

Reduced basis methods: From low-rank matrices to low-rank tensors

Daniel Kressner

Chair of Numerical Algorithms and HPC

MATH / FSB / EPF Lausanne

daniel.kressner@epfl.ch

<http://anchp.epfl.ch>

Based on joint work with J. Ballani (EPFL, Akselos).

MoRePaS 2018

10.4.2018



ÉCOLE POLYTECHNIQUE
FÉDÉRALE DE LAUSANNE

Parameter-dependent linear systems

Linear system

$$A(\mu)x(\mu) = f(\mu)$$

depending on parameters $\mu \in \mathcal{D} \subset \mathbb{R}^p$.

- ▶ $\mathcal{D}$ compact
- ▶ $A(\cdot)$ invertible on $\mathcal{D}$

Applications:

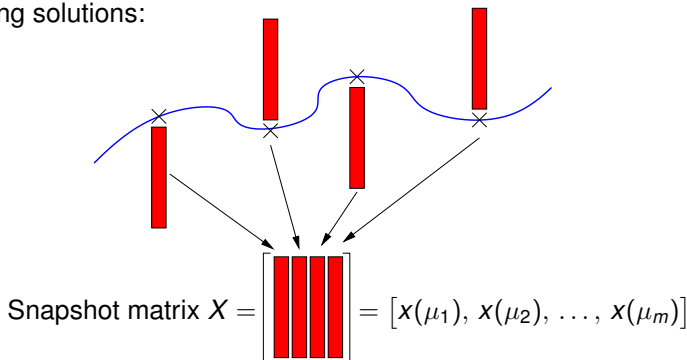
- ▶ PDEs with varying or unknown coefficients
- ▶ PDEs with varying geometry
- ▶ Uncertainty quantification
- ▶ Fuzzy systems

From matrices ...

$$A(\mu)x(\mu) = b(\mu)$$

- ▶ $A(\cdot) : [\mu_{\min}, \mu_{\max}] \rightarrow \mathbb{R}^{n \times n}$ invertible
- ▶ $f(\cdot) : [\mu_{\min}, \mu_{\max}] \rightarrow \mathbb{R}^n$

Sampling solutions:



... to tensors

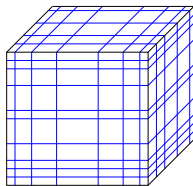
$$A(\mu)x(\mu) = b(\mu), \quad \mu \in [-1, 1]^p$$

with

$$A : [-1, 1]^p \rightarrow \mathbb{R}^{n \times n}$$

$$x, b : [-1, 1]^p \rightarrow \mathbb{R}^n$$

and $A(\mu)$ invertible.

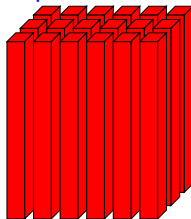


Sample $\mu = (\mu^{(1)}, \dots, \mu^{(p)})^T$ on tensorized grid:

$$\mu_{i_1, \dots, i_p} = (\mu_{i_1}^{(1)}, \dots, \mu_{i_p}^{(p)})^T$$

where $-1 \leq \mu_1^{(j)} \leq \dots \leq \mu_m^{(j)} \leq 1$.

Snapshot tensor $\mathcal{X}$



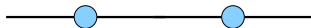
Overview

- ▶ Low-rank tensor formats
- ▶ Low-rank iterative methods
- ▶ Reduced basis methods and low-rank tensors

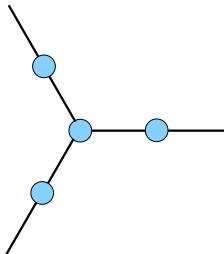
Low-rank tensor formats

Tensor network diagrams

- ▶ Introduced by Roger Penrose.
- ▶ Heavily used in quantum mechanics (spin networks).
- ▶ This is the matrix product $C = AB$:


$$C_{ij} = \sum_{k=1}^r A_{ik} B_{kj}$$

Tensor of order 3 in Tucker format

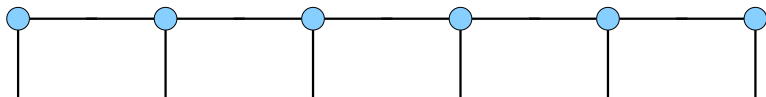


$$\mathcal{X}_{ijk} = \sum_{\ell_1=1}^{r_1} \sum_{\ell_2=1}^{r_2} \sum_{\ell_3=1}^{r_3} \mathcal{C}_{\ell_1 \ell_2 \ell_3} U_{i\ell_1} V_{j\ell_2} W_{k\ell_3}$$

Implicit representation of $\mathcal{X}$ via

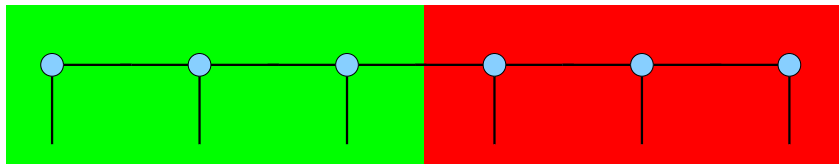
- ▶ $r_1 \times r_2 \times r_3$ core tensor $\mathcal{C}$
- ▶ $n_1 \times r_1$ matrix U spans first mode
- ▶ $n_2 \times r_2$ matrix V spans second mode
- ▶ $n_3 \times r_3$ matrix W spans third mode.

Tensor of order 6 in tensor train (TT) format



- ▶ $\mathcal{X}$ implicitly represented by four $r \times n \times r$ tensors and two $n \times r$ matrices
- ▶ Quantum mechanics: MPS (matrix product states)
- ▶ Matrix-based tensor formats introduced in numerical analysis by Grasedyck, Hackbusch, Kühn, Oseledets, Tyrtishnikov.

Tensor of order 6 in tensor train (TT) format



This partition corresponds to low-rank factorization

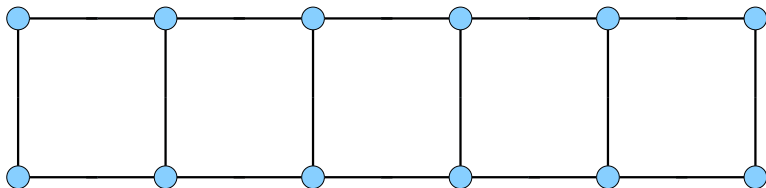
$$X^{(1,2,3)} = UV^T, \quad X^{(1,2,3)} \in \mathbb{R}^{n_1 n_2 n_3 \times n_4 n_5 n_6}, \quad U \in \mathbb{R}^{n_1 n_2 n_3 \times r_3}, \quad V \in \mathbb{R}^{n_4 n_5 n_6 \times r_3}$$

$X^{(1,2,3)}$ is matricization of $\mathcal{X}$:

Merge multi-indices (1, 2, 3) into row indices and
multi-indices (4, 5, 6) into column indices

The ranks of $X^{(1,\dots,\mu)}$ for $\mu = 1, \dots, d-1$ are the TT ranks of $\mathcal{X}$.

Inner product of two tensors in TT format



- ▶ Carrying out contractions requires $O(dnr^4)$ instead of $O(n^d)$ operations for tensors of order d .

Linear operators in TT format

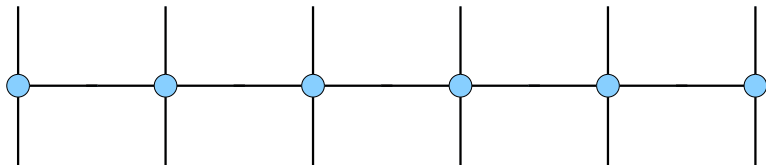
Consider linear operator

$$\mathcal{A} : \mathbb{R}^{n_1 \times \dots \times n_d} \rightarrow \mathbb{R}^{n_1 \times \dots \times n_d}$$

Reinterpret matrix representation of $\mathcal{A}$ as tensor

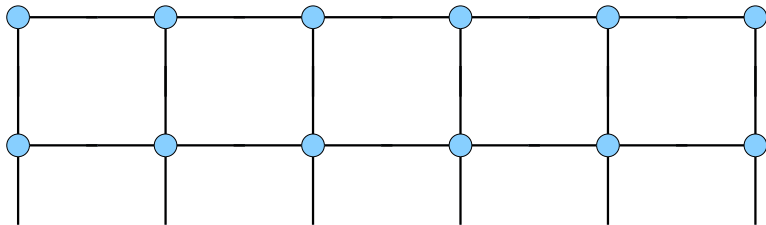
$$\text{vec}(\mathcal{A}) \in \mathbb{R}^{n_1^2 \times \dots \times n_d^2}$$

and apply TT format:



In practice: Perfect shuffle permutation of modes.

Multiplication with linear operators in TT format



- Carrying out contractions requires $O(dnr^4)$ operations for tensors of order d .

Operations with tensors in TT format

Easy:

- ▶ (partial) contractions
- ▶ multiplication with TT operators
- ▶ recompression/truncation

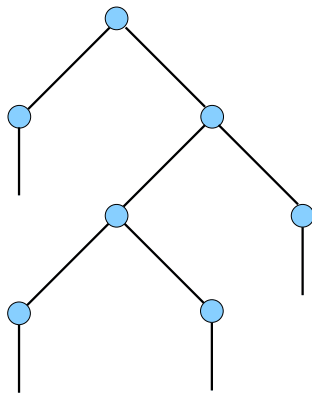
Medium:

- ▶ entrywise products

Hard:

- ▶ almost everything else

Tensor of order 4 in hierarchical Tucker format



- ▶ Proposed in [Grasedyck'2010], [Hackbusch/Kühn'2009]
- ▶ Defined for general binary trees $\rightsquigarrow$ extension of TT format
- ▶ Linear operators in hierarchical Tucker format analogous to TT format

Back to snapshot tensors

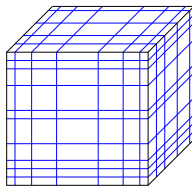
$$A(\mu)x(\mu) = b(\mu), \quad \mu \in [-1, 1]^p$$

with

$$A : [-1, 1]^p \rightarrow \mathbb{R}^{n \times n}$$

$$x, b : [-1, 1]^p \rightarrow \mathbb{R}^n$$

and $A(\mu)$ invertible.

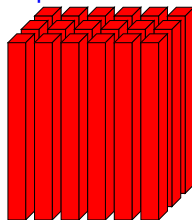


Sample $\mu = (\mu^{(1)}, \dots, \mu^{(p)})^T$ on tensorized grid:

$$\mu_{i_1, \dots, i_p} = (\mu_{i_1}^{(1)}, \dots, \mu_{i_p}^{(p)})^T$$

where $-1 \leq \mu_1^{(j)} \leq \dots \leq \mu_m^{(j)} \leq 1$.

Snapshot tensor $\mathcal{X}$



Low rank approximation of snapshot tensor

Approximation of snapshot tensor $\mathcal{X}$:

$$\mathcal{X} = \text{tensor of rank } k + \text{error}.$$

- ▶ If $A(\cdot), b(\cdot)$ analytic in each parameter $\rightsquigarrow x(\cdot)$ analytic.
- ▶ For any choice of $s \geq 1$:

$$\text{error} \leq C(s) k^{-s},$$

$\rightsquigarrow$ super-polynomial convergence.

(But not exponential, $C(s) \rightarrow \infty$ as $s \rightarrow \infty$.)

- ▶ Follows from approximation results in [Bieri/Andreev/Schwab'09], [Cohen/DeVore/Schwab'10].
- ▶ More can be said for special situations: [Tobler'2012], [Bachmayr/Cohen'2015].

Example

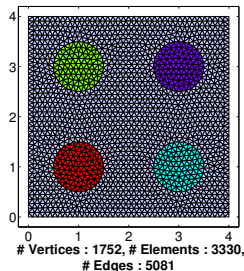
Stationary heat equation with pw constant heat conductivity $\sigma(x, \mu)$:

$$\begin{aligned} -\nabla(\sigma(x, \mu)\nabla u) &= f && \text{in } \Omega = [-1, 1]^2 \\ u &= 0 && \text{on } \partial\Omega, \end{aligned}$$

- ▶ $\sigma(\text{baking tray}) = 1$
- ▶ $\sigma(\text{cookie}_j) = 1 + \mu^{(j)}, \quad \mu^{(j)} = 0, 1, \dots, 100$

FE discretization $\rightsquigarrow$

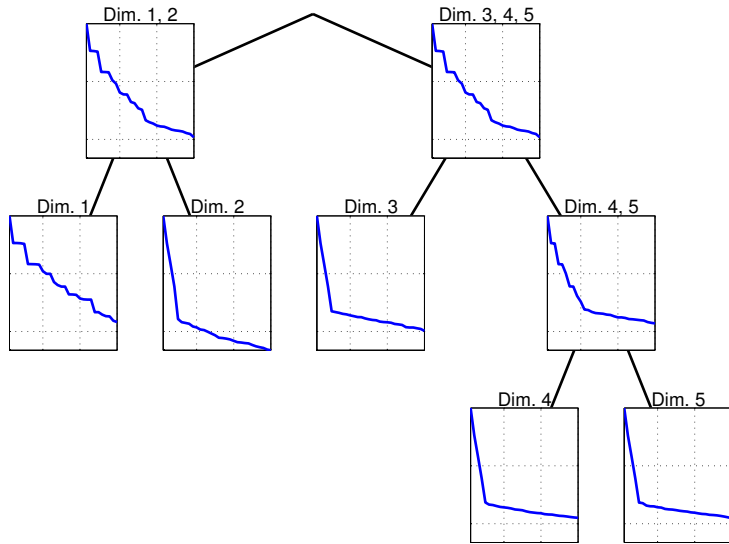
$$(A_0 + \mu^{(1)}A_1 + \mu^{(2)}A_2 + \mu^{(3)}A_3 + \mu^{(4)}A_4)x = b.$$



4 cookies – singular value decay

- ▶ $m = 100$ samples/parameter

Singular value decays in hierarchical Tucker decomposition

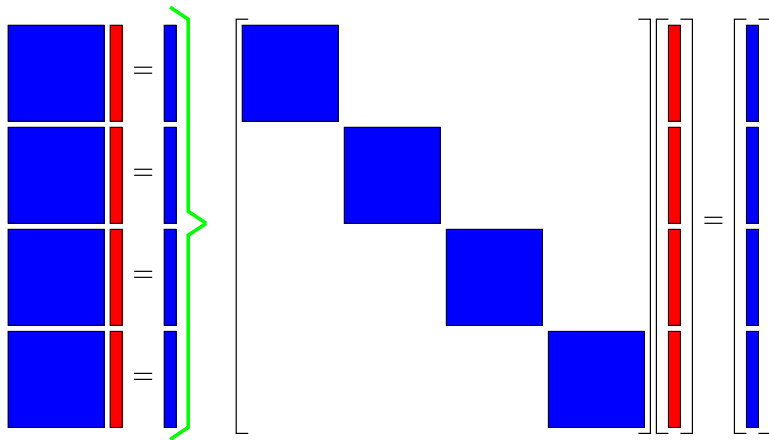


Tensor methods for parametrized linear systems

- ▶ **Combination of standard iterative schemes with low-rank truncation** [Bachmayr/Cohen/Dahmen'2016], [Ballani/Grasedyck'2015], [Benner/Onwunta/Stoll'2015], [Dolgov'2013], [Khoromskij/Oseledets'2011], [Khoromskij/Schwab'2011], [Kressner/Tobler'2011], [Matthies/Zander'2012].
- ▶ **Optimization-based approaches** [Dolgov/Savostyanov'2014], [Holtz/Rohwedder/Schneider'2012].
- ▶ **Greedy low-rank approximation** [Billaud-Friess/Nouy/Zahm'2014], [Nouy'2015]
- ▶ **Sampling techniques** [Ballani/Grasedyck'2015], [Ballani/Kressner'2016], [Dolgov/Khoromskij/Litvinenko/Matthies'2015], [Grasedyck/Kluge/Krämer'2013], [Oseledets/Tyrtysnikov'2009].
- ▶ **Multilevel approximation** [Ballani/Kressner/Peters'2017],
- ▶ Comparison between sparse and low-rank approaches [Bachmayr/Cohen/Dahmen'2016]
- ▶ ...

Low-rank iterative methods

Couple the uncoupled



m linear systems $\Rightarrow$

1 BIG linear system

Couple the uncoupled

$$\begin{array}{l} A(\mu_1)x(\mu_1) = b(\mu_1) \\ \vdots \\ A(\mu_m)x(\mu_m) = b(\mu_m) \end{array} \Rightarrow \underbrace{\begin{bmatrix} A(\mu_1) & & \\ & \ddots & \\ & & A(\mu_m) \end{bmatrix}}_{\mathcal{A}} \underbrace{\begin{bmatrix} x(\mu_1) \\ \vdots \\ x(\mu_m) \end{bmatrix}}_{\mathbf{x}} = \underbrace{\begin{bmatrix} b(\mu_1) \\ \vdots \\ b(\mu_m) \end{bmatrix}}_{\mathbf{b}}$$

► Linear case, $A(\mu) = A_0 + \mu A_1$:

$$\mathcal{A} = I \otimes A_0 + D \otimes A_1, \quad D = \begin{bmatrix} \mu_1 & & \\ & \ddots & \\ & & \mu_m \end{bmatrix}$$

(General: Kronecker structure if $A(\mu) = A_0 + f_1(\mu)A_1 + f_2(\mu)A_2 + \dots$)

Couple the uncoupled

$$\begin{array}{l} A(\mu_1)x(\mu_1) = b(\mu_1) \\ \vdots \\ A(\mu_m)x(\mu_m) = b(\mu_m) \end{array} \Rightarrow \underbrace{\begin{bmatrix} A(\mu_1) & & \\ & \ddots & \\ & & A(\mu_m) \end{bmatrix}}_{\mathcal{A}} \underbrace{\begin{bmatrix} x(\mu_1) \\ \vdots \\ x(\mu_m) \end{bmatrix}}_{\mathbf{x}} = \underbrace{\begin{bmatrix} b(\mu_1) \\ \vdots \\ b(\mu_m) \end{bmatrix}}_{\mathbf{b}}$$

- ▶ Linear case, $A(\mu) = A_0 + \mu A_1$:

$$\mathcal{A} = I \otimes A_0 + D \otimes A_1, \quad D = \begin{bmatrix} \mu_1 & & \\ & \ddots & \\ & & \mu_m \end{bmatrix}$$

- ▶ De-vectorize:

$$\mathcal{A} : X \mapsto A_0 X + A_1 X D, \quad X = [x(\mu_1), \dots, x(\mu_m)]$$

Connection to matrix Sylvester equations exploited in [Simoncini'10].

Applying PCG to BIG linear system

Preconditioner $\mathcal{P}$.

$$x_0 = 0, r_0 = b, z_0 = \mathcal{P}^{-1} r_0, \gamma_0 = \langle r_0, z_0 \rangle$$
$$p_0 = z_0, q_0 = \mathcal{A}p_0, \xi_0 = \langle p_0, q_0 \rangle, k = 0$$

iterate:

$$x_{k+1} = x_k + \frac{\gamma_k}{\xi_k} p_k$$
$$r_{k+1} = b - \mathcal{A}x_k \quad \% \text{ non-standard}$$

$$z_{k+1} = \mathcal{P}^{-1} r_{k+1}$$

$$\gamma_{k+1} = \langle r_{k+1}, z_{k+1} \rangle$$

$$p_{k+1} = z_{k+1} + \frac{\gamma_{k+1}}{\gamma_k} p_k$$

$$q_{k+1} = \mathcal{A}p_{k+1}$$

$$\xi_{k+1} = \langle p_{k+1}, q_{k+1} \rangle$$

$$k = k + 1$$

... de-vectorizing ...

Applying PCG to BIG linear system

Preconditioner $\mathcal{P}$.

$$X_0 = 0, R_0 = B, Z_0 = \mathcal{P}^{-1}(R_0), \gamma_0 = \langle R_0, Z_0 \rangle \\ P_0 = Z_0, Q_0 = \mathcal{A}(P_0), \xi_0 = \langle P_0, Q_0 \rangle, k = 0$$

iterate:

$$\begin{aligned} X_{k+1} &= X_k + \frac{\gamma_k}{\xi_k} P_k \\ R_{k+1} &= B - \mathcal{A}(X_k) \quad \% \text{ non-standard} \\ Z_{k+1} &= \mathcal{P}^{-1}(R_{k+1}) \\ \gamma_{k+1} &= \langle R_{k+1}, Z_{k+1} \rangle \\ P_{k+1} &= Z_{k+1} + \frac{\gamma_{k+1}}{\gamma_k} P_k \\ Q_{k+1} &= \mathcal{A}(P_{k+1}) \\ \xi_{k+1} &= \langle P_{k+1}, Q_{k+1} \rangle \\ k &= k + 1 \end{aligned}$$

$\langle X, Y \rangle = \text{tr}(Y^T X)$ denotes matrix inner product.

Remark: Coincides with a [global Krylov subspace method](#) [Jbilou/Messaoudi/Sadok'99] in the case of *no* parameter dependence.

PCG with low rank truncations

Preconditioner $\mathcal{P}$.

$$X_0 = 0, R_0 = B, Z_0 = \mathcal{P}^{-1}(R_0), \gamma_0 = \langle R_0, Z_0 \rangle \\ P_0 = Z_0, Q_0 = \mathcal{A}(P_0), \xi_0 = \langle P_0, Q_0 \rangle, k = 0$$

iterate:

$$\begin{aligned} X_{k+1} &= X_k + \frac{\gamma_k}{\xi_k} P_k, & X_{k+1} &= \mathcal{T}(X_{k+1}) \\ R_{k+1} &= B - \mathcal{A}(X_k), & R_{k+1} &= \mathcal{T}(R_{k+1}) \\ Z_{k+1} &= \mathcal{P}^{-1}(R_{k+1}) \\ \gamma_{k+1} &= \langle R_{k+1}, Z_{k+1} \rangle \\ P_{k+1} &= Z_{k+1} + \frac{\gamma_{k+1}}{\gamma_k} P_k, & P_{k+1} &= \mathcal{T}(P_{k+1}) \\ Q_{k+1} &= \mathcal{A}(P_{k+1}), & Q_{k+1} &= \mathcal{T}(Q_{k+1}) \\ \xi_{k+1} &= \langle P_{k+1}, Q_{k+1} \rangle \\ k &= k + 1 \end{aligned}$$

$\mathcal{T}$: low rank truncation with relative error $\leq \epsilon$.

Low rank r decreases cost of CG.

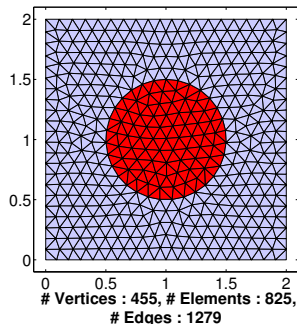
Example: Cost of inner product is $O(mr^2 + nr^2)$ instead of $O(mn)$.

Example

Stationary heat equation with pw constant heat conductivity $\sigma(x, \mu)$:

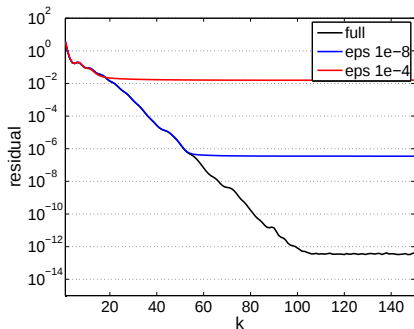
$$\begin{aligned} -\nabla(\sigma(x, \mu)\nabla u) &= f \quad \text{in } \Omega = [-1, 1]^2 \\ u &= 0 \quad \text{on } \partial\Omega, \end{aligned}$$

- ▶ $(A_0 + \mu A_1)x(\mu) = b.$
- ▶ Const. preconditioner
 $M = A_0 + \bar{\mu}A_1$ with optimal $\bar{\mu} > 0.$
(spectral equivalence $\rightsquigarrow$
 h -independent convergence)
- ▶ Samples: $\mu = 0, 1, 2, \dots, 100.$
- ▶ $\mathcal{A} : X \mapsto A_0X + A_1XD.$
- ▶ $\mathcal{P}^{-1} : X \mapsto (A_0 + \bar{\mu}A_1)^{-1}X$

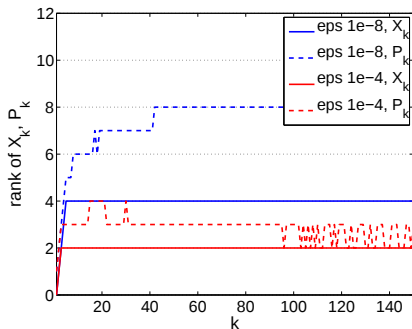


Example: PCG with low rank truncations

Convergence



Rank profile



Extension to several parameters

Stationary heat equation with pw constant heat conductivity $\sigma(x, \mu)$:

$$\begin{aligned} -\nabla(\sigma(x, \mu)\nabla u) &= f \quad \text{in } \Omega = [-1, 1]^2 \\ u &= 0 \quad \text{on } \partial\Omega, \end{aligned}$$

- ▶ $\sigma(\text{baking tray}) = 1$
- ▶ $\sigma(\text{cookie}_j) = 1 + \mu^{(j)}, \quad \mu^{(j)} = 0, 1, \dots, 100$

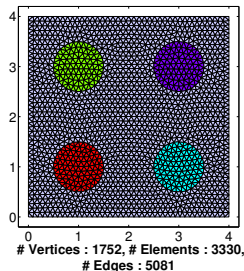
FE discretization $\rightsquigarrow$

$$(A_0 + \mu^{(1)}A_1 + \mu^{(2)}A_2 + \mu^{(3)}A_3 + \mu^{(4)}A_4)x = b.$$

Parameter sampling + assembly $\rightsquigarrow \mathcal{A}x = b \otimes \mathbf{1} \otimes \mathbf{1} \otimes \mathbf{1} \otimes \mathbf{1}$,

$$\mathcal{A} = I \otimes I \otimes I \otimes I \otimes A_0 + D_1 \otimes I \otimes I \otimes I \otimes A_1 + I \otimes D_2 \otimes I \otimes I \otimes A_2 + \dots$$

with $D_j = \text{diag}(\mu_1^{(j)}, \dots, \mu_m^{(j)})$.



PCG with low rank truncations

Idea: Store Hierarchical Tucker format of all vectors.

Preconditioner $\mathcal{P}$.

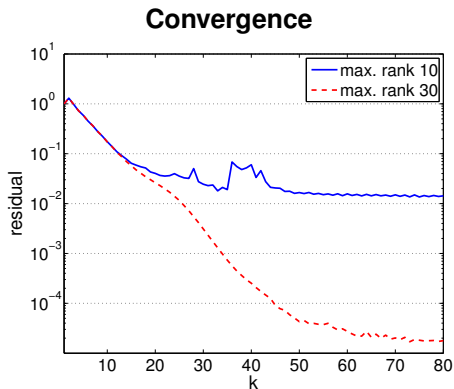
$$\begin{aligned} X_0 &= 0, R_0 = B, Z_0 = \mathcal{P}^{-1}(R_0), \gamma_0 = \langle R_0, Z_0 \rangle \\ P_0 &= Z_0, Q_0 = \mathcal{A}(P_0), \xi_0 = \langle P_0, Q_0 \rangle, k = 0 \end{aligned}$$

iterate:

$$\begin{aligned} X_{k+1} &= X_k + \frac{\gamma_k}{\xi_k} P_k, & X_{k+1} &= \mathcal{T}(X_{k+1}) \\ R_{k+1} &= B - \mathcal{A}(X_k), & R_{k+1} &= \mathcal{T}(R_{k+1}) \\ Z_{k+1} &= \mathcal{P}^{-1}(R_{k+1}) \\ \gamma_{k+1} &= \langle R_{k+1}, Z_{k+1} \rangle \\ P_{k+1} &= Z_{k+1} + \frac{\gamma_{k+1}}{\gamma_k} P_k, & P_{k+1} &= \mathcal{T}(P_{k+1}) \\ Q_{k+1} &= \mathcal{A}(P_{k+1}), & Q_{k+1} &= \mathcal{T}(Q_{k+1}) \\ \xi_{k+1} &= \langle P_{k+1}, Q_{k+1} \rangle \\ k &= k + 1 \end{aligned}$$

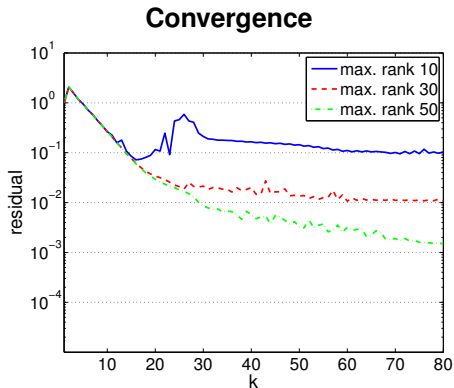
$\mathcal{T}$: low hierarchical Tucker rank truncation with max rank r .

4 cookies: PCG with low rank truncations



- ▶ $m = 100$ samples/parameter
- ▶ Const. preconditioner: $A_0 + \overline{\mu^{(1)}}A_1 + \overline{\mu^{(2)}}A_2 + \overline{\mu^{(3)}}A_3 + \overline{\mu^{(4)}}A_4$.

9 cookies: PCG with low rank truncations



- ▶ $m = 100$ samples/parameter
- ▶ Const. preconditioner: $\mathcal{P} = A_0 + \overline{\mu^{(1)}}A_1 + \dots + \overline{\mu^{(9)}}A_9$.
- ▶ Compt. for $r = 50$ took ≈ 1 h.

Requirements for low-rank iterative methods

- ▶ $\mathcal{A}$ needs to have low-rank operator representation.
- ▶ Availability of good low-rank preconditioner.
- ▶ Difficult to control transient rank growth [Bachmayr/Dahmen'2015/2016].

Reduced basis methods and low-rank tensors

Low-rank tensor MDEIM

Goal. Approximation of $A(\mu)$:

$$A(\mu^i) \approx \sum_{q=1}^Q \Theta_{(i,q)} A_q, \quad \mu^i \in \mathcal{D}_{\mathcal{I}}.$$

- ▶ Parameter-independent matrices $A_q \in \mathbb{R}^{n \times n}$.
- ▶ Tensor product index set $\mathcal{D}_{\mathcal{I}} \subset \mathcal{D}$.
- ▶ Tensor $\Theta \in \mathbb{R}^{n_{\mathcal{I}} \times Q}$ represented in hierarchical tensor format.

Produced in two steps:

Step 1. find an orthogonal matrix $W \in \mathbb{R}^{n^2 \times Q}$ such that $\text{span}\{\text{vec}(A(\mu)) : \mu \in \mathcal{D}_{\mathcal{I}}\} \approx \text{range}(W)$,

Step 2. construct hierarchical Tucker Θ such that $\Theta_{(i,\cdot)} \approx W^T \text{vec}(A(\mu^i))$ for all $\mu^i \in \mathcal{D}_{\mathcal{I}}$.

Low-rank tensor MDEIM

Step 1: MDEIM [Benner/Gugercin/Willcox'2105],
[Negri/Manzoni/Quarteroni'2015]

- 1: $W := []$
- 2: **repeat**
- 3: $\mu^* := \arg \max_{\mu \in \mathcal{D}_{\text{op}}} \|(I - WW^\top) \text{vec}(\mathbf{A}(\mu))\|_2$
- 4: $W := \text{orth}[W, \text{vec}(\mathbf{A}(\mu^*))]$
- 5: **until** $\|(I - WW^\top) \text{vec}(\mathbf{A}(\mu^*))\|_2 \leq \varepsilon_{\text{op}}$

Step 2:

- ▶ Perform cross approximation of tensor

$$\Theta_{(i,\cdot)} = W^\top \text{vec}(\mathbf{A}(\mu^i)), \quad \mu^i \in \mathcal{D}_{\mathcal{I}}$$

in hierarchical Tucker format [Ballani/Grasedyck/Kluge'2013].

Greedy RB Low-rank Tensor Approximation

- 1: $m := 0, V_0 := []$
- 2: **while** $\max_{\mu \in \mathcal{D}_{\text{train}}} \Delta_m(\mu) > \text{tolerance}$ **do**
- 3: $m := m + 1$
- 4: $\mu_m := \arg \max_{\mu \in \mathcal{D}_{\text{train}}} \Delta_{m-1}(\mu)$
- 5: Solve $A(\mu_m)u(\mu_m) = f(\mu_m)$
- 6: $V_m := \text{orth}[V_{m-1}, u(\mu_m)]$
- 7: Solve $\mathcal{A}_m \mathcal{X}_m = \mathcal{B}_m$ in hierarchical Tucker format to update error estimate.
- 8: **end while**

Remarks:

- ▶ $\Delta_m(\mu)$ suitable error bound for approximate solution $x_m(\mu)$ from online phase. Well suited for low-rank tensor formats: residual norm.
- ▶ $\mathcal{A}_m \mathcal{X}_m = \mathcal{B}_m$ is global reduced system in RB method!
- ▶ Can use preconditioned iterative solver or cross approximation.

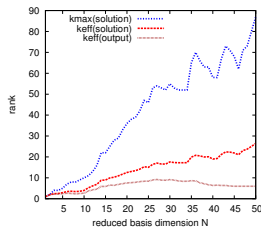
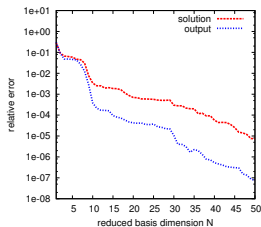
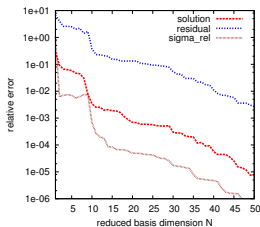
Output: Tensor $V_m \cdot \mathcal{X}_m$ approximating solutions on entire product index set $\mathcal{D}_{\mathcal{I}}$. Online phase replaced by tensor evaluation.

Numerical results: Cookie problem

MDEIM recovers maximal rank but yields lower effective rank:

p	ε_{op}	Q	kmax	keff
4	1e-08	5	5	2.69
9	1e-08	10	10	3.81
16	1e-08	17	17	5.14

Results for $p = 9$ cookies:

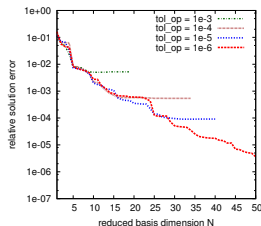
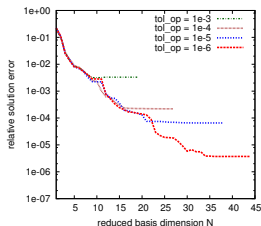
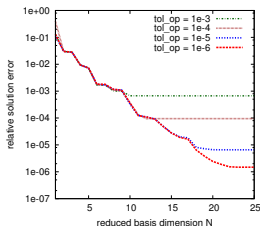


Numerical results: Nonlinear problem

Problem from [Grepl et al.'2007]:

- Square root nonlinearities in parameter dependence.

	$p = 2$			$p = 4$			$p = 8$		
ε_{op}	Q	kmax	keff	Q	kmax	keff	Q	kmax	keff
1e-03	10	10	5.78	15	15	7.44	21	21	8.76
1e-04	14	14	7.36	23	23	11.33	35	35	14.45
1e-05	20	20	9.83	34	34	15.76	49	49	19.84
1e-06	25	25	10.88	44	44	19.81	66	66	26.50



Conclusions

Conclusions and references

- ▶ Low-rank tensor techniques allow to access solutions without need for solving reduced problems $\rightsquigarrow$ computing statistics of solutions reduces to tensor operations.
- ▶ Use of MDEIM+adaptive cross approximation lowers requirements for use of low-rank tensor techniques.
- ▶ Talk based on [Ballani/Kressner. Reduced basis methods: from low-rank matrices to low-rank tensors. SISC'2016].
- ▶ Multilevel ideas from [Teckentrup et al.'2015] for stochastic collocation can be used to reduce ranks; see [Ballani/Kressner/Peters'2017].