



Gaussian Process Variational Autoencoder

Xinxing Shi

Sept. 16th, 2026

Contents

❖ Background

- GP

- VAE

❖ GPVAE Models

- GPVAE with Inducing Points

- Nearest-Neighbour GPVAE

Contents

❖ Background

- GP

- VAE

❖ GPVAE Models

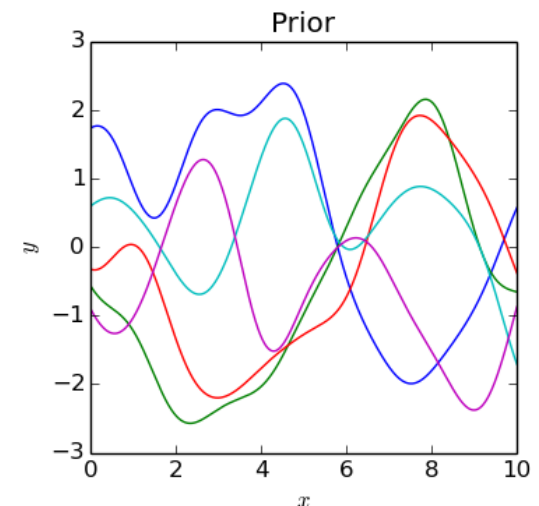
- GPVAE with Inducing Points

- Nearest-Neighbour GPVAE

Gaussian Process (GP)

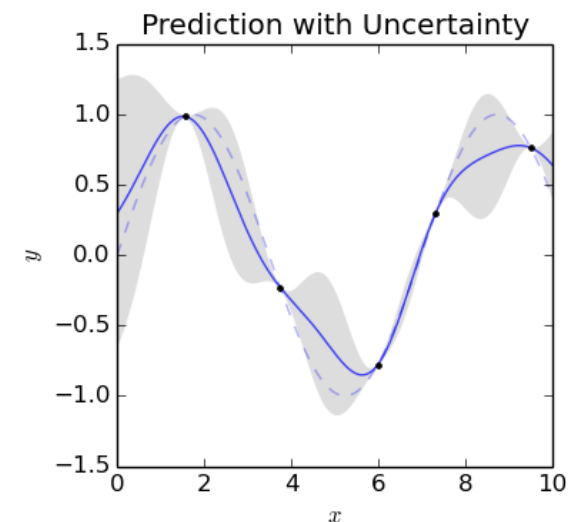
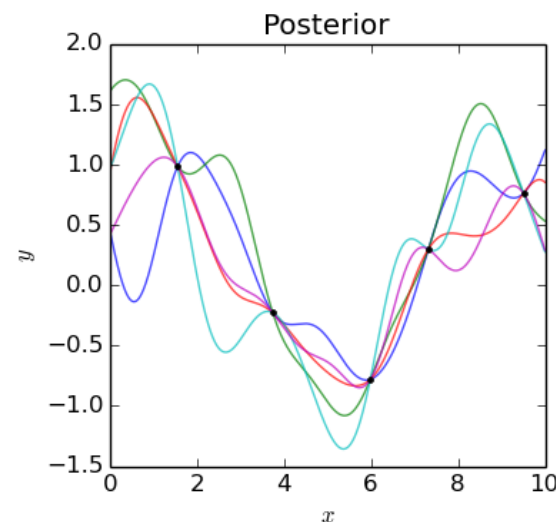
- GP builds correlations and calibrates uncertainty for variables
 - Prior: Any finite collection of random variables is Gaussian

$$\begin{bmatrix} f(\mathbf{x}_1) \\ f(\mathbf{x}_2) \\ \vdots \\ f(\mathbf{x}_i) \\ \vdots \end{bmatrix} \sim \mathcal{N} \left(\begin{bmatrix} \mu(\mathbf{x}_1) \\ \mu(\mathbf{x}_2) \\ \vdots \\ \mu(\mathbf{x}_i) \\ \vdots \end{bmatrix}, \begin{bmatrix} k(\mathbf{x}_1, \mathbf{x}_1) & k(\mathbf{x}_1, \mathbf{x}_2) & \cdots & k(\mathbf{x}_1, \mathbf{x}_i) & \cdots \\ k(\mathbf{x}_2, \mathbf{x}_1) & k(\mathbf{x}_2, \mathbf{x}_2) & \cdots & k(\mathbf{x}_2, \mathbf{x}_i) & \cdots \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ k(\mathbf{x}_i, \mathbf{x}_1) & k(\mathbf{x}_i, \mathbf{x}_2) & \cdots & k(\mathbf{x}_i, \mathbf{x}_i) & \cdots \\ \vdots & \vdots & \vdots & \vdots & \ddots \end{bmatrix} \right)$$



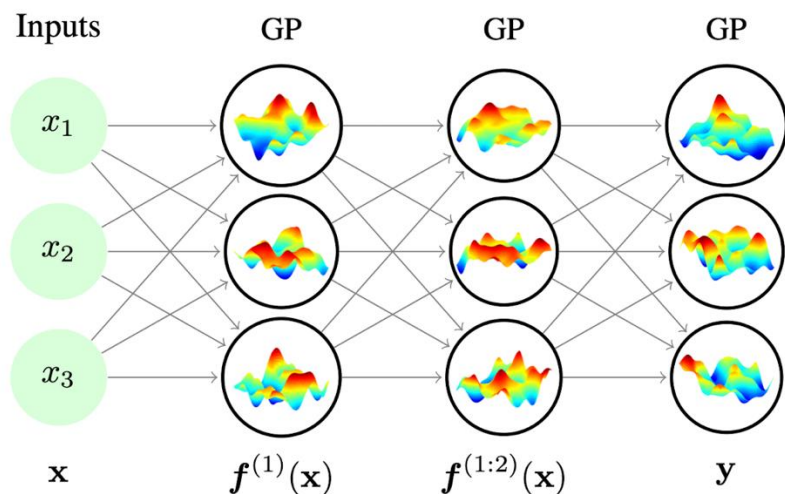
- Posterior: Conditioned on observations through Bayes' theorem

$$p(\mathbf{f} \mid \mathbf{y}) = \frac{p(\mathbf{y} \mid \mathbf{f})p(\mathbf{f})}{p(\mathbf{y})}$$

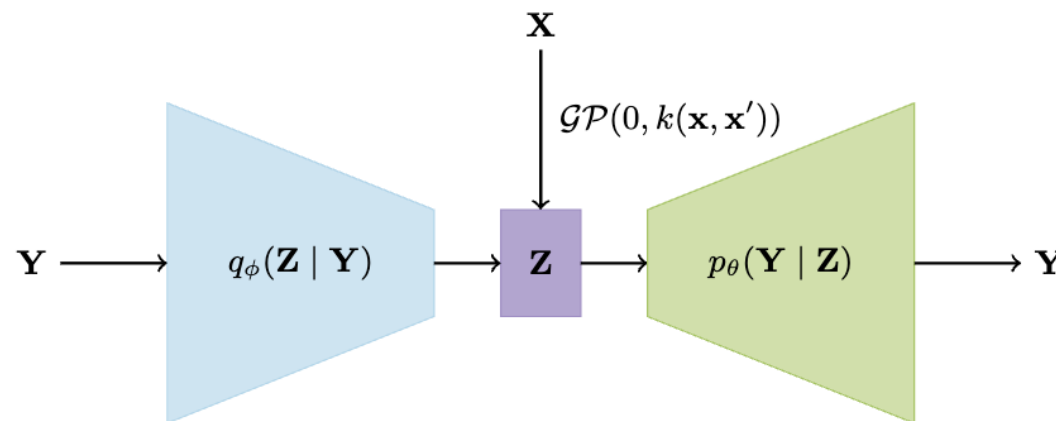


GP as Building Block

- More complex GP models:
 - (supervised) Deep GP



- (unsupervised) GP Variational AutoEncoder (GPVAE)



- Problems of GPs are inherited:
 - Intractability of posteriors
 - Scalability issues due to $N \times N$ kernel matrix

$$\log p(\mathbf{y}) = \log \int p(\mathbf{y} \mid \mathbf{f}) p(\mathbf{f}) d\mathbf{f}$$

- Intractability
- Scalability issues due to *the* $N \times N$ covariance matrix

$$p(\mathbf{y}) = \mathcal{N}(\mathbf{y} \mid \mathbf{0}, \mathbf{K}_{\mathbf{ff}} + \sigma^2 \mathbf{I})$$

$$p(\mathbf{f}_* \mid \mathbf{y}) = \mathcal{N}(\mathbf{f}_* \mid \mathbf{K}_{\mathbf{f}_* \mathbf{f}} (\mathbf{K}_{\mathbf{ff}} + \sigma^2 \mathbf{I})^{-1} \mathbf{y}, \mathbf{K}_{\mathbf{f}_* \mathbf{f}_*} - \mathbf{K}_{\mathbf{f}_* \mathbf{f}} (\mathbf{K}_{\mathbf{ff}} + \sigma^2 \mathbf{I})^{-1} \mathbf{K}_{\mathbf{ff}_*})$$

Contents

❖ Background

- GP

- VAE

❖ GPVAE Models

- GPVAE with Inducing Points

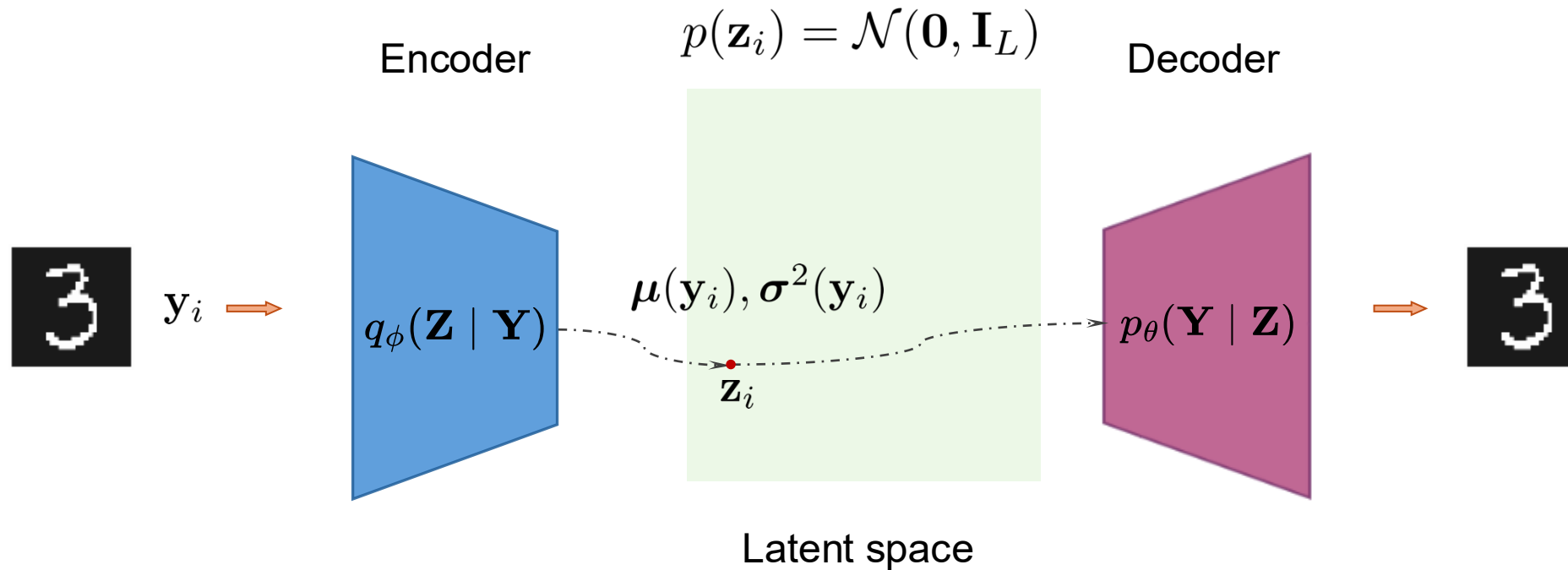
- Nearest-Neighbour GPVAE

Variational Autoencoder (VAE)

- Represent data points $\mathbf{Y}$ with Gaussian latent variables $\mathbf{Z}$

$$\mathbf{Y} = [\mathbf{y}_i]_{i=1}^N \in \mathbb{R}^{N \times K} \quad \mathbf{Z} = [\mathbf{z}_i]_{i=1}^N \in \mathbb{R}^{N \times L}$$

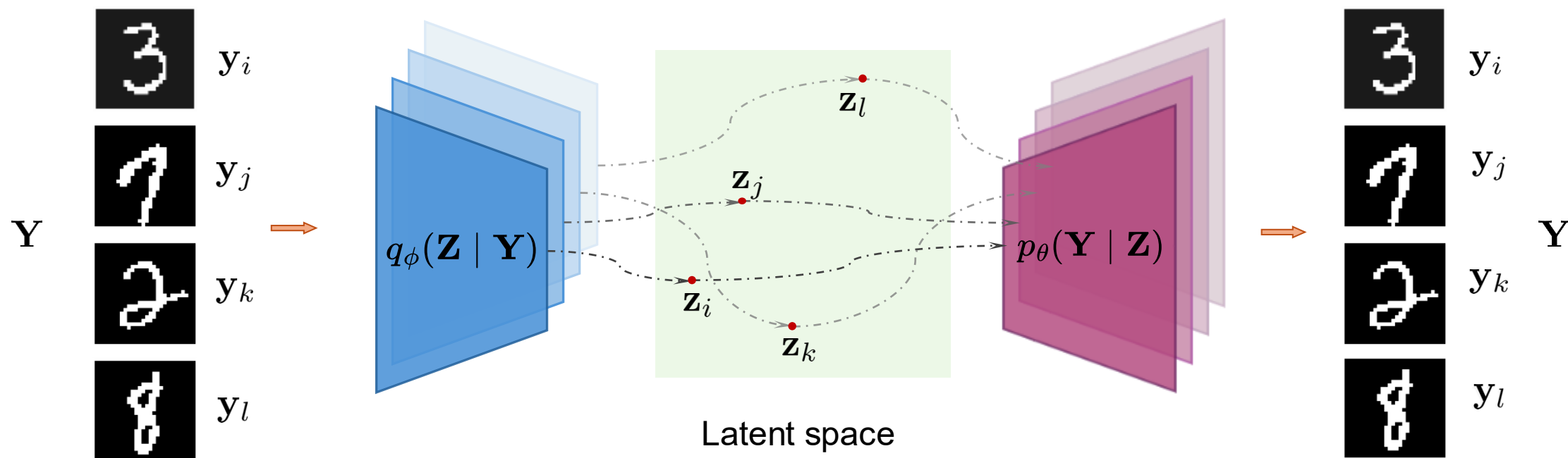
- Assumes a fully factorised Gaussian prior over the latent space



VAE: Factorised Prior

- Assumes a fully factorised Gaussian prior over the latent space

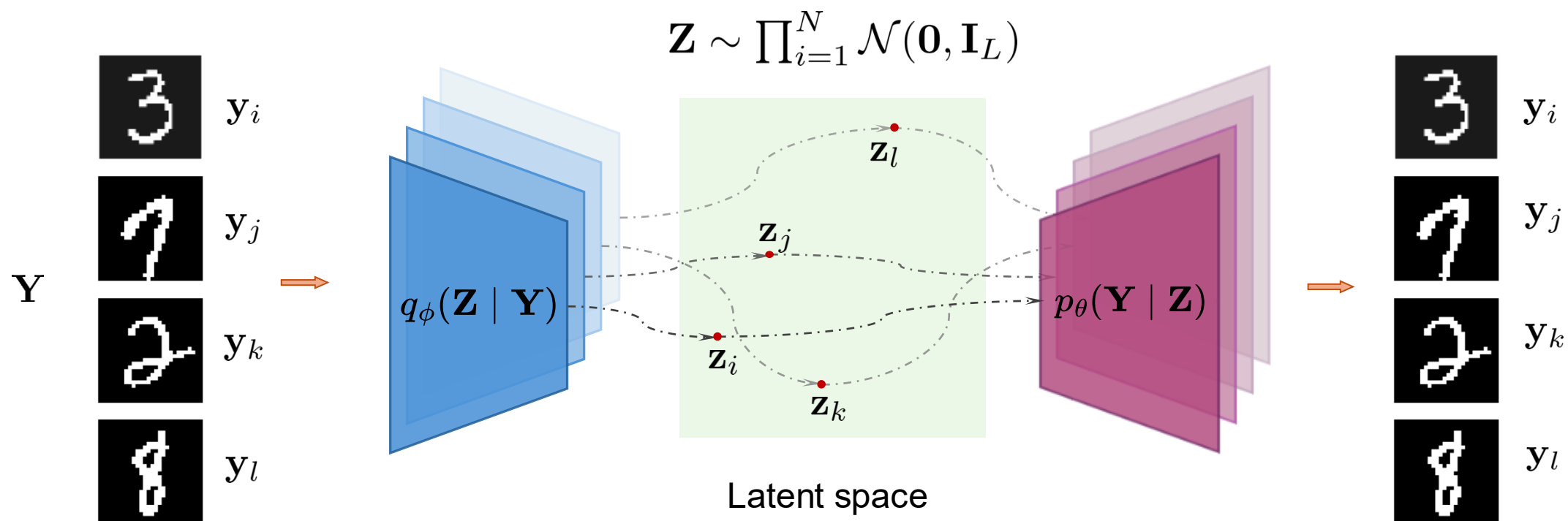
$$\mathbf{Z} \sim \prod_{i=1}^N p(\mathbf{z}_i) = \prod_{i=1}^N \mathcal{N}(\mathbf{z}_i \mid \mathbf{0}, \mathbf{I}_L)$$



VAE: Training

- A standard VAE assumes a fully factorised Gaussian prior over the latent space for amortised variational inference

$$\mathcal{L}_{\text{VAE}} = \frac{1}{|\mathcal{B}|} \sum_{i \in \mathcal{B}} \left\{ \mathbb{E}_{q(\mathbf{z}_i | \mathbf{y}_i)} [\log p(\mathbf{y}_i | \mathbf{z}_i)] - \text{KL}[q(\mathbf{z}_i | \mathbf{y}_i) \| p(\mathbf{z}_i)] \right\}$$

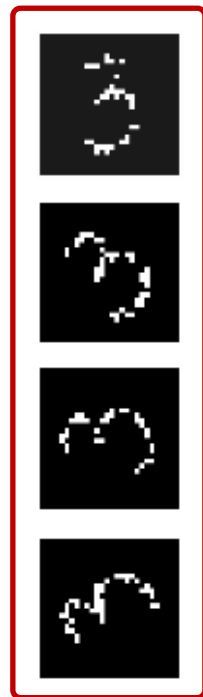


VAE: Independent Latent Variables

- A fully factorised Gaussian prior in a standard VAE ignores correlations of structured data $\mathbf{Y}$ with “timestamps” X

$$\mathbf{Y} = [\mathbf{y}_i]_{i=1}^N \in \mathbb{R}^{N \times K}, \quad \mathbf{X} = [\mathbf{x}_i]_{i=1}^N \in \mathbb{R}^{N \times D}$$

*correlated
samples*

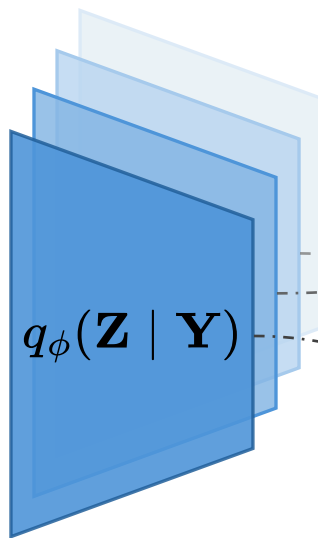


$(\mathbf{y}_{i-2}, \mathbf{x}_{i-2})$

$(\mathbf{y}_{i-1}, \mathbf{x}_{i-1})$

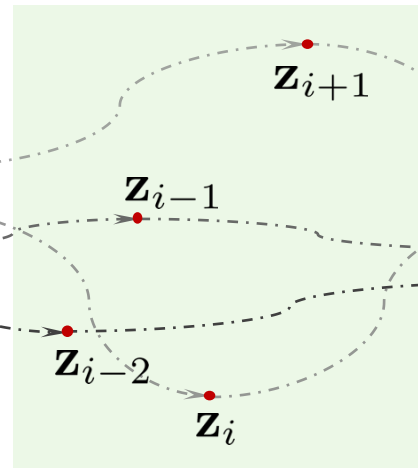
$(\mathbf{y}_i, \mathbf{x}_i)$

$(\mathbf{y}_{i+1}, \mathbf{x}_{i+1})$



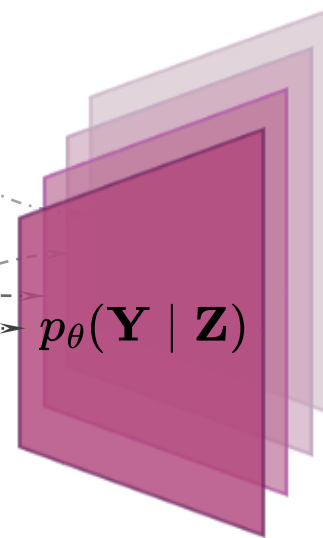
$q_\phi(\mathbf{Z} | \mathbf{Y})$

$$\mathbf{Z} \sim \prod_{i=1}^N \mathcal{N}(\mathbf{0}, \mathbf{I}_L)$$



Latent space

Introduce GPs for explicit, structured latent modelling



$p_\theta(\mathbf{Y} | \mathbf{Z})$



Contents

❖ Background

- GP

- From VAE to GPVAE

❖ GPVAE Models

- GPVAE with Inducing Points

- Nearest-Neighbour GPVAE

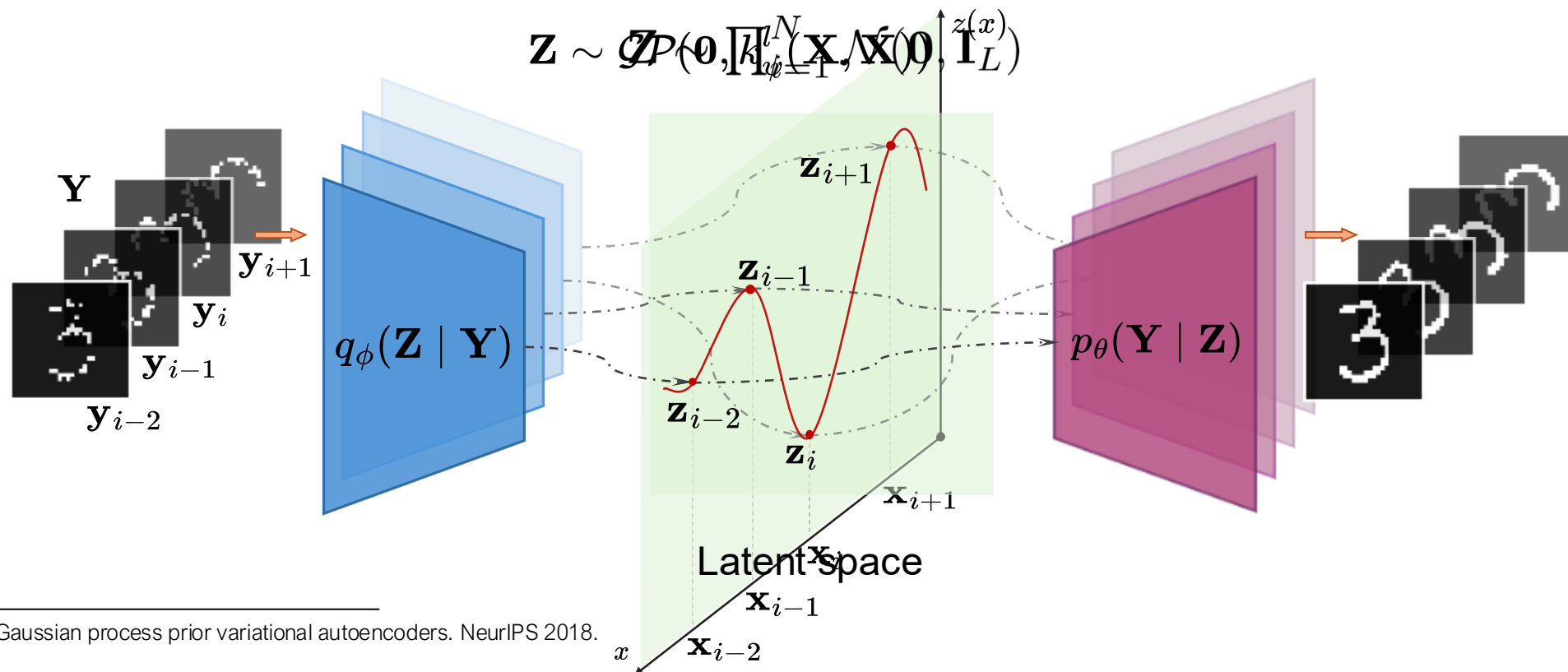
GPVAE: GP Prior in VAE but with Poor Scalability

- Build correlations via GP prior on data pairs $(\mathbf{X}, \mathbf{Y})$

$$\mathbf{Y} = [\mathbf{y}_i]_{i=1}^N, \quad \mathbf{X} = [\mathbf{x}_i]_{i=1}^N$$

- Evidence lower bound (ELBO) scales poorly with $\mathcal{O}(N^3)$

$$\mathcal{L} = \mathbb{E}_{q(\mathbf{Z}|\mathbf{Y})} [\log p(\mathbf{Y} | \mathbf{Z})] - \text{KL}[q(\mathbf{Z} | \mathbf{Y}) \parallel p(\mathbf{Z})]$$



GPVAE: ELBO Derivation

- Data pairs $(\mathbf{X}, \mathbf{Y})$

$$\mathbf{Y} = [\mathbf{y}_i]_{i=1}^N \in \mathbb{R}^{N \times K}, \quad \mathbf{X} = [\mathbf{x}_i]_{i=1}^N \in \mathbb{R}^{N \times D}$$

- Latent variables $\mathbf{Z}$

$$\text{GP prior: } p(\mathbf{Z}) = \mathcal{N}(\mathbf{Z} \mid \mathbf{0}, \mathbf{K}_{\mathbf{X}\mathbf{X}})$$

$$\text{Variational posterior: } q(\mathbf{Z} \mid \mathbf{Y}) = \mathcal{N}(\mathbf{Z} \mid \mu(\mathbf{Y}), \sigma^2(\mathbf{Y}))$$

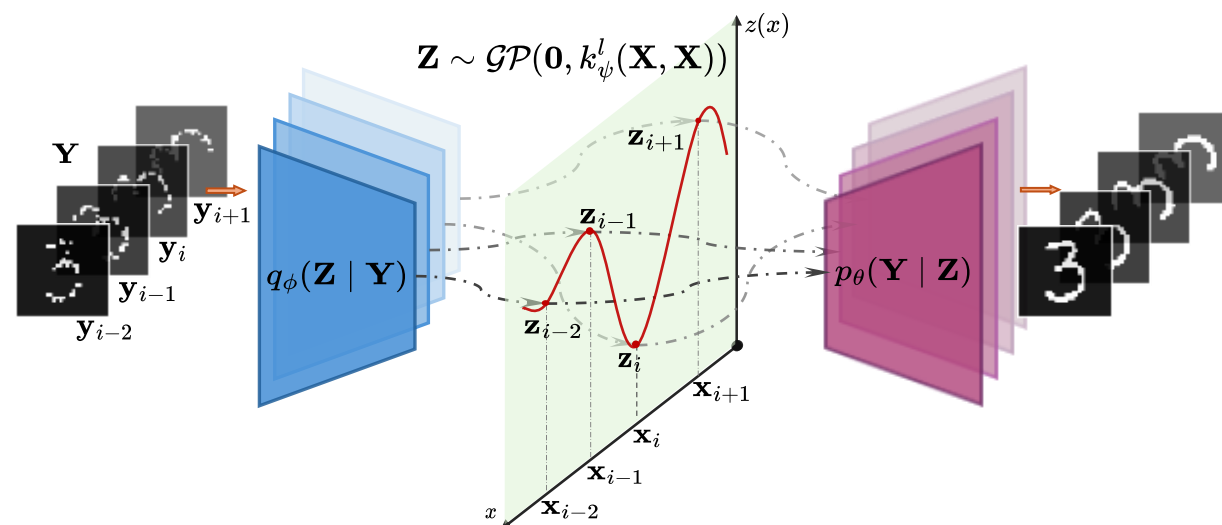
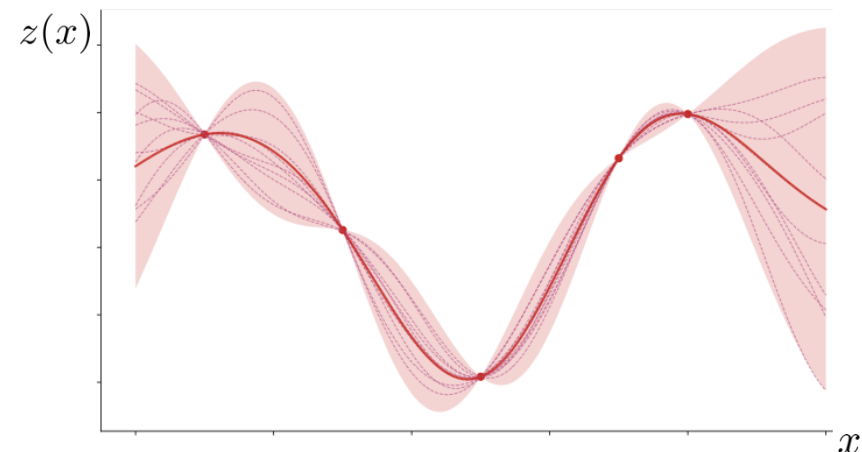
- ELBO

$$\begin{aligned} \log p(\mathbf{Y}) &= \log \int p(\mathbf{Z}) p(\mathbf{Y} \mid \mathbf{Z}) d\mathbf{Z} \\ &\geq \int q(\mathbf{Z} \mid \mathbf{Y}) \log \frac{p(\mathbf{Z}) p(\mathbf{Y} \mid \mathbf{Z})}{q(\mathbf{Z} \mid \mathbf{Y})} d\mathbf{Z} \\ &= \mathbb{E}_{q(\mathbf{Z} \mid \mathbf{Y})} [\log p(\mathbf{Y} \mid \mathbf{Z})] - \text{KL}[q(\mathbf{Z} \mid \mathbf{Y}) \parallel p(\mathbf{Z})] \end{aligned}$$

$O(N^3)$ complexity!

Mindset: Find Solutions from Scalable GP

- Desiderata:
 - Maintain the VAE architecture
 - Enable mini-batch training
 - Support various kernel choices
 - ...
- Draw inspirations from scalable GPs
 - Sparse GP with inducing point methods
 - Nearest-neighbour GP
 - Markovian GP with state-space representation
 - Sampling-based GP
 - ...



Jazbec, M., et al. Scalable Gaussian process variational autoencoders. AISTATS 2021.

Shi, X., et al. Neighbour-driven Gaussian process variational autoencoders for scalable structured latent modelling. ICML 2025.

Zhu, H., et al. Markovian Gaussian process variational autoencoders. ICML 2023.

Tran, B.-H., et al. Fully Bayesian autoencoders with latent sparse Gaussian processes. ICML 2023.

Contents

❖ Background

- GP

- From VAE to GPVAE

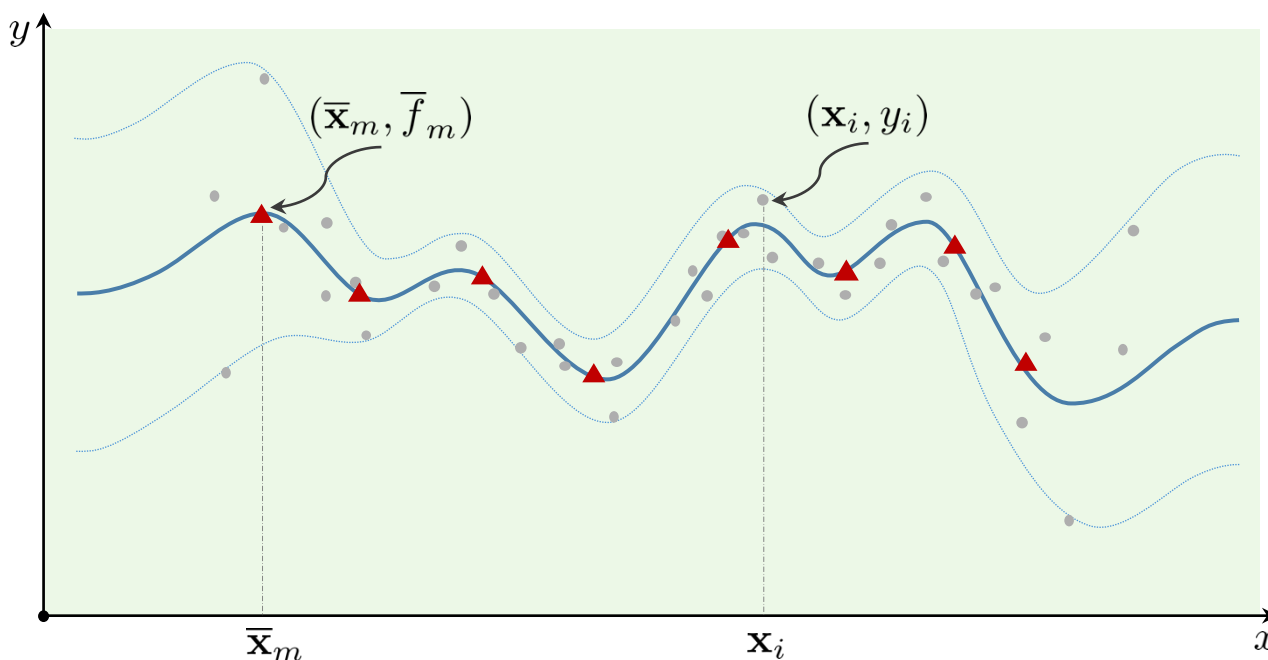
❖ GPVAE Models

- GPVAE with Inducing Points

- Nearest-Neighbour GPVAE

Sparse Variational GP (SVGP)

- Use a small set of M ($\ll N$) representative points to summarise data



Maximise:

$$\log p(\mathbf{y}) \geq \mathbb{E}_{p(\mathbf{f}|\bar{\mathbf{f}})q(\bar{\mathbf{f}})}[\log p(\mathbf{y} | \mathbf{f})] - \text{KL}[q(\bar{\mathbf{f}}) \| p(\bar{\mathbf{f}})]$$

Evidence Lower Bound (ELBO)

- Observations $\mathbf{y} = [y_i]_{i=1}^N$, $\mathbf{X} = [\mathbf{x}_i]_{i=1}^N$

GP prior: $p(\mathbf{f}) = \mathcal{N}(\mathbf{f} | \mathbf{0}, \mathbf{K}_{\mathbf{ff}})$

Likelihood: $p(\mathbf{y} | \mathbf{f}) = \prod_i p(y_i | f_i)$

- ▲ Inducing points $\bar{\mathbf{f}} = [\bar{f}_m]_{m=1}^M$, $\bar{\mathbf{X}} = [\bar{\mathbf{x}}_m]_{m=1}^M$

Augmented GP prior $p(\mathbf{f}, \bar{\mathbf{f}})$:

$$\begin{bmatrix} \mathbf{f} \\ \bar{\mathbf{f}} \end{bmatrix} \sim \mathcal{N} \left(\mathbf{0}, \begin{bmatrix} \mathbf{K}_{\mathbf{ff}} & \mathbf{K}_{\mathbf{f}\bar{\mathbf{f}}} \\ \mathbf{K}_{\bar{\mathbf{f}}\mathbf{f}} & \mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}} \end{bmatrix} \right)$$

Variational posterior:

$$q(\mathbf{f}, \bar{\mathbf{f}}) = p(\mathbf{f} | \bar{\mathbf{f}})q(\bar{\mathbf{f}}), \quad q(\bar{\mathbf{f}}) = \mathcal{N}(\bar{\mathbf{f}} | \mathbf{m}, \mathbf{S})$$

SVGP: ELBO

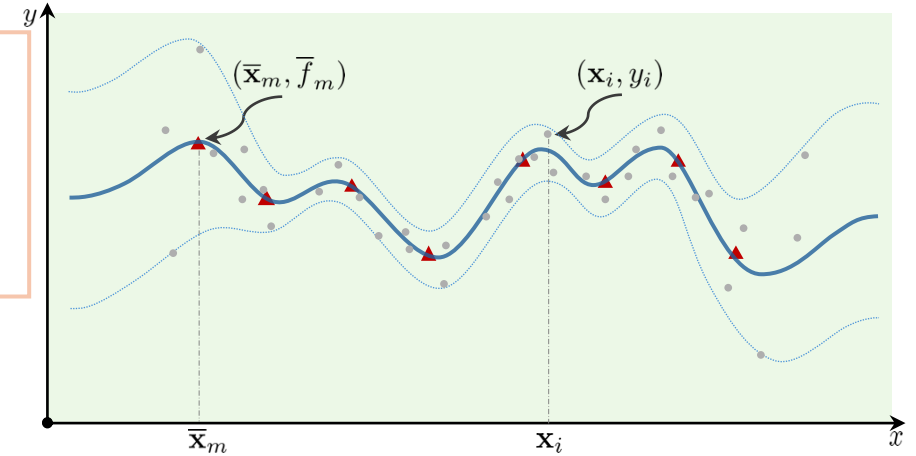
- A generic likelihood $p(y_i | f_i)$ may be non-Gaussian

$$\mathbb{E}_{p(\mathbf{f}|\bar{\mathbf{f}})q(\bar{\mathbf{f}})}[\log p(\mathbf{y} | \mathbf{f})] - \text{KL}[q(\bar{\mathbf{f}})||p(\bar{\mathbf{f}})]$$

- Closed-form KL

$$p(\bar{\mathbf{f}}) = \mathcal{N}(\bar{\mathbf{f}} | \mathbf{0}, \mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}}) \quad q(\bar{\mathbf{f}}) = \mathcal{N}(\bar{\mathbf{f}} | \mathbf{m}, \mathbf{S})$$

$$\text{KL}[q||p] = \frac{1}{2} [\mathbf{m}^\top \mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}}^{-1} \mathbf{m} + \text{tr}\{\mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}}^{-1} \mathbf{S}\} - \log \frac{|\mathbf{S}|}{|\mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}}|} - M]$$



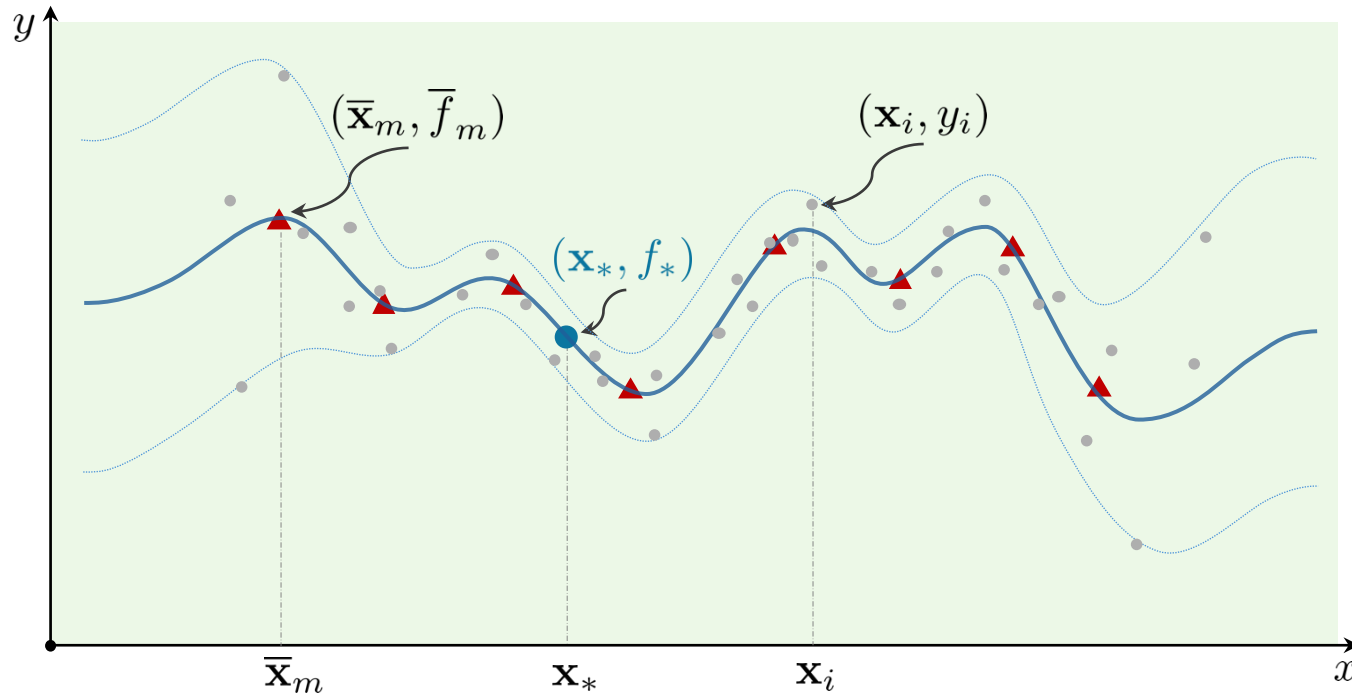
- Decomposable expected log-likelihood

$$\mathbb{E}_{p(\mathbf{f}|\bar{\mathbf{f}})q(\bar{\mathbf{f}})}[\log p(\mathbf{y} | \mathbf{f})] = \sum_{i=1}^N \mathbb{E}_{q(f_i)}[\log p(y_i | f_i)] \approx \frac{N}{|\mathcal{B}|} \sum_{i \in \mathcal{B}} \mathbb{E}_{q(f_i)}[\log p(y_i | f_i)] \quad \text{Marginalising + Mini-batching}$$

$$q(f_i) = \int p(f_i | \bar{\mathbf{f}})q(\bar{\mathbf{f}}) d\bar{\mathbf{f}} = \mathcal{N}(f_i | \mathbf{K}_{f_i \bar{\mathbf{f}}} \mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}}^{-1} \mathbf{m}, \underbrace{K_{f_i f_i} + \mathbf{K}_{f_i \bar{\mathbf{f}}} \mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}}^{-1} (\mathbf{S} - \mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}}) \mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}}^{-1} \mathbf{K}_{\bar{\mathbf{f}} f_i}}_{\mathbf{K}_{q(\mathbf{f})}})$$

SVGP: Prediction

- Compute posterior predictive distribution via $q(f_*, \bar{\mathbf{f}})$



- GP Prediction $f_*(\mathbf{x}_*)$

$$q(f_*, \bar{\mathbf{f}}) = p(f_* | \bar{\mathbf{f}})q(\bar{\mathbf{f}})$$

$$\begin{aligned} q(f_*) &= \int p(f_* | \bar{\mathbf{f}})q(\bar{\mathbf{f}}) d\bar{\mathbf{f}} \\ &= \mathcal{N}(f_* | \mathbf{K}_{f_* \bar{\mathbf{f}}} \mathbf{K}_{\bar{\mathbf{f}} \bar{\mathbf{f}}}^{-1} \mathbf{m}, \\ &\quad K_{f_* f_*} + \mathbf{K}_{f_* \bar{\mathbf{f}}} \mathbf{K}_{\bar{\mathbf{f}} \bar{\mathbf{f}}}^{-1} (\mathbf{S} - \mathbf{K}_{\bar{\mathbf{f}} \bar{\mathbf{f}}}) \mathbf{K}_{\bar{\mathbf{f}} \bar{\mathbf{f}}}^{-1} \mathbf{K}_{\bar{\mathbf{f}} f_*}) \end{aligned}$$

- Prediction $y_*(\mathbf{x}_*)$

$$q(y_* | \mathbf{y}) = \int p(y_* | f_*)q(f_*) df_*$$

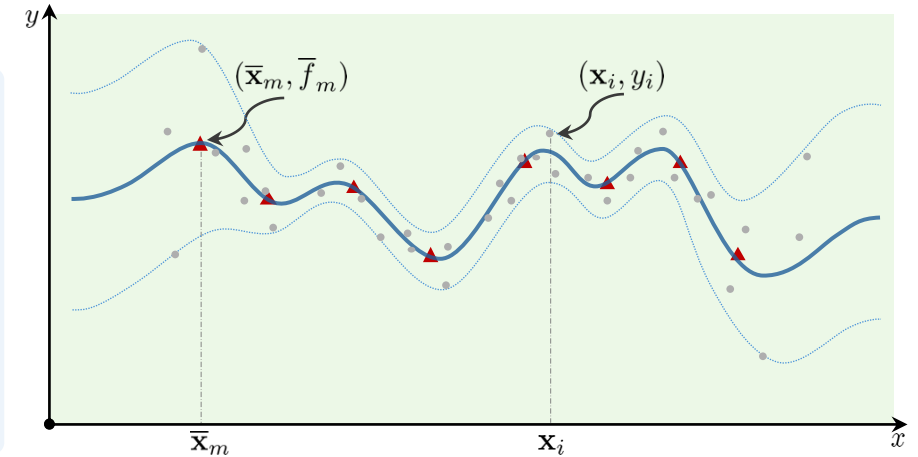
SVGP with Gaussian Likelihood

- When the likelihood is Gaussian $\mathcal{N}(y_i | f_i, \sigma^2)$

$$\mathbb{E}_{p(\mathbf{f}|\bar{\mathbf{f}})q(\bar{\mathbf{f}})}[\log p(\mathbf{y} | \mathbf{f})] - \text{KL}[q(\bar{\mathbf{f}})||p(\bar{\mathbf{f}})]$$

- Uncollapsed bound with mini-batching

$$\begin{aligned} \log \mathcal{N}(\mathbf{y} | \mathbf{K}_{\mathbf{ff}}\mathbf{K}_{\mathbf{ff}}^{-1}\mathbf{m}, \sigma^2\mathbf{I}) - \frac{\text{tr}\{\mathbf{K}_{q(\mathbf{f})}\}}{2\sigma^2} \\ \approx \frac{N}{|\mathcal{B}|} \sum_{i \in \mathcal{B}} \left\{ \log \mathcal{N}(y_i | \mathbf{K}_{f_i\bar{\mathbf{f}}}\mathbf{K}_{\mathbf{ff}}^{-1}\mathbf{m}, \sigma^2) - \frac{\text{tr}\{\mathbf{K}_{q(f_i)}\}}{2\sigma^2} \right\} \end{aligned}$$



- Collapsed bound with analytical $q(\bar{\mathbf{f}})$

$$q^*(\bar{\mathbf{f}}) = \mathcal{N}(\bar{\mathbf{f}} | \sigma^{-2}\mathbf{K}_{\mathbf{ff}}\hat{\Sigma}^{-1}\mathbf{K}_{\mathbf{ff}}\mathbf{y}, \mathbf{K}_{\mathbf{ff}}\hat{\Sigma}^{-1}\mathbf{K}_{\mathbf{ff}}) \quad \hat{\Sigma} = \mathbf{K}_{\mathbf{ff}} + \sigma^{-2}\mathbf{K}_{\mathbf{ff}}\mathbf{K}_{\mathbf{ff}} \quad \text{No mini-batching}$$

$$\text{ELBO} = \mathcal{N}(\mathbf{y} | \mathbf{0}, \mathbf{K}_{\mathbf{ff}}\mathbf{K}_{\mathbf{ff}}^{-1}\mathbf{K}_{\mathbf{ff}} + \sigma^2\mathbf{I}) - \frac{1}{2\sigma^2} \text{tr}\{\mathbf{K}_{\mathbf{ff}} - \mathbf{K}_{\mathbf{ff}}\mathbf{K}_{\mathbf{ff}}^{-1}\mathbf{K}_{\mathbf{ff}}\}$$

SVGPVAE: Amortised Inducing Points

- Start from the inducing-point distribution

$$p(\bar{\mathbf{f}}) = \mathcal{N}(\bar{\mathbf{f}} \mid \mathbf{0}, \mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}})$$

$$q^*(\bar{\mathbf{f}}) = \mathcal{N}(\bar{\mathbf{f}} \mid \sigma^{-2} \mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}} \hat{\Sigma}^{-1} \mathbf{K}_{\bar{\mathbf{f}}\mathbf{y}}, \mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}} \hat{\Sigma}^{-1} \mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}}) \quad \hat{\Sigma} = \mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}} + \sigma^{-2} \mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}} \mathbf{K}_{\bar{\mathbf{f}}\bar{\mathbf{f}}}$$

- Introduce inducing points in latent space

$$p(\bar{\mathbf{Z}}) = \mathcal{N}(\mathbf{0}, \mathbf{K}_{\bar{\mathbf{Z}}\bar{\mathbf{Z}}}), \quad k(\bar{\mathbf{X}}, \bar{\mathbf{X}}) = \mathbf{K}_{\bar{\mathbf{Z}}\bar{\mathbf{Z}}}$$

- Design variational posterior

- Encoder outputs pseudo-observations / variances for a mini-batch $(\mathbf{X}_{\mathcal{B}}, \mathbf{Y}_{\mathcal{B}}, \mathbf{Z}_{\mathcal{B}})$

$$\mathbf{y} \longrightarrow \boldsymbol{\mu}_{\mathcal{B}}, \quad \sigma^2 \longrightarrow \boldsymbol{\sigma}_{\mathcal{B}}^2$$

- Mini-batch version of “ $q^*(\bar{\mathbf{f}})$ ”: $q(\bar{\mathbf{Z}} \mid \mathbf{Y}_{\mathcal{B}}) = \mathcal{N}(\bar{\mathbf{Z}} \mid \mathbf{m}_{\mathcal{B}}, \mathbf{S}_{\mathcal{B}})$

$$\mathbf{m}_{\mathcal{B}} = \frac{N}{|\mathcal{B}|} \mathbf{K}_{\bar{\mathbf{Z}}\bar{\mathbf{Z}}} \boldsymbol{\Sigma}_{\mathcal{B}}^{-1} \mathbf{K}_{\bar{\mathbf{Z}}\mathbf{z}_{\mathcal{B}}} \boldsymbol{\sigma}_{\mathcal{B}}^{-2} \boldsymbol{\mu}_{\mathcal{B}} \quad \mathbf{S}_{\mathcal{B}} = \mathbf{K}_{\bar{\mathbf{Z}}\bar{\mathbf{Z}}} \boldsymbol{\Sigma}_{\mathcal{B}}^{-1} \mathbf{K}_{\bar{\mathbf{Z}}\bar{\mathbf{Z}}} \quad \boldsymbol{\Sigma}_{\mathcal{B}} = \mathbf{K}_{\bar{\mathbf{Z}}\bar{\mathbf{Z}}} + \frac{N}{|\mathcal{B}|} \mathbf{K}_{\bar{\mathbf{Z}}\mathbf{z}_{\mathcal{B}}} \boldsymbol{\sigma}_{\mathcal{B}}^{-2} \mathbf{K}_{\mathbf{z}_{\mathcal{B}}\bar{\mathbf{Z}}}$$

- Variational posterior: $q(\mathbf{z}_n \mid \mathbf{Y}_{\mathcal{B}}) = \int p(\mathbf{z}_n \mid \bar{\mathbf{Z}}_{\mathcal{B}}) q(\bar{\mathbf{Z}}_{\mathcal{B}} \mid \mathbf{Y}_{\mathcal{B}}) d\bar{\mathbf{Z}}_{\mathcal{B}}$

GPVAE with Inducing Points

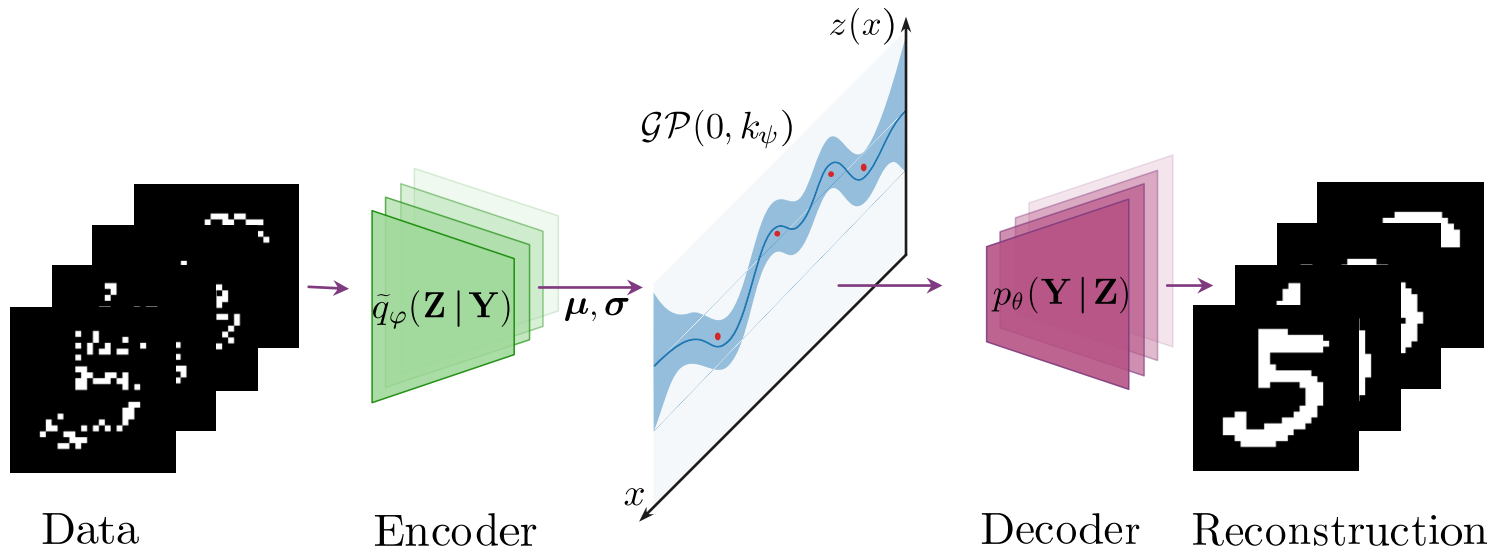
- Training objective based on VAE and SVGP formulation

$$\mathcal{L}_{\text{SVGPVAE}} = \frac{N}{|\mathcal{B}|} \sum_{n \in \mathcal{B}} \mathbb{E}_{q(\mathbf{z}_n | \mathbf{Y}_{\mathcal{B}})} [\log p(\mathbf{y}_n | \mathbf{z}_n)] - \text{KL}[q(\bar{\mathbf{Z}} | \mathbf{Y}_{\mathcal{B}}) \| p(\bar{\mathbf{Z}})]$$

Simplified reconstruction-minus-KL training objective

- Prediction through Gaussian conditional

$$q(\mathbf{z}_* | \mathbf{Y}) = \int p(\mathbf{z}_* | \bar{\mathbf{Z}}) q(\bar{\mathbf{Z}} | \mathbf{Y}) d\bar{\mathbf{Z}} \quad \text{Use full-batch or mini-batch}$$



Contents

❖ Background

- GP

- From VAE to GPVAE

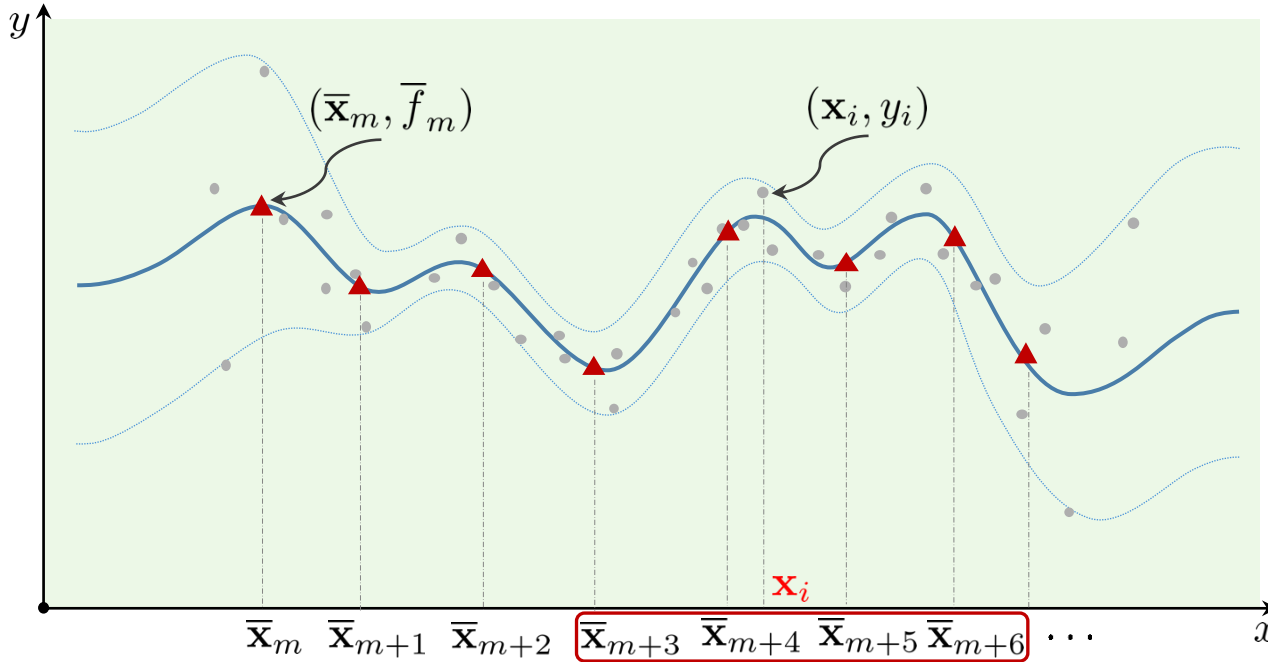
❖ GPVAE Models

- GPVAE with Inducing Points

- Nearest-Neighbour GPVAE

Nearest-Neighbour GP (NNGP)

- Consider the H -nearest inducing points of SVGP



neighbours of $\mathbf{x}_i$:

$$n(i) = \{m + 3, m + 4, m + 5, m + 6\}$$

- NNGP: “sparse within sparse”

$$p(f_i | \bar{\mathbf{f}}) \longrightarrow p(f_i | \bar{\mathbf{f}}_{n(i)})$$

$$\mathcal{N}(f_i | \mathbf{K}_{f_i \bar{\mathbf{f}}} \mathbf{K}_{\bar{\mathbf{f}} \bar{\mathbf{f}}}^{-1} \bar{\mathbf{f}}, K_{f_i f_i} - \mathbf{K}_{f_i \bar{\mathbf{f}}} \mathbf{K}_{\bar{\mathbf{f}} \bar{\mathbf{f}}}^{-1} \mathbf{K}_{\bar{\mathbf{f}} f_i})$$

$$q(\bar{\mathbf{f}}) \longrightarrow q(\bar{\mathbf{f}}_{n(i)})$$

$$p(\bar{\mathbf{f}}) \longrightarrow p(\bar{\mathbf{f}}_{n(i)})$$

- ELBO:

$$\frac{N}{|\mathcal{B}|} \sum_{i \in \mathcal{B}} \mathbb{E}_{p(f_i | \bar{\mathbf{f}}) q(\bar{\mathbf{f}})} [\log p(y_i | f_i)] - \text{KL}[q(\bar{\mathbf{f}}) \| p(\bar{\mathbf{f}})]$$

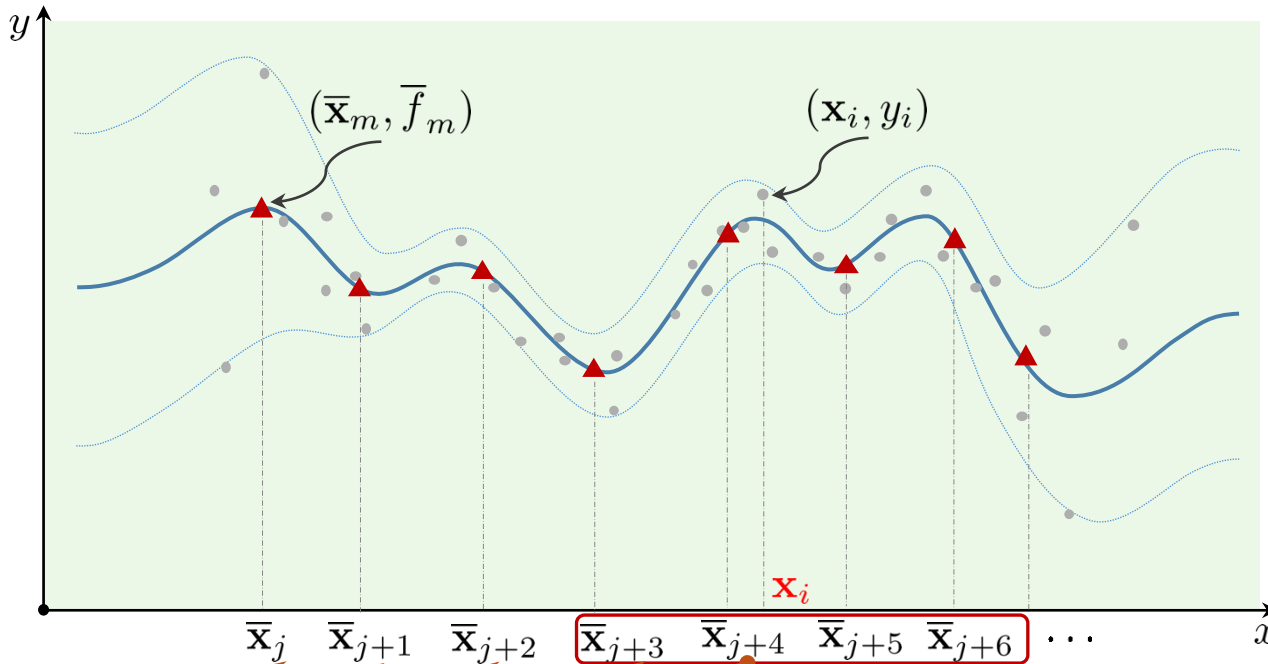


$$\frac{N}{|\mathcal{B}|} \sum_{i \in \mathcal{B}} \mathbb{E}_{p(f_i | \bar{\mathbf{f}}_{n(i)}) q(\bar{\mathbf{f}}_{n(i)})} [\log p(y_i | f_i)]$$

$$- \frac{1}{|\mathcal{B}|} \sum_{i \in \mathcal{B}} \text{KL}[q(\bar{\mathbf{f}}_{n(i)}) \| p(\bar{\mathbf{f}}_{n(i)})]$$

NNGP: Vecchia Approximation

- Apply the Vecchia approximation to prior $p(\bar{\mathbf{f}})$



neighbours of $\mathbf{x}_i$:

$$n(i) = \{j + 3, j + 4, j + 5, j + 6\}$$

neighbours of $\bar{\mathbf{x}}_{j+4}$ (selected by “looking backward”):

$$n(j + 4) = \{j, j + 1, j + 2, j + 3\}$$

- NNGP based on Vecchia approximation

$$p(f_i | \bar{\mathbf{f}}) \longrightarrow p(f_i | \bar{\mathbf{f}}_{n(i)})$$

$$q(\bar{\mathbf{f}}) \longrightarrow \text{mean-field } q(\bar{\mathbf{f}})$$

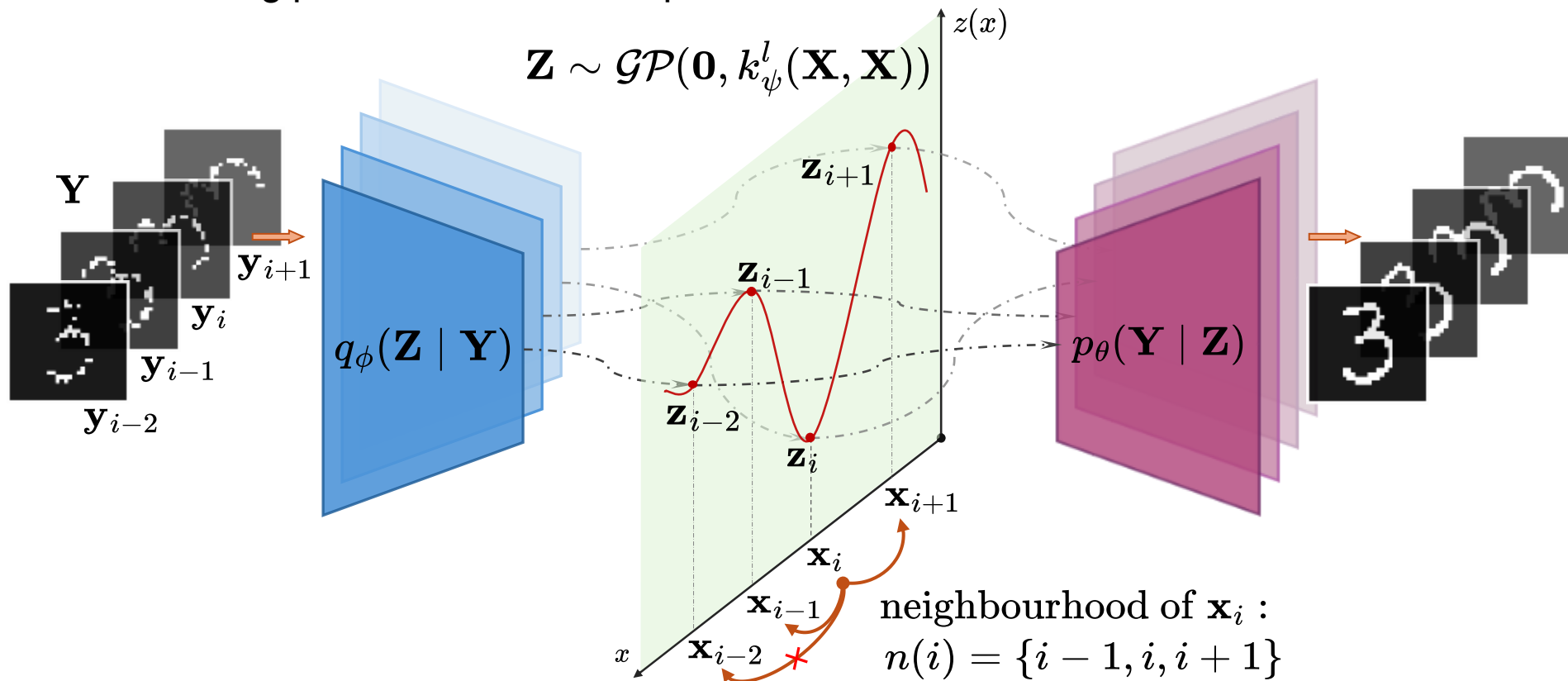
$$p(\bar{\mathbf{f}}) \longrightarrow p(\bar{f}_1) \prod_{j=2}^M p(\bar{f}_m | \bar{\mathbf{f}}_{n(j)})$$

- ELBO:

$$\begin{aligned} & \frac{N}{|\mathcal{I}|} \sum_{i \in \mathcal{I}} [\mathbb{E}_{q(f_i)} \log p(y_i | f_i)] \\ & - \frac{M}{|\mathcal{J}|} \sum_{j \in \mathcal{J}} \mathbb{E}_{q(\bar{\mathbf{f}}_{n(j)})} \{ \text{KL}[q(\bar{f}_j) \| p(\bar{f}_j | \bar{\mathbf{f}}_{n(j)})] \} \end{aligned}$$

NNGPVAE: Talk Only to Nearest Neighbours

- Inspired by the “first law of geography” in recent NNGP
- Focusing on a small set of nearest neighbours captures most of the essential correlation structures
- Place each inducing point over each data point



Two Neighbour-Driven Prior Approximations

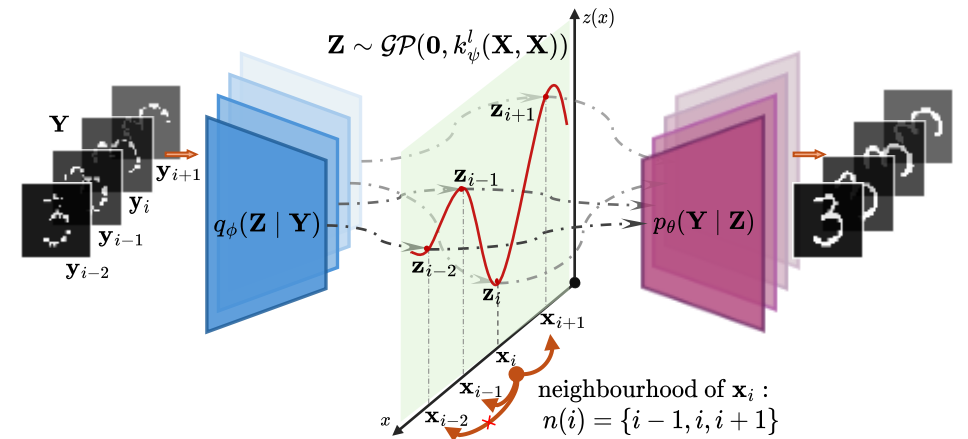
- Hierarchical Prior Approximation (HPA): sparsify covariance via hierarchical masks

$$\mathcal{L}_{\text{HPA}} \approx \frac{N}{|\mathcal{I}|} \sum_{i \in \mathcal{I}} \left\{ \mathbb{E}_{q(\mathbf{z}_i | \mathbf{y}_i)} [\log p(\mathbf{y}_i | \mathbf{z}_i)] - \frac{1}{N} \text{KL} \left[q(\mathbf{Z}_{n(i)}) \parallel p(\mathbf{Z}_{n(i)}) \right] \right\}$$

- Sparse Precision Approximation (SPA): sparsify precision matrix via chained conditionals

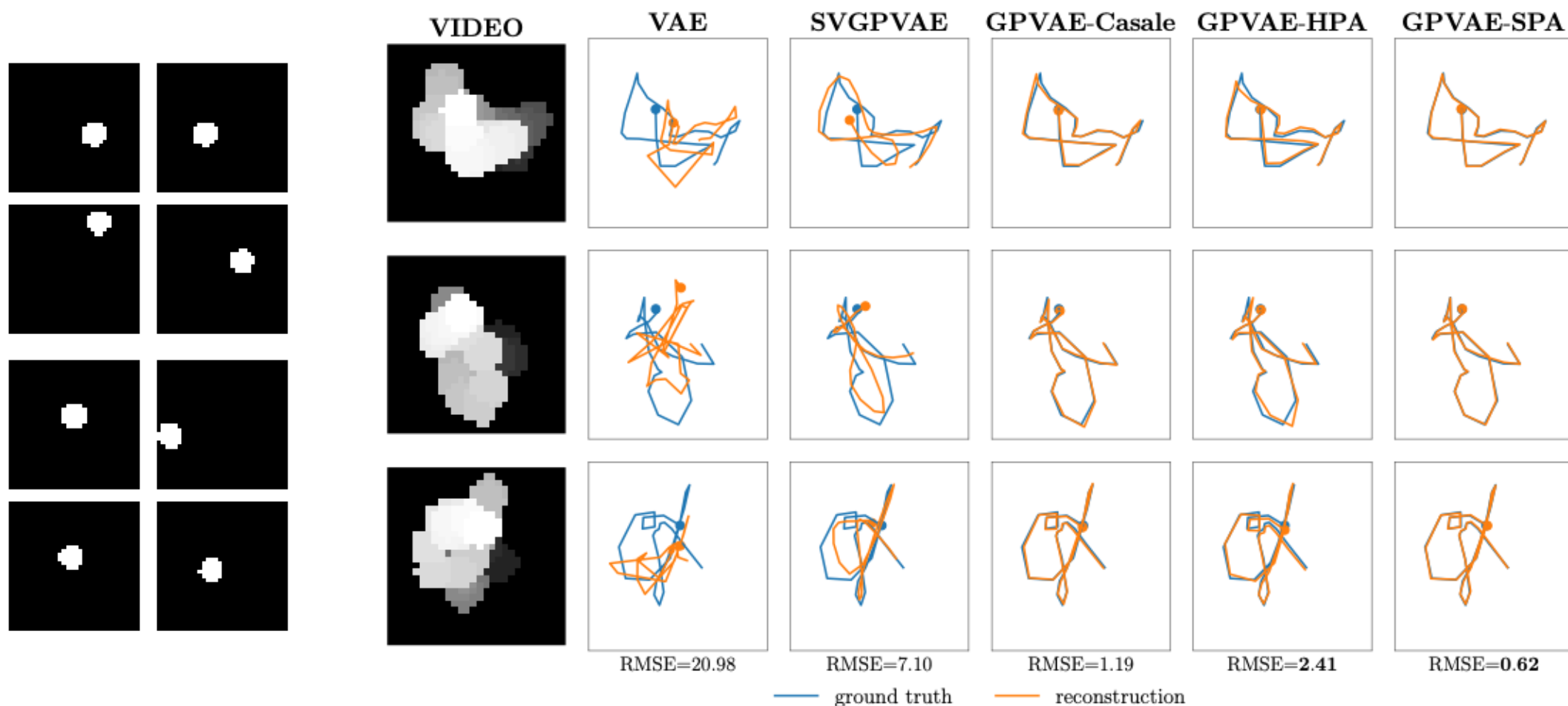
$$\mathcal{L}_{\text{SPA}} \approx \frac{N}{|\mathcal{I}|} \sum_{i \in \mathcal{I}} \mathbb{E}_{q(\mathbf{z}_i | \mathbf{y}_i)} [\log p(\mathbf{y}_i | \mathbf{z}_i)] - \frac{N}{|\mathcal{J}|} \sum_{j \in \mathcal{J}} \mathbb{E}_{q(\mathbf{z}_{n(j)})} \text{KL} \left[q(\mathbf{z}_j) \parallel p(\mathbf{z}_j | \mathbf{Z}_{n(j)}) \right]$$

KL only involves H ($\ll N$) nearest neighbouring variables!



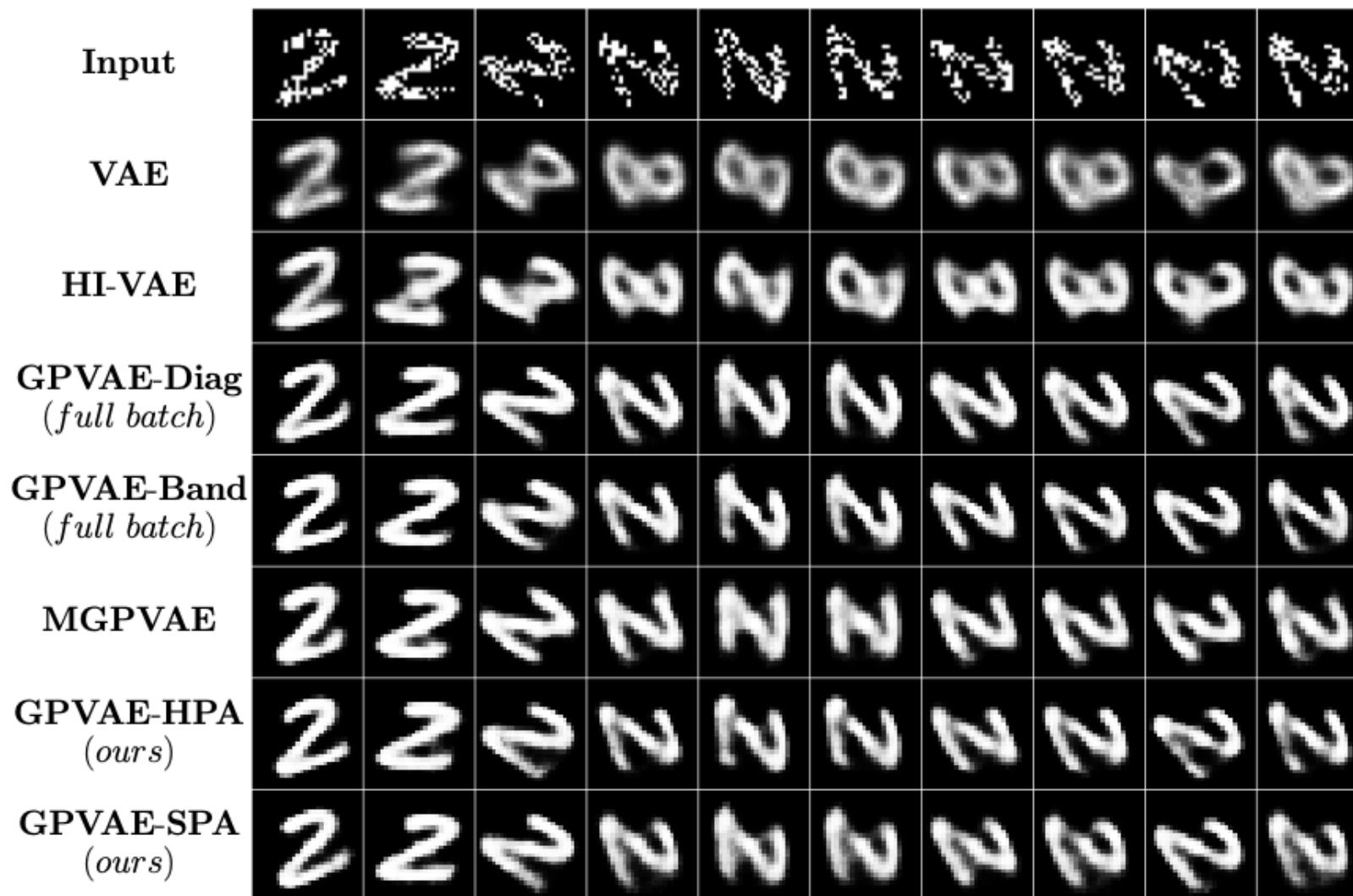
Experiments: Latent Representation in Moving Ball

- Thousands of black-and-white videos of 30 frames



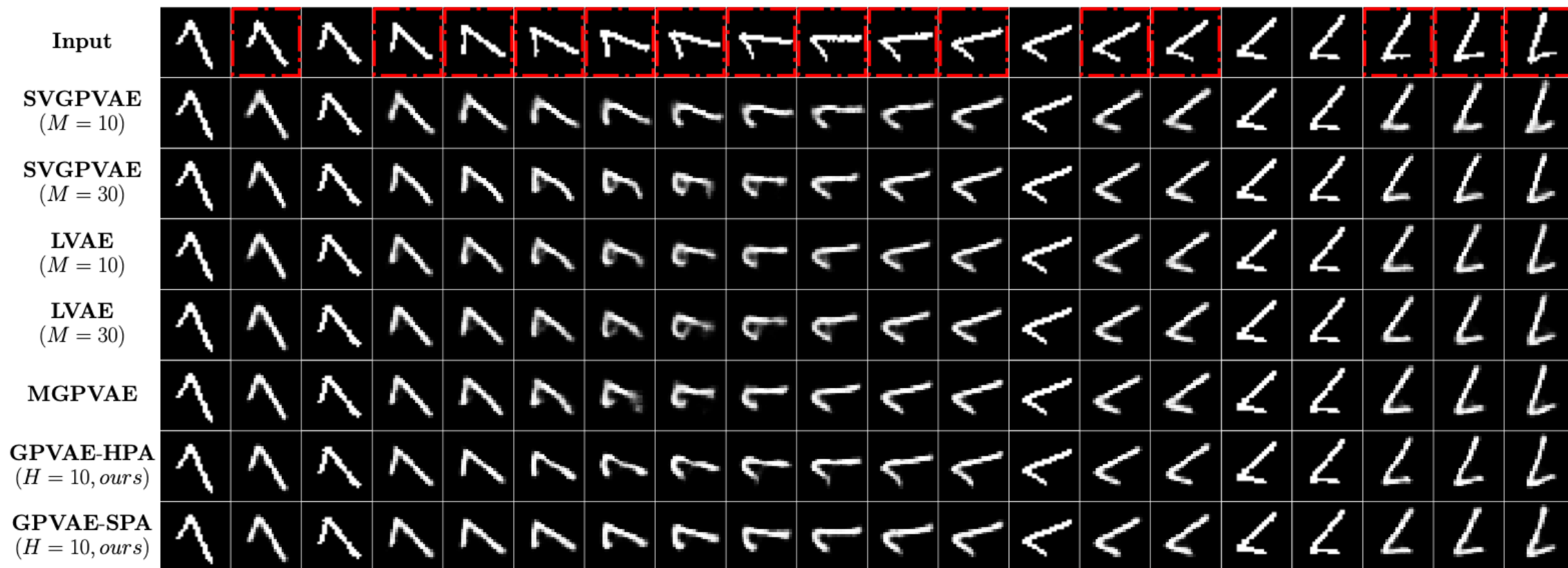
Experiments: Impute Corrupted Frames in Rotated MNIST

- 50,000 training sequences of 10 frames with 60% pixels absent



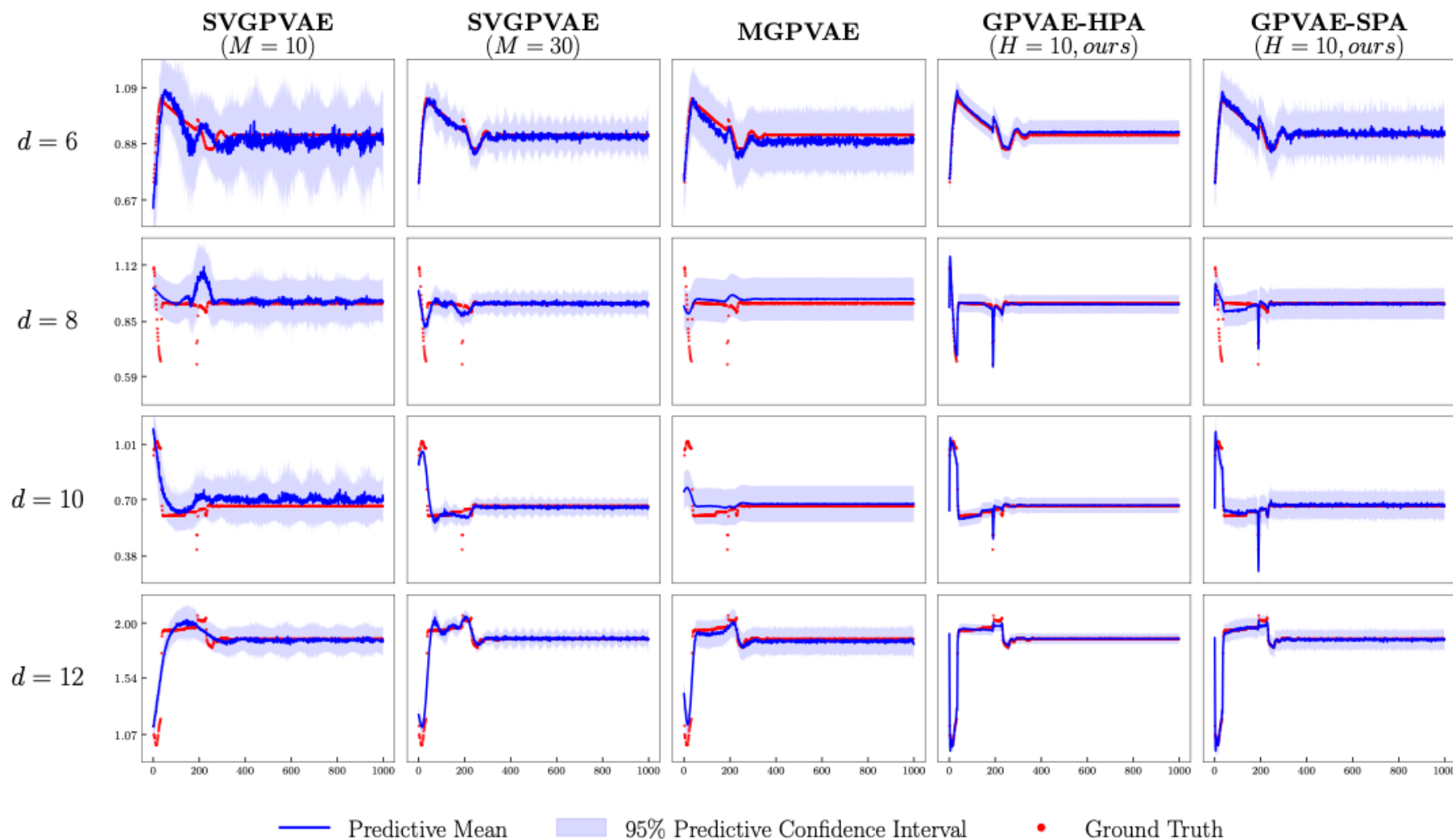
Experiments: Generate Missing Frames in Rotated MNIST

- 100 frames, each with a 0.6 probability of being dropped.



Experiments: Generate Missing States in MuJoCo

- 1000 timesteps, each with a 0.6 probability of being dropped



Experiments: Impute Missing Locations in SPE10

- $\sim 10^5$ locations of reservoir properties, high spatial heterogeneity

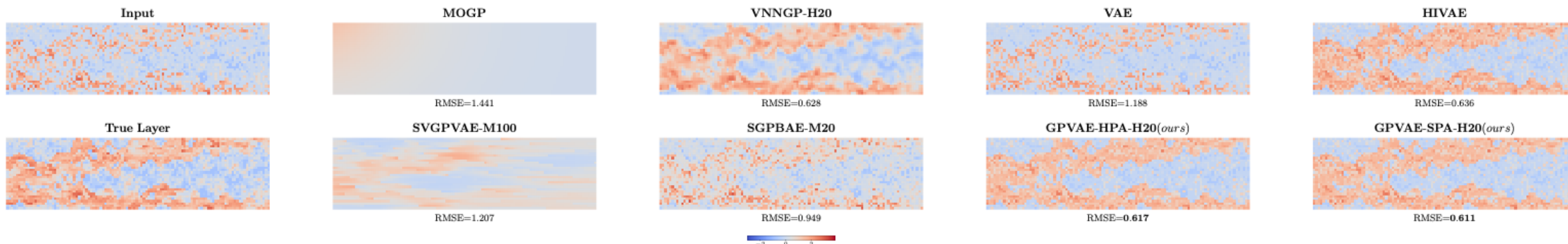


Table 3: Imputation results of various models on SPE10 dataset.

Models	NLL	RMSE	Training Time (s/epoch)
MOGP-M1000 (Hensman et al., 2013)	6.034 ± 0.003	1.230 ± 0.000	3.820 ± 0.021
VNNGP-H20 (Wu et al., 2022)	1.052 ± 0.000	0.683 ± 0.000	2.490 ± 0.084
VAE (Kingma & Welling, 2014)	3.436 ± 0.015	1.298 ± 0.002	1.416 ± 0.020
HI-VAE (Nazabal et al., 2020)	0.752 ± 0.034	0.647 ± 0.008	1.448 ± 0.011
SVGPVAE-M100 (Jazbec et al., 2021)	2.025 ± 0.020	0.979 ± 0.006	4.755 ± 0.097
SGPBAE-M20 (Tran et al., 2023)	1.335 ± 0.106	0.898 ± 0.125	1585.6 ± 21.3
GPVAE-HPA-H20	0.735 ± 0.007	0.638 ± 0.007	4.019 ± 0.036
GPVAE-SPA-H20	0.736 ± 0.004	0.638 ± 0.004	3.365 ± 0.117

Conclusion

- **SVGP $\rightarrow$ SVGPVAE**
 - Inducing variables for latent GPs
 - A reconstruction-minus-KL objective for scalable training
 - No special kernel structures
- **NNGP $\rightarrow$ NNGPVAE**
 - HPA & SPA: covariance & precision sparsity in GPVAE
 - No special kernel structures
 - Do not rely on large inducing points