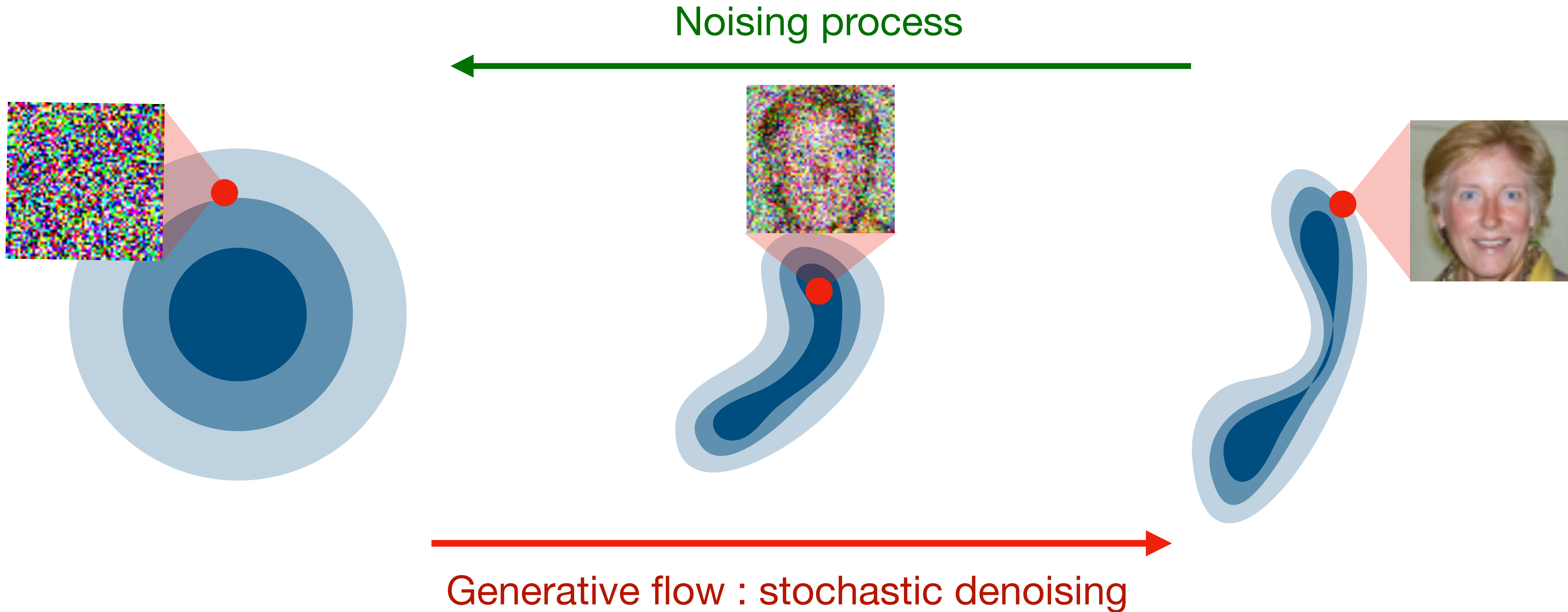
✓ Continuous transformation
noise $\rightarrow$ data

✓ Easier sub-training at each time

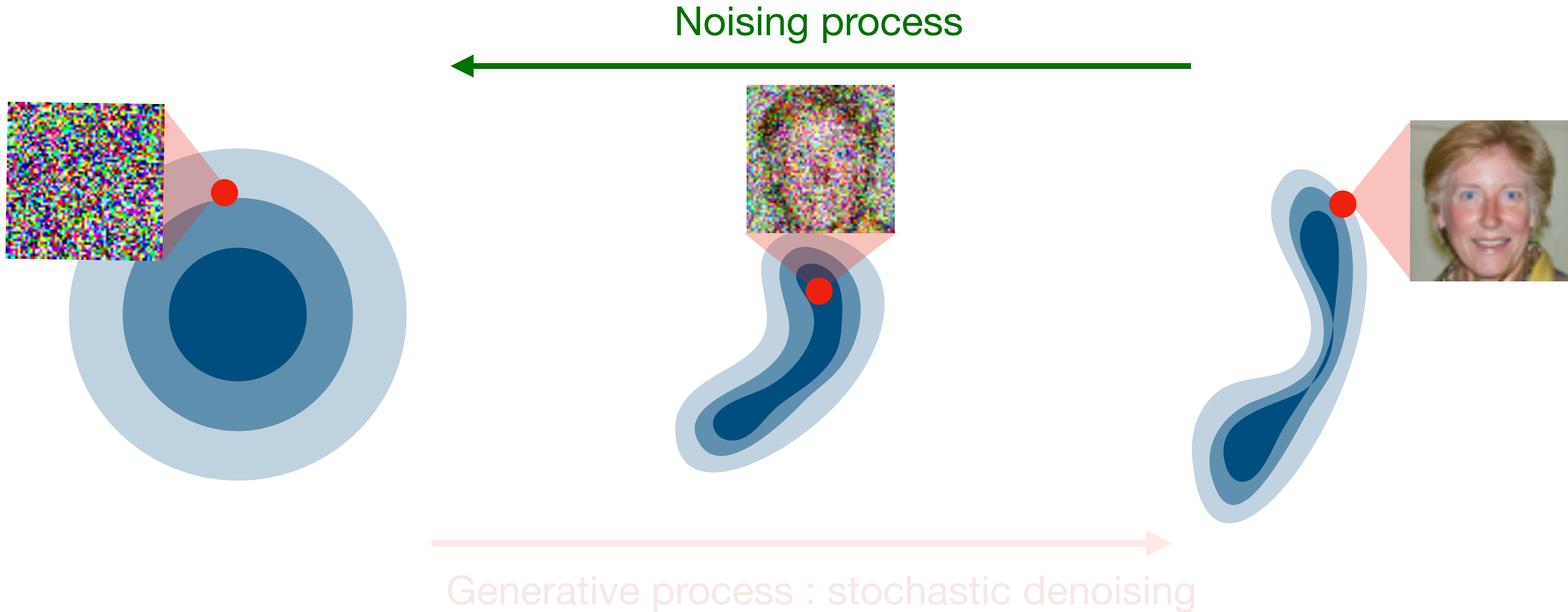
$$\min_{\theta} \int_0^T \mathbb{E}_{x \sim p_{data}} \|v_\theta(x_t, t) - v^*(x_t, t)\|^2 dt$$

✗ Discretization error at
inference

Diffusion models



Diffusion models

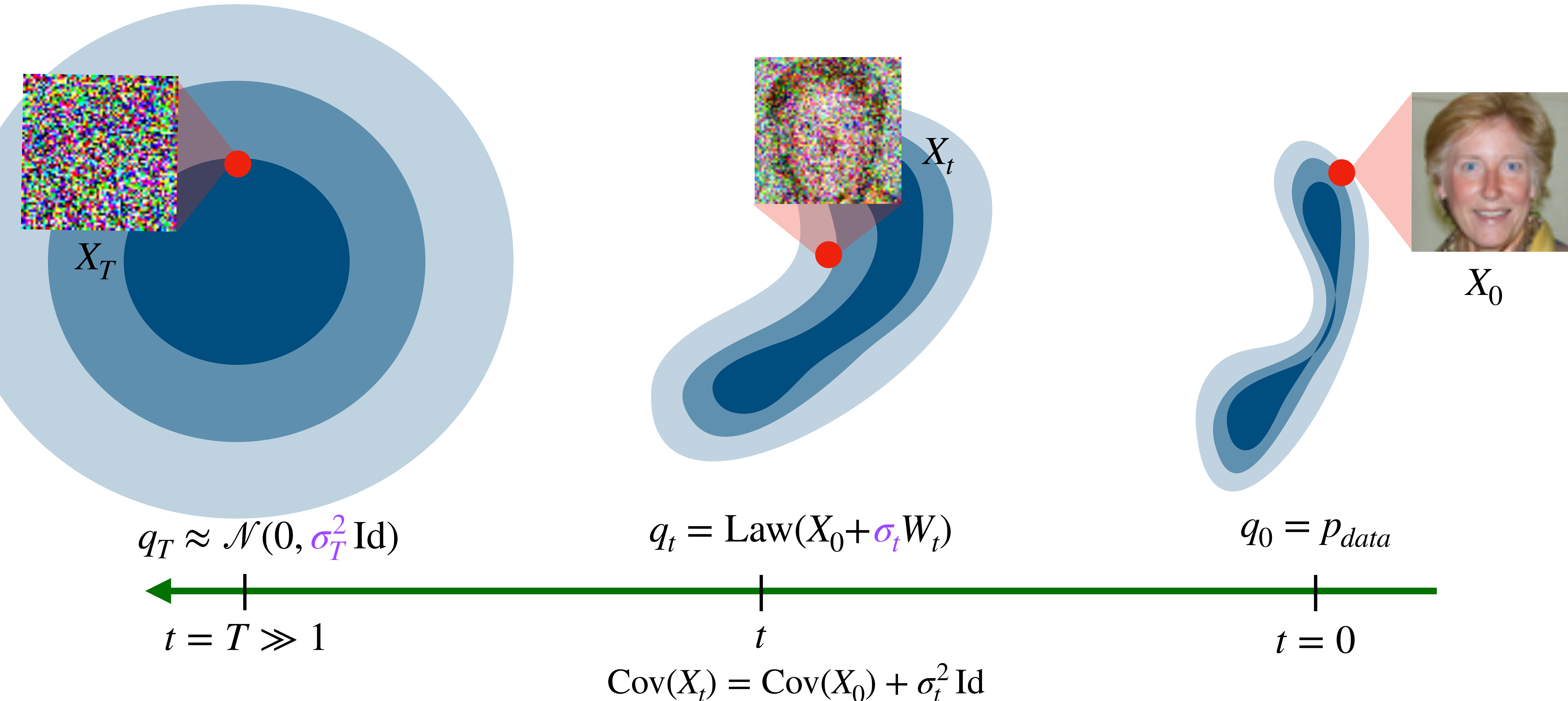


Diffusion models **Noising process** (from data to noise)

Variance Exploding (VE)

$$X_0 \sim p_{\text{data}}, \quad W_t \sim \mathcal{N}(0, \text{Id}), \quad X_t \stackrel{d}{=} X_0 + \sigma_t W_t$$

$$\sigma_t > 0, \quad \sigma_t \rightarrow \infty$$

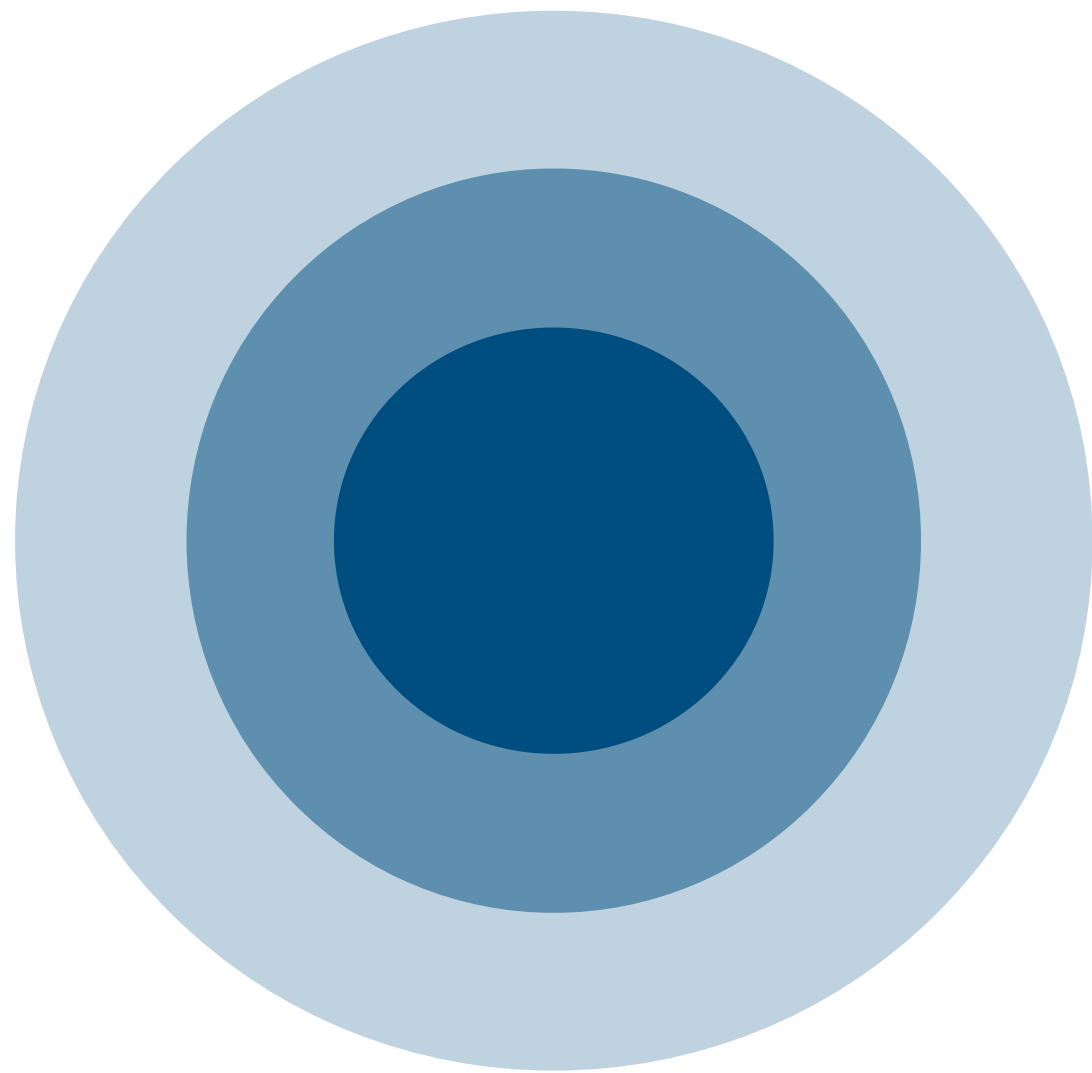


Diffusion models **Noising process** (from data to noise)

Variance Preserving (VP)

$$X_0 \sim p_{\text{data}}, W_t \sim \mathcal{N}(0, \text{Id}), X_t \stackrel{d}{=} \eta_t(X_0 + \sigma_t W_t)$$

$$\sigma_t > 0, \quad \eta_t = \frac{1}{\sqrt{\sigma_t^2 + 1}}$$



$$q_T \approx \mathcal{N}(0, \text{Id})$$



$$t = T \gg 1$$



$$q_t = \text{Law}(\eta_t(X_0 + \sigma_t w_t))$$

$$t$$

$$\text{Cov}(X_t) = \eta_t^2 \text{Cov}(X_0) + (1 - \eta_t^2) \text{Id}$$



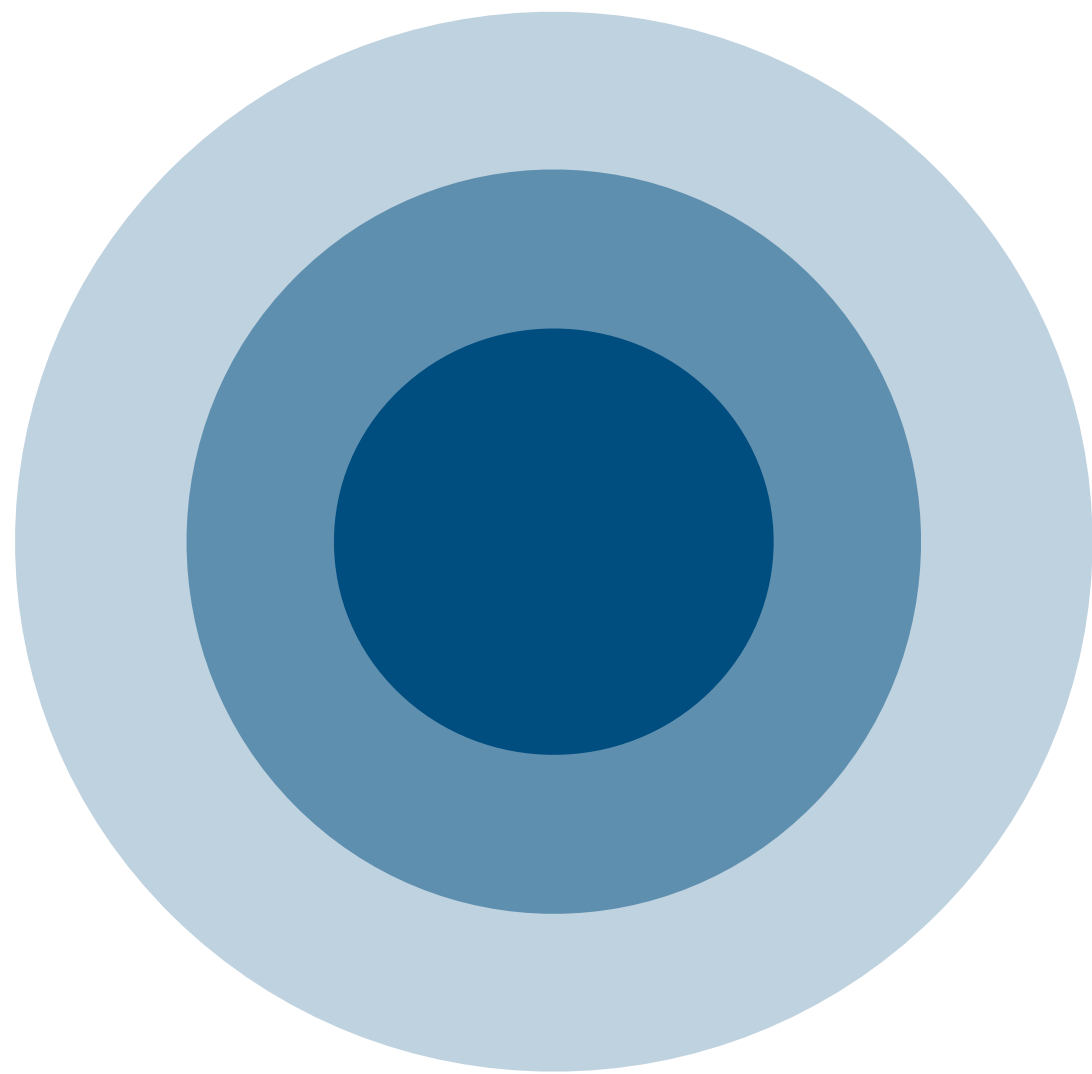
$$q_0 = p_{\text{data}}$$

$$t = 0$$

Diffusion models **Noising process** (from data to noise)

General schedule (η_t, σ_t)

$$X_0 \sim p_{\text{data}}, W_t \sim \mathcal{N}(0, \text{Id}), X_t \stackrel{d}{=} \eta_t(X_0 + \sigma_t W_t)$$



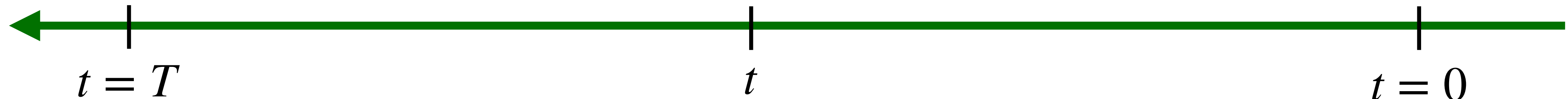
q_T



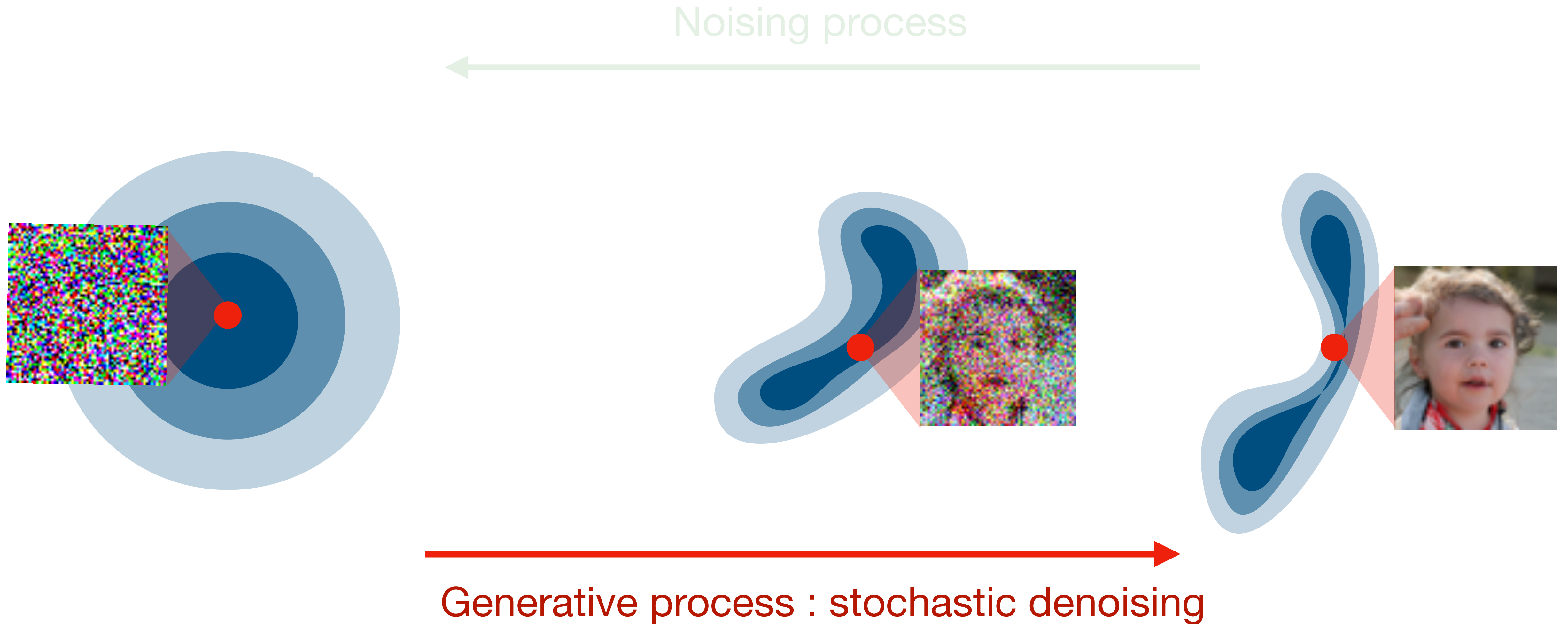
$q_t = \text{Law}(\eta_t(X_0 + \sigma_t W_t))$



$q_0 = p_{\text{data}}$



Diffusion models



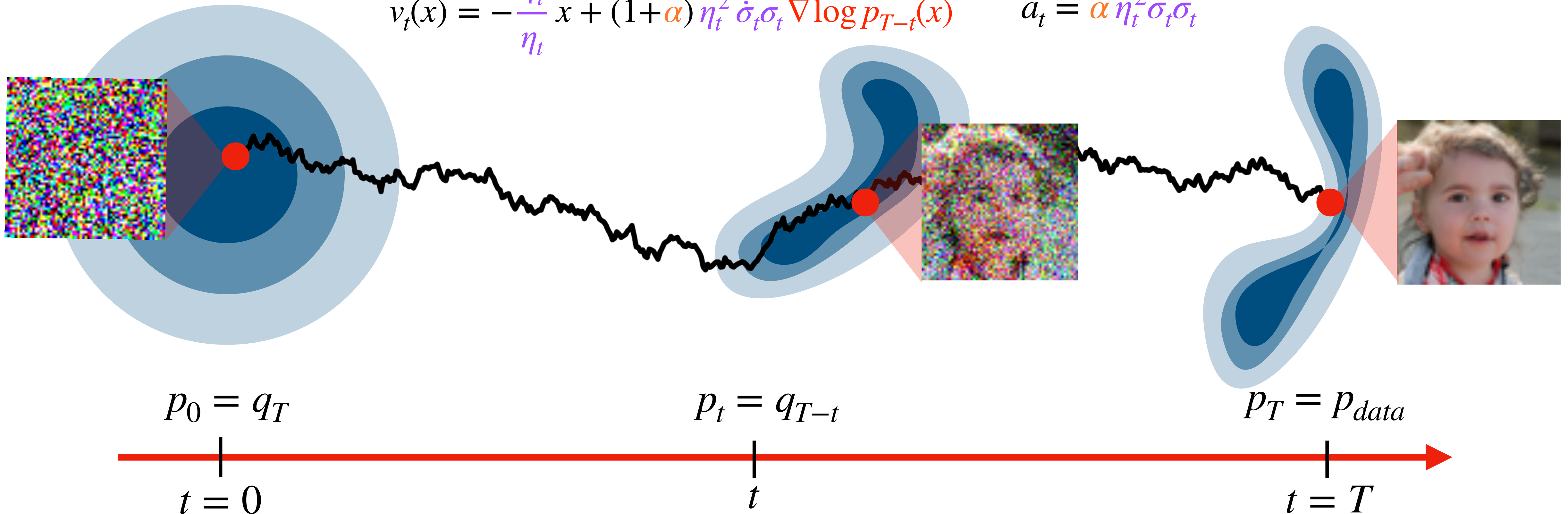
Diffusion models **Generative process** (from noise to data)

$p_t = q_{T-t}$ is sampled from the **SDE**

For $\alpha \geq 0$,

$$Y_0 \sim p_0, \quad dY_t = v_t(Y_t) dt + \sqrt{2a_t} dW_t$$

$$v_t(x) = -\frac{\dot{\eta}_t}{\eta_t} x + (1+\alpha) \eta_t^2 \dot{\sigma}_t \sigma_t \nabla \log p_{T-t}(x) \quad a_t = \alpha \eta_t^2 \dot{\sigma}_t \sigma_t$$



Parameters of the algorithm: diffusion coefficient α , time schedulers η_t, σ_t

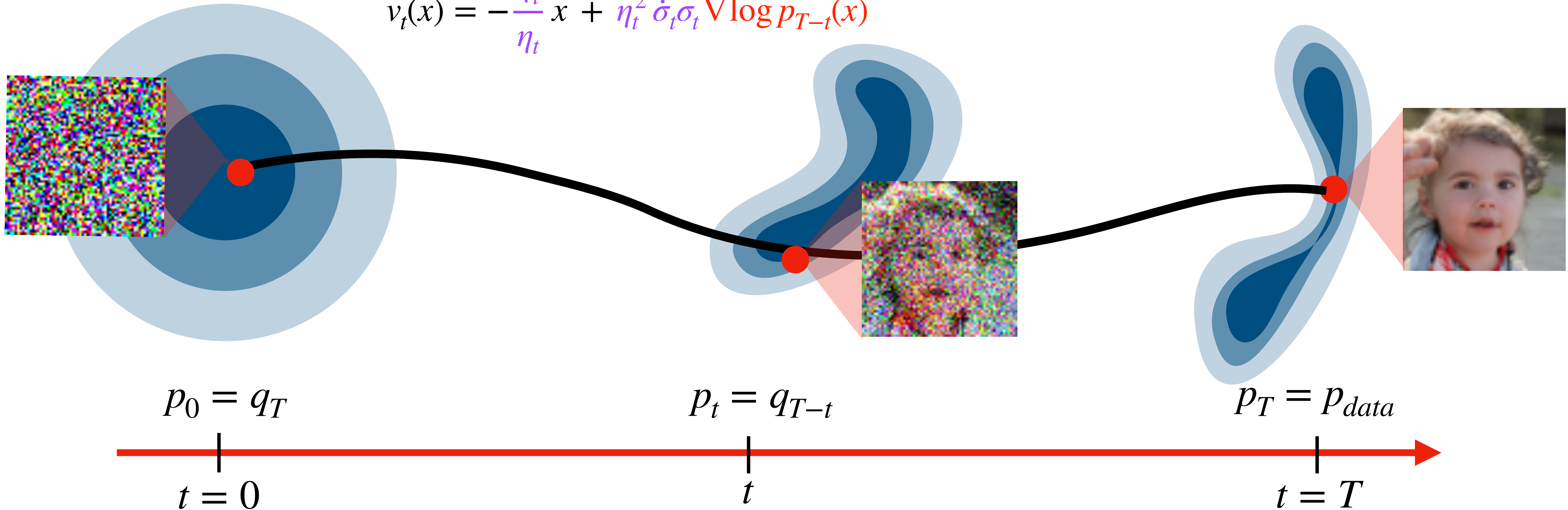
Diffusion models **Generative process** (from noise to data)

$p_t = q_{T-t}$ is sampled from the **ODE**

For $\alpha \geq 0$,

$$Y_0 \sim p_0, \quad dY_t = v_t(Y_t) dt$$

$$v_t(x) = -\frac{\dot{\eta}_t}{\eta_t} x + \eta_t^2 \dot{\sigma}_t \sigma_t \nabla \log p_{T-t}(x)$$

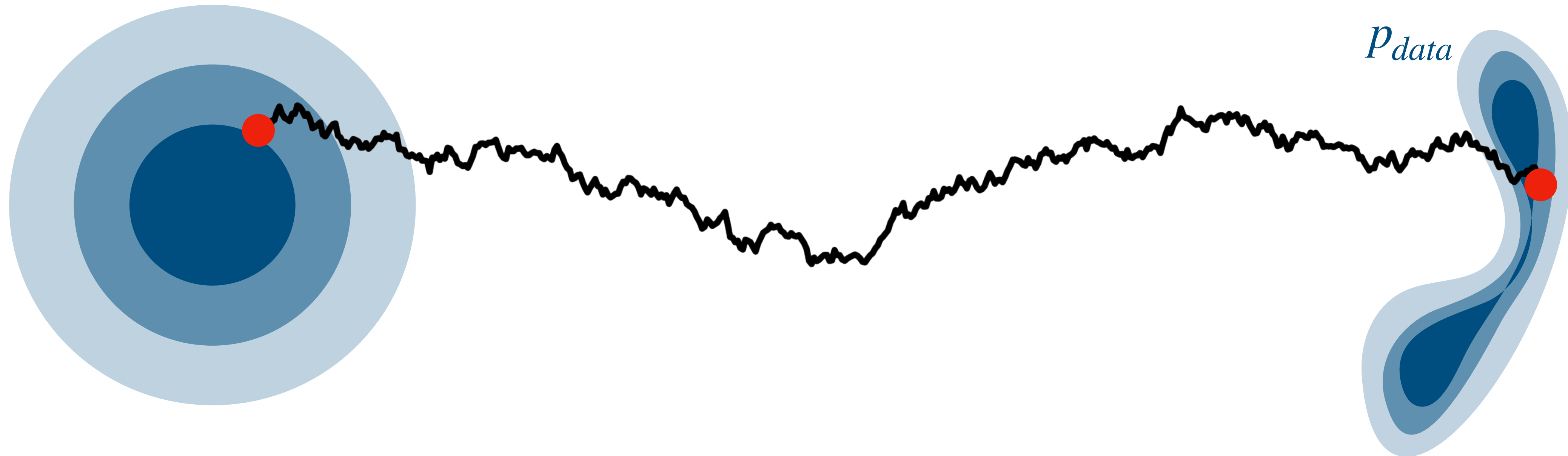


Parameters of the algorithm: diffusion coefficient α , time schedulers η_t, σ_t

Diffusion models errors

$$Y_0 \sim p_0, \quad dY_t = v_t(Y_t) dt + \sqrt{2a_t} dW_t$$

$$v_t(x) = -\frac{\dot{\eta}_{T-t}}{\eta_{T-t}} x + (1+\alpha) \eta_{T-t}^2 \dot{\sigma}_{T-t} \sigma_{T-t} \nabla \log p_{T-t}(x)$$



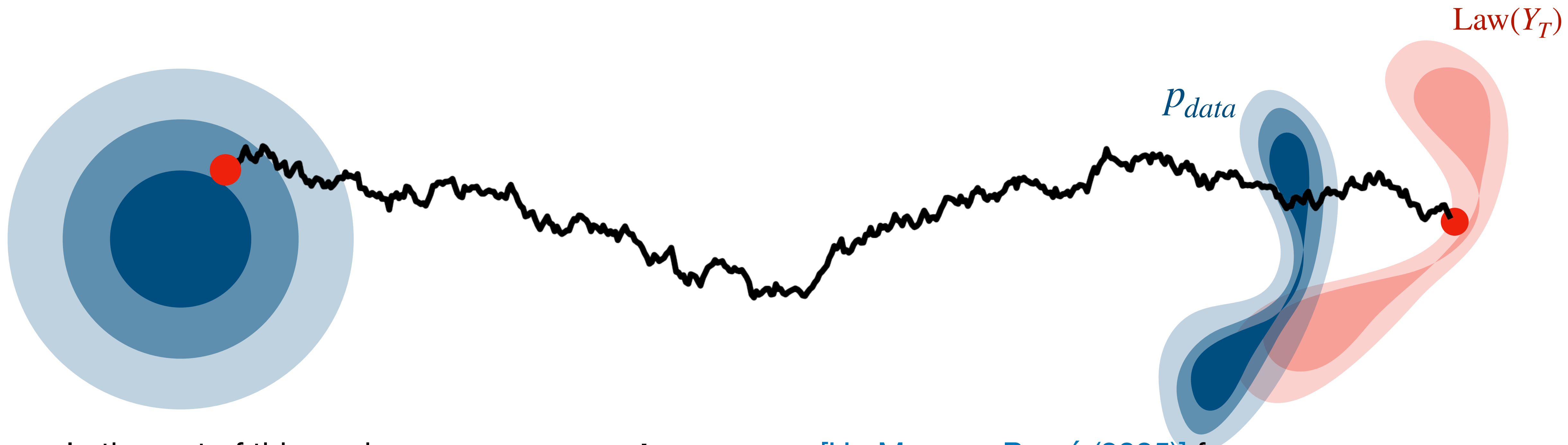
Diffusion models errors: **score training error**

$$Y_0 \sim p_0, \quad dY_t = v_t(Y_t) dt + \sqrt{2a_t} dW_t$$

$$v_t(x) = -\frac{\dot{\eta}_t}{\eta_t} x + (1+\alpha) \eta_t^2 \dot{\sigma}_t \nabla \log p_{T-t}(x)$$

$$s_t^\theta \approx \nabla \log p_t$$

Trained by denoising Score Matching
on a finite dataset $\{x_i \sim p\}_{i \leq N_{data}}$

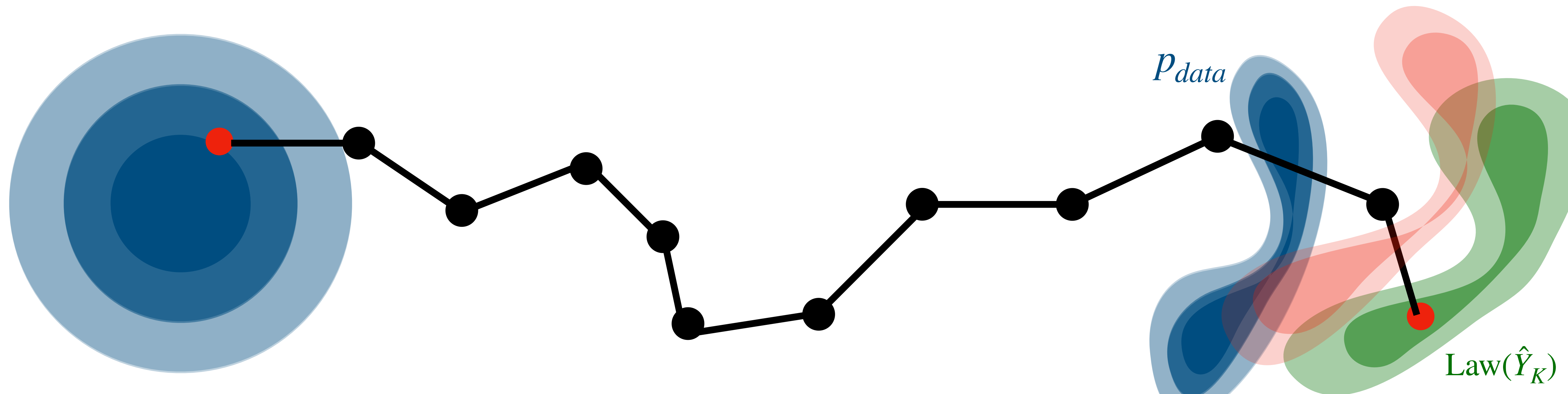


In the rest of this work, we assume **exact score**, see [H., Moreau, Peyré (2025)] for **score error**.

Diffusion models errors: discretization error

Euler discretization: for $\gamma = \frac{T}{K}$,

$$\hat{Y}_0 \sim p_0, \quad \hat{Y}_{k+1} = \hat{Y}_k + \gamma v_{t_k}(\hat{Y}_k) + \sqrt{2\gamma a_{t_k}} W_k$$



In the rest of this work, we assume **exact score**, see [H., Moreau, Peyré (2025)] for **score error**.

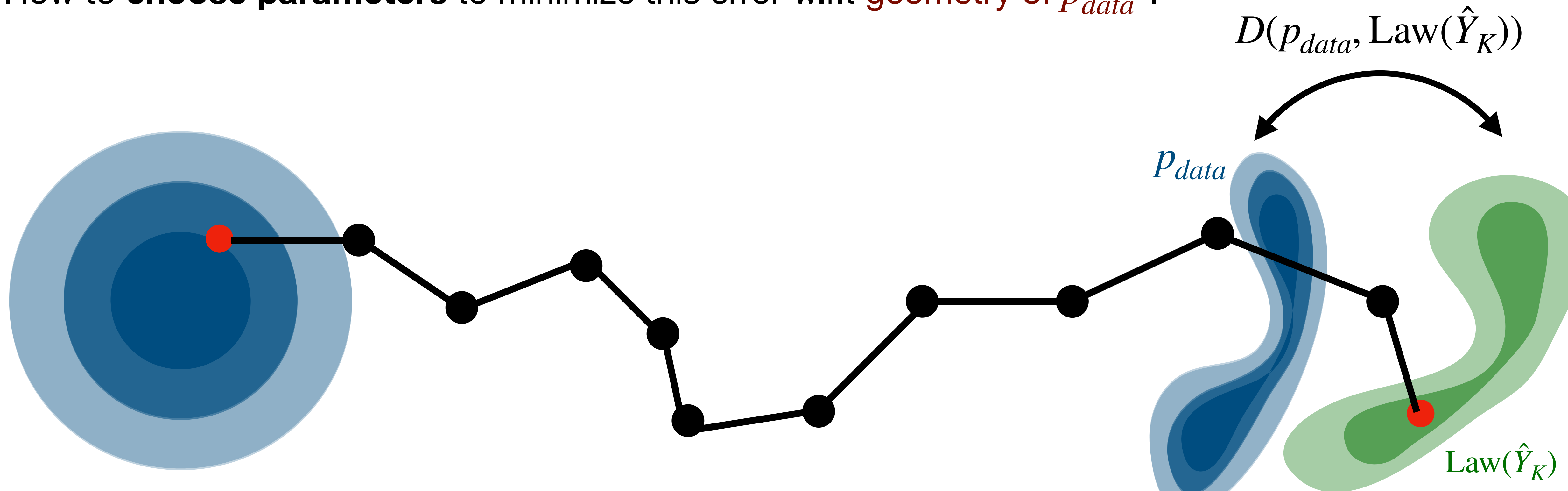
Objective of this work

$$\hat{Y}_0 \sim p_0, \quad \hat{Y}_{k+1} = \hat{Y}_k + \gamma v_{t_k}(\hat{Y}_k) + \sqrt{2\gamma a_{t_k}} W_k$$

We calculate explicitly how **discretization error** affects sampling:

$D(p_{data}, \text{law}(\hat{Y}_K))$ w.r.t. **diffusion parameters** and **geometry of p_{data}**

- How to **choose parameters** to minimize this error w.r.t **geometry of p_{data}** ?



In the rest of this work, we assume **exact score**, see [H., Moreau, Peyré (2025)] for **score error**.

Objective of this work

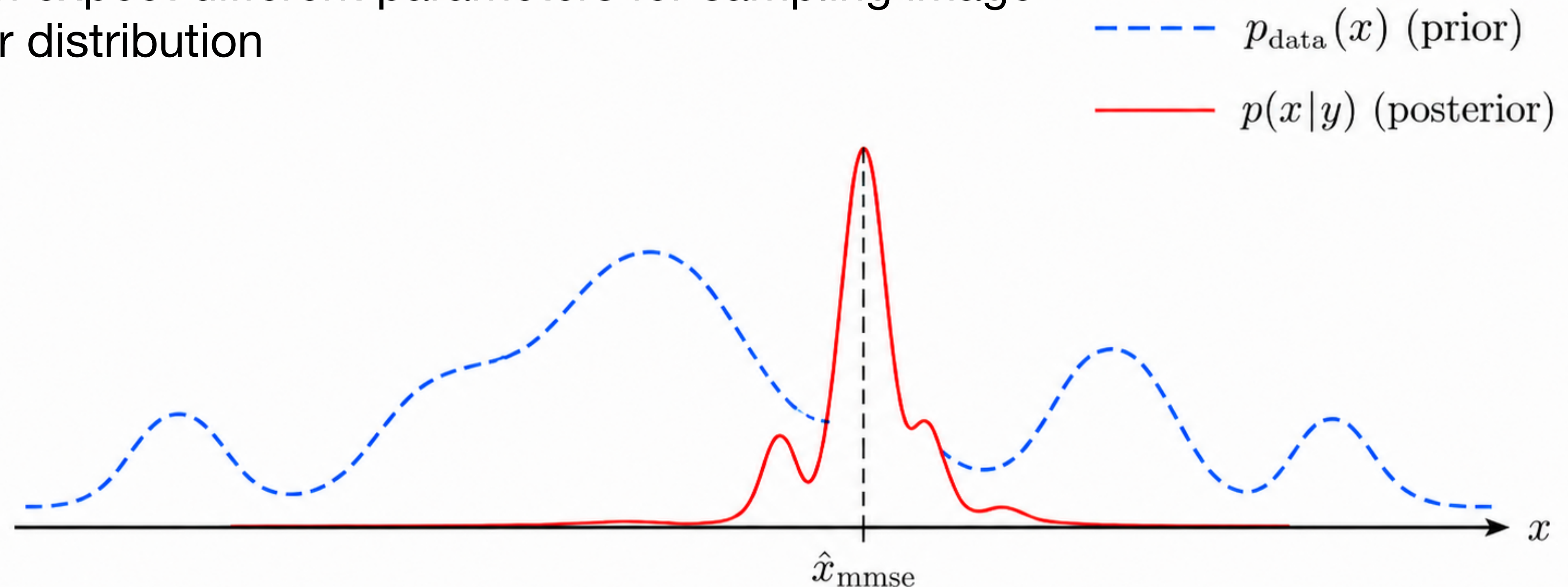
$$\hat{Y}_0 \sim p_0, \quad \hat{Y}_{k+1} = \hat{Y}_k + \gamma v_{t_k}(\hat{Y}_k) + \sqrt{2\gamma a_{t_k}} W_k$$

We calculate explicitly how **discretization error** affects sampling:

$D(p_{data}, \text{law}(\hat{Y}_K))$ w.r.t. **diffusion parameters** and **geometry of p_{data}**

- How to **choose parameters** to minimize this error w.r.t **geometry of p_{data}** ?

Example: we can expect different parameters for sampling image prior or posterior distribution



Objective of this work

$$\hat{Y}_0 \sim p_0, \quad \hat{Y}_{k+1} = \hat{Y}_k + \gamma v_{t_k}(\hat{Y}_k) + \sqrt{2\gamma a_{t_k}} W_k$$

We calculate explicitly how **discretization error** affects sampling:

$$D(p_{data}, \text{law}(\hat{Y}_K)) \text{ w.r.t. diffusion parameters and geometry of } p_{data}$$

- How to **choose parameters** to minimize this error w.r.t **geometry of p_{data}** ?

Existing works:

In a **general setting**, upper bounds e.g. [Benton et al.]

$$D(p_{data}, \text{Law}(\hat{Y}_K)) \lesssim \gamma \mathcal{E}_{disc}$$

$\mathcal{E}_{disc}$ w.r.t. dim, Lipschitz const. of score, bound on p_{data} ...

- Not tight
- No influence of **geometry of the data**.

Not adapted for parameter design

This work:

Assuming **Gaussian data** $p_{data} \sim \mathcal{N}(\mu_{data}, \Sigma_{data})$

$$W_2(p_{data}, \text{law}(\hat{Y}_K)) = \gamma \mathcal{E}_{disc} + o(\gamma)$$

with $\mathcal{E}_{disc}$ explicit w.r.t.

- sampling parameters (α, σ_t, η_t)
- **geometry of the data** : spectrum of Σ_{data}

Adapted for parameter

Discretization error: Gaussian case

Assuming $p_{\text{data}} = \mathcal{N}(0, \Sigma_{\text{data}} = U \text{Diag}(\lambda_i) U^\top)$,

$$\hat{Y}_0 \sim p_0, \quad \hat{Y}_{k+1} = \hat{Y}_k + \gamma v_{t_k}(\hat{Y}_k) + \sqrt{2\gamma a_{t_k}} W_k$$

$$v_t(x) = H_t x \quad \text{with} \quad H_t = U \text{Diag}\left(\Delta(\lambda_i)\right)_i U^\top$$

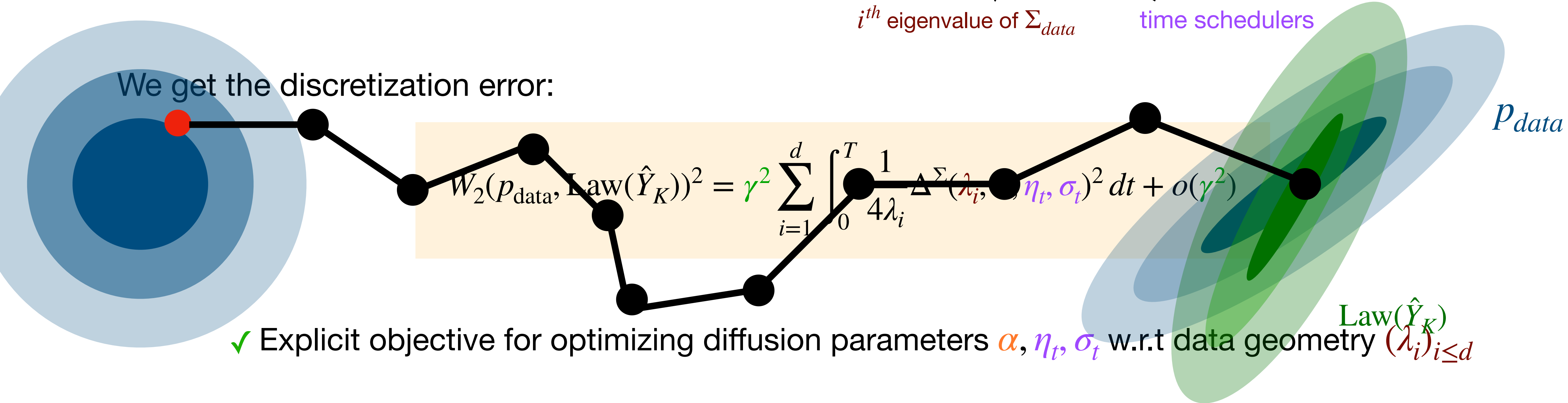
- Linear velocity $v_t \rightarrow$ samples a Gaussian $\text{Law}(\hat{Y}_K) = \mathcal{N}(\mathbb{E}[\hat{Y}_K], \text{Cov}[\hat{Y}_K])$
- v_t decomposes per eigendirections

$$\mathbb{E}[\hat{Y}_K] = \mu_{\text{data}} \qquad \text{Cov}[\hat{Y}_K] - \Sigma_{\text{data}} = \gamma U \text{Diag}\left(\int_0^T \Delta^\Sigma(\lambda_i, \alpha, \eta_t, \sigma_t) dt\right)_i U^\top + o(\gamma)$$

diffusion coefficient

i^{th} eigenvalue of Σ_{data}

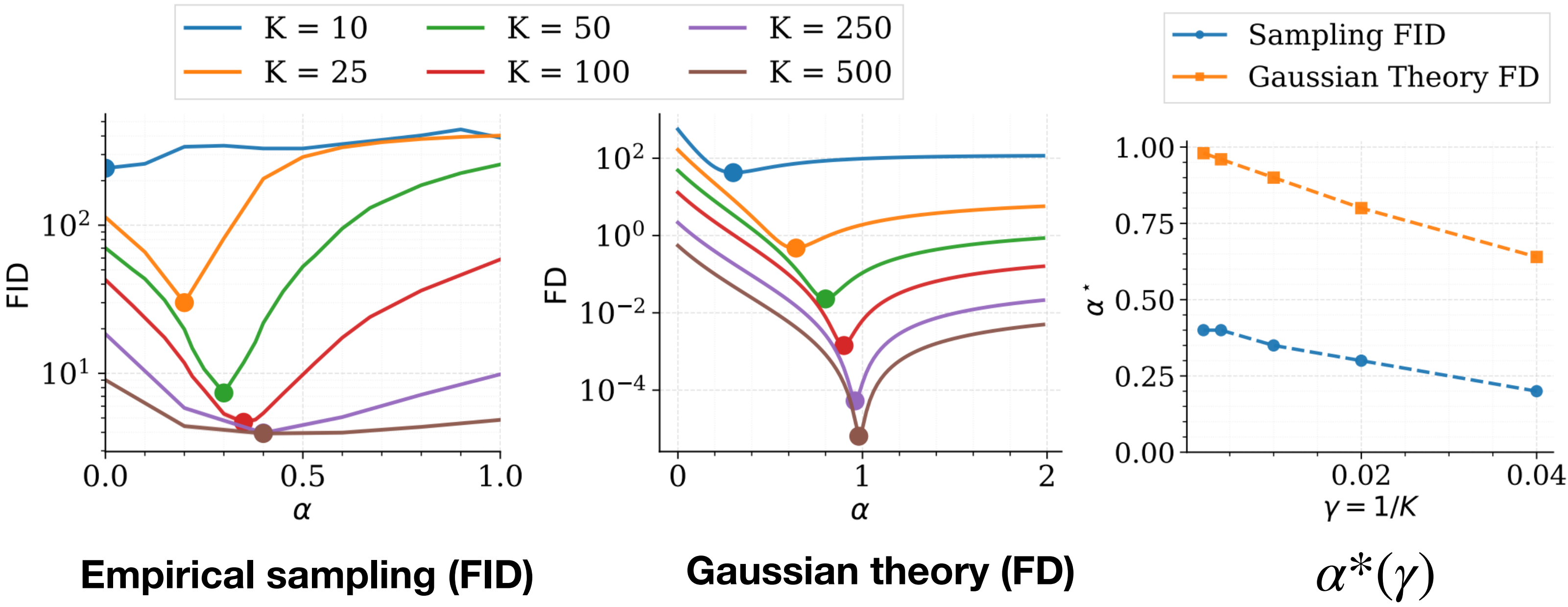
time schedulers



Optimization of the diffusion parameter α

Image dataset sampling (FFHQ) : EDM schedule [Karras et al., '23] : *Variance Exploding (VE)* $\eta_t = 1$, noise $\sigma_t = \sigma_{\max} \left(\frac{t}{T} \right)^\beta$.

Sampling error w.r.t number of discretization steps K (stepsize $\gamma = T/K$) :



$\alpha^\star$ decreases (linearly) with the stepsize γ

For *Variance Exploding (VE)* $\eta_t = 1$, and noise $\sigma_t = \sigma_{\max} \left(\frac{t}{T} \right)^\beta$ with $\sigma_{\max} \gg 1$

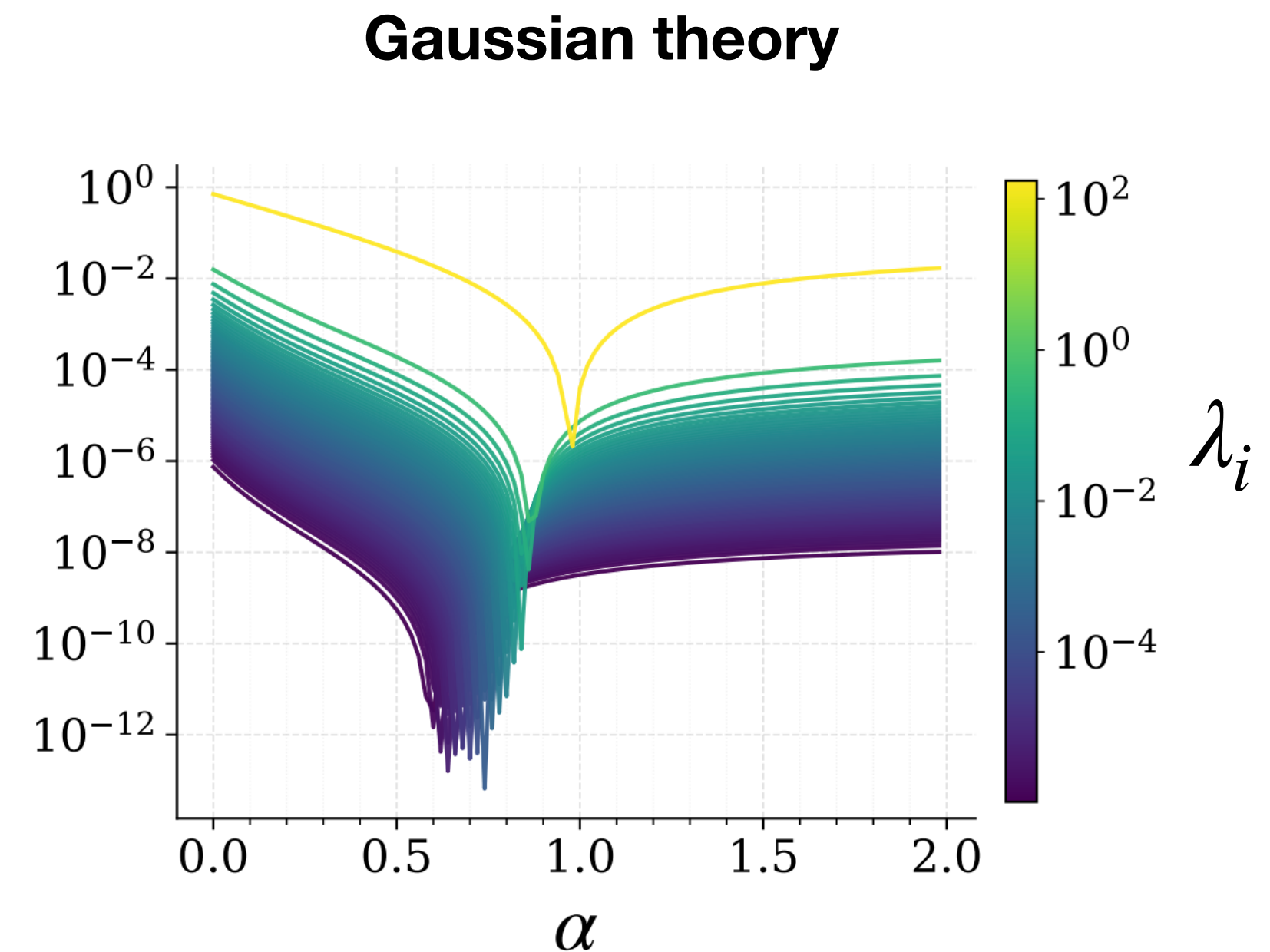
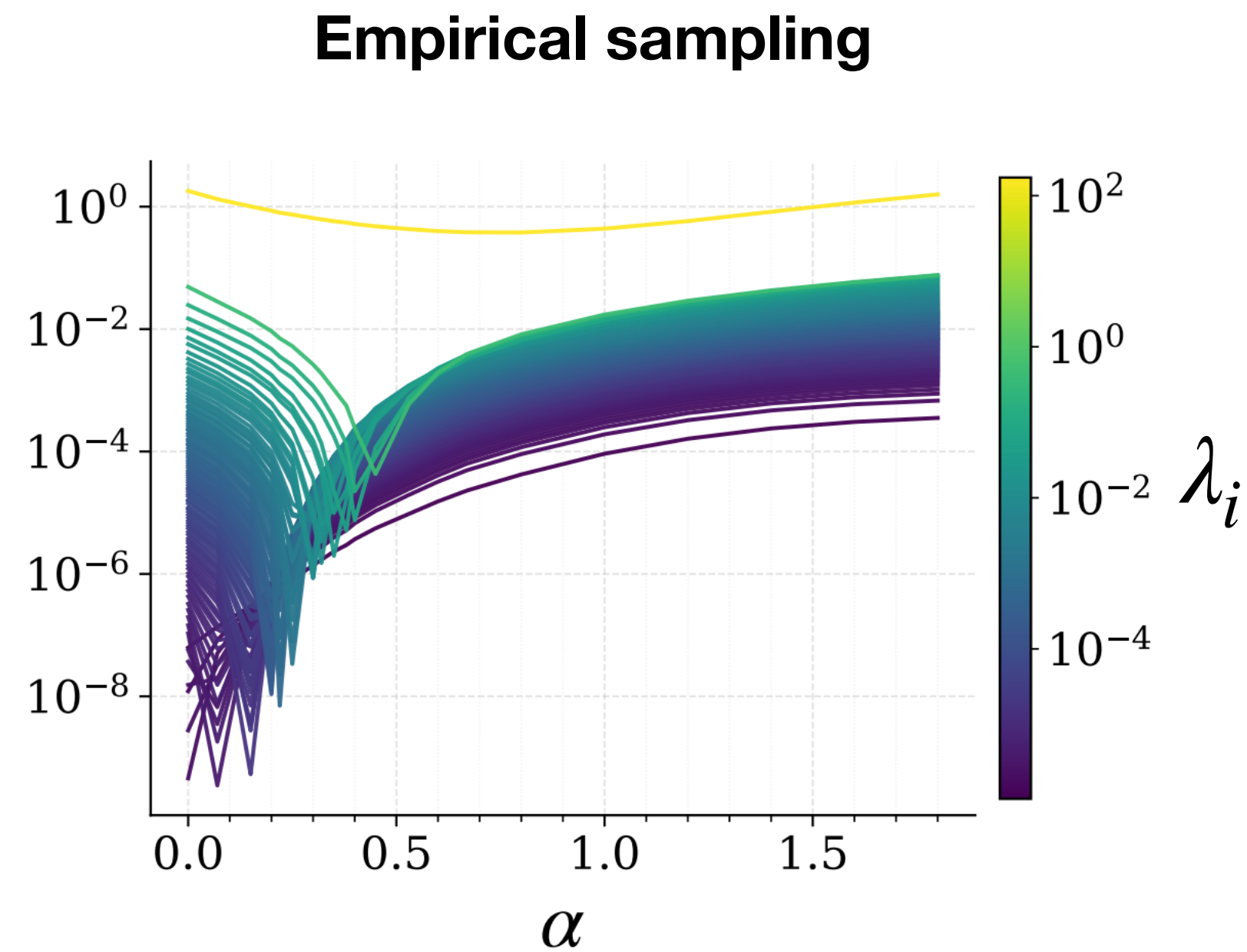
$$\alpha^\star(\gamma) = 1 - \gamma \sum_{i=1}^d w_i \lambda_i^{\frac{-1}{2\beta}} + o(\gamma)$$

The decrease slope depends on $(\lambda_i)_{i \leq d}$:
directions with smaller variance λ_i favor
smaller values of $\alpha^\star$

Optimization of the diffusion parameter α

Image dataset sampling (FFHQ) : EDM schedule [Karras et al., '23] : *Variance Exploding (VE)* $\eta_t = 1$, noise $\sigma_t = \sigma_{\max} \left(\frac{t}{T} \right)^\beta$.

Error decomposed per eigendirection of Σ_{data} :



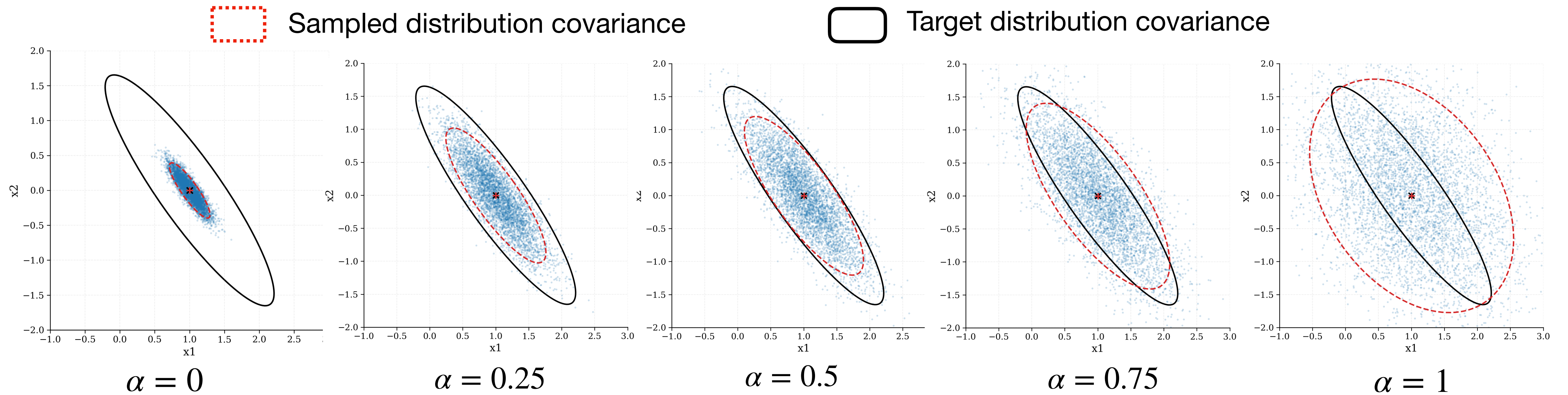
For *Variance Exploding (VE)* $\eta_t = 1$, and noise $\sigma_t = \sigma_{\max} \left(\frac{t}{T} \right)^\beta$ with $\sigma_{\max} \gg 1$

$$\alpha^\star(\gamma) = 1 - \gamma \sum_{i=1}^d w_i \lambda_i^{\frac{-1}{2\beta}} + o(\gamma)$$

The decrease slope depends on $(\lambda_i)_{i \leq d}$:
directions with smaller variance λ_i favor
smaller values of $\alpha^\star$

Optimization of the diffusion parameter α

Gaussian sampling in 10 steps



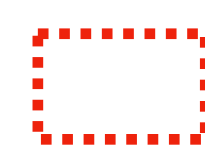
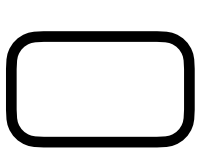
For *Variance Exploding (VE)* $\eta_t = 1$, and noise $\sigma_t = \sigma_{\max} \left(\frac{t}{T} \right)^\beta$ with $\sigma_{\max} \gg 1$

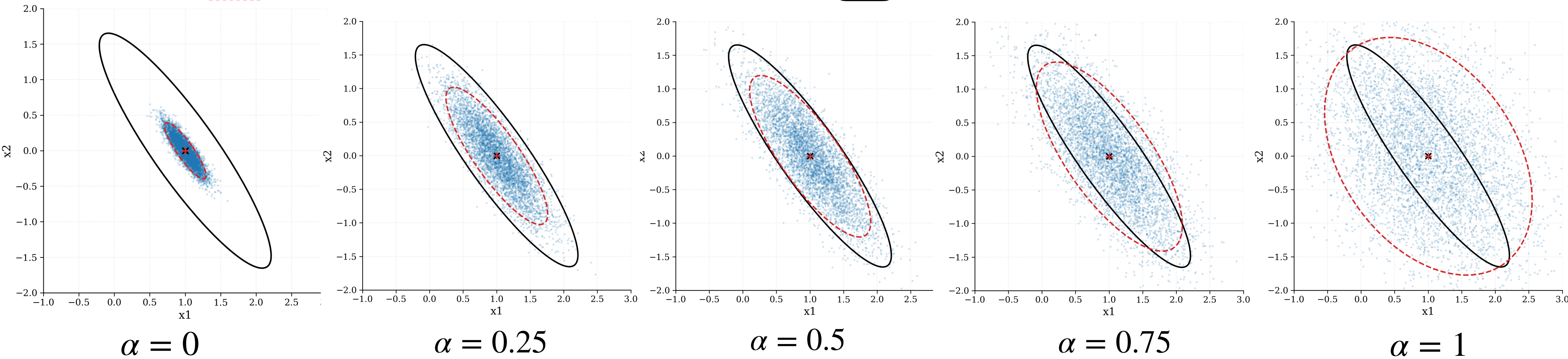
$$\alpha^\star(\gamma) = 1 - \gamma \sum_{i=1}^d w_i \lambda_i^{\frac{-1}{2\beta}} + o(\gamma)$$

The decrease slope depends on $(\lambda_i)_{i \leq d}$:
directions with smaller variance λ_i favor
smaller values of $\alpha^\star$

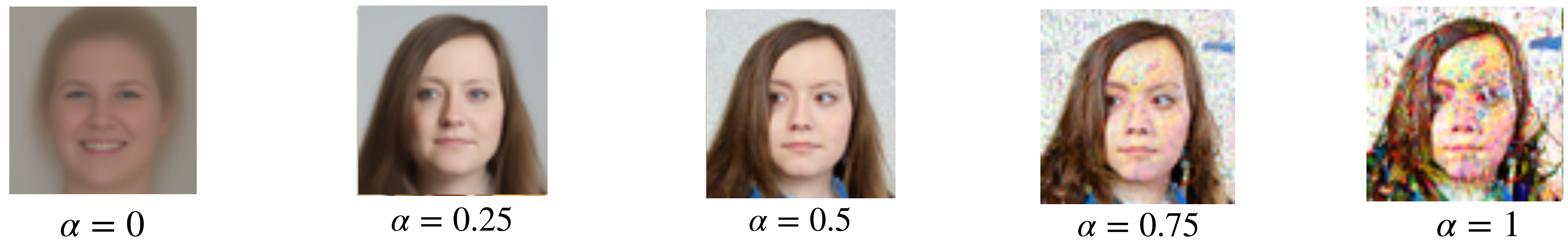
Optimization of the diffusion parameter α

Gaussian sampling in 10 steps

 Sampled distribution covariance  Target distribution covariance



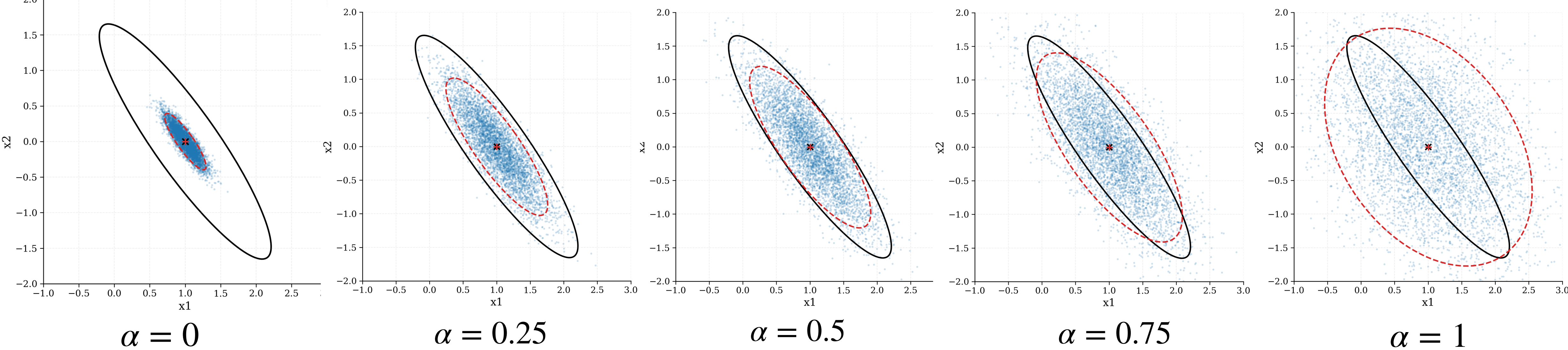
FFHQ sampling:



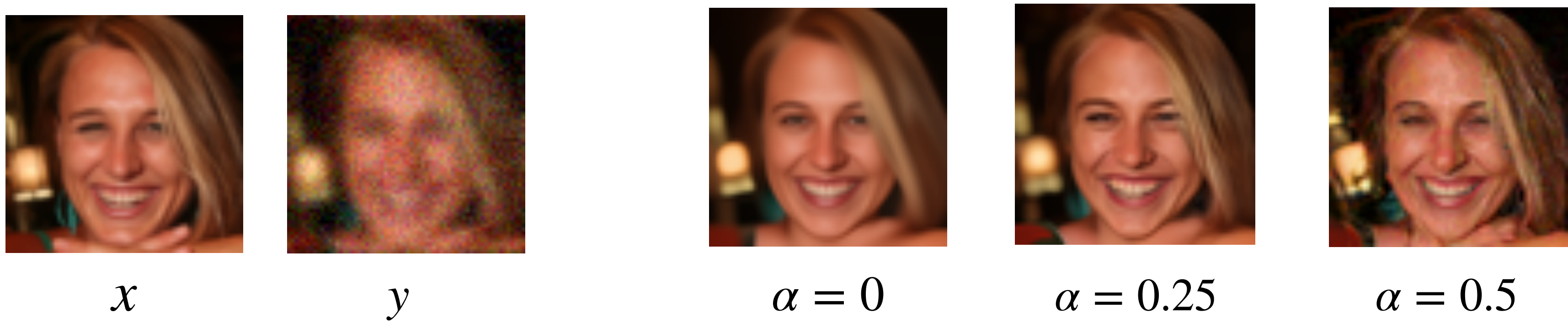
► For small number of steps, choosing $\alpha = 0$ samples close to the data mean

Optimization of the diffusion parameter α

Gaussian sampling in 10 steps



Posterior sampling: sample the posterior $p(x | y)$



Discretization error allows **MMSE estimation** using diffusion sampling

► For small number of steps, choosing $\alpha = 0$ samples close to the **MMSE**

Optimization of the rescaling schedule η_t

Classical choices of rescaling schedules:

We choose the parametrized

$$\eta_t(c) = \sqrt{\frac{c}{c + \sigma_t^2}}$$

$$\eta_t^{VE} = 1, \quad \eta_t^{VP} = \sqrt{\frac{1}{1 + \sigma_t^2}}$$

- The forward noising process converges to $\mathcal{N}(0, c \mathbf{I})$
- Variance Preserving (VP) : $c = 1$
- Variance Exploding (VE) : $c = \infty$

And look for $c^\star = \operatorname{argmin}_{c>0} W_2(p_{\text{data}}, \text{Law}(\hat{Y}_K))$

Assuming a power spectrum $\lambda_i = \lambda_{\max} i^{-p}$ and noise $\sigma_t = \sigma_{\max} \left(\frac{t}{T}\right)^\beta$ with large $\sigma_{\max}$ and β

$$c^\star = \lambda_{\max} \tau(p) \quad \text{with } \tau(p) > 0 \text{ decreasing for } p \lesssim 1$$

Optimization of the rescaling schedule η_t

$$\eta_t(c) = \sqrt{\frac{c}{c + \sigma_t^2}}$$

Assuming a power spectrum $\lambda_i = \lambda_{\max} i^{-p}$ and noise $\sigma_t = \sigma_{\max} \left(\frac{t}{T}\right)^\beta$ with large $\sigma_{\max}$ and β

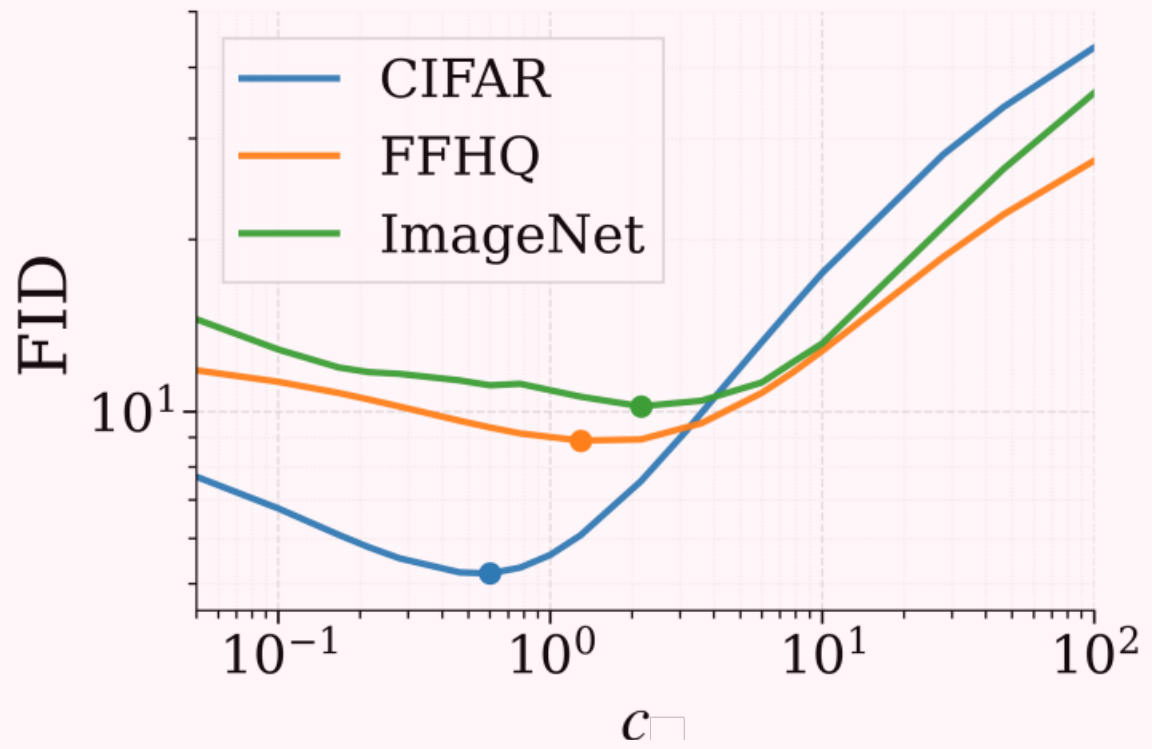
$c^\star = \lambda_{\max} \tau(p)$ with $\tau(p) > 0$ decreasing for $p \lesssim 1$

Full image datasets sampling

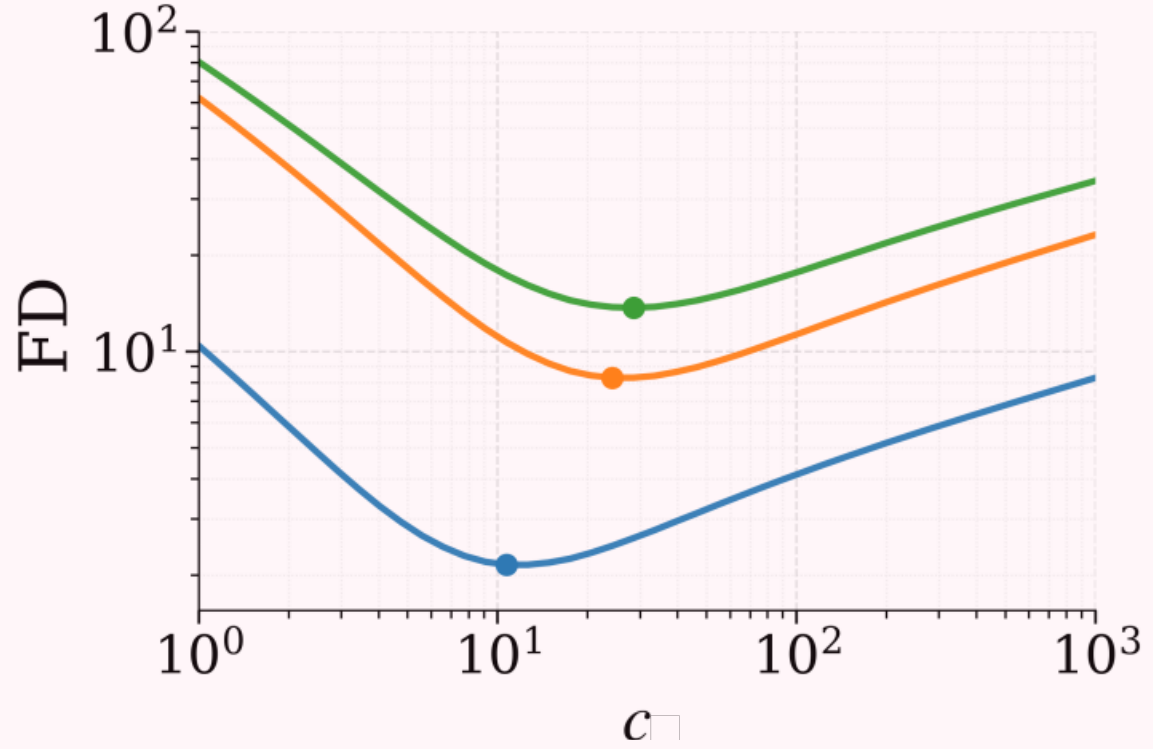
Small relative change of p across datasets

Influence of $\lambda_{\max}$

Empirical FID



Gaussian theory



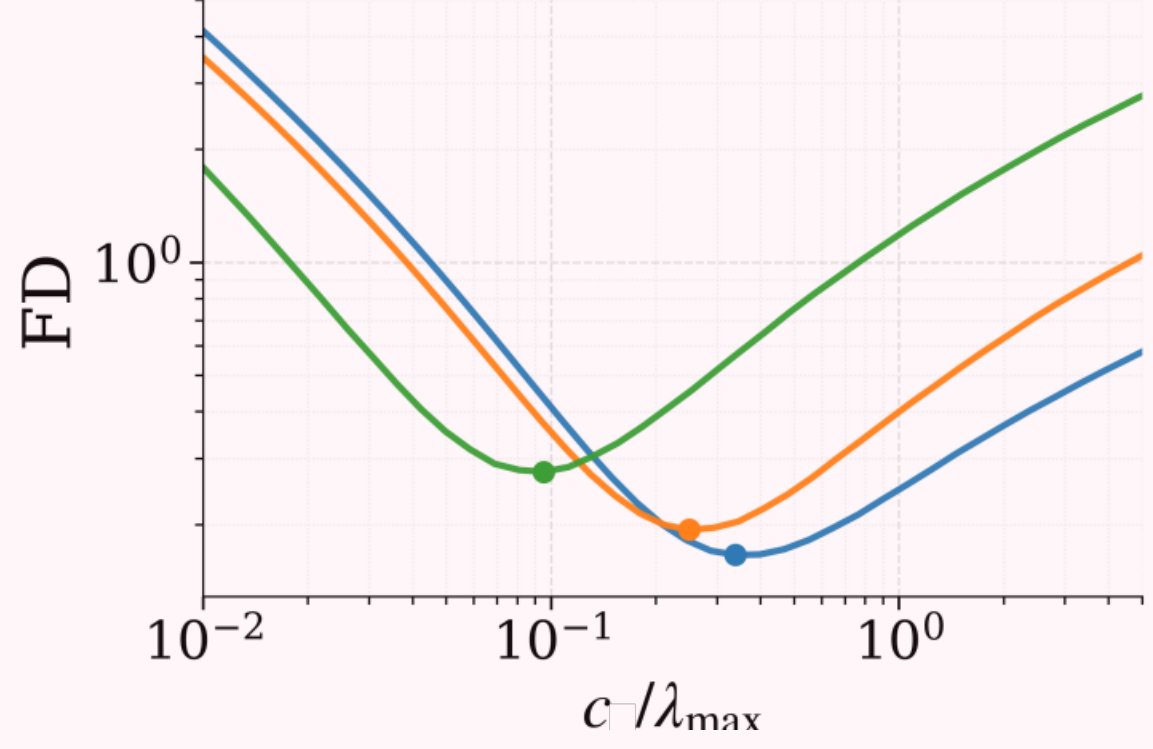
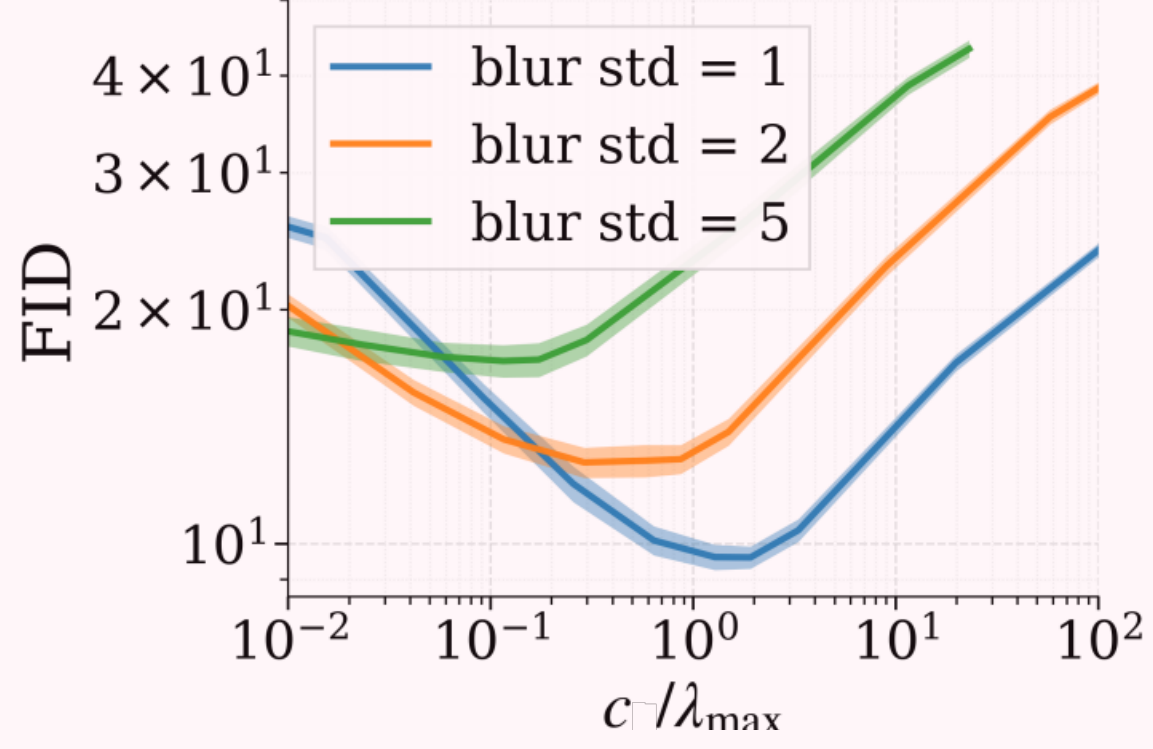
Dataset	$\lambda_{\max}$	p
CIFAR	55.3	1.43
FFHQ	173	1.47
ImageNet	225	1.35

Deblurring Posterior sampling

Sampling of $p(x | y = k \star x + w)$

c-axis rescaled by $\lambda_{\max}$

Influence of p



blur std	$\lambda_{\max}$	p
1	0.39	0.57
2	0.86	0.61
5	0.78	0.83

Conclusion

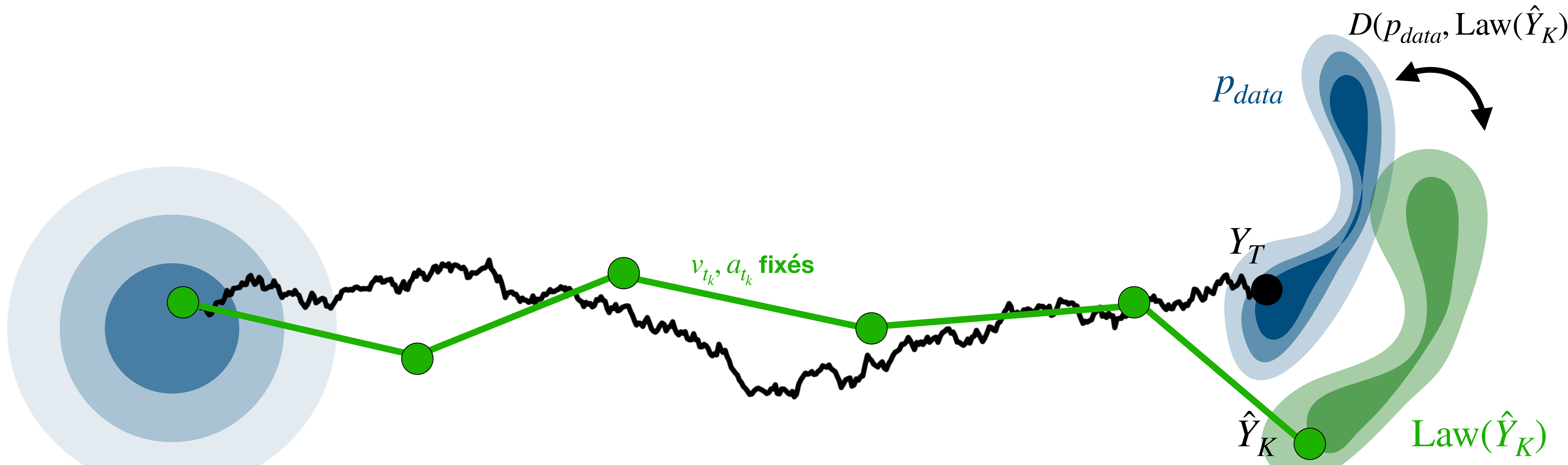
Data geometry matters for optimizing parameters w.r.t **discretization error**:

- Derived explicit first-order discretization error of diffusion sampling.
- In Gaussian case, gives a tractable objective for tuning diffusion parameters.
- Study how the optimal diffusion parameter α and schedulers η_t, σ_t change with the shape of the data covariance spectrum $(\lambda_i)_{i \leq d}$.
- Gaussian theory predicts trends observed for (non-Gaussian) real image sampling.

Discretization error

Euler discretization: for $\gamma = \frac{T}{K}$,

$$\hat{Y}_0 \sim p_0, \quad \hat{Y}_{k+1} = \hat{Y}_k + \gamma v_{t_k}(\hat{Y}_k) + \sqrt{2\gamma a_{t_k}} W_k$$



Discretization error: general weak error

Euler discretization: for $\gamma = \frac{T}{K}$,

$$\hat{Y}_0 \sim p_0, \quad \hat{Y}_{k+1} = \hat{Y}_k + \gamma v_{t_k}(\hat{Y}_k) + \sqrt{2\gamma a_{t_k}} W_k$$

For a test function $f \in \mathcal{C}^\infty(\mathbb{R}^d, \mathbb{R})$, under mild assumptions on p_{data}

End of discretized process End of exact (continuous) process

$$\mathbb{E}[f(\hat{Y}_K)] - \mathbb{E}[f(Y_T)] = \gamma \int_0^T \mathbb{E}[\psi_s(Y_s)] ds + o(\gamma)$$

Euler error due to freezing v_s, a_s on $[s, s + \gamma]$

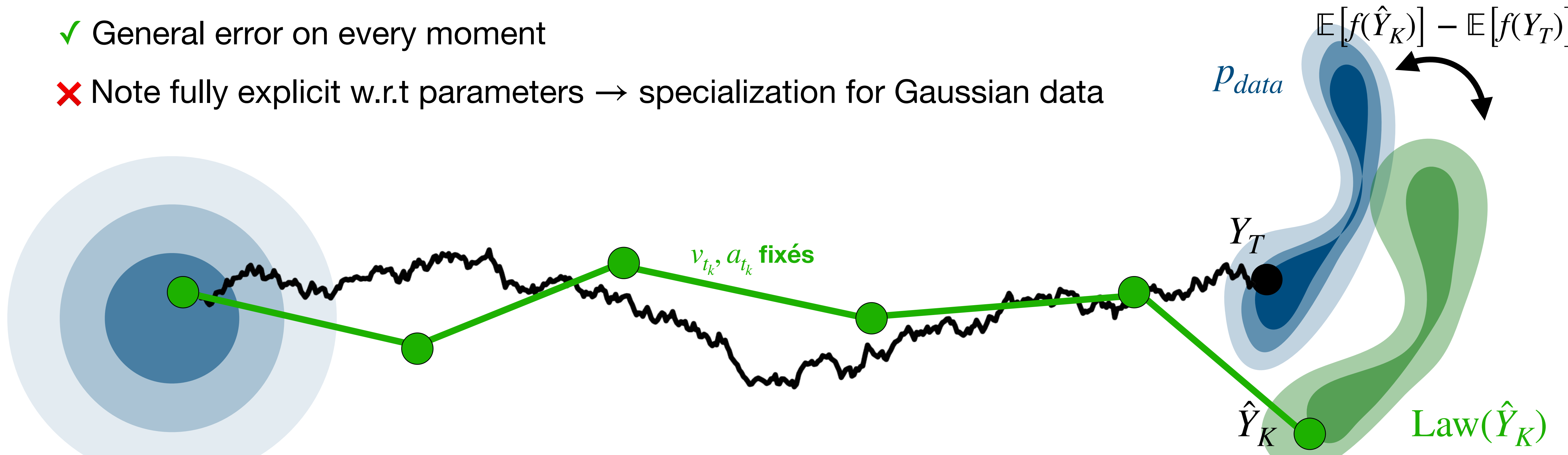
For $a_t = 0$

$$\psi_s(Y_s) = -\frac{1}{2} J_s(Y) D_s v_s(Y_s) \cdot \nabla f(Y_T)$$

Propagation $s \rightarrow T$ (flow Jacobian) local acceleration (material derivative)

✓ General error on every moment

✗ Not fully explicit w.r.t parameters → specialization for Gaussian data



Diffusion models for posterior sampling $p(x | y = Ax + w)$

Apply the same algorithm with p_{data} replaced by $p(x | y)$

For $\alpha \geq 0$,

$$Y_0 \sim p_0, \quad dY_t = v_t(Y_t) dt + \sqrt{2a_t} dW_t$$

$$v_t(x) = -\frac{\dot{\eta}_{T-t}}{\eta_{T-t}} x + (1+\alpha) \eta_{T-t}^2 \dot{\sigma}_{T-t} \sigma_{T-t} \nabla \log p_{T-t}(x | y)$$

Using Baye's formula, for $x_t \sim p_t$ $\nabla \log p_t(x_t | y) = \nabla \log p_t(x_t) + \nabla \log p_t(y | x_t)$

✓ Unconditional score training

$$= \int p(y | x_0) p(x_0 | x_t) dx_0$$

✗ Untractable

DPS [Chung et. al, '23], PiGDM [Song et. al, '23], Moment Matching [Rozet et. al, '24] :

Use the Gaussian approximation $p(x_0 | x_t) \approx \mathcal{N}(\mathbb{E}[x_0 | x_t], \mathbb{V}[x_0 | x_t])$

Related to a pretrained denoiser by Tweedie's formulas

Diffusion models

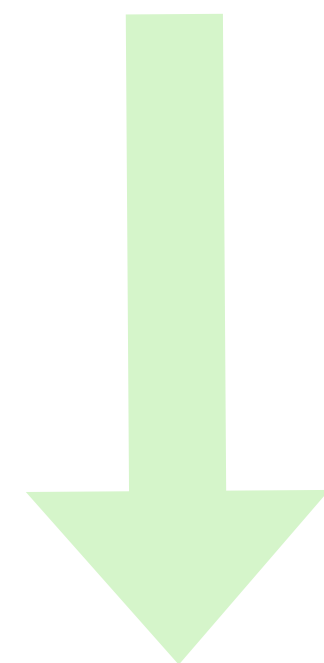
In the **continuous** setting, for fixed σ_t : VE $\eta_t = 1 \Leftrightarrow$ VP $\eta_t = \frac{1}{\sqrt{1 + \sigma_t^2}} \Leftrightarrow$ any rescaling η_t

VE schedule ($\sigma_t, \eta_t = 1$)

$$y_0 \sim p_T, \quad dy_t = v_t(y_t) dt + \sqrt{2\alpha\xi_t} dW_t$$

$$v_t = \xi_t \nabla \log p_{T-t}$$

$$\begin{cases} \beta_t = 0 \\ \xi_t = \dot{\sigma}_{T-t} \sigma_{T-t} \end{cases}$$



$$\hat{y}_t = \eta_t y_t$$

VP schedule (σ_t, η_t)

$$y_0 \sim \hat{p}_T, \quad d\hat{y}_t = v_t(\hat{y}_t) dt + \sqrt{2\alpha\xi_t} dW_t$$

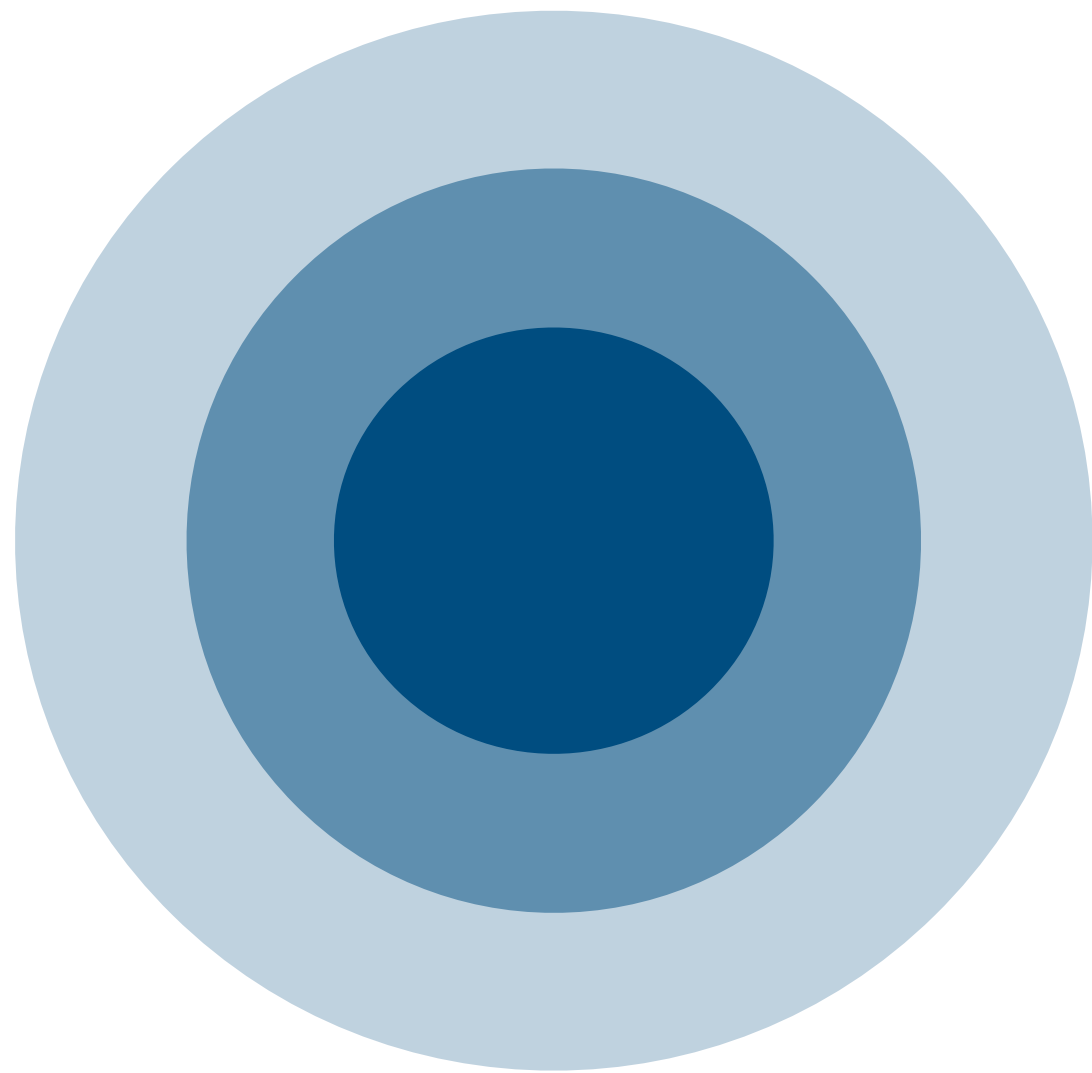
$$v_t = \beta_t \text{Id} + \xi_t \nabla \log \hat{p}_{T-t} \left(\frac{\dot{}}{\eta_t} \right)$$

$$\begin{cases} \beta_t = -\frac{\dot{\eta}_{T-t}}{\eta_{T-t}} \\ \xi_t = \eta_{T-t}^2 \dot{\sigma}_{T-t} \sigma_{T-t} \end{cases}$$

Diffusion models Other Noising Process (from data to noise)

Flow Matching (FM)

$$X_0 \sim p_{\text{data}}, W_t \sim \mathcal{N}(0, \text{Id}), \quad X_t \stackrel{d}{=} (1-t)X_0 + tW_t \quad \sigma_t = \frac{t}{1-t}, \quad \eta_t = \frac{1}{1+\sigma_t}$$



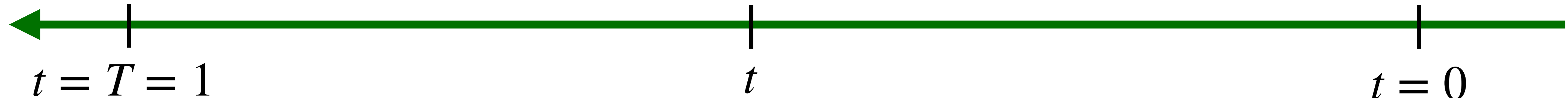
$$q_1 = \mathcal{N}(0, \text{Id})$$



$$q_t = \text{Law}(\eta_t(X_0 + \sigma_t X_t))$$



$$q_0 = p_{\text{data}}$$



$$\text{Cov}(X_t) = (1-t)^2 \text{Cov}(X_0) + t^2 \text{Id}$$

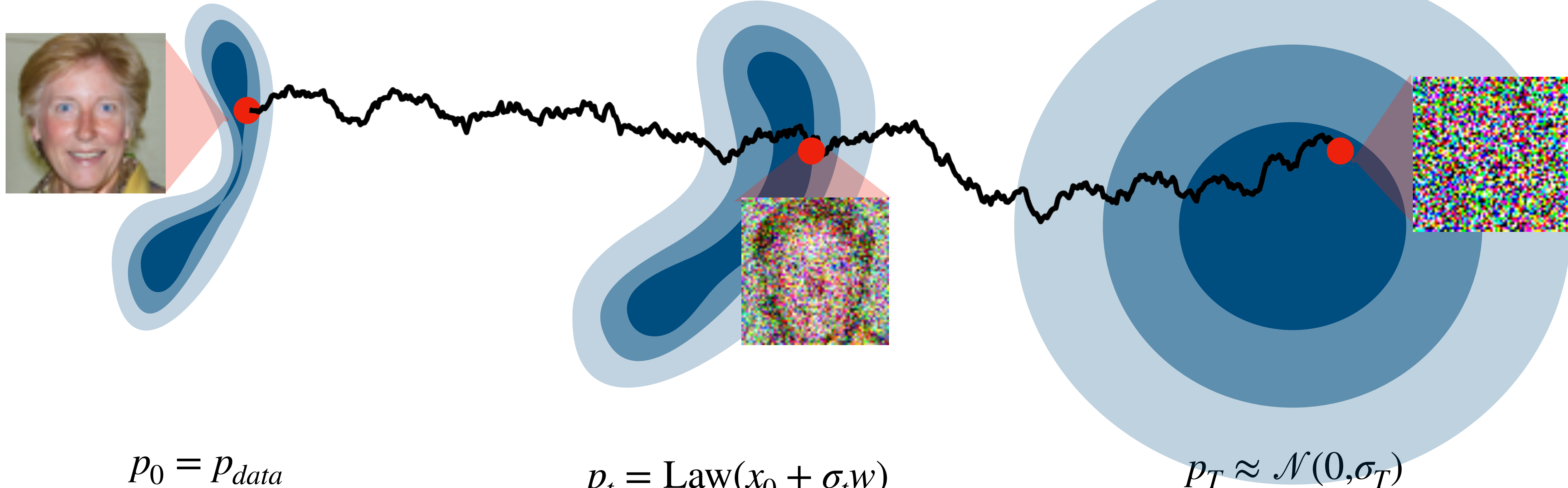
Diffusion models

Forward flow : from data to noise

Variance Exploding (VE)

$\sigma_t > 0$ e.g. [Song et al.] $\sigma_t = t$

$$x_0 \sim p_{\text{data}}, \quad dx_t = \sqrt{2\dot{\sigma}_t \sigma_t} dw_t \quad \Leftrightarrow \quad w_t \sim \mathcal{N}(0, \text{Id}), \quad x_t = x_0 + \sigma_t w_t$$



$$p_0 = p_{\text{data}}$$

$$p_t = \text{Law}(x_0 + \sigma_t w)$$

$$p_T \approx \mathcal{N}(0, \sigma_T)$$

$$t = 0$$

$$t$$

$$t = T \gg 1$$

$$\text{Cov}(x_0)$$

$$\text{Cov}(x_t) = \text{Cov}(x_0) + \sigma_t^2 \text{Id}$$

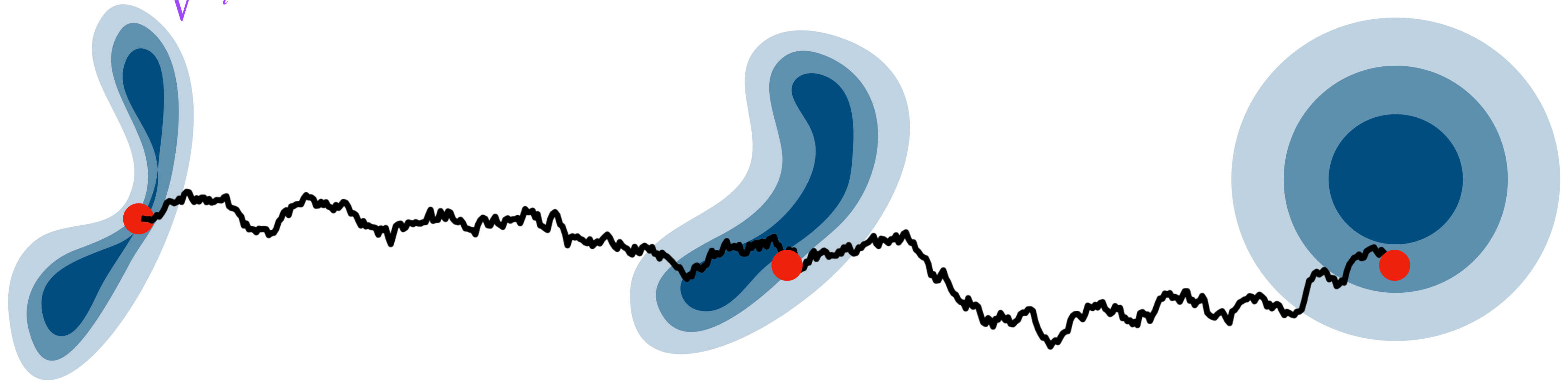
$$\text{Cov}(x_T) \approx \sigma_T^2 \text{Id}$$

Diffusion models

Forward flow : from data to noise

Variance Preserving (VP) $x_0 \sim p_{\text{data}}, \quad dx_t = \frac{\dot{\eta}_t}{\eta_t} x_t dt + \eta_t \sqrt{2\dot{\sigma}_t \sigma_t} dw_t \Leftrightarrow \boxed{w_t \sim \mathcal{N}(0, \text{Id}), \quad x_t = \eta_t(x_0 + \sigma_t w_t)}$

$$\sigma_t > 0, \quad \eta_t = \frac{1}{\sqrt{\sigma_t^2 + 1}}$$



$$p_0 = p_{\text{data}}$$

$$p_t = \text{Law}(\eta_t(x_0 + \sigma_t w))$$

$$p_T \approx \mathcal{N}(0, \sigma_T)$$

$$t = 0$$

$$t$$

$$t = T \gg 1$$

$$\text{Cov}(x_0)$$

$$\text{Cov}(x_t) = \eta_t^2 \text{Cov}(x_0) + (1 - \eta_t^2) \text{Id}$$

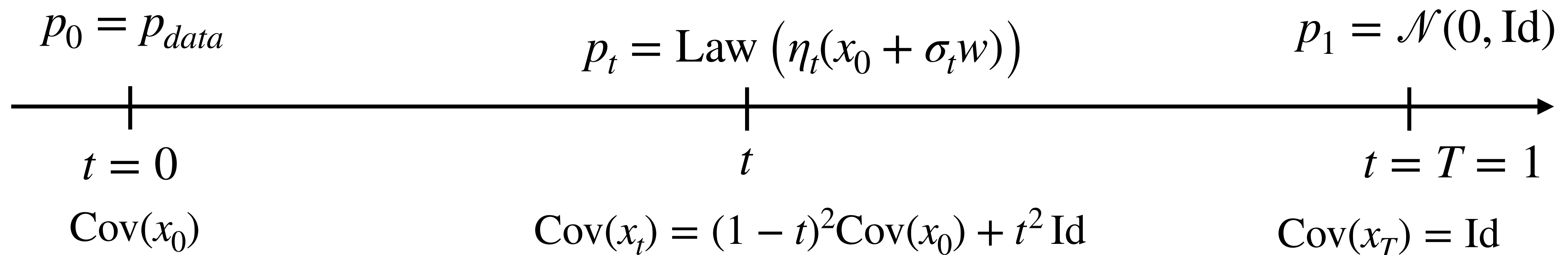
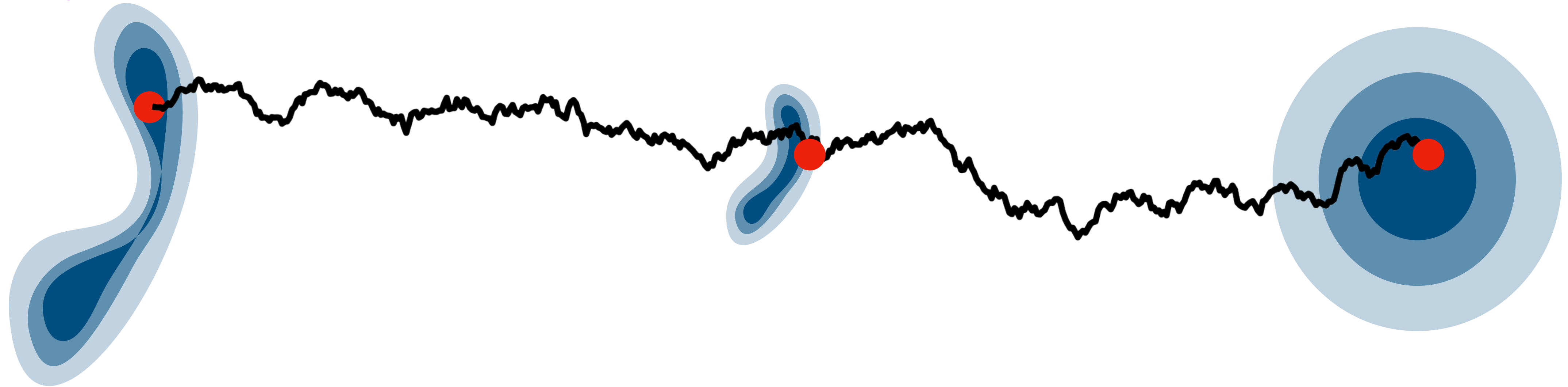
$$\text{Cov}(x_T) \approx \text{Id}$$

Diffusion models **Forward flow** : from data to noise

Flow Matching (FM)

$$\sigma_t = \frac{t}{1-t}, \quad \eta_t = 1-t \quad (T=1)$$

$$w_t \sim \mathcal{N}(0, \text{Id}), \quad x_t = t x_0 + (1-t) w_t$$



Discretization error: : general weak error

Euler discretization: for $\gamma = \frac{T}{K}$,

$$\hat{Y}_0 \sim p_0, \quad \hat{Y}_{k+1} = \hat{Y}_k + \gamma v_{t_k}(\hat{Y}_k) + a_{t_k} W_k$$

Under mild assumptions on p_{data}

End discretized process End continuous process

$$\mathbb{E}[f(\hat{Y}_K)] - \mathbb{E}[f(Y_T)] = \gamma \int_0^T \mathbb{E}[\psi_s(Y_s)] ds + o(\gamma)$$

Euler error due to freezing v_s, a_s on $[s, s + \gamma]$

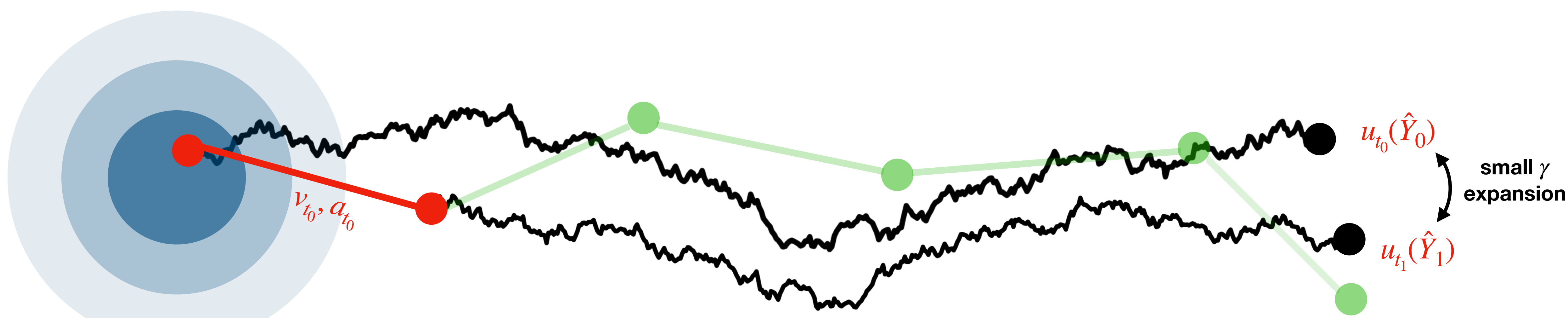
For $a_t = 0$

$$\psi_s(Y_s) = -\frac{1}{2} J_s(Y) D_s v_s(Y_s) \cdot \nabla f(Y_T)$$

Propagation $s \rightarrow T$ (flow Jacobian) local acceleration (material derivative)

Sketch of the proof: Let $u_t(x) = \mathbb{E}[f(Y_T) | Y_t = x]$. We have:

$$\mathbb{E}[f(\hat{Y}_K)] - \mathbb{E}[f(Y_T)] = \mathbb{E}[u_{t_k}(\hat{Y}_k)] - \mathbb{E}[u_0(\hat{Y}_0)] \overset{\text{Telescopic sum}}{=} \sum_{k=0}^K u_{t_{k+1}}(\hat{Y}_{k+1}) - u_{t_k}(\hat{Y}_k) \overset{\text{Expansion for small } \gamma}{=} \sum_{k=0}^K \left(\gamma^2 \psi_{t_k}(\hat{Y}_k) + o(\gamma^2) \right) \overset{\text{Riemann-sum error}}{=} \gamma \int_0^T \psi_s(Y_s) ds + o(\gamma)$$



Discretization error: : general weak error

Euler discretization: for $\gamma = \frac{T}{K}$,

$$\hat{Y}_0 \sim p_0, \quad \hat{Y}_{k+1} = \hat{Y}_k + \gamma v_{t_k}(\hat{Y}_k) + a_{t_k} W_k$$

Under mild assumptions on p_{data}

End discretized process End continuous process

$$\mathbb{E}[f(\hat{Y}_K)] - \mathbb{E}[f(Y_T)] = \gamma \int_0^T \mathbb{E}[\psi_s(Y_s)] ds + o(\gamma)$$

Euler error due to freezing v_s, a_s on $[s, s + \gamma]$

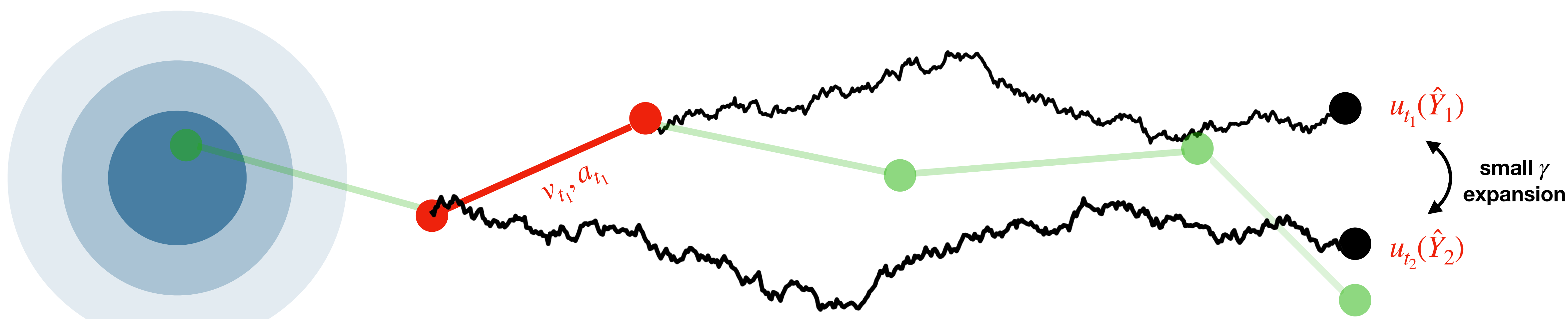
For $a_t = 0$

$$\psi_s(Y_s) = -\frac{1}{2} J_s(Y) D_s v_s(Y_s) \cdot \nabla f(Y_T)$$

Propagation $s \rightarrow T$ (flow Jacobian) local acceleration (material derivative)

Sketch of the proof: Let $u_t(x) = \mathbb{E}[f(Y_T) | Y_t = x]$. We have:

$$\mathbb{E}[f(\hat{Y}_K)] - \mathbb{E}[f(Y_T)] = \mathbb{E}[u_{t_K}(\hat{Y}_K)] - \mathbb{E}[u_0(\hat{Y}_0)] \underset{\text{Telescopic sum}}{=} \sum_{k=0}^K u_{t_{k+1}}(\hat{Y}_{k+1}) - u_{t_k}(\hat{Y}_k) \underset{\text{Expansion for small } \gamma}{=} \sum_{k=0}^K \left(\gamma^2 \psi_{t_k}(\hat{Y}_k) + o(\gamma^2) \right) \underset{\text{Riemann-sum error}}{=} \gamma \int_0^T \psi_s(Y_s) ds + o(\gamma)$$



Discretization error: : general weak error

Euler discretization: for $\gamma = \frac{T}{K}$,

$$\hat{Y}_0 \sim p_0, \quad \hat{Y}_{k+1} = \hat{Y}_k + \gamma v_{t_k}(\hat{Y}_k) + a_{t_k} W_k$$

Under mild assumptions on p_{data}

End discretized process End continuous process

$$\mathbb{E}[f(\hat{Y}_K)] - \mathbb{E}[f(Y_T)] = \gamma \int_0^T \mathbb{E}[\psi_s(Y_s)] ds + o(\gamma)$$

Euler error due to freezing v_s, a_s on $[s, s + \gamma]$

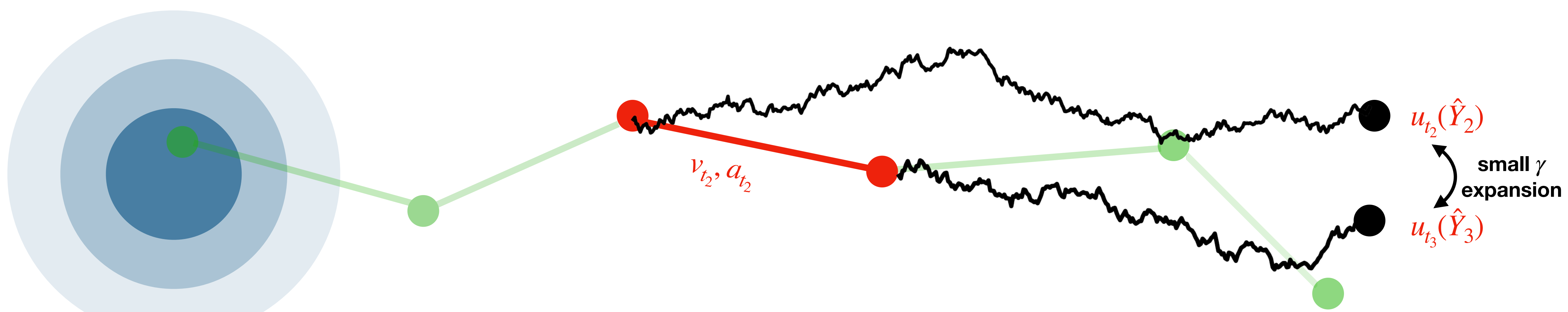
For $a_t = 0$

$$\psi_s(Y_s) = -\frac{1}{2} J_s(Y) D_s v_s(Y_s) \cdot \nabla f(Y_T)$$

Propagation $s \rightarrow T$ (flow Jacobian) local acceleration (material derivative)

Sketch of the proof: Let $u_t(x) = \mathbb{E}[f(Y_T) | Y_t = x]$. We have:

$$\mathbb{E}[f(\hat{Y}_K)] - \mathbb{E}[f(Y_T)] = \mathbb{E}[u_{t_k}(\hat{Y}_k)] - \mathbb{E}[u_0(\hat{Y}_0)] \underset{\text{Telescopic sum}}{=} \sum_{k=0}^K u_{t_{k+1}}(\hat{Y}_{k+1}) - u_{t_k}(\hat{Y}_k) \underset{\text{Expansion for small } \gamma}{=} \sum_{k=0}^K \left(\gamma^2 \psi_{t_k}(\hat{Y}_k) + o(\gamma^2) \right) \underset{\text{Riemann-sum error}}{=} \gamma \int_0^T \psi_s(Y_s) ds + o(\gamma)$$



Discretization error: : general weak error

Euler discretization: for $\gamma = \frac{T}{K}$,

$$\hat{Y}_0 \sim p_0, \quad \hat{Y}_{k+1} = \hat{Y}_k + \gamma v_{t_k}(\hat{Y}_k) + a_{t_k} W_k$$

Under mild assumptions on p_{data}

End discretized process End continuous process

$$\mathbb{E}[f(\hat{Y}_K)] - \mathbb{E}[f(Y_T)] = \gamma \int_0^T \mathbb{E}[\psi_s(Y_s)] ds + o(\gamma)$$

Euler error due to freezing v_s, a_s on $[s, s + \gamma]$

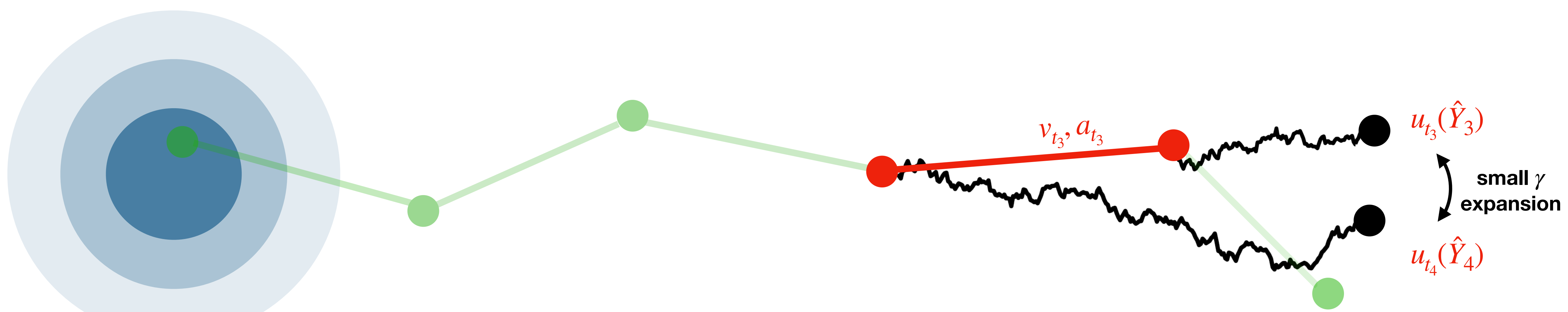
For $a_t = 0$

$$\psi_s(Y_s) = -\frac{1}{2} J_s(Y) D_s v_s(Y_s) \cdot \nabla f(Y_T)$$

Propagation $s \rightarrow T$ (flow Jacobian) local acceleration (material derivative)

Sketch of the proof: Let $u_t(x) = \mathbb{E}[f(Y_T) | Y_t = x]$. We have:

$$\mathbb{E}[f(\hat{Y}_K)] - \mathbb{E}[f(Y_T)] = \mathbb{E}[u_{t_k}(\hat{Y}_k)] - \mathbb{E}[u_0(\hat{Y}_0)] \underset{\text{Telescopic sum}}{=} \sum_{k=0}^K u_{t_{k+1}}(\hat{Y}_{k+1}) - u_{t_k}(\hat{Y}_k) \underset{\text{Expansion for small } \gamma}{=} \sum_{k=0}^K \left(\gamma^2 \psi_{t_k}(\hat{Y}_k) + o(\gamma^2) \right) \underset{\text{Riemann-sum error}}{=} \gamma \int_0^T \psi_s(Y_s) ds + o(\gamma)$$



Discretization error: : general weak error

Euler discretization: for $\gamma = \frac{T}{K}$,

$$\hat{Y}_0 \sim p_0, \quad \hat{Y}_{k+1} = \hat{Y}_k + \gamma v_{t_k}(\hat{Y}_k) + a_{t_k} W_k$$

Under mild assumptions on p_{data}

End discretized process End continuous process

$$\mathbb{E}[f(\hat{Y}_K)] - \mathbb{E}[f(Y_T)] = \gamma \int_0^T \mathbb{E}[\psi_s(Y_s)] ds + o(\gamma)$$

Euler error due to freezing v_s, a_s on $[s, s + \gamma]$

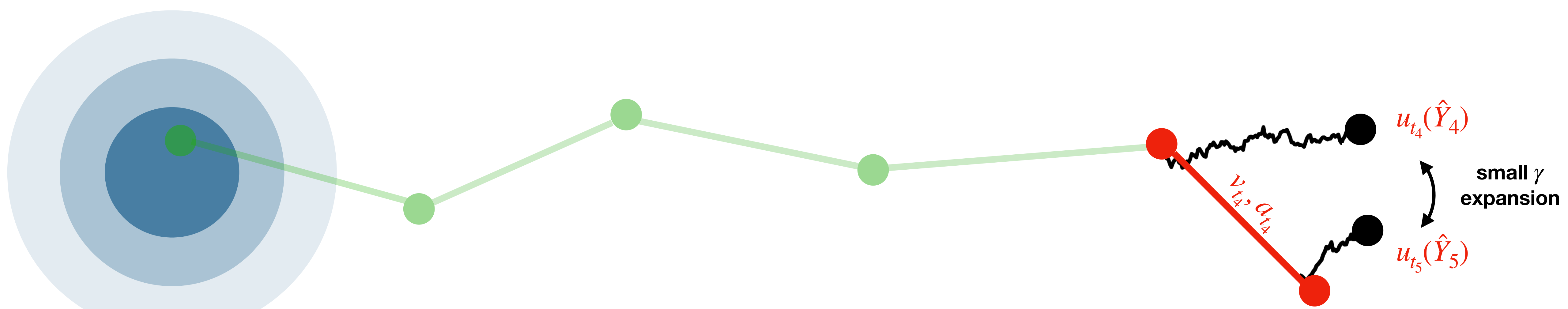
For $a_t = 0$

$$\psi_s(Y_s) = -\frac{1}{2} J_s(Y) D_s v_s(Y_s) \cdot \nabla f(Y_T)$$

Propagation $s \rightarrow T$ (flow Jacobian) local acceleration (material derivative)

Sketch of the proof: Let $u_t(x) = \mathbb{E}[f(Y_T) | Y_t = x]$. We have:

$$\mathbb{E}[f(\hat{Y}_K)] - \mathbb{E}[f(Y_T)] = \mathbb{E}[u_{t_k}(\hat{Y}_k)] - \mathbb{E}[u_0(\hat{Y}_0)] \overset{\text{Telescopic sum}}{=} \sum_{k=0}^K u_{t_{k+1}}(\hat{Y}_{k+1}) - u_{t_k}(\hat{Y}_k) \overset{\text{Expansion for small } \gamma}{=} \sum_{k=0}^K \left(\gamma^2 \psi_{t_k}(\hat{Y}_k) + o(\gamma^2) \right) \overset{\text{Riemann-sum error}}{=} \gamma \int_0^T \psi_s(Y_s) ds + o(\gamma)$$



Optimization of the rescaling schedule η_t

Classical choices of rescaling schedules:

$$\eta_t^{VE} = 1, \quad \eta_t^{VP} = \sqrt{\frac{1}{1+\sigma_t^2}}, \quad \eta_t^{FM} = \frac{1}{1+\sigma_t^2}$$

Gaussian sampling covariance error:

$$\text{Cov}[\hat{Y}_K] = \Sigma_{\text{data}} + \gamma U \text{Diag}\left(\int_0^T \Delta^\Sigma(\lambda_i, \eta_t, \sigma_t) dt\right) U^\top + o(\gamma)$$

$$\Delta^\Sigma(\lambda_i, \eta_t^\star, \sigma_t) = 0 \Leftrightarrow \eta_t^\star(\lambda_i) = \sqrt{\frac{\lambda_i}{\lambda_i + \sigma_t^2}} \Rightarrow \text{Var}_{\lambda_i}[X_t] = \text{Var}_{\lambda_i}[X_0]$$

Variance preserving condition on eigendirection λ_i

This is optimal **for a single** eigendirection λ_i .

For multidimensional and anisotropic data, we set

$$\eta_t(c) = \sqrt{\frac{c}{c + \sigma_t^2}}$$

Comes back to VP initialized at $\mathcal{N}(0, c \mathbf{I})$ instead of $\mathcal{N}(0, \mathbf{I})$

$$c^\star = \operatorname{argmin}_{c>0} \text{FD}(p_{\text{data}}, \text{Law}(\hat{Y}_K)) = \operatorname{argmin}_{c>0} \sum_{i=1}^d \int_0^T f(\lambda_i, \eta_t(c), \sigma_t) dt$$

Assuming a power spectrum $\lambda_i = \lambda_{\max} i^{-p}$ and noise $\sigma_t = \sigma_{\max} \left(\frac{t}{T}\right)^\beta$ with large $\sigma_{\max}$ and β

$$c^\star = \lambda_{\max} \tau(p) \quad \text{with } \tau(p) > 0 \text{ and decreasing for } p \lesssim 1$$