

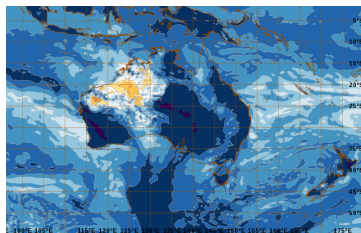
Tensor trains for sequential state and parameter estimation in state-space models

In collaboration with Yiran Zhao

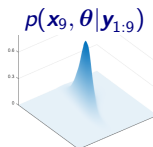
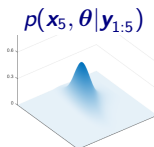
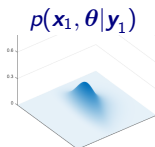
Tiangang Cui
University of Sydney

`tiangang.cui@sydney.edu.au`
`https://www.fastfins.org/`

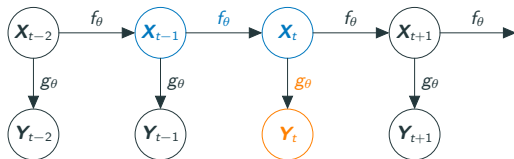
Institute for Mathematical and Statistical Innovation
October 6, 2025



- Aim to **sequentially** estimate hidden states $\{\mathbf{x}_t\}_{t=0}^T$ and parameters θ of stochastic dynamical systems from indirect data $\{\mathbf{y}_t\}_{t=1}^T$
- Need to **recursively** learn posterior random variables



State-space models



- State evolution and observation follow the **hidden Markov property**

$$\text{State transition} \quad X_t | (X_{t-1} = x) \sim f(\cdot | x, \theta)$$

$$\text{Likelihood function} \quad Y_t | (X_t = x) \sim g(\cdot | x, \theta)$$

- Filtering** estimates the current state using observations up to time t

$$p(x_t, \theta | y_{1:t}) \propto \int p(x_{0:t}, \theta | y_{1:t}) dx_{0:t-1}$$

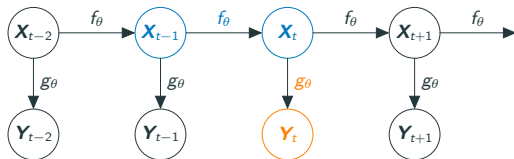
- Smoothing** estimates the state using future observations

$$p(x_t, \theta | y_{1:T}) \propto \int p(x_{1:T}, \theta | y_{0:T}) dx_{0:t-1} dx_{t+1:T}$$

- Parameter estimation:**

$$p(\theta | y_{1:t}) \propto \int p(x_{0:t}, \theta | y_{1:t}) dx_{0:t}$$

State-space models



- State evolution and observation follow the **hidden Markov property**

$$\text{State transition} \quad X_t | (X_{t-1} = x) \sim f(\cdot | x, \theta)$$

$$\text{Likelihood function} \quad Y_t | (X_t = x) \sim g(\cdot | x, \theta)$$

- Filtering** estimates the current state using observations up to time t

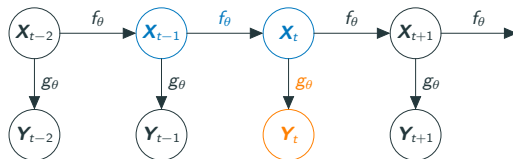
$$p(x_t, \theta | y_{1:t}) \propto \int p(x_{0:t}, \theta | y_{1:t}) dx_{0:t-1}$$

- Smoothing** estimates the state using future observations

$$p(x_t, \theta | y_{1:T}) \propto \int p(x_{1:T}, \theta | y_{0:T}) dx_{0:t-1} dx_{t+1:T}$$

- Parameter estimation:**

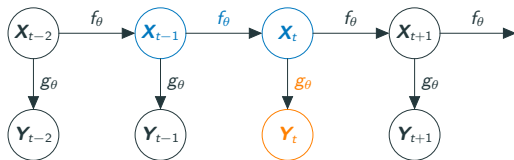
$$p(\theta | y_{1:t}) \propto \int p(x_{0:t}, \theta | y_{1:t}) dx_{0:t}$$



- The joint posterior density is recursively defined:

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \boldsymbol{\theta} | \mathbf{y}_{1:t}) \propto p(\mathbf{x}_{t-1}, \boldsymbol{\theta} | \mathbf{y}_{1:t-1}) f(\mathbf{x}_t | \mathbf{x}_{t-1}, \boldsymbol{\theta}) g(\mathbf{y}_t | \mathbf{x}_t, \boldsymbol{\theta})$$

- Existing methods:
 - **Without parameter:** particle filters [Del Moral, Doucet, Liu, ...], TT-Zakai [Li et al. 2022], Frobenius–Perron [Fox et al. 2021], (optimal) transport maps [Gottwald, Reich, Baptista, ...]
 - **With parameter:** SMC² and nested filters (filters within filters) [Chopin et al. 2013; Crisan et al. 2018; etc.], Knothe–Rosenblatt rearrangement [Spantini et al. 2018]
- We propose one method for all sequential estimation problems:
 1. Recursive posterior approximation and marginalization using tensor trains
 2. Filtering and smoothing via accompanying Knothe–Rosenblatt rearrangement
 3. Examples



- The joint posterior density is recursively defined:

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \boldsymbol{\theta} | \mathbf{y}_{1:t}) \propto p(\mathbf{x}_{t-1}, \boldsymbol{\theta} | \mathbf{y}_{1:t-1}) f(\mathbf{x}_t | \mathbf{x}_{t-1}, \boldsymbol{\theta}) g(\mathbf{y}_t | \mathbf{x}_t, \boldsymbol{\theta})$$

- Existing methods:

- **Without parameter:** particle filters [Del Moral, Doucet, Liu, ...], TT-Zakai [Li et al. 2022], Frobenius–Perron [Fox et al. 2021], (optimal) transport maps [Gottwald, Reich, Baptista, ...]
- **With parameter:** SMC² and nested filters (filters within filters) [Chopin et al. 2013; Crisan et al. 2018; etc.], Knothe–Rosenblatt rearrangement [Spantini et al. 2018]

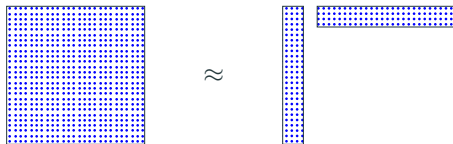
- **We propose one method for all sequential estimation problems:**

1. Recursive posterior approximation and marginalization using tensor trains
2. Filtering and smoothing via accompanying Knothe–Rosenblatt rearrangement
3. Examples

Recursive approximation

Functional tensor train: what happened when $d = 2$

- 2D: discretisation of $\pi(x_1, x_2)$ gives a matrix $A(i, j)$



- Function/matrix decomposition

$$\pi(x_1, x_2) = \sum_{\alpha_1=1}^r \varphi_1(x_1, \alpha_1) \varphi_2(\alpha_1, x_2) + \mathcal{O}(\varepsilon)$$

$$A(i, j) = \sum_{\alpha_1=1}^r V(i, \alpha_1) W(\alpha_1, j) + \mathcal{O}(\varepsilon)$$

- Rank $r \ll n$ to save storage and computing cost. Separates x_1 and x_2

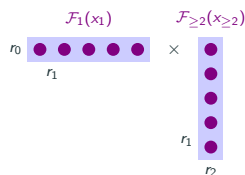
$$\pi(x_2) = \int \pi(\mathbf{x}_1, x_2) d\mathbf{x}_1 \approx \sum_{\alpha_1=1}^r \left(\int \varphi_1(\mathbf{x}_1, \alpha_1) d\mathbf{x}_1 \right) \varphi_2(\alpha_1, x_2)$$

$$\pi(x_1) = \int \pi(x_1, \mathbf{x}_2) d\mathbf{x}_2 \approx \sum_{\alpha_1=1}^r \varphi_1(x_1, \alpha_1) \left(\int \varphi_2(\alpha_1, \mathbf{x}_2) d\mathbf{x}_2 \right)$$

How to build a functional tensor train $d > 2$ (on pen and paper)?

- Group variables and apply, e.g., SVD, we obtain

$$\pi(x_1, \underbrace{x_2, \dots, x_d}_{x_{\geq 2}}) = \pi(x_1, x_{\geq 2}) \approx \sum_{\alpha_1=1}^{r_1} \varphi_1(x_1, \alpha_1) \varphi_{\geq 2}(\alpha_1, x_{\geq 2})$$

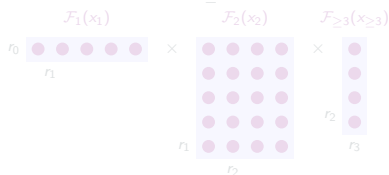


$$\mathcal{F}_1(x_1) = [\varphi_1(x_1, 1), \dots, \varphi_1(x_1, r_1)]$$

$$\mathcal{F}_{\geq 2}(x_{\geq 2}) = [\varphi_{\geq 2}(1, x_{\geq 2}), \dots, \varphi_{\geq 2}(r_1, x_{\geq 2})]^T$$

- Decompose $\mathcal{F}_{\geq 2}(x_{\geq 2})$:

$$\varphi_{\geq 2}(\underbrace{\alpha_1, x_2, x_{\geq 3}}_{\tilde{x}_{\leq 2}}) = \varphi_{\geq 2}(\tilde{x}_{\leq 2}, x_{\geq 3}) \approx \sum_{\alpha_2=1}^{r_2} \varphi_2(\alpha_1, x_2, \alpha_2) \varphi_{\geq 3}(\alpha_2, x_{\geq 3})$$



$$\mathcal{F}_2(x_2) = \begin{bmatrix} \varphi_2(1, x_2, 1) & \cdots & \varphi_2(1, x_2, r_2) \\ \vdots & \ddots & \vdots \\ \varphi_2(r_1, x_2, 1) & \cdots & \varphi_2(r_1, x_2, r_2) \end{bmatrix}$$

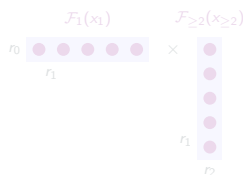
$$\mathcal{F}_{\geq 3}(x_{\geq 3}) = [\varphi_{\geq 3}(1, x_{\geq 3}), \dots, \varphi_{\geq 3}(r_2, x_{\geq 3})]^T$$

- Decompose $\mathcal{F}_{\geq 3}(x_{\geq 3})$ and cont ...

How to build a functional tensor train $d > 2$ (on pen and paper)?

- Group variables and apply, e.g., SVD, we obtain

$$\pi(x_1, \underbrace{x_2, \dots, x_d}_{x_{\geq 2}}) = \pi(x_1, x_{\geq 2}) \approx \sum_{\alpha_1=1}^{r_1} \varphi_1(x_1, \alpha_1) \varphi_{\geq 2}(\alpha_1, x_{\geq 2})$$

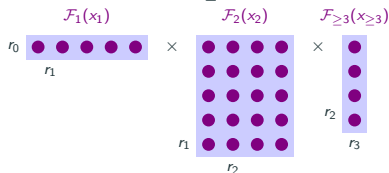


$$\mathcal{F}_1(x_1) = [\varphi_1(x_1, 1), \dots, \varphi_1(x_1, r_1)]$$

$$\mathcal{F}_{\geq 2}(x_{\geq 2}) = [\varphi_{\geq 2}(1, x_{\geq 2}), \dots, \varphi_{\geq 2}(r_1, x_{\geq 2})]^\top$$

- Decompose $\mathcal{F}_{\geq 2}(x_{\geq 2})$:

$$\varphi_{\geq 2}(\underbrace{\alpha_1, x_2}_{\tilde{x}_{\leq 2}}, x_{\geq 3}) = \varphi_{\geq 2}(\tilde{x}_{\leq 2}, x_{\geq 3}) \approx \sum_{\alpha_2=1}^{r_2} \varphi_2(\alpha_1, x_2, \alpha_2) \varphi_{\geq 3}(\alpha_2, x_{\geq 3})$$



$$\mathcal{F}_2(x_2) = \begin{bmatrix} \varphi_2(1, x_2, 1) & \cdots & \varphi_2(1, x_2, r_2) \\ \vdots & \ddots & \vdots \\ \varphi_2(r_1, x_2, 1) & \cdots & \varphi_2(r_1, x_2, r_2) \end{bmatrix}$$

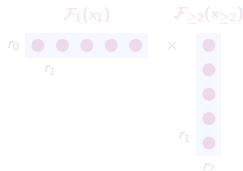
$$\mathcal{F}_{\geq 3}(x_{\geq 3}) = [\varphi_{\geq 3}(1, x_{\geq 3}), \dots, \varphi_{\geq 3}(r_2, x_{\geq 3})]^\top$$

- Decompose $\mathcal{F}_{\geq 3}(x_{\geq 3})$ and cont ...

How to build a functional tensor train $d > 2$ (on pen and paper)?

- Group variables and apply, e.g., SVD, we obtain

$$\pi(x_1, \underbrace{x_2, \dots, x_d}_{x_{\geq 2}}) = \pi(x_1, x_{\geq 2}) \approx \sum_{\alpha_1=1}^{r_1} \varphi_1(x_1, \alpha_1) \varphi_{\geq 2}(\alpha_1, x_{\geq 2})$$

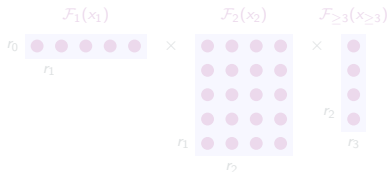


$$\mathcal{F}_1(x_1) = [\varphi_1(x_1, 1), \dots, \varphi_1(x_1, r_1)]$$

$$\mathcal{F}_{\geq 2}(x_{\geq 2}) = [\varphi_{\geq 2}(1, x_{\geq 2}), \dots, \varphi_{\geq 2}(r_1, x_{\geq 2})]^T$$

- Decompose $\mathcal{F}_{\geq 2}(x_{\geq 2})$:

$$\varphi_{\geq 2}(\underbrace{\alpha_1, x_2, x_{\geq 3}}_{\tilde{x}_{\leq 2}}) = \varphi_{\geq 2}(\tilde{x}_{\leq 2}, x_{\geq 3}) \approx \sum_{\alpha_2=1}^{r_2} \varphi_2(\alpha_1, x_2, \alpha_2) \varphi_{\geq 3}(\alpha_2, x_{\geq 3})$$



$$\mathcal{F}_2(x_2) = \begin{bmatrix} \varphi_2(1, x_2, 1) & \cdots & \varphi_2(1, x_2, r_2) \\ \vdots & \ddots & \vdots \\ \varphi_2(r_1, x_2, 1) & \cdots & \varphi_2(r_1, x_2, r_2) \end{bmatrix}$$

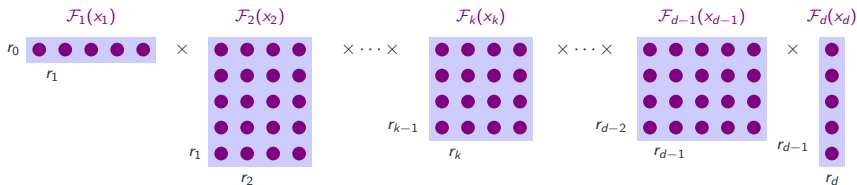
$$\mathcal{F}_{\geq 3}(x_{\geq 3}) = [\varphi_{\geq 3}(1, x_{\geq 3}), \dots, \varphi_{\geq 3}(r_2, x_{\geq 3})]^T$$

- Decompose $\mathcal{F}_{\geq 3}(x_{\geq 3})$ and cont ...

TT factorization $\implies$ separation of variable

$$\pi(x_1, \dots, x_d) \approx \mathcal{F}_1(x_1) \cdots \mathcal{F}_k(x_k) \cdots \mathcal{F}_d(x_d)$$

Each $\mathcal{F}_k : \mathbb{R} \rightarrow \mathbb{R}^{r_{k-1} \times r_k}$ is a *matrix-valued univariate function*



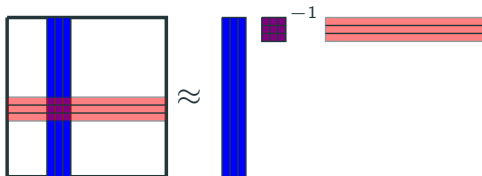
Integration of a matrix-valued function $\mathcal{F}_k(x_k) \implies$ integration over x_k

Cross interpolation: how we build it in practice

- [How to build this?](#) SVD/Schmidt decomposition (unpractical for large d) needs high-dimensional integration
- A rank- r matrix can be recovered by its r indep. rows and cols:

$$A = \hat{A} \equiv A(:, \mathcal{J})A^{-1}(\mathcal{I}, \mathcal{J})A(\mathcal{I}, :),$$

for some *index sets* $\mathcal{I}, \mathcal{J} \subset \{1, \dots, n\}$ of cardinality r .



Tricks are in details:

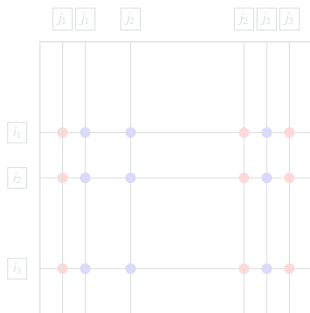
- what are the *best* and *practically realizable* indices $\mathcal{I}, \mathcal{J}$?
- how to generalize to $d > 2$ dimensions?

Tensor train: MaxVol and alternating iteration

Maximum Volume (MaxVol): *best* indices in Chebyshev norm:

- If $|\det A(\mathcal{I}, \mathcal{J})| = \max_{\hat{\mathcal{I}}, \hat{\mathcal{J}}} |\det A(\hat{\mathcal{I}}, \hat{\mathcal{J}})|$, then $\|A - \hat{A}\|_{\infty} \leq (r+1)\sigma_{r+1}(A)$. *
- NP-hard to look through all submatrices, so apply alternating iterations

Given an *initial set* $\mathcal{J} \subset \{1 \dots n\}$, repeat the following



1. $V = A(:, \mathcal{J})$
2. $\mathcal{I} = \text{pivots}(V)$
3. $W = A(\mathcal{I}, :)$.
4. $\mathcal{J} = \text{pivots}(W)$.

pivots via LU or QR with $\mathcal{O}(nr^2)$ flops [Goreinov & Tyrtyshnikov 2001]. Easily generalisable to $d > 2$.

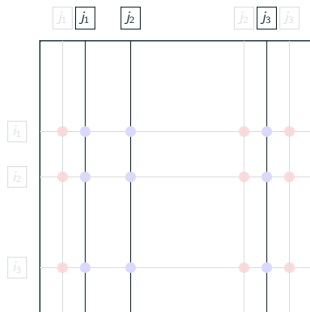
*Frank de Hoog and Markus Hegland further optimised this bound [de Hoog & Hegland 2023]

Tensor train: MaxVol and alternating iteration

Maximum Volume (MaxVol): *best* indices in Chebyshev norm:

- If $|\det A(\mathcal{I}, \mathcal{J})| = \max_{\hat{\mathcal{I}}, \hat{\mathcal{J}}} |\det A(\hat{\mathcal{I}}, \hat{\mathcal{J}})|$, then $\|A - \hat{A}\|_{\infty} \leq (r+1)\sigma_{r+1}(A)$. *
- NP-hard to look through all submatrices, so apply alternating iterations

Given an *initial set* $\mathcal{J} \subset \{1 \dots n\}$, repeat the following



1. $V = A(:, \mathcal{J})$
2. $\mathcal{I} = \text{pivots}(V)$
3. $W = A(\mathcal{I}, :)$.
4. $\mathcal{J} = \text{pivots}(W)$.

pivots via LU or QR with $\mathcal{O}(nr^2)$ flops [Goreinov & Tyrtyshnikov 2001]. Easily generalisable to $d > 2$.

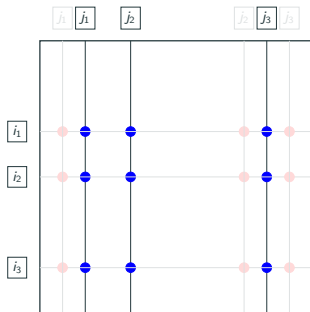
*Frank de Hoog and Markus Hegland further optimised this bound [de Hoog & Hegland 2023]

Tensor train: MaxVol and alternating iteration

Maximum Volume (MaxVol): *best* indices in Chebyshev norm:

- If $|\det A(\mathcal{I}, \mathcal{J})| = \max_{\hat{\mathcal{I}}, \hat{\mathcal{J}}} |\det A(\hat{\mathcal{I}}, \hat{\mathcal{J}})|$, then $\|A - \hat{A}\|_{\infty} \leq (r+1)\sigma_{r+1}(A)$. *
- NP-hard to look through all submatrices, so apply alternating iterations

Given an *initial set* $\mathcal{J} \subset \{1 \dots n\}$, repeat the following



1. $V = A(:, \mathcal{J})$
2. $\mathcal{I} = \text{pivots}(V)$
3. $W = A(\mathcal{I}, :)$.
4. $\mathcal{J} = \text{pivots}(W)$.

pivots via LU or QR with $\mathcal{O}(nr^2)$ flops [Goreinov & Tyrtyshnikov 2001]. Easily generalisable to $d > 2$.

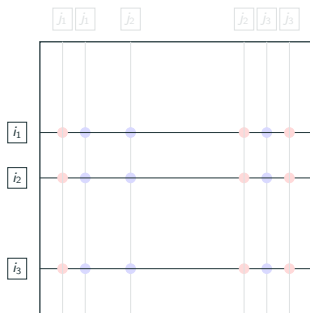
*Frank de Hoog and Markus Hegland further optimised this bound [de Hoog & Hegland 2023]

Tensor train: MaxVol and alternating iteration

Maximum Volume (MaxVol): *best* indices in Chebyshev norm:

- If $|\det A(\mathcal{I}, \mathcal{J})| = \max_{\hat{\mathcal{I}}, \hat{\mathcal{J}}} |\det A(\hat{\mathcal{I}}, \hat{\mathcal{J}})|$, then $\|A - \hat{A}\|_{\infty} \leq (r+1)\sigma_{r+1}(A)$. *
- NP-hard to look through all submatrices, so apply alternating iterations

Given an *initial set* $\mathcal{J} \subset \{1 \dots n\}$, repeat the following



1. $V = A(:, \mathcal{J})$
2. $\mathcal{I} = \text{pivots}(V)$
3. $W = A(\mathcal{I}, :)$.
4. $\mathcal{J} = \text{pivots}(W)$.

pivots via LU or QR with $\mathcal{O}(nr^2)$ flops [Goreinov & Tyrtyshnikov 2001]. Easily generalisable to $d > 2$.

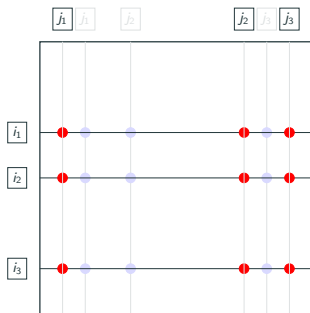
*Frank de Hoog and Markus Hegland further optimised this bound [de Hoog & Hegland 2023]

Tensor train: MaxVol and alternating iteration

Maximum Volume (MaxVol): *best* indices in Chebyshev norm:

- If $|\det A(\mathcal{I}, \mathcal{J})| = \max_{\hat{\mathcal{I}}, \hat{\mathcal{J}}} |\det A(\hat{\mathcal{I}}, \hat{\mathcal{J}})|$, then $\|A - \hat{A}\|_{\infty} \leq (r+1)\sigma_{r+1}(A)$. *
- NP-hard to look through all submatrices, so apply alternating iterations

Given an *initial set* $\mathcal{J} \subset \{1 \dots n\}$, repeat the following



1. $V = A(:, \mathcal{J})$
2. $\mathcal{I} = \text{pivots}(V)$
3. $W = A(\mathcal{I}, :)$.
4. $\mathcal{J} = \text{pivots}(W)$.

pivots via LU or QR with $\mathcal{O}(nr^2)$ flops [Goreinov & Tyrtyshnikov 2001]. Easily generalisable to $d > 2$.

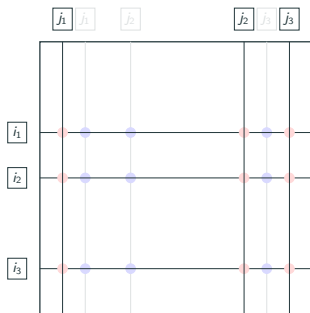
*Frank de Hoog and Markus Hegland further optimised this bound [de Hoog & Hegland 2023]

Tensor train: MaxVol and alternating iteration

Maximum Volume (MaxVol): *best* indices in Chebyshev norm:

- If $|\det A(\mathcal{I}, \mathcal{J})| = \max_{\hat{\mathcal{I}}, \hat{\mathcal{J}}} |\det A(\hat{\mathcal{I}}, \hat{\mathcal{J}})|$, then $\|A - \hat{A}\|_{\infty} \leq (r+1)\sigma_{r+1}(A)$. *
- NP-hard to look through all submatrices, so apply alternating iterations

Given an *initial set* $\mathcal{J} \subset \{1 \dots n\}$, repeat the following



1. $V = A(:, \mathcal{J})$
2. $\mathcal{I} = \text{pivots}(V)$
3. $W = A(\mathcal{I}, :)$.
4. $\mathcal{J} = \text{pivots}(W)$.

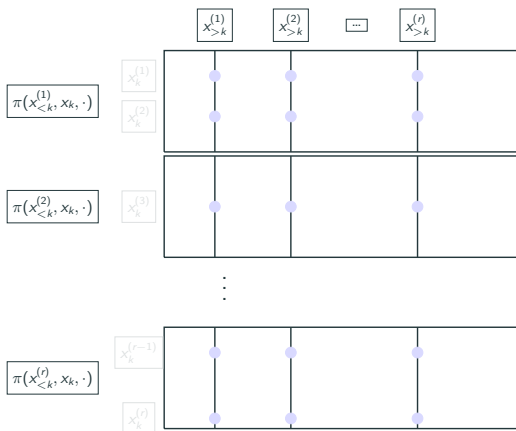
pivots via LU or QR with $\mathcal{O}(nr^2)$ flops [Goreinov & Tyrtyshnikov 2001]. Easily generalisable to $d > 2$.

*Frank de Hoog and Markus Hegland further optimised this bound [de Hoog & Hegland 2023]

Functional TT cross

Grouped coordinates: $x_{<k} = (x_1, x_2, \dots, x_{k-1})$, $x_{>k} = (x_{k+1}, \dots, x_{d-1}, x_d)$

Left and right interpolation sets: $\mathcal{I}_{<k} = \{x_{<k}^{(i)}\}_{i=1}^r$, $\mathcal{J}_{>k} = \{x_{>k}^{(j)}\}_{j=1}^r$



Iterate to next sets:

- Discretize x_k : $\pi(\mathcal{I}_{<k}, x_k, \mathcal{J}_{>k})$ for all $x_k \in \mathcal{I}_k$
- $\mathcal{I}_{<k+1} = \text{MaxVol} \subset \mathcal{I}_{<k} \otimes \mathcal{I}_k$
- $\mathcal{J}_{>k+1} = \{x_{>k+1}^{(j)}\}_{j=1}^r$
- $k \rightarrow k+1$, move to the x_{k+1} .
- $k = d$, switch direction and update $\mathcal{J}_{>k}$.
- Stop if converged, repeat otherwise.
- $O(dr^2)$ function evaluations.

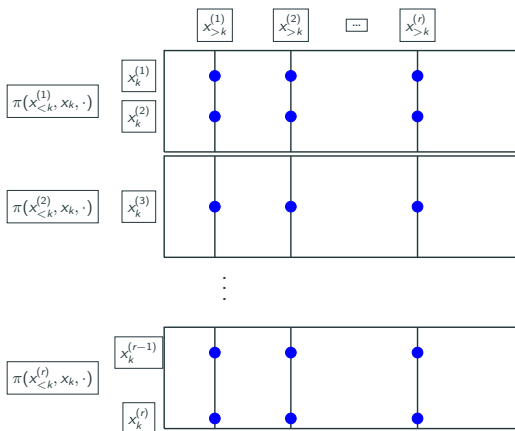
TT cross: [\[Oseledets & Tyrtshnikov 2010\]](#), [\[Gorodetsky, Karaman & Marzouk 2018\]](#)

Empirical interpolation: [\[Mada, Chaturantabut, Sorensen ...\]](#)

Functional TT cross

Grouped coordinates: $x_{<k} = (x_1, x_2, \dots, x_{k-1})$, $x_{>k} = (x_{k+1}, \dots, x_{d-1}, x_d)$

Left and right interpolation sets: $\mathcal{I}_{<k} = \{x_{<k}^{(i)}\}_{i=1}^r$, $\mathcal{J}_{>k} = \{x_{>k}^{(j)}\}_{j=1}^r$



Iterate to next sets:

- Discretize x_k : $\pi(\mathcal{I}_{<k}, \underline{x}_k, \mathcal{J}_{>k})$ for all $x_k \in \mathcal{I}_k$
- $\mathcal{I}_{<k+1} = \text{MaxVol} \subset \mathcal{I}_{<k} \otimes \mathcal{I}_k$
- $\mathcal{J}_{>k+1} = \{x_{>k+1}^{(j)}\}_{j=1}^r$
- $k \rightarrow k+1$, move to the x_{k+1} .
- $k = d$, switch direction and update $\mathcal{J}_{>k}$.
- Stop if converged, repeat otherwise.
- $O(d^2)$ function evaluations.

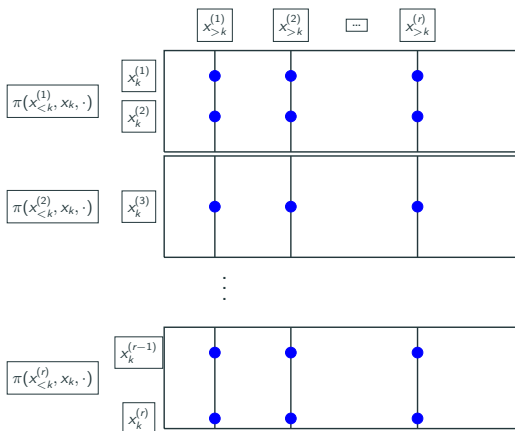
TT cross: [\[Oseledets & Tyrtshnikov 2010\]](#), [\[Gorodetsky, Karaman & Marzouk 2018\]](#)

Empirical interpolation: [\[Madaay, Chaturantabut, Sorensen ...\]](#)

Functional TT cross

Grouped coordinates: $x_{<k} = (x_1, x_2, \dots, x_{k-1})$, $x_{>k} = (x_{k+1}, \dots, x_{d-1}, x_d)$

Left and right interpolation sets: $\mathcal{I}_{<k} = \{x_{<k}^{(i)}\}_{i=1}^r$, $\mathcal{J}_{>k} = \{x_{>k}^{(j)}\}_{j=1}^r$



Iterate to next sets:

- Discretize x_k : $\pi(\mathcal{I}_{<k}, \underline{x}_k, \mathcal{J}_{>k})$ for all $x_k \in \mathcal{I}_k$
- $\mathcal{I}_{<k+1} = \text{MaxVol} \subset \mathcal{I}_{<k} \otimes \mathcal{I}_k$
- $\mathcal{J}_{>k+1} = \{x_{>k+1}^{(j)}\}_{j=1}^r$
- $k \rightarrow k+1$, move to the x_{k+1} .
- $k = d$, switch direction and update $\mathcal{J}_{>k}$.
- Stop if converged, repeat otherwise.
- $O(d^2)$ function evaluations.

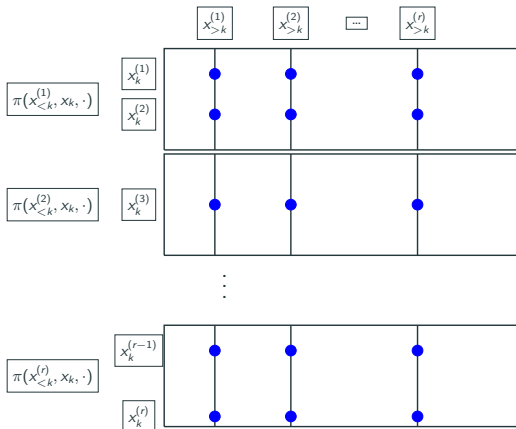
TT cross: [\[Oseledets & Tyrtshnikov 2010\]](#), [\[Gorodetsky, Karaman & Marzouk 2018\]](#)

Empirical interpolation: [\[Madaay, Chaturantabut, Sorensen ...\]](#)

Functional TT cross

Grouped coordinates: $x_{<k} = (x_1, x_2, \dots, x_{k-1})$, $x_{>k} = (x_{k+1}, \dots, x_{d-1}, x_d)$

Left and right interpolation sets: $\mathcal{I}_{<k} = \{x_{<k}^{(i)}\}_{i=1}^r$, $\mathcal{J}_{>k} = \{x_{>k}^{(j)}\}_{j=1}^r$



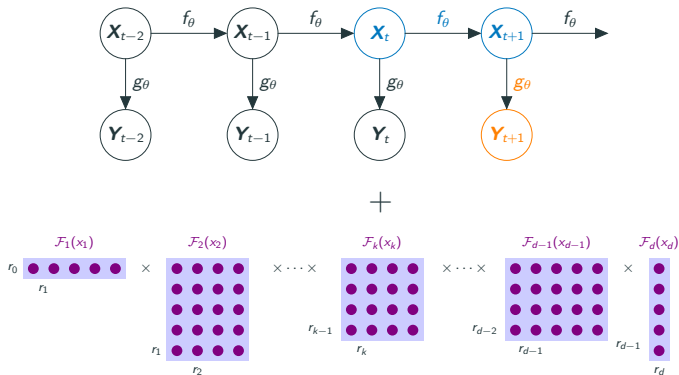
Iterate to next sets:

- Discretize x_k : $\pi(\mathcal{I}_{<k}, \underline{x}_k, \mathcal{J}_{>k})$ for all $x_k \in \mathcal{I}_k$
- $\mathcal{I}_{<k+1} = \text{MaxVol} \subset \mathcal{I}_{<k} \otimes \mathcal{I}_k$
- $\mathcal{J}_{>k+1} = \{x_{>k+1}^{(j)}\}_{j=1}^r$
- $k \rightarrow k+1$, move to the x_{k+1} .
- $k = d$, switch direction and update $\mathcal{J}_{>k}$.
- Stop if converged, repeat otherwise.
- $O(dr^2)$ function evaluations.

TT cross: [\[Oseledets & Tyrtshnikov 2010\]](#), [\[Gorodetsky, Karaman & Marzouk 2018\]](#)

Empirical interpolation: [\[Mada, Chaturantabut, Sorensen ...\]](#)

Recursive approximation and marginalization



Recursive approximation and marginalization

- Previous filtering density:

$$p_{t-1} := p(\mathbf{x}_{t-1}, \boldsymbol{\theta} | \mathbf{y}_{1:t-1})$$

$$p_{t-1}(\mathbf{x}_{t-1}, \boldsymbol{\theta}) \approx \tilde{p}_{t-1}(\mathbf{x}_{t-1}, \boldsymbol{\theta})$$

- Propagate via the joint:

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \boldsymbol{\theta} | \mathbf{y}_{1:t})$$

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \boldsymbol{\theta} | \mathbf{y}_{1:t}) \approx q_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1})$$

$$\propto p_{t-1}(\mathbf{x}_{t-1}, \boldsymbol{\theta}) \ell(\mathbf{x}_t | \mathbf{x}_{t-1}, \boldsymbol{\theta}) g(\mathbf{y}_t | \mathbf{x}_t, \boldsymbol{\theta})$$

$$\propto \tilde{p}_{t-1}(\mathbf{x}_{t-1}, \boldsymbol{\theta}) \ell(\mathbf{x}_t | \mathbf{x}_{t-1}, \boldsymbol{\theta}) g(\mathbf{y}_t | \mathbf{x}_t, \boldsymbol{\theta})$$

- Build ϕ_{TT} such that

$$\|\phi_{\text{TT}}(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) - \sqrt{q_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1})}\|_{L^2} \leq \epsilon_t$$

and then

$$\tilde{\pi}_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = \phi_{\text{TT}}(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1})^2 + \tau \lambda(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}), \quad \tau \leq \epsilon_t$$

- Marginalization for new filtering density

$$p(\mathbf{x}_t, \boldsymbol{\theta} | \mathbf{y}_{1:t}) \propto \int p_t(\mathbf{x}_t, \mathbf{x}_{t-1}, \boldsymbol{\theta}) d\mathbf{x}_{t-1}$$

$$p(\mathbf{x}_t, \boldsymbol{\theta} | \mathbf{y}_{1:t}) \approx \tilde{p}_t(\mathbf{x}_t, \boldsymbol{\theta})$$

$$\propto \int \tilde{\pi}_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) d\mathbf{x}_{t-1}$$

Approximation targets

Computation

Existing approximations

Defensive term

Recursive approximation and marginalization

- Previous filtering density:

$$p_{t-1} := p(\mathbf{x}_{t-1}, \boldsymbol{\theta} | \mathbf{y}_{1:t-1})$$

$$p_{t-1}(\mathbf{x}_{t-1}, \boldsymbol{\theta}) \approx \tilde{p}_{t-1}(\mathbf{x}_{t-1}, \boldsymbol{\theta})$$

- Propagate via the joint:

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \boldsymbol{\theta} | \mathbf{y}_{1:t})$$

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \boldsymbol{\theta} | \mathbf{y}_{1:t}) \approx q_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1})$$

$$\propto p_{t-1}(\mathbf{x}_{t-1}, \boldsymbol{\theta}) f(\mathbf{x}_t | \mathbf{x}_{t-1}, \boldsymbol{\theta}) g(y_t | \mathbf{x}_t, \boldsymbol{\theta})$$

$$\propto \tilde{p}_{t-1}(\mathbf{x}_{t-1}, \boldsymbol{\theta}) f(\mathbf{x}_t | \mathbf{x}_{t-1}, \boldsymbol{\theta}) g(y_t | \mathbf{x}_t, \boldsymbol{\theta})$$

- Build ϕ_{TT} such that

$$\|\phi_{\text{TT}}(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) - \sqrt{q_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1})}\|_{L^2} \leq \epsilon_t$$

and then

$$\tilde{\pi}_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = \phi_{\text{TT}}(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1})^2 + \tau \lambda(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}), \quad \tau \leq \epsilon_t$$

- Marginalization for new filtering density

$$p(\mathbf{x}_t, \boldsymbol{\theta} | \mathbf{y}_{1:t}) \propto \int p_t(\mathbf{x}_t, \mathbf{x}_{t-1}, \boldsymbol{\theta}) d\mathbf{x}_{t-1}$$

$$p(\mathbf{x}_t, \boldsymbol{\theta} | \mathbf{y}_{1:t}) \approx \tilde{p}_t(\mathbf{x}_t, \boldsymbol{\theta})$$

$$\propto \int \tilde{\pi}_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) d\mathbf{x}_{t-1}$$

Approximation targets

Computation

Existing approximations

Defensive term

Recursive approximation and marginalization

- Previous filtering density:

$$p_{t-1} := p(\mathbf{x}_{t-1}, \boldsymbol{\theta} | \mathbf{y}_{1:t-1})$$

$$p_{t-1}(\mathbf{x}_{t-1}, \boldsymbol{\theta}) \approx \tilde{p}_{t-1}(\mathbf{x}_{t-1}, \boldsymbol{\theta})$$

- Propagate via the joint:

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \boldsymbol{\theta} | \mathbf{y}_{1:t})$$

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \boldsymbol{\theta} | \mathbf{y}_{1:t}) \approx q_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1})$$

$$\propto p_{t-1}(\mathbf{x}_{t-1}, \boldsymbol{\theta}) f(\mathbf{x}_t | \mathbf{x}_{t-1}, \boldsymbol{\theta}) g(\mathbf{y}_t | \mathbf{x}_t, \boldsymbol{\theta})$$

$$\propto \tilde{p}_{t-1}(\mathbf{x}_{t-1}, \boldsymbol{\theta}) f(\mathbf{x}_t | \mathbf{x}_{t-1}, \boldsymbol{\theta}) g(\mathbf{y}_t | \mathbf{x}_t, \boldsymbol{\theta})$$

- Build ϕ_{TT} such that

$$\|\phi_{\text{TT}}(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) - \sqrt{q_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1})}\|_{L^2} \leq \epsilon_t$$

and then

$$\tilde{\pi}_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = \phi_{\text{TT}}(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1})^2 + \tau \lambda(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}), \quad \tau \leq \epsilon_t$$

- Marginalization for new filtering density

$$p(\mathbf{x}_t, \boldsymbol{\theta} | \mathbf{y}_{1:t}) \propto \int p_t(\mathbf{x}_t, \mathbf{x}_{t-1}, \boldsymbol{\theta}) d\mathbf{x}_{t-1}$$

$$p(\mathbf{x}_t, \boldsymbol{\theta} | \mathbf{y}_{1:t}) \approx \tilde{p}_t(\mathbf{x}_t, \boldsymbol{\theta})$$

$$\propto \int \tilde{\pi}_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) d\mathbf{x}_{t-1}$$

Approximation targets

Computation

Existing approximations

Defensive term

Recursive approximation and marginalization

- Previous filtering density:

$$p_{t-1} := p(\mathbf{x}_{t-1}, \theta | \mathbf{y}_{1:t-1})$$

$$p_{t-1}(\mathbf{x}_{t-1}, \theta) \approx \tilde{p}_{t-1}(\mathbf{x}_{t-1}, \theta)$$

- Propagate via the joint:

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \theta | \mathbf{y}_{1:t})$$

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \theta | \mathbf{y}_{1:t}) \approx q_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1})$$

$$\propto p_{t-1}(\mathbf{x}_{t-1}, \theta) \ell(\mathbf{x}_t | \mathbf{x}_{t-1}, \theta) g(\mathbf{y}_t | \mathbf{x}_t, \theta)$$

$$\propto \tilde{p}_{t-1}(\mathbf{x}_{t-1}, \theta) \ell(\mathbf{x}_t | \mathbf{x}_{t-1}, \theta) g(\mathbf{y}_t | \mathbf{x}_t, \theta)$$

- Build ϕ_{TT} such that

$$\|\phi_{\text{TT}}(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}) - \sqrt{q_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1})}\|_{L^2} \leq \epsilon_t$$

and then

$$\tilde{\pi}_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}) = \phi_{\text{TT}}(\mathbf{x}_t, \theta, \mathbf{x}_{t-1})^2 + \tau \lambda(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}), \quad \tau \leq \epsilon_t$$

- Marginalization for new filtering density

$$p(\mathbf{x}_t, \theta | \mathbf{y}_{1:t}) \propto \int p_t(\mathbf{x}_t, \mathbf{x}_{t-1}, \theta) d\mathbf{x}_{t-1}$$

$$p(\mathbf{x}_t, \theta | \mathbf{y}_{1:t}) \approx \tilde{p}_t(\mathbf{x}_t, \theta)$$

$$\propto \int \tilde{\pi}_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}) d\mathbf{x}_{t-1}$$

Approximation targets

Computation

Existing approximations

Defensive term

Recursive approximation and marginalization

- Previous filtering density:

$$p_{t-1} := p(\mathbf{x}_{t-1}, \theta | \mathbf{y}_{1:t-1})$$

$$p_{t-1}(\mathbf{x}_{t-1}, \theta) \approx \tilde{p}_{t-1}(\mathbf{x}_{t-1}, \theta)$$

- Propagate via the joint:

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \theta | \mathbf{y}_{1:t})$$

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \theta | \mathbf{y}_{1:t}) \approx q_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1})$$

$$\propto p_{t-1}(\mathbf{x}_{t-1}, \theta) f(\mathbf{x}_t | \mathbf{x}_{t-1}, \theta) g(y_t | \mathbf{x}_t, \theta)$$

$$\propto \tilde{p}_{t-1}(\mathbf{x}_{t-1}, \theta) f(\mathbf{x}_t | \mathbf{x}_{t-1}, \theta) g(y_t | \mathbf{x}_t, \theta)$$

- Build ϕ_{TT} such that

$$\|\phi_{\text{TT}}(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}) - \sqrt{q_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1})}\|_{L^2} \leq \epsilon_t$$

and then

$$\tilde{\pi}_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}) = \phi_{\text{TT}}(\mathbf{x}_t, \theta, \mathbf{x}_{t-1})^2 + \tau \lambda(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}), \quad \tau \leq \epsilon_t$$

- Marginalization for new filtering density

$$p(\mathbf{x}_t, \theta | \mathbf{y}_{1:t}) \propto \int p_t(\mathbf{x}_t, \mathbf{x}_{t-1}, \theta) d\mathbf{x}_{t-1}$$

$$p(\mathbf{x}_t, \theta | \mathbf{y}_{1:t}) \approx \tilde{p}_t(\mathbf{x}_t, \theta)$$

$$\propto \int \tilde{\pi}_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}) d\mathbf{x}_{t-1}$$

Approximation targets

Computation

Existing approximations

Defensive term

Recursive approximation and marginalization

- Previous filtering density:

$$p_{t-1} := p(\mathbf{x}_{t-1}, \theta | \mathbf{y}_{1:t-1})$$

$$p_{t-1}(\mathbf{x}_{t-1}, \theta) \approx \tilde{p}_{t-1}(\mathbf{x}_{t-1}, \theta)$$

- Propagate via the joint:

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \theta | \mathbf{y}_{1:t})$$

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \theta | \mathbf{y}_{1:t}) \approx q_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1})$$

$$\propto p_{t-1}(\mathbf{x}_{t-1}, \theta) f(\mathbf{x}_t | \mathbf{x}_{t-1}, \theta) g(\mathbf{y}_t | \mathbf{x}_t, \theta)$$

$$\propto \tilde{p}_{t-1}(\mathbf{x}_{t-1}, \theta) f(\mathbf{x}_t | \mathbf{x}_{t-1}, \theta) g(\mathbf{y}_t | \mathbf{x}_t, \theta)$$

- Build ϕ_{TT} such that

$$\|\phi_{\text{TT}}(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}) - \sqrt{q_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1})}\|_{L^2} \leq \epsilon_t$$

and then

$$\tilde{\pi}_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}) = \phi_{\text{TT}}(\mathbf{x}_t, \theta, \mathbf{x}_{t-1})^2 + \tau \lambda(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}), \quad \tau \leq \epsilon_t$$

- Marginalization for new filtering density

$$p(\mathbf{x}_t, \theta | \mathbf{y}_{1:t}) \propto \int p_t(\mathbf{x}_t, \mathbf{x}_{t-1}, \theta) d\mathbf{x}_{t-1}$$

$$p(\mathbf{x}_t, \theta | \mathbf{y}_{1:t}) \approx \tilde{p}_t(\mathbf{x}_t, \theta)$$

$$\propto \int \tilde{\pi}_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}) d\mathbf{x}_{t-1}$$

Approximation targets

Computation

Existing approximations

Defensive term

Recursive approximation and marginalization

- Previous filtering density:

$$p_{t-1} := p(\mathbf{x}_{t-1}, \theta | \mathbf{y}_{1:t-1})$$

$$p_{t-1}(\mathbf{x}_{t-1}, \theta) \approx \tilde{p}_{t-1}(\mathbf{x}_{t-1}, \theta)$$

- Propagate via the joint:

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \theta | \mathbf{y}_{1:t})$$

$$p(\mathbf{x}_t, \mathbf{x}_{t-1}, \theta | \mathbf{y}_{1:t}) \approx q_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1})$$

$$\propto p_{t-1}(\mathbf{x}_{t-1}, \theta) f(\mathbf{x}_t | \mathbf{x}_{t-1}, \theta) g(y_t | \mathbf{x}_t, \theta)$$

$$\propto \tilde{p}_{t-1}(\mathbf{x}_{t-1}, \theta) f(\mathbf{x}_t | \mathbf{x}_{t-1}, \theta) g(y_t | \mathbf{x}_t, \theta)$$

- Build ϕ_{TT} such that

$$\|\phi_{\text{TT}}(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}) - \sqrt{q_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1})}\|_{L^2} \leq \epsilon_t$$

and then

$$\tilde{\pi}_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}) = \phi_{\text{TT}}(\mathbf{x}_t, \theta, \mathbf{x}_{t-1})^2 + \tau \lambda(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}), \quad \tau \leq \epsilon_t$$

- Marginalization for new filtering density

$$p(\mathbf{x}_t, \theta | \mathbf{y}_{1:t}) \propto \int p_t(\mathbf{x}_t, \mathbf{x}_{t-1}, \theta) d\mathbf{x}_{t-1}$$

$$p(\mathbf{x}_t, \theta | \mathbf{y}_{1:t}) \approx \tilde{p}_t(\mathbf{x}_t, \theta)$$

$$\propto \int \tilde{\pi}_t(\mathbf{x}_t, \theta, \mathbf{x}_{t-1}) d\mathbf{x}_{t-1}$$

Approximation targets

Computation

Existing approximations

Defensive term

- Cost: $\mathcal{O}((2d_x + d_\theta)r^2)$ densities + $\mathcal{O}(d_x r^3)$ flops (for marginalization)
- Assume either of the following bounds holds

$$C_t^{(g)} = \sup_{x_t \in \mathcal{X}, \theta \in \Theta} g(y_t | x_t, \theta) < \infty, \quad (1)$$

$$C_t^{(f)} = \sup_{x_{t-1} \in \mathcal{X}, \theta \in \Theta} \int f(x_t | x_{t-1}, \theta) g(y_t | x_t, \theta) dx_t < \infty. \quad (2)$$

- At time t , approximation satisfies $\|\sqrt{q_t} - \phi_{\text{TT}}\|_{L^2} \leq \epsilon_t$. See  [Cui & Dolgov 2023] and  [Griebel & Harbrecht 2023].
- Then the Hellinger distance between the joint density p_t and its (normalized) recursive tensor-train approximation $\tilde{p}_t \propto \tilde{\pi}_t$ is bounded by

$$D_H(p_t, \tilde{p}_t) \leq \frac{2}{\sqrt{p(y_{1:t})}} \sum_{k=1}^t C^{t-k} \epsilon_k,$$

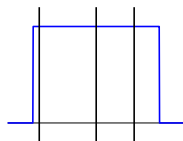
where $p(y_{1:t})$ is the evidence and $C = \sup_t C_t^{(h)} < \infty$ for $h \in \{f, g\}$.

Filtering and smoothing

Knothe–Rosenblatt (KR) rearrangement

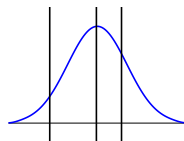
A prob. density has factorization: $\pi(x_1, \dots, x_d) = \pi_1(x_1) \cdots \pi_k(x_k | x_{<k}) \cdots \pi_d(x_d | x_{<d})$

Couple $X \sim \pi$ with $Z \sim \text{uniform}(0, 1]^d$

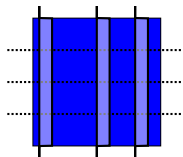


(a) $u_1(z_1)$

$$x_1^j = F_1^{-1}(z_1^j)$$

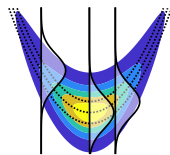


(b) $\pi_{\leq 1}(x_1)$



(c) $u(z_1, z_2) = u_1(z_1)u_2(z_2)$

$$x_2^j | (x_1 = x_1^j) = F_2^{-1}(z_2^j | x_1^j)$$



(d) $\pi(x_1, x_2) = \pi_{\leq 1}(x_1)\pi_2(x_2|x_1)$

- In each iteration, unnormalized (joint) filtering density π_t has an approximation

$$\pi_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) \approx \tilde{\pi}_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) := \phi_{\text{TT}}^2(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) + \tau \lambda(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1})$$

- Obtain marginal densities and conditional CDFs using TT, e.g.,

$$\tilde{\pi}_{t, \leq k}(\mathbf{x}_{t, \leq k}) = \int \phi_{\text{TT}}^2(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) d\mathbf{x}_{t, > k} d\boldsymbol{\theta} d\mathbf{x}_{t-1} + \tau \lambda(\mathbf{x}_{t, \leq k})$$

$$F_{t,k}(x_{t,k} | \mathbf{x}_{t, < k}, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = \int_{-\infty}^{x_{t,k}} \frac{\tilde{\pi}_{t, \leq k}(\mathbf{x}_{t, < k}, x'_{t,k})}{\tilde{\pi}_{t, < k}(\mathbf{x}_{t, < k})} dx'_{t,k}$$

- Integrate $\hat{\pi}_t$ from right to left to obtain **lower triangular KR rearrangement**

$$\mathcal{F}_t^l(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = \begin{bmatrix} \mathcal{F}_{t,t}^l & (\mathbf{x}_t) \\ \mathcal{F}_{t,\theta}^l & (\boldsymbol{\theta} \mid \mathbf{x}_t) \\ \mathcal{F}_{t,t-1}^l & (\mathbf{x}_{t-1} \mid \boldsymbol{\theta}, \mathbf{x}_t) \end{bmatrix}$$

- (Conditional) sampling from $\tilde{\pi}_t$ is easy—a sequence of 1D inverse CDF

- Use the last block of the lower triangular KR rearrangement

$$\mathcal{F}_t^l(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = \begin{bmatrix} \mathcal{F}_{t,t}^l & (\mathbf{x}_t) \\ \mathcal{F}_{t,\theta}^l & (\boldsymbol{\theta} \mid \mathbf{x}_t) \\ \mathcal{F}_{t,t-1}^l & (\mathbf{x}_{t-1} \mid \boldsymbol{\theta}, \mathbf{x}_t) \end{bmatrix}$$

for generating (backward) conditional random variable $\widehat{\mathbf{X}}_{t-1} \mid \widehat{\boldsymbol{\theta}}, \widehat{\mathbf{X}}_t$

- Integrate $\tilde{\pi}_t$ from left to right to obtain upper triangular KR rearrangement

$$\mathcal{F}_t^u(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = \begin{bmatrix} \mathcal{F}_{t,t}^u & (\mathbf{x}_t \mid \boldsymbol{\theta}, \mathbf{x}_{t-1}) \\ \mathcal{F}_{t,\theta}^u & (\boldsymbol{\theta} \mid \mathbf{x}_{t-1}) \\ \mathcal{F}_{t,t-1}^u & (\mathbf{x}_{t-1}) \end{bmatrix}$$

for generating (forward) conditional random variable $\widehat{\mathbf{X}}_t \mid \widehat{\boldsymbol{\theta}}, \widehat{\mathbf{X}}_{t-1}$

Use the **upper-triangular** $\mathcal{F}_{t,t}^u$

$$\mathcal{F}_t^u(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = \begin{bmatrix} \mathcal{F}_{t,t}^u & (\mathbf{x}_t | \boldsymbol{\theta}, \mathbf{x}_{t-1}) \\ \mathcal{F}_{t,\theta}^u & (\boldsymbol{\theta} | \mathbf{x}_{t-1}) \\ \mathcal{F}_{t,t-1}^u & (\mathbf{x}_{t-1}) \end{bmatrix}$$

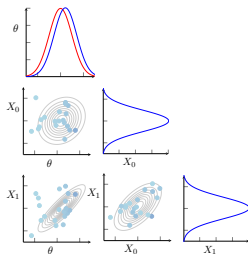
$$\{\hat{\mathbf{x}}_{0:t-1}^{(i)}, \hat{\boldsymbol{\theta}}^{(i)}, w_{t-1}^{(i)}\}$$

$$(\mathcal{F}_{t,t}^u)^{-1}$$

$$\{\hat{\mathbf{x}}_{0:t}^{(i)}, \hat{\boldsymbol{\theta}}^{(i)}, w_{t-1}^{(i)}\}$$

$$w_t^{(i)} = w_{t-1}^{(i)} \times \omega_F^{(i)}$$

$$\{\hat{\mathbf{x}}_{0:t}^{(i)}, \hat{\boldsymbol{\theta}}^{(i)}\}$$



Use the **upper-triangular** $\mathcal{F}_{t,t}^u$

$$\mathcal{F}_t^u(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = \begin{bmatrix} \mathcal{F}_{t,t}^u & (\mathbf{x}_t | \boldsymbol{\theta}, \mathbf{x}_{t-1}) \\ \mathcal{F}_{t,\theta}^u & (\boldsymbol{\theta} | \mathbf{x}_{t-1}) \\ \mathcal{F}_{t,t-1}^u & (\mathbf{x}_{t-1}) \end{bmatrix}$$

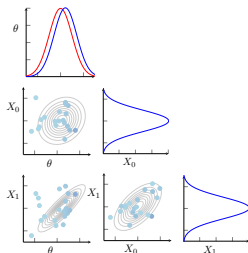
$$\{\hat{\mathbf{x}}_{0:t-1}^{(i)}, \hat{\boldsymbol{\theta}}^{(i)}, w_{t-1}^{(i)}\}$$

$$(\mathcal{F}_{t,t}^u)^{-1}$$

$$\{\hat{\mathbf{x}}_{0:t}^{(i)}, \hat{\boldsymbol{\theta}}^{(i)}, w_{t-1}^{(i)}\}$$

$$w_t^{(i)} = w_{t-1}^{(i)} \times \omega_F^{(i)}$$

$$\{\hat{\mathbf{x}}_{0:t}^{(i)}, \hat{\boldsymbol{\theta}}^{(i)}, w_t^{(i)}\}$$



Use the **upper-triangular** $\mathcal{F}_{t,t}^u$

$$\mathcal{F}_t^u(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = \begin{bmatrix} \mathcal{F}_{t,t}^u & (\mathbf{x}_t | \boldsymbol{\theta}, \mathbf{x}_{t-1}) \\ \mathcal{F}_{t,\theta}^u & (\boldsymbol{\theta} | \mathbf{x}_{t-1}) \\ \mathcal{F}_{t,t-1}^u & (\mathbf{x}_{t-1}) \end{bmatrix}$$

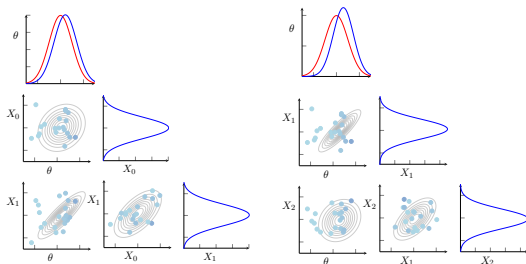
$$\{\hat{\mathbf{x}}_{0:t-1}^{(i)}, \hat{\boldsymbol{\theta}}^{(i)}, w_{t-1}^{(i)}\}$$

$$(\mathcal{F}_{t,t}^u)^{-1}$$

$$\{\hat{\mathbf{x}}_{0:t}^{(i)}, \hat{\boldsymbol{\theta}}^{(i)}, w_{t-1}^{(i)}\}$$

$$w_t^{(i)} = w_{t-1}^{(i)} \times \omega_F^{(i)}$$

$$\{\hat{\mathbf{x}}_{0:t}^{(i)}, \hat{\boldsymbol{\theta}}^{(i)}, w_t^{(i)}\}$$



Accompanying particle filter (forward sampling) $\hat{\mathbf{x}}_t^{(i)} | \hat{\boldsymbol{\theta}}^{(i)}, \hat{\mathbf{x}}_{t-1}^{(i)}$

Use the **upper-triangular** $\mathcal{F}_{t,t}^u$

$$\mathcal{F}_t^u(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = \begin{bmatrix} \mathcal{F}_{t,t}^u & (\mathbf{x}_t | \boldsymbol{\theta}, \mathbf{x}_{t-1}) \\ \mathcal{F}_{t,\theta}^u & (\boldsymbol{\theta} | \mathbf{x}_{t-1}) \\ \mathcal{F}_{t,t-1}^u & (\mathbf{x}_{t-1}) \end{bmatrix}$$

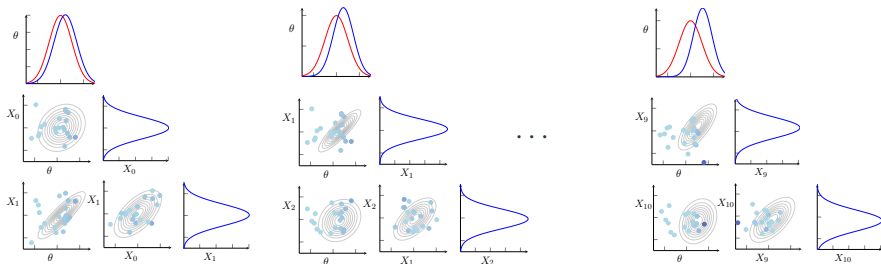
$$\{\hat{\mathbf{x}}_{0:t-1}^{(i)}, \hat{\boldsymbol{\theta}}^{(i)}, w_{t-1}^{(i)}\}$$

$$(\mathcal{F}_{t,t}^u)^{-1}$$

$$\{\hat{\mathbf{x}}_{0:t}^{(i)}, \hat{\boldsymbol{\theta}}^{(i)}, w_{t-1}^{(i)}\}$$

$$w_t^{(i)} = w_{t-1}^{(i)} \times \omega_F^{(i)}$$

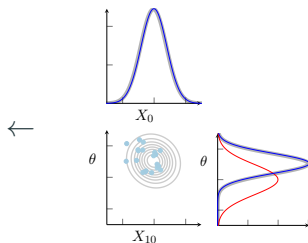
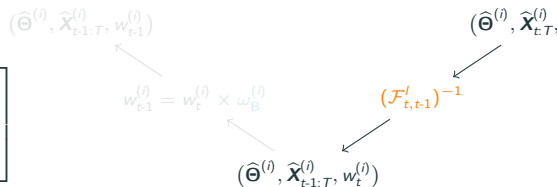
$$\{\hat{\mathbf{x}}_{0:t}^{(i)}, \hat{\boldsymbol{\theta}}^{(i)}, w_t^{(i)}\}$$



Accompanying smoother (backward sampling) $\hat{\mathbf{x}}_{t-1} | \hat{\Theta}^{(i)}, \hat{\mathbf{x}}_t^{(i)}$

Use the lower-triangular $\mathcal{F}'_{t,t-1}$

$$\mathcal{F}'_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = \begin{bmatrix} \mathcal{F}'_{t,t} & (\mathbf{x}_t) \\ \mathcal{F}'_{t,\theta} & (\boldsymbol{\theta} | \mathbf{x}_t) \\ \mathcal{F}'_{t,t-1} & (\mathbf{x}_{t-1} | \boldsymbol{\theta}, \mathbf{x}_t) \end{bmatrix}$$



Accompanying smoother (backward sampling) $\hat{\mathbf{x}}_{t-1} | \hat{\Theta}^{(i)}, \hat{\mathbf{x}}_t^{(i)}$

Use the lower-triangular $\mathcal{F}'_{t,t-1}$

$$\mathcal{F}'_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = \begin{bmatrix} \mathcal{F}'_{t,t} & (\mathbf{x}_t) \\ \mathcal{F}'_{t,\theta} & (\boldsymbol{\theta} | \mathbf{x}_t) \\ \mathcal{F}'_{t,t-1} & (\mathbf{x}_{t-1} | \boldsymbol{\theta}, \mathbf{x}_t) \end{bmatrix}$$

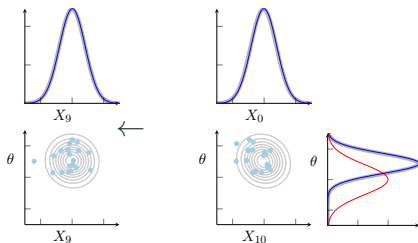
$$(\hat{\Theta}^{(i)}, \hat{\mathbf{x}}_{t-1:T}^{(i)}, w_{t-1}^{(i)})$$

$$w_{t-1}^{(i)} = w_t^{(i)} \times \omega_B^{(i)}$$

$$(\hat{\Theta}^{(i)}, \hat{\mathbf{x}}_{t-1:T}^{(i)}, w_t^{(i)})$$

$$(\mathcal{F}'_{t,t-1})^{-1}$$

$$(\hat{\Theta}^{(i)}, \hat{\mathbf{x}}_{t:T}^{(i)})$$



Accompanying smoother (backward sampling) $\hat{\mathbf{x}}_{t-1} | \hat{\boldsymbol{\theta}}^{(i)}, \hat{\mathbf{x}}_t^{(i)}$

Use the lower-triangular $\mathcal{F}'_{t,t-1}$

$$\mathcal{F}'_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = \begin{bmatrix} \mathcal{F}'_{t,t} & (\mathbf{x}_t) \\ \mathcal{F}'_{t,\theta} & (\boldsymbol{\theta} | \mathbf{x}_t) \\ \mathcal{F}'_{t,t-1} & (\mathbf{x}_{t-1} | \boldsymbol{\theta}, \mathbf{x}_t) \end{bmatrix}$$

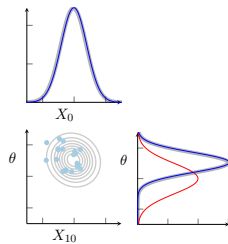
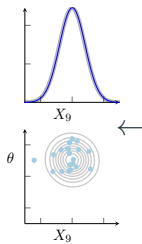
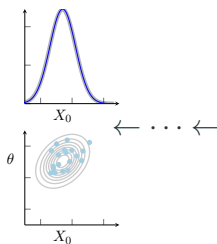
$$(\hat{\boldsymbol{\theta}}^{(i)}, \hat{\mathbf{x}}_{t-1:T}^{(i)}, w_{t-1}^{(i)})$$

$$w_{t-1}^{(i)} = w_t^{(i)} \times \omega_B^{(i)}$$

$$(\hat{\boldsymbol{\theta}}^{(i)}, \hat{\mathbf{x}}_{t-1:T}^{(i)}, w_t^{(i)})$$

$$(\hat{\boldsymbol{\theta}}^{(i)}, \hat{\mathbf{x}}_t^{(i)})$$

$$(\mathcal{F}'_{t,t-1})^{-1}$$



Block sparse structure of the joint KR rearrangement

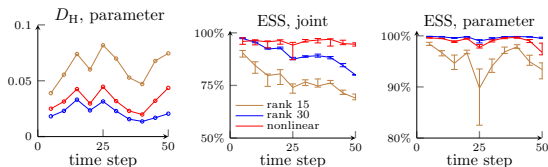
	$\hat{\Theta}$	$\hat{X}_T$	$\hat{X}_{T-1}$	$\hat{X}_{T-2}$	$\hat{X}_{T-3}$	...	$\hat{X}_2$	$\hat{X}_1$	$\hat{X}_0$	
Ξ_θ	Orange	Blue	Dashed	Dashed	Dashed	...	Dashed	Dashed	Dashed	$\Xi_\theta = \mathcal{F}_{T,\theta}^l \quad (\hat{\Theta} \hat{X}_T)$
Ξ_T	Dashed	Blue	Dashed	Dashed	Dashed	...	Dashed	Dashed	Dashed	$\Xi_T = \mathcal{F}_{T,T}^l \quad (\hat{X}_T)$
Ξ_{T-1}	Orange	Blue	Blue	Dashed	Dashed	...	Dashed	Dashed	Dashed	$\Xi_{T-1} = \mathcal{F}_{T,T-1}^l \quad (\hat{X}_{T-1} \hat{X}_T, \hat{\Theta})$
Ξ_{T-2}	Orange	Dashed	Blue	Blue	Dashed	...	Dashed	Dashed	Dashed	$\Xi_{T-2} = \mathcal{F}_{T-1,T-2}^l (\hat{X}_{T-2} \hat{X}_{T-1}, \hat{\Theta})$
Ξ_{T-3}	Orange	Dashed	Dashed	Blue	Blue	...	Dashed	Dashed	Dashed	$\Xi_{T-3} = \mathcal{F}_{T-2,T-3}^l (\hat{X}_{T-3} \hat{X}_{T-2}, \hat{\Theta})$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	$\vdots$	$\vdots$	
Ξ_1	Orange	Dashed	Dashed	Dashed	Dashed	...	Blue	Blue	Dashed	$\Xi_1 = \mathcal{F}_{2,1}^l \quad (\hat{X}_1 \hat{X}_2, \hat{\Theta})$
Ξ_0	Orange	Dashed	Dashed	Dashed	Dashed	...	Dashed	Blue	Blue	$\Xi_0 = \mathcal{F}_{1,0}^l \quad (\hat{X}_0 \hat{X}_1, \hat{\Theta})$

Examples

Example 1: Kalman filter (with analytical solutions)

$$\begin{cases} \mathbf{X}_t - \mu &= \sqrt{1 - a^2} (\mathbf{X}_{t-1} - \mu) + a \varepsilon_t^{(x)} \\ \mathbf{Y}_t &= \mathbf{C} \mathbf{X}_t + d \varepsilon_t^{(y)} \end{cases}$$

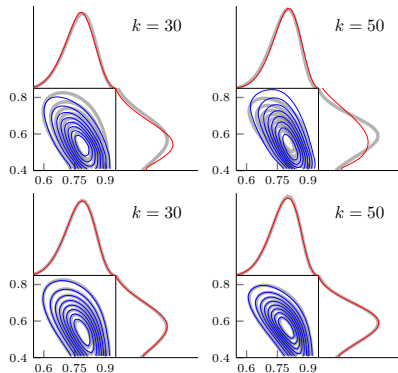
- $\mathbf{X}_t, \mathbf{Y}_t \in \mathbb{R}^3$ with fixed $\mu = 0, \mathbf{C} \in \mathbb{R}^{3 \times 3}$
- Unknown parameter $\theta = (a, d)$
- Synthetic data are generated using $\theta = (0.8, 0.5)$, to $T = 50$



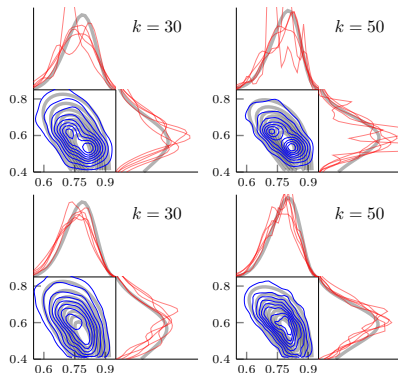
- Non-linear precondition: composition of two rank-15 tensor trains[†]

Example 1: parameter posterior densities

TT (top: rank 15; bottom: rank 30)



SMC² (top: 500; bottom: 5000 particles)



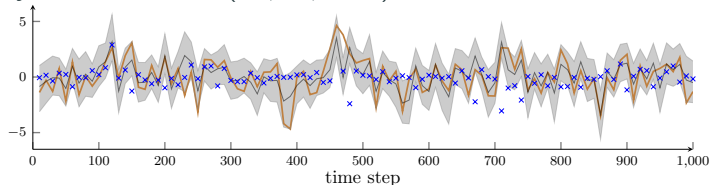
- The cost of running SMC² with 5000 particles is comparable to TT with rank 30.

Example 2: stochastic volatility

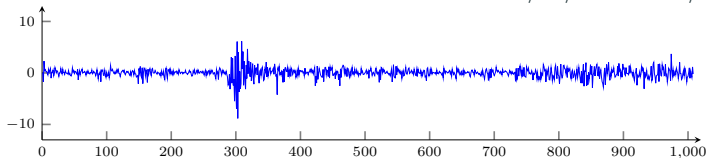
- The unknown parameter $\theta = (\gamma, \beta, \sigma)$

$$\begin{cases} \mathbf{X}_t = \gamma \mathbf{X}_{t-1} + \sigma \varepsilon_t^{(x)} \\ \mathbf{Y}_t = \varepsilon_t^{(y)} \beta \exp(\frac{1}{2} \mathbf{X}_t) \end{cases}$$

- Synthetic data:** $\theta = (0.6, 0.4, \sigma = 1)$ to $T = 1000$

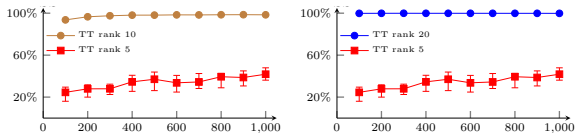


- Real data:** returns of the S&P 500 index from 31/12/2019 to 29/12/2023

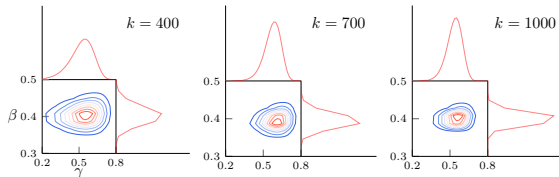


Example 2: synthetic data

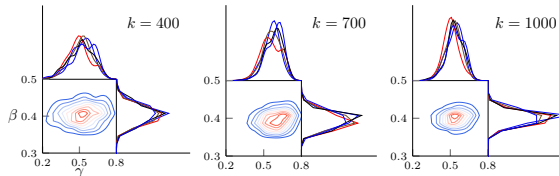
TT ESS



TT density



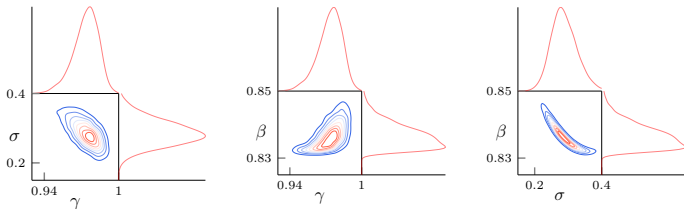
SMC²



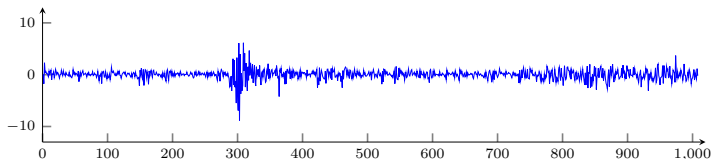
- Computation time of TT is 10% of that of SMC² (1000 parameter particles)

Example 2: S&P 500 returns

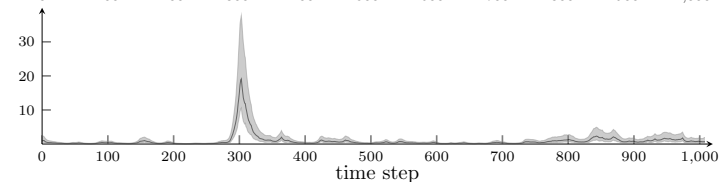
Estimated parameters



S&P 500 returns



Estimated volatilities



- The peak of estimated volatilities matches the volatile returns at $t = 300$

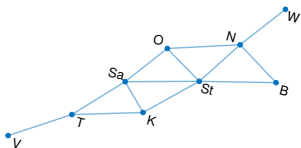
Example 3: susceptible-infectious-removed (SIR) model

$$\begin{cases} dS_j &= (-\kappa_j S_j I_j + \frac{1}{2} \sum_{i \in \mathcal{I}_j} (S_i - S_j)) dt + \sigma d\mathcal{W} \\ dl_j &= (\kappa_j S_j I_j - \nu_j I_j + \frac{1}{2} \sum_{i \in \mathcal{I}_j} (I_i - I_j)) dt + \sigma d\mathcal{W} \\ dR_j &= (\nu_j I_j + \frac{1}{2} \sum_{i \in \mathcal{I}_j} (R_i - R_j)) dt + \sigma d\mathcal{W} \end{cases}, \quad j = 1, \dots, J$$

- $J \in \mathbb{N}$: number of spatially dependent demographic compartments
- $\mathcal{I}_j$: geographical connections
- Collect observed data from $\{I_j\}^J$

$$\mathbf{Y}_{t,j} = \mathbf{I}_{t,j} + \varepsilon_{t,j}^{(y)}, \quad k = 1, 2, \dots, T, \quad j = 1, 2, \dots, J.$$

- Parameters: $\theta = (\kappa_j, \mu_j, \sigma)$ with $\kappa_j = 0.1$, $\nu_j = 18$, $\sigma^2 = 50$ for simulating synthetic data



V - Vorarlberg

T - Tyrol

Sa - Salzburg

K - Carinthia

St - Styria

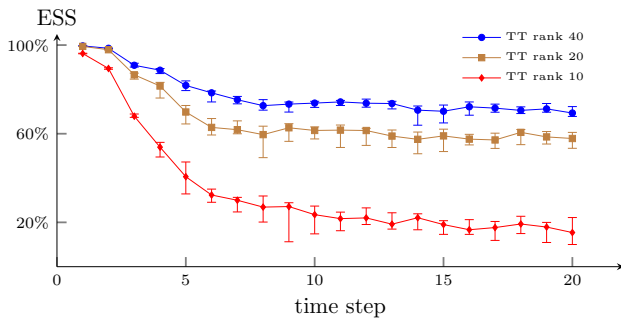
O - Upper Austria

N - Lower Austria

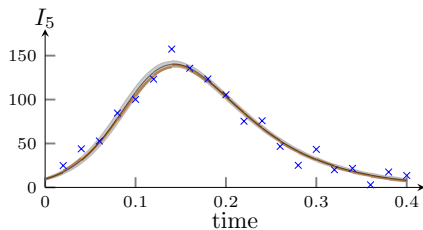
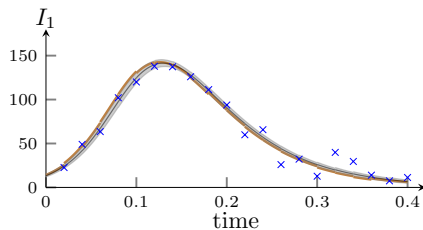
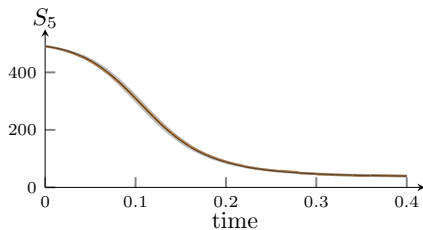
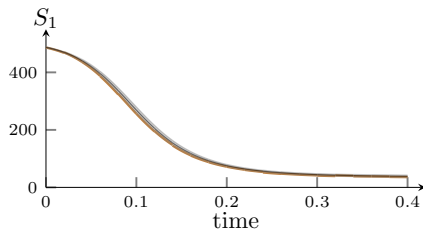
W - Vienna

B - Burgenland

Example 3: effective sample size



Example 3: state trajectories



- The estimated paths of Vorarlberg and Styria

- Zhao, Y., & Cui, T. (2024). Tensor-train methods for sequential state and parameter learning in state-space models. *Journal of Machine Learning Research*, 25 (244), 1–51.
- Matlab code: <https://github.com/DeepTransport/tensor-ssm-paper-demo>
- Dolgov, S., Anaya-Izquierdo, K., Fox, C., & Scheichl, R. (2020). Approximation and sampling of multivariate probability distributions in the tensor train decomposition. *Statistics and Computing*, 30, 603–625.
- Cui, T., & Dolgov, S. (2022). Deep composition of Tensor-Trains using squared inverse Rosenblatt transports. *Foundations of Computational Mathematics*, 22(6), 1863–1922.
- DIRT package (Matlab): <https://github.com/DeepTransport/deep-tensor>
- Spantini, A., Bigoni, D., & Marzouk, Y. (2018). Inference via low-dimensional couplings. *Journal of Machine Learning Research*, 19(66), 1–71.

Preconditioning

Preconditioning

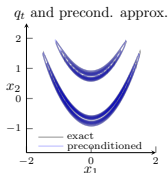
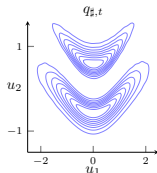
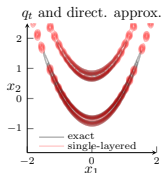
Consider a function $q(x_1, x_2)$ with high tensor ranks and a transport map M . The pushforward function $q_{\#}(u_1, u_2) \propto M_{\#}q(u_1, u_2)$ may be easier to approximate.

Example: linear preconditioning

For $\mathbf{X} \sim N(\mu, LL^T)$ with density $q(\mathbf{x})$, define the linear transport map $M(\mathbf{x}) = L^{-1}(\mathbf{x} - \mu)$, then the pushforward function $q_{\#}(\mathbf{u}) \propto M_{\#}q(\mathbf{u})$ is rank-one.

Example: Non-linear preconditioning via tempering

- Tempered density: $q_t^{\alpha}(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1}) = q_t(\mathbf{x}_t, \boldsymbol{\theta}, \mathbf{x}_{t-1})^{\alpha}$
- Build TT approximation $\phi_{\rho,t} \approx \sqrt{q^{\alpha}}$, and hence the KR map $M_{\#}q_t^{\alpha} \approx \rho$
- The pushforward density is given by $q_{\#,t} = M_{\#}q_t \approx q_t^{1-\alpha}$



More examples

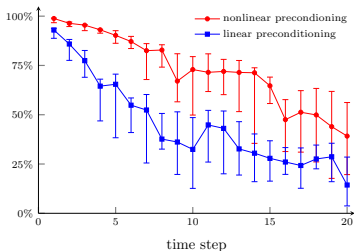
Example 4: predator prey

- State transition: $\mathbf{x}_t = (P_t, Q_t)$ following

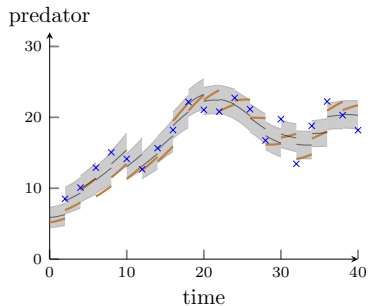
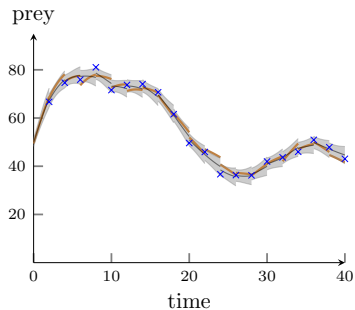
$$\begin{cases} dP &= \left(rP \left(1 - \frac{P}{K} \right) - s \left(\frac{sPQ}{a+P} \right) \right) dt + d\mathcal{W} \\ dQ &= \left(u \left(\frac{PQ}{a+P} - vQ \right) \right) dt + d\mathcal{W} \end{cases}$$

- Observation: $\mathbf{Y}_t = \mathbf{X}_t + \varepsilon_t^{(y)}$
- Parameters: $\theta = (r, K, a, s, u, v)$

TT ESS



Example 4: state trajectories



- Trajectories of the states estimated from tensor-train smoother
- Orange curves: discretized true trajectory