

Tensor Decompositions and Structured Randomization

Lecture 2 of 2: what replaces the SVD when there are more than two modes

Hussam Al Daas

Computational Mathematics Theme, STFC-RAL

What to do with multi-way data

Exploratory analysis of *matrix* data starts with the SVD, or PCA.

But many datasets have more than two modes:

- ▶ if an image is 2D, a video is 3D
- ▶ a single social network is 2D; several modalities makes it 3D or higher
- ▶ a parametrized PDE has one mode per parameter

What to do with multi-way data

Exploratory analysis of *matrix* data starts with the SVD, or PCA.

But many datasets have more than two modes:

- ▶ if an image is 2D, a video is 3D
- ▶ a single social network is 2D; several modalities makes it 3D or higher
- ▶ a parametrized PDE has one mode per parameter

Three options:

- ▶ **extract or average** into a matrix — **losing information**
- ▶ **reshape/unfold** into a matrix — **losing relationships between modes**
- ▶ **keep the multi-way structure** and use tensor methods

Warmup: the matrix SVD, and the four things we want from it

The **SVD** of an $M \times N$ matrix A :

$$A = U\Sigma V^T = \sum_{i=1}^N \sigma_i u_i v_i^T \quad \text{and the truncated SVD, a rank-}R \text{ approximation,}$$
$$A \approx A_R = U_R \Sigma_R V_R^T = \sum_{i=1}^R \sigma_i u_i v_i^T$$

We are going to meet three tensor generalizations. Rather than list them, let us score each one on what the SVD gives us for free:

unique / interpretable factors	✓
solves the approximation problem	✓
compresses (storage linear in dimensions)	✓
stable, predictable algorithms	✓

No tensor format scores four out of four. Which boxes you need decides which format you use — that is the real content of the next forty minutes.

Notation

Basic tensor notation

Vector
 $N = 1$



$\mathbf{x}$

Matrix
 $N = 2$



$\mathbf{X}$

3rd-Order Tensor
 $N = 3$



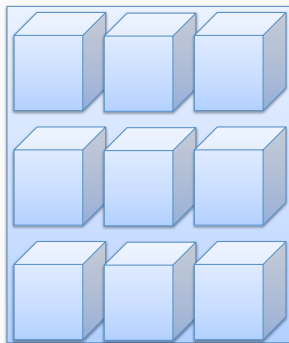
$\mathcal{X}$

4th-Order Tensor
 $N = 4$



$\mathcal{X}$

5th-Order Tensor
 $N = 5$

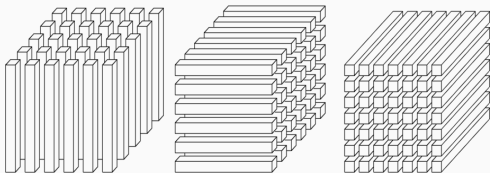


$\mathcal{X}$

An N^{th} -order tensor has N modes

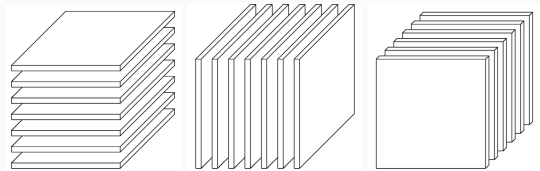
Notation convention: vector $\mathbf{v}$, matrix $\mathbf{M}$, tensor $\mathcal{T}$

Fibers and slices



Mode-1 Fibers Mode-2 Fibers Mode-3 Fibers

fibers: fix all indices but one



(a) Horizontal slices: $\mathbf{X}_{i,:}$

(b) Lateral slices: $\mathbf{X}_{:,j}$

(c) Frontal slices: $\mathbf{X}_{:,:k}$ (or $\mathbf{X}_k$)

slices: fix one index

A tensor can be partitioned into the fibers, or the slices, of each mode.

Modal unfoldings / matricizations

$$\mathbf{x} = \begin{array}{|c|c|c|} \hline & 5 & 7 \\ \hline 1 & 3 & \\ \hline 2 & 6 & 8 \\ \hline & 4 & \\ \hline \end{array}$$
$$\mathbf{X}_{(1)} = \begin{bmatrix} 1 & 3 & 5 & 7 \\ 2 & 4 & 6 & 8 \end{bmatrix}$$
$$\mathbf{X}_{(2)} = \begin{bmatrix} 1 & 2 & 5 & 6 \\ 3 & 4 & 7 & 8 \end{bmatrix}$$
$$\mathbf{X}_{(3)} = \begin{bmatrix} 1 & 2 & 3 & 4 \\ 5 & 6 & 7 & 8 \end{bmatrix}$$

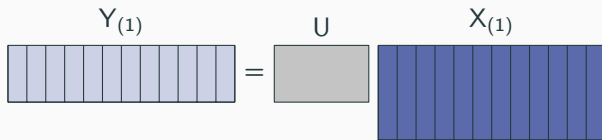
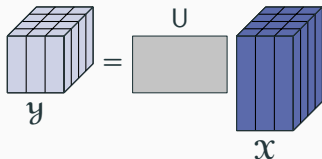
A tensor can be reshaped into matrices, called *unfoldings* or *matricizations*

For *modal unfoldings*, fibers form columns and slices form rows

This is the bridge to everything you already know — but note the shape: $\mathbf{X}_{(n)}$ is $I_n \times (I_1 \cdots I_{n-1} I_{n+1} \cdots I_N)$, catastrophically wide.

Tensor times matrix (TTM)

Notation: $\mathcal{Y} = \mathcal{X} \times_1 U$



- ▶ TTM is mode specific
- ▶ the input matrix multiplies all fibers in that mode
- ▶ unfolded, it is **just matrix multiplication**
- ▶ multiple TTMs compose, in any order

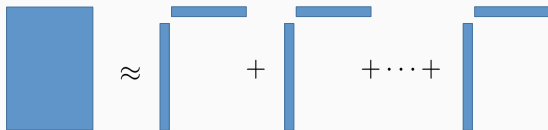
$$\mathcal{T} = \mathcal{G} \times_1 U \times_2 V \times_3 W$$

This one operation is most of what we do all lecture. Remember it.

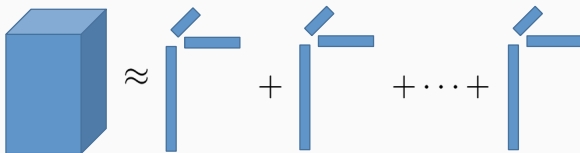
CP

CP: think of the SVD as a sum of outer products

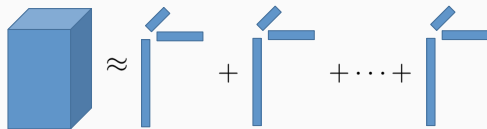
Matrix:



Tensor:



Canonical Polyadic, a.k.a. CANDECOMP/PARAFAC.



$$\mathcal{X} \approx \mathbf{a}_1 \circ \mathbf{b}_1 \circ \mathbf{c}_1 + \cdots + \mathbf{a}_r \circ \mathbf{b}_r \circ \mathbf{c}_r,$$

$$\mathcal{X} \in \mathbb{R}^{m \times n \times p}$$

$$\mathcal{X} \approx [\mathbf{A}, \mathbf{B}, \mathbf{C}],$$

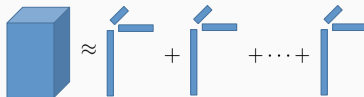
$$\mathbf{A} \in \mathbb{R}^{m \times r}, \mathbf{B} \in \mathbb{R}^{n \times r}, \mathbf{C} \in \mathbb{R}^{p \times r}$$

are *factor matrices*

$$x_{ijk} \approx \sum_{\ell=1}^r a_{i\ell} b_{j\ell} c_{k\ell},$$

$$1 \leq i \leq m, 1 \leq j \leq n, 1 \leq k \leq p$$

The CP optimization problem



For fixed rank r , we want to solve

$$\min_{A,B,C} \left\| \mathcal{X} - \sum_{\ell=1}^r \mathbf{a}_{\ell} \circ \mathbf{b}_{\ell} \circ \mathbf{c}_{\ell} \right\|$$

which is a nonlinear, **nonconvex** optimization problem

- ▶ in the matrix case, the SVD hands us the optimal solution
- ▶ in the tensor case, **there is no such thing** — no HOSVD-style shortcut, and the best rank- r approximation may not even exist

Alternating least squares — and the first structured matrix

Fix all but one factor matrix and the problem becomes *linear* least squares:

$$\min_{\mathbf{B}} \left\| \mathcal{X} - \sum_{\ell=1}^r \hat{\mathbf{a}}_{\ell} \circ \mathbf{b}_{\ell} \circ \hat{\mathbf{c}}_{\ell} \right\|$$

Alternating least squares — and the first structured matrix

Fix all but one factor matrix and the problem becomes *linear* least squares:

$$\min_{\mathbf{B}} \left\| \mathcal{X} - \sum_{\ell=1}^r \hat{\mathbf{a}}_{\ell} \circ \mathbf{b}_{\ell} \circ \hat{\mathbf{c}}_{\ell} \right\|$$

To see why, *matricize*:

$$\min_{\mathbf{B}} \left\| \underbrace{\mathbf{X}_{(2)}}_{n \times mp} - \underbrace{\mathbf{B}}_{n \times r} \underbrace{(\hat{\mathbf{C}} \odot \hat{\mathbf{A}})^{\mathsf{T}}}_{r \times mp} \right\|_F$$

where $\odot$ is the **Khatri-Rao product**, a column-wise Kronecker product:

$$\hat{\mathbf{C}} \odot \hat{\mathbf{A}} = \left[\hat{\mathbf{c}}_1 \otimes \hat{\mathbf{a}}_1 \cdots \hat{\mathbf{c}}_r \otimes \hat{\mathbf{a}}_r \right]$$

Note what just happened. The coefficient matrix is $mp \times r$ — enormous — but it is stored as two small factors. *We will never form it.*

The least squares subproblem, and where the time goes

The subproblem

$$\min_{\mathbf{B}} \left\| \underbrace{\mathbf{X}_{(2)}}_{n \times mp} - \underbrace{\mathbf{B}}_{n \times r} \underbrace{(\hat{\mathbf{C}} \odot \hat{\mathbf{A}})^T}_{r \times mp} \right\|_F$$

has normal equations

$$\underbrace{(\hat{\mathbf{C}} \odot \hat{\mathbf{A}})^T (\hat{\mathbf{C}} \odot \hat{\mathbf{A}})}_{r \times r} \cdot \underbrace{\mathbf{B}^T}_{r \times n} = \underbrace{(\mathbf{X}_{(2)})}_{n \times mp} \underbrace{(\hat{\mathbf{C}} \odot \hat{\mathbf{A}})^T}_{mp \times r}$$

- ▶ $(\hat{\mathbf{C}} \odot \hat{\mathbf{A}})^T (\hat{\mathbf{C}} \odot \hat{\mathbf{A}})$ is **cheap**: it equals $\hat{\mathbf{C}}^T \hat{\mathbf{C}} * \hat{\mathbf{A}}^T \hat{\mathbf{A}}$, an $r \times r$ elementwise product — structure paying off immediately
- ▶ $\mathbf{X}_{(2)}(\hat{\mathbf{C}} \odot \hat{\mathbf{A}})$ is the **Matricized-Tensor Times Khatri-Rao Product (MTTKRP)**, and it touches every entry of $\mathcal{X}$

MTTKRP is bottleneck number one. $O(I^3 r)$ per mode per iteration.

Repeat

1. Solve $(B^T B * C^T C) A^T = (X_{(1)}(C \odot B))^T$ for A
2. Solve $(A^T A * C^T C) B^T = (X_{(2)}(C \odot A))^T$ for B
3. Solve $(A^T A * B^T B) C^T = (X_{(3)}(B \odot A))^T$ for C

Repeat

1. Solve $(B^T B * C^T C) A^T = (X_{(1)}(C \odot B))^T$ for A
2. Solve $(A^T A * C^T C) B^T = (X_{(2)}(C \odot A))^T$ for B
3. Solve $(A^T A * B^T B) C^T = (X_{(3)}(B \odot A))^T$ for C

Other details

- ▶ initialization can be random, or more intelligent
- ▶ handling the scaling ambiguity within rank-one components
- ▶ overlapping computations across the three steps

Repeat

1. Solve $(B^T B * C^T C) A^T = (X_{(1)}(C \odot B))^T$ for A
2. Solve $(A^T A * C^T C) B^T = (X_{(2)}(C \odot A))^T$ for B
3. Solve $(A^T A * B^T B) C^T = (X_{(3)}(B \odot A))^T$ for C

Other details

- ▶ initialization can be random, or more intelligent
- ▶ handling the scaling ambiguity within rank-one components
- ▶ overlapping computations across the three steps

Simple loop; the entire cost is three MTTKRPCs per iteration.

Addition in CP format

Let $\mathcal{X} = \llbracket A, B, C \rrbracket$ have rank r and $\mathcal{Y} = \llbracket \tilde{A}, \tilde{B}, \tilde{C} \rrbracket$ have rank $\tilde{r}$. The sum of two sums of rank-one terms is again a sum of rank-one terms, so we simply **concatenate the factor matrices column-wise**:

$$\mathcal{X} + \mathcal{Y} = \llbracket \begin{bmatrix} A & \tilde{A} \end{bmatrix}, \begin{bmatrix} B & \tilde{B} \end{bmatrix}, \begin{bmatrix} C & \tilde{C} \end{bmatrix} \rrbracket, \quad \text{factors now } I \times (r + \tilde{r}), \text{ etc.}$$

Addition in CP format

Let $\mathcal{X} = \llbracket A, B, C \rrbracket$ have rank r and $\mathcal{Y} = \llbracket \tilde{A}, \tilde{B}, \tilde{C} \rrbracket$ have rank $\tilde{r}$. The sum of two sums of rank-one terms is again a sum of rank-one terms, so we simply **concatenate the factor matrices column-wise**:

$$\mathcal{X} + \mathcal{Y} = \llbracket \begin{bmatrix} A & \tilde{A} \end{bmatrix}, \begin{bmatrix} B & \tilde{B} \end{bmatrix}, \begin{bmatrix} C & \tilde{C} \end{bmatrix} \rrbracket, \quad \text{factors now } I \times (r + \tilde{r}), \text{ etc.}$$

- ▶ **zero flops** — the operation is a copy; storage is the sum of the two storages
- ▶ scalar multiples are just as cheap: $\alpha\mathcal{X} = \llbracket \alpha A, B, C \rrbracket$, absorb α into any one factor
- ▶ the rank of the result is *at most* $r + \tilde{r}$, and generically equal to it

Addition in CP format

Let $\mathcal{X} = \llbracket A, B, C \rrbracket$ have rank r and $\mathcal{Y} = \llbracket \tilde{A}, \tilde{B}, \tilde{C} \rrbracket$ have rank $\tilde{r}$. The sum of two sums of rank-one terms is again a sum of rank-one terms, so we simply **concatenate the factor matrices column-wise**:

$$\mathcal{X} + \mathcal{Y} = \llbracket \begin{bmatrix} A & \tilde{A} \end{bmatrix}, \begin{bmatrix} B & \tilde{B} \end{bmatrix}, \begin{bmatrix} C & \tilde{C} \end{bmatrix} \rrbracket, \quad \text{factors now } I \times (r + \tilde{r}), \text{ etc.}$$

- ▶ **zero flops** — the operation is a copy; storage is the sum of the two storages
- ▶ scalar multiples are just as cheap: $\alpha\mathcal{X} = \llbracket \alpha A, B, C \rrbracket$, absorb α into any one factor
- ▶ the rank of the result is *at most* $r + \tilde{r}$, and generically equal to it

The catch: k additions give rank $O(kr)$, and CP has *no* cheap rounding — truncating back to rank r means solving another nonconvex problem with CP-ALS, at $O(I^3 r)$ per mode per iteration, with no guarantee of a good answer.

Inner products in CP format

Everything follows from one fact about rank-one tensors:

$$\langle a \circ b \circ c, \tilde{a} \circ \tilde{b} \circ \tilde{c} \rangle = (a^T \tilde{a})(b^T \tilde{b})(c^T \tilde{c})$$

Inner products in CP format

Everything follows from one fact about rank-one tensors:

$$\langle \mathbf{a} \circ \mathbf{b} \circ \mathbf{c}, \tilde{\mathbf{a}} \circ \tilde{\mathbf{b}} \circ \tilde{\mathbf{c}} \rangle = (\mathbf{a}^T \tilde{\mathbf{a}})(\mathbf{b}^T \tilde{\mathbf{b}})(\mathbf{c}^T \tilde{\mathbf{c}})$$

Expanding bilinearly over all $r\tilde{r}$ pairs of terms,

$$\langle \mathcal{X}, \mathcal{Y} \rangle = \sum_{\ell=1}^r \sum_{\ell'=1}^{\tilde{r}} (\mathbf{a}_{\ell}^T \tilde{\mathbf{a}}_{\ell'}) (\mathbf{b}_{\ell}^T \tilde{\mathbf{b}}_{\ell'}) (\mathbf{c}_{\ell}^T \tilde{\mathbf{c}}_{\ell'}) = \mathbf{1}^T \left[(\mathbf{A}^T \tilde{\mathbf{A}}) * (\mathbf{B}^T \tilde{\mathbf{B}}) * (\mathbf{C}^T \tilde{\mathbf{C}}) \right] \mathbf{1}$$

- ▶ $O((I + J + K)r\tilde{r})$: three thin Gram products, then an $r \times \tilde{r}$ elementwise product — **the tensors are never formed**
- ▶ the same $*$ identity behind the ALS normal equations: $(\mathbf{C} \odot \mathbf{B})^T (\tilde{\mathbf{C}} \odot \tilde{\mathbf{B}}) = \mathbf{C}^T \tilde{\mathbf{C}} * \mathbf{B}^T \tilde{\mathbf{B}}$

Inner products in CP format

Everything follows from one fact about rank-one tensors:

$$\langle \mathbf{a} \circ \mathbf{b} \circ \mathbf{c}, \tilde{\mathbf{a}} \circ \tilde{\mathbf{b}} \circ \tilde{\mathbf{c}} \rangle = (\mathbf{a}^T \tilde{\mathbf{a}})(\mathbf{b}^T \tilde{\mathbf{b}})(\mathbf{c}^T \tilde{\mathbf{c}})$$

Expanding bilinearly over all $r\tilde{r}$ pairs of terms,

$$\langle \mathbf{x}, \mathbf{y} \rangle = \sum_{\ell=1}^r \sum_{\ell'=1}^{\tilde{r}} (\mathbf{a}_{\ell}^T \tilde{\mathbf{a}}_{\ell'}) (\mathbf{b}_{\ell}^T \tilde{\mathbf{b}}_{\ell'}) (\mathbf{c}_{\ell}^T \tilde{\mathbf{c}}_{\ell'}) = \mathbf{1}^T \left[(\mathbf{A}^T \tilde{\mathbf{A}}) * (\mathbf{B}^T \tilde{\mathbf{B}}) * (\mathbf{C}^T \tilde{\mathbf{C}}) \right] \mathbf{1}$$

- ▶ $O((I + J + K)r\tilde{r})$: three thin Gram products, then an $r \times \tilde{r}$ elementwise product — **the tensors are never formed**
- ▶ the same $*$ identity behind the ALS normal equations: $(\mathbf{C} \odot \mathbf{B})^T (\tilde{\mathbf{C}} \odot \tilde{\mathbf{B}}) = \mathbf{C}^T \tilde{\mathbf{C}} * \mathbf{B}^T \tilde{\mathbf{B}}$

Two consequences, letting ALS monitor its *fit* with no residual tensor:

$$\|\mathbf{x}\|^2 = \mathbf{1}^T (\mathbf{A}^T \mathbf{A} * \mathbf{B}^T \mathbf{B} * \mathbf{C}^T \mathbf{C}) \mathbf{1}, \quad \|\mathbf{x} - \mathbf{y}\|^2 = \|\mathbf{x}\|^2 - 2\langle \mathbf{x}, \mathbf{y} \rangle + \|\mathbf{y}\|^2$$

$$\mathcal{X} \approx \sum_{i=1}^R \mathbf{u}_i \circ \mathbf{v}_i \circ \mathbf{w}_i$$

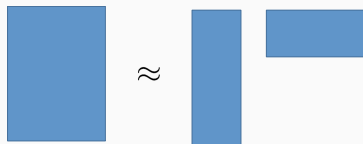
unique / interpretable factors	✓
solves the approximation problem	✗
compresses (storage linear in dimensions)	✓
stable, predictable algorithms	✗

- ▶ CP is almost always unique — so you can interpret the components. This is the property the SVD does *not* give you for tensors of order ≥ 3 , and it is CP's whole reason to exist
- ▶ but there is no exact solve: you cannot reliably hit a prescribed ϵ
- ▶ storage is linear in the dimensions: $(I + J + K)R$
- ▶ algorithms are iterative and not predictable — run from several starting points and hope

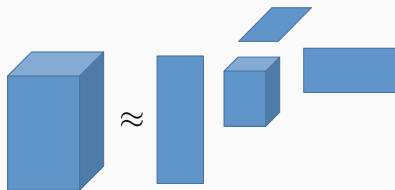
Tucker

Tucker: think of the SVD as identifying subspaces

Matrix:

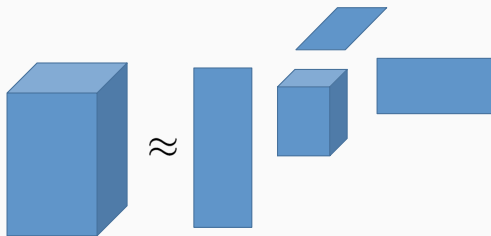


Tensor:



One subspace per mode, plus a small core that couples them.

Tucker notation



$$\mathcal{X} \approx \mathcal{G} \times_1 \mathbf{U} \times_2 \mathbf{V} \times_3 \mathbf{W}$$

$$\mathcal{X} \approx [\![\mathcal{G}; \mathbf{U}, \mathbf{V}, \mathbf{W}]\!],$$

$$\mathcal{X} \in \mathbb{R}^{I \times J \times K}, \mathcal{G} \in \mathbb{R}^{P \times Q \times R}$$

is the core tensor

$$\mathbf{U} \in \mathbb{R}^{I \times P}, \mathbf{V} \in \mathbb{R}^{J \times Q}, \mathbf{W} \in \mathbb{R}^{K \times R}$$

are factor matrices

$$x_{ijk} \approx \sum_{p=1}^P \sum_{q=1}^Q \sum_{r=1}^R g_{pqr} u_{ip} v_{jq} w_{kr}$$

Note: this is exactly a chain of TTMs. Everything we do to it will be TTMs.

The Tucker optimization problem, and why it is easier

For fixed ranks P, Q, R :

$$\min_{\mathcal{G}, \mathbf{U}, \mathbf{V}, \mathbf{W}} \|\mathcal{X} - \mathcal{G} \times_1 \mathbf{U} \times_2 \mathbf{V} \times_3 \mathbf{W}\|$$

— again nonlinear and nonconvex. **But** matricize it:

$$\min_{\mathcal{G}, \mathbf{U}, \mathbf{V}, \mathbf{W}} \left\| \underbrace{\mathbf{X}_{(1)}}_{I \times JK} - \underbrace{\mathbf{U}}_{I \times P} \underbrace{\mathbf{G}_{(1)}(\mathbf{W} \otimes \mathbf{V})^T}_{P \times JK} \right\|_F$$

The Tucker optimization problem, and why it is easier

For fixed ranks P, Q, R :

$$\min_{\mathcal{G}, \mathbf{U}, \mathbf{V}, \mathbf{W}} \|\mathcal{X} - \mathcal{G} \times_1 \mathbf{U} \times_2 \mathbf{V} \times_3 \mathbf{W}\|$$

— again nonlinear and nonconvex. **But** matricize it:

$$\min_{\mathcal{G}, \mathbf{U}, \mathbf{V}, \mathbf{W}} \left\| \underbrace{\mathbf{X}_{(1)}}_{I \times JK} - \underbrace{\mathbf{U}}_{I \times P} \underbrace{\mathbf{G}_{(1)}(\mathbf{W} \otimes \mathbf{V})^T}_{P \times JK} \right\|_F$$

If we ignore the structure imposed on the $P \times JK$ factor and simply ask for the best rank- P approximation of $\mathbf{X}_{(1)}$, we would take $\mathbf{U}$ to be its leading left singular vectors.

That turns out to be good enough — and it decouples the modes completely.

Also note the **Kronecker product** $\mathbf{W} \otimes \mathbf{V}$ appearing on its own. Remember it too.

Truncated higher-order SVD (T-HOSVD)

1. Compute U = leading P left singular vectors of $X_{(1)}$
2. Compute V = leading Q left singular vectors of $X_{(2)}$
3. Compute W = leading R left singular vectors of $X_{(3)}$
4. $\mathcal{G} = \mathcal{X} \times_1 U^T \times_2 V^T \times_3 W^T$

Quasi-optimality: the T-HOSVD error is within a factor of $\sqrt{N}$ of optimal for the given ranks.

No iteration, no local minima, and a computable error bound.

Sequentially truncated variant (ST-HOSVD): truncate the tensor after each factor matrix is computed — cheaper, and equally accurate.

Truncated higher-order SVD (T-HOSVD)

1. Compute U = leading P left singular vectors of $X_{(1)}$
2. Compute V = leading Q left singular vectors of $X_{(2)}$
3. Compute W = leading R left singular vectors of $X_{(3)}$
4. $\mathcal{G} = \mathcal{X} \times_1 U^T \times_2 V^T \times_3 W^T$

Quasi-optimality: the T-HOSVD error is within a factor of $\sqrt{N}$ of optimal for the given ranks.
No iteration, no local minima, and a computable error bound.

Sequentially truncated variant (ST-HOSVD): truncate the tensor after each factor matrix is computed — cheaper, and equally accurate.

Bottleneck number two: those SVDs are of $I \times I^{N-1}$ unfoldings.

Addition in Tucker format

Let $\mathcal{X} = \llbracket \mathcal{G}; \mathbf{U}, \mathbf{V}, \mathbf{W} \rrbracket$ and $\mathcal{Y} = \llbracket \tilde{\mathcal{G}}; \tilde{\mathbf{U}}, \tilde{\mathbf{V}}, \tilde{\mathbf{W}} \rrbracket$. Concatenate the factors and stack the cores **block-diagonally**:

$$\mathcal{X} + \mathcal{Y} = \llbracket \begin{bmatrix} \mathcal{G} & \\ & \tilde{\mathcal{G}} \end{bmatrix}; [\mathbf{U} \quad \tilde{\mathbf{U}}], [\mathbf{V} \quad \tilde{\mathbf{V}}], [\mathbf{W} \quad \tilde{\mathbf{W}}] \rrbracket$$

The new core is $(P+\tilde{P}) \times (Q+\tilde{Q}) \times (R+\tilde{R})$: $\mathcal{G}$ in the leading block, $\tilde{\mathcal{G}}$ in the trailing one, *zeros elsewhere* — keeping the two sets of subspaces from mixing.

Addition in Tucker format

Let $\mathcal{X} = \llbracket \mathcal{G}; \mathbf{U}, \mathbf{V}, \mathbf{W} \rrbracket$ and $\mathcal{Y} = \llbracket \tilde{\mathcal{G}}; \tilde{\mathbf{U}}, \tilde{\mathbf{V}}, \tilde{\mathbf{W}} \rrbracket$. Concatenate the factors and stack the cores **block-diagonally**:

$$\mathcal{X} + \mathcal{Y} = \llbracket \begin{bmatrix} \mathcal{G} & \\ & \tilde{\mathcal{G}} \end{bmatrix}; [\mathbf{U} \quad \tilde{\mathbf{U}}], [\mathbf{V} \quad \tilde{\mathbf{V}}], [\mathbf{W} \quad \tilde{\mathbf{W}}] \rrbracket$$

The new core is $(P+\tilde{P}) \times (Q+\tilde{Q}) \times (R+\tilde{R})$: $\mathcal{G}$ in the leading block, $\tilde{\mathcal{G}}$ in the trailing one, *zeros elsewhere* — keeping the two sets of subspaces from mixing.

- ▶ again **zero flops**: ranks *add* in every mode, so the core *multiplies* — 2^N times bigger for two order- N tensors, mostly zeros
- ▶ so, unlike CP, a sum cannot be left unrounded for long

Addition in Tucker format

Let $\mathcal{X} = \llbracket \mathcal{G}; \mathbf{U}, \mathbf{V}, \mathbf{W} \rrbracket$ and $\mathcal{Y} = \llbracket \tilde{\mathcal{G}}; \tilde{\mathbf{U}}, \tilde{\mathbf{V}}, \tilde{\mathbf{W}} \rrbracket$. Concatenate the factors and stack the cores **block-diagonally**:

$$\mathcal{X} + \mathcal{Y} = \llbracket \begin{bmatrix} \mathcal{G} & \\ & \tilde{\mathcal{G}} \end{bmatrix}; \begin{bmatrix} \mathbf{U} & \tilde{\mathbf{U}} \end{bmatrix}, \begin{bmatrix} \mathbf{V} & \tilde{\mathbf{V}} \end{bmatrix}, \begin{bmatrix} \mathbf{W} & \tilde{\mathbf{W}} \end{bmatrix} \rrbracket$$

The new core is $(P+\tilde{P}) \times (Q+\tilde{Q}) \times (R+\tilde{R})$: $\mathcal{G}$ in the leading block, $\tilde{\mathcal{G}}$ in the trailing one, *zeros elsewhere* — keeping the two sets of subspaces from mixing.

- ▶ again **zero flops**: ranks *add* in every mode, so the core *multiplies* — 2^N times bigger for two order- N tensors, mostly zeros
- ▶ so, unlike CP, a sum cannot be left unrounded for long

But rounding is cheap here, and touches nothing of size I^N :

1. $\begin{bmatrix} \mathbf{U} & \tilde{\mathbf{U}} \end{bmatrix} = \mathbf{Q}_1 \mathbf{R}_1$ and likewise per mode — thin QR, $O(I(P+\tilde{P})^2)$
2. apply the triangular factors to the core, $\mathcal{G} \times_1 \mathbf{R}_1 \times_2 \mathbf{R}_2 \times_3 \mathbf{R}_3$ — TTMs on a small tensor
3. ST-HOSVD that small core and fold the result back into $\mathbf{Q}_1, \mathbf{Q}_2, \mathbf{Q}_3$

Inner products in Tucker format

Vectorize: $\text{vec}(\mathcal{X}) = (W \otimes V \otimes U) \text{vec}(\mathcal{G})$, then use the mixed-product rule

$$(W \otimes V \otimes U)^T (\tilde{W} \otimes \tilde{V} \otimes \tilde{U}) = (W^T \tilde{W}) \otimes (V^T \tilde{V}) \otimes (U^T \tilde{U})$$

Inner products in Tucker format

Vectorize: $\text{vec}(\mathcal{X}) = (W \otimes V \otimes U) \text{vec}(\mathcal{G})$, then use the mixed-product rule

$$(W \otimes V \otimes U)^T (\tilde{W} \otimes \tilde{V} \otimes \tilde{U}) = (W^T \tilde{W}) \otimes (V^T \tilde{V}) \otimes (U^T \tilde{U})$$

which, back in tensor notation, says: **push all six factor matrices onto one small core**

$$\langle \mathcal{X}, \mathcal{Y} \rangle = \langle \mathcal{G}, \tilde{\mathcal{G}} \times_1 (U^T \tilde{U}) \times_2 (V^T \tilde{V}) \times_3 (W^T \tilde{W}) \rangle$$

Inner products in Tucker format

Vectorize: $\text{vec}(\mathcal{X}) = (W \otimes V \otimes U) \text{vec}(\mathcal{G})$, then use the mixed-product rule

$$(W \otimes V \otimes U)^T (\tilde{W} \otimes \tilde{V} \otimes \tilde{U}) = (W^T \tilde{W}) \otimes (V^T \tilde{V}) \otimes (U^T \tilde{U})$$

which, back in tensor notation, says: **push all six factor matrices onto one small core**

$$\langle \mathcal{X}, \mathcal{Y} \rangle = \langle \mathcal{G}, \tilde{\mathcal{G}} \times_1 (U^T \tilde{U}) \times_2 (V^T \tilde{V}) \times_3 (W^T \tilde{W}) \rangle$$

- ▶ three Gram products between the factor matrices: $O(IP\tilde{P} + JQ\tilde{Q} + KR\tilde{R})$
- ▶ a chain of TTMs on the small core, $O(\tilde{P}\tilde{Q}\tilde{R}P + P\tilde{Q}\tilde{R}Q + PQ\tilde{R}R)$, then a dot product
- ▶ **the mode sizes I, J, K enter only linearly**, through those Gram products

Inner products in Tucker format

Vectorize: $\text{vec}(\mathcal{X}) = (W \otimes V \otimes U) \text{vec}(\mathcal{G})$, then use the mixed-product rule

$$(W \otimes V \otimes U)^T (\tilde{W} \otimes \tilde{V} \otimes \tilde{U}) = (W^T \tilde{W}) \otimes (V^T \tilde{V}) \otimes (U^T \tilde{U})$$

which, back in tensor notation, says: **push all six factor matrices onto one small core**

$$\langle \mathcal{X}, \mathcal{Y} \rangle = \langle \mathcal{G}, \tilde{\mathcal{G}} \times_1 (U^T \tilde{U}) \times_2 (V^T \tilde{V}) \times_3 (W^T \tilde{W}) \rangle$$

- ▶ three Gram products between the factor matrices: $O(IP\tilde{P} + JQ\tilde{Q} + KR\tilde{R})$
- ▶ a chain of TTMs on the small core, $O(\tilde{P}\tilde{Q}\tilde{R}P + P\tilde{Q}\tilde{R}Q + PQ\tilde{R}R)$, then a dot product
- ▶ **the mode sizes I, J, K enter only linearly**, through those Gram products

With orthonormal factors — as T-HOSVD always leaves them — $U^T U = I$ and this collapses to $\|\mathcal{X}\| = \|\mathcal{G}\|$: **the norm of a Tucker tensor is the norm of its core.**

$$\mathcal{X} \approx \mathcal{G} \times_1 \mathbf{U} \times_2 \mathbf{V} \times_3 \mathbf{W}$$

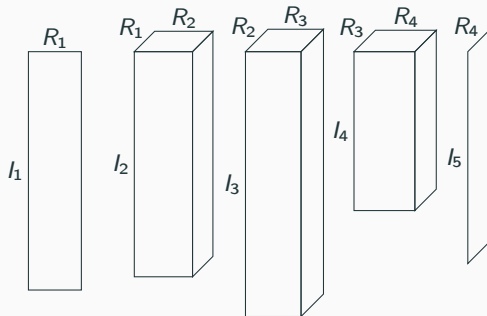
unique / interpretable factors	✗
solves the approximation problem	✓
compresses (storage linear in dimensions)	✓
stable, predictable algorithms	✓

- ▶ Tucker is *not* unique — the factors are subspaces, and you cannot interpret individual columns
- ▶ but it solves the approximation problem **quasi-optimally**, and you can hit a prescribed ϵ
- ▶ compression is real but **not linear in the order**: the core is $R_1 \cdots R_N$, so storage still grows exponentially in N
- ▶ algorithms are direct, stable and predictable — running time depends only on the dimensions

That third point is the catch, and it is why we need a third format.

Tensor Train

Tensor Train (TT) notation



$$\mathcal{X} \approx \{\mathcal{T}_{\mathcal{X},k}\}, \mathcal{X} \in \mathbb{R}^{l_1 \times l_2 \times l_3 \times l_4 \times l_5}$$

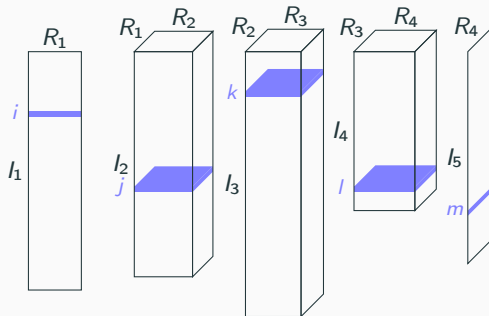
$$\mathcal{T}_{\mathcal{X},n} \in \mathbb{R}^{R_{n-1} \times l_n \times R_n}$$

are *TT cores*

$$x_{ijklm} \approx \sum_{\alpha=1}^{R_1} \sum_{\beta=1}^{R_2} \sum_{\gamma=1}^{R_3} \sum_{\delta=1}^{R_4} \mathcal{T}_{\mathcal{X},1}(i, \alpha) \mathcal{T}_{\mathcal{X},2}(\alpha, j, \beta) \mathcal{T}_{\mathcal{X},3}(\beta, k, \gamma) \mathcal{T}_{\mathcal{X},4}(\gamma, l, \delta) \mathcal{T}_{\mathcal{X},5}(\delta, m)$$

Storage $(l_1 + \dots + l_N)R^2$: linear in the order. The format for the cookies problem.

Tensor Train (TT) notation



$$\mathcal{X} \approx \{\mathcal{T}_{\mathcal{X},k}\}, \mathcal{X} \in \mathbb{R}^{l_1 \times l_2 \times l_3 \times l_4 \times l_5}$$

$$\mathcal{T}_{\mathcal{X},n} \in \mathbb{R}^{R_{n-1} \times l_n \times R_n}$$

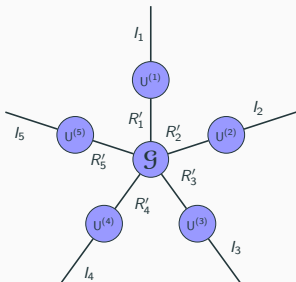
are *TT cores*

$$x_{ijklm} \approx \sum_{\alpha=1}^{R_1} \sum_{\beta=1}^{R_2} \sum_{\gamma=1}^{R_3} \sum_{\delta=1}^{R_4} \mathcal{T}_{\mathcal{X},1}(i, \alpha) \mathcal{T}_{\mathcal{X},2}(\alpha, j, \beta) \mathcal{T}_{\mathcal{X},3}(\beta, k, \gamma) \mathcal{T}_{\mathcal{X},4}(\gamma, l, \delta) \mathcal{T}_{\mathcal{X},5}(\delta, m)$$

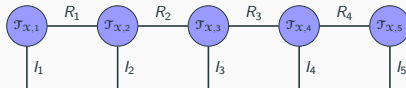
Storage $(l_1 + \dots + l_N)R^2$: linear in the order. The format for the cookies problem.

Tensor network diagrams

Tucker



Tensor Train



Vertices are tensors, edges are modes, degree is dimension, connected edges are contractions

Tucker has one high-degree vertex (the core, R^N entries); TT has a chain of degree-3 vertices (NIR^2).

Intuition for TT ranks

Tucker: the n th Tucker rank approximates the numerical rank of the n th modal unfolding

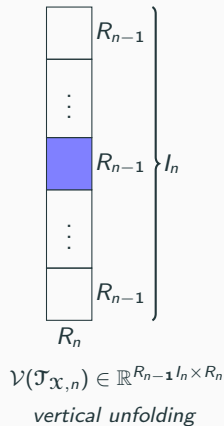
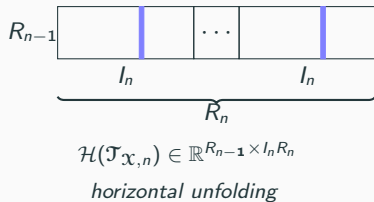
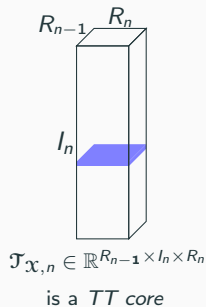
$$X_{(n)} \text{ is } I_n \times (I_1 \cdots I_{n-1} I_{n+1} \cdots I_N)$$

Tensor Train: the n th TT rank approximates the numerical rank of the unfolding that maps the *first* n modes to rows

$$X_{(1:n)} \text{ is } (I_1 \cdots I_n) \times (I_{n+1} \cdots I_N)$$

- ▶ so TT asks about a *sequence of splits* of the mode list, not about each mode against all the others
- ▶ the full-format tensor would be far too big to store, so $X_{(1:n)}$ is only ever available **implicitly**, through the cores
- ▶ truncating the ranks of a tensor already in TT format is called **TT rounding** (Oseledets 2011)

Vertical and horizontal unfoldings of TT cores



Example algorithm: TT rounding

function $\{\mathcal{T}_{\mathcal{Z},n}\} = \text{TT-ROUNDING}(\{\mathcal{T}_{\mathcal{X},n}\})$

▷ Orthogonalization Phase

for $n = N$ down to 2 **do**

$[\mathbf{Y}_n, \mathbf{R}_n] = \text{QR}(\mathcal{H}(\mathcal{T}_{\mathcal{X},n})^T)$

$\mathcal{V}(\mathcal{T}_{\mathcal{X},n-1}) = \mathcal{V}(\mathcal{T}_{\mathcal{X},n-1}) \cdot \mathbf{R}^T$

▷ (tall-skinny) QR factorization

▷ Apply R to previous core

▷ Truncation Phase

$\mathcal{Z} = \mathcal{X}$

for $n = 1$ to $N - 1$ **do**

$[\mathbf{Y}_n, \mathbf{R}_n] = \text{QR}(\mathcal{V}(\mathcal{T}_{\mathcal{Z},n}))$

$[\hat{\mathbf{U}}_R, \hat{\Sigma}, \hat{\mathbf{V}}] \approx \text{TSVD}(\mathbf{R}_n)$

$\mathcal{V}(\mathcal{T}_{\mathcal{Z},n}) = \text{APPLY-Q}(\mathbf{Y}_n, \hat{\mathbf{U}}_R)$

$\mathcal{H}(\mathcal{T}_{\mathcal{Z},n+1})^T = \text{APPLY-Q}(\mathbf{Y}_{n+1}, \hat{\mathbf{V}}\hat{\Sigma})$

▷ (tall-skinny) QR factorization

▷ Truncated SVD of R

Both core unfoldings are column major, so BLAS/LAPACK apply directly.

Bottleneck number three: that orthogonalization sweep, $O(NIR^3)$, dominates.

Addition in TT format

Each entry of a TT tensor is a product of matrices, one per mode. To add two of them, make those matrices **block-diagonal**, so the two chains run side by side without interacting:

$$\mathcal{T}_{\mathbf{z},n}(:, i_n, :) = \begin{bmatrix} \mathcal{T}_{\mathbf{x},n}(:, i_n, :) & \\ & \mathcal{T}_{\mathbf{y},n}(:, i_n, :) \end{bmatrix} \in \mathbb{R}^{(R_{n-1} + \tilde{R}_{n-1}) \times (R_n + \tilde{R}_n)}, \quad 1 < n < N$$

closing the chain at both ends so the two products are *summed*, not kept apart:

$$\mathcal{T}_{\mathbf{z},1}(i_1, :) = \begin{bmatrix} \mathcal{T}_{\mathbf{x},1}(i_1, :) & \mathcal{T}_{\mathbf{y},1}(i_1, :) \end{bmatrix}, \quad \mathcal{T}_{\mathbf{z},N}(:, i_N) = \begin{bmatrix} \mathcal{T}_{\mathbf{x},N}(:, i_N) \\ \mathcal{T}_{\mathbf{y},N}(:, i_N) \end{bmatrix}$$

Addition in TT format

Each entry of a TT tensor is a product of matrices, one per mode. To add two of them, make those matrices **block-diagonal**, so the two chains run side by side without interacting:

$$\mathcal{T}_{\mathbf{z},n}(:, i_n, :) = \begin{bmatrix} \mathcal{T}_{\mathbf{x},n}(:, i_n, :) & \\ & \mathcal{T}_{\mathbf{y},n}(:, i_n, :) \end{bmatrix} \in \mathbb{R}^{(R_{n-1} + \tilde{R}_{n-1}) \times (R_n + \tilde{R}_n)}, \quad 1 < n < N$$

closing the chain at both ends so the two products are *summed*, not kept apart:

$$\mathcal{T}_{\mathbf{z},1}(i_1, :) = \begin{bmatrix} \mathcal{T}_{\mathbf{x},1}(i_1, :) & \mathcal{T}_{\mathbf{y},1}(i_1, :) \end{bmatrix}, \quad \mathcal{T}_{\mathbf{z},N}(:, i_N) = \begin{bmatrix} \mathcal{T}_{\mathbf{x},N}(:, i_N) \\ \mathcal{T}_{\mathbf{y},N}(:, i_N) \end{bmatrix}$$

- ▶ multiplying out the block-diagonal chain gives $x_{i_1 \dots i_N} + y_{i_1 \dots i_N}$, as required
- ▶ **zero flops, pure data movement** — but every internal rank adds: $R_n + \tilde{R}_n$

Addition in TT format

Each entry of a TT tensor is a product of matrices, one per mode. To add two of them, make those matrices **block-diagonal**, so the two chains run side by side without interacting:

$$\mathcal{T}_{\mathbf{z},n}(:, i_n, :) = \begin{bmatrix} \mathcal{T}_{\mathbf{x},n}(:, i_n, :) & \\ & \mathcal{T}_{\mathbf{y},n}(:, i_n, :) \end{bmatrix} \in \mathbb{R}^{(R_{n-1} + \tilde{R}_{n-1}) \times (R_n + \tilde{R}_n)}, \quad 1 < n < N$$

closing the chain at both ends so the two products are *summed*, not kept apart:

$$\mathcal{T}_{\mathbf{z},1}(i_1, :) = \begin{bmatrix} \mathcal{T}_{\mathbf{x},1}(i_1, :) & \mathcal{T}_{\mathbf{y},1}(i_1, :) \end{bmatrix}, \quad \mathcal{T}_{\mathbf{z},N}(:, i_N) = \begin{bmatrix} \mathcal{T}_{\mathbf{x},N}(:, i_N) \\ \mathcal{T}_{\mathbf{y},N}(:, i_N) \end{bmatrix}$$

- ▶ multiplying out the block-diagonal chain gives $x_{i_1 \dots i_N} + y_{i_1 \dots i_N}$, as required
- ▶ **zero flops, pure data movement** — but every internal rank adds: $R_n + \tilde{R}_n$

This is why TT rounding is the workhorse kernel. Sums are everywhere — time stepping, Krylov methods, residuals — and each doubles the ranks. Rounding is the inner loop.

Inner products in TT format

Contract the two trains together **one mode at a time**, carrying a small $R_n \times \tilde{R}_n$ matrix M_n from left to right:

$$M_0 = 1, \quad M_n = \sum_{i_n=1}^{I_n} \mathcal{T}_{\mathbf{x},n}(:, i_n, :)^T M_{n-1} \mathcal{T}_{\mathbf{y},n}(:, i_n, :), \quad \langle \mathbf{x}, \mathbf{y} \rangle = M_N$$

Inner products in TT format

Contract the two trains together **one mode at a time**, carrying a small $R_n \times \tilde{R}_n$ matrix M_n from left to right:

$$M_0 = 1, \quad M_n = \sum_{i_n=1}^{I_n} \mathcal{T}_{\mathcal{X},n}(:, i_n, :)^T M_{n-1} \mathcal{T}_{\mathcal{Y},n}(:, i_n, :), \quad \langle \mathcal{X}, \mathcal{Y} \rangle = M_N$$

Each step is two dense matrix products on reshaped cores:

1. $M_{n-1} \cdot \mathcal{V}(\mathcal{T}_{\mathcal{Y},n})$ costs $O(R_{n-1} \tilde{R}_{n-1} I_n \tilde{R}_n)$
2. contract the result against $\mathcal{T}_{\mathcal{X},n}$, costing $O(R_{n-1} I_n \tilde{R}_n R_n)$

for a total of $O(NIR\tilde{R}(R + \tilde{R}))$ — **linear in the order and in the mode sizes**, and BLAS-3 throughout.

Inner products in TT format

Contract the two trains together **one mode at a time**, carrying a small $R_n \times \tilde{R}_n$ matrix M_n from left to right:

$$M_0 = 1, \quad M_n = \sum_{i_n=1}^{I_n} \mathcal{T}_{\mathcal{X},n}(:, i_n, :)^T M_{n-1} \mathcal{T}_{\mathcal{Y},n}(:, i_n, :), \quad \langle \mathcal{X}, \mathcal{Y} \rangle = M_N$$

Each step is two dense matrix products on reshaped cores:

1. $M_{n-1} \cdot \mathcal{V}(\mathcal{T}_{\mathcal{Y},n})$ costs $O(R_{n-1} \tilde{R}_{n-1} I_n \tilde{R}_n)$
2. contract the result against $\mathcal{T}_{\mathcal{X},n}$, costing $O(R_{n-1} I_n \tilde{R}_n R_n)$

for a total of $O(NIR\tilde{R}(R + \tilde{R}))$ — **linear in the order and in the mode sizes**, and BLAS-3 throughout.

Never form the block-diagonal sum and then contract: that would square the ranks. Sweep instead.

Inner products in TT format

Contract the two trains together **one mode at a time**, carrying a small $R_n \times \tilde{R}_n$ matrix M_n from left to right:

$$M_0 = 1, \quad M_n = \sum_{i_n=1}^{I_n} \mathcal{T}_{\mathcal{X},n}(:, i_n, :)^T M_{n-1} \mathcal{T}_{\mathcal{Y},n}(:, i_n, :), \quad \langle \mathcal{X}, \mathcal{Y} \rangle = M_N$$

Each step is two dense matrix products on reshaped cores:

1. $M_{n-1} \cdot \mathcal{V}(\mathcal{T}_{\mathcal{Y},n})$ costs $O(R_{n-1} \tilde{R}_{n-1} I_n \tilde{R}_n)$
2. contract the result against $\mathcal{T}_{\mathcal{X},n}$, costing $O(R_{n-1} I_n \tilde{R}_n R_n)$

for a total of $O(NIR\tilde{R}(R + \tilde{R}))$ — **linear in the order and in the mode sizes**, and BLAS-3 throughout.

Never form the block-diagonal sum and then contract: that would square the ranks. Sweep instead.

Norms: orthogonalize the cores first (the sweep of QRs from the TT rounding algorithm); then all but one core cancels and $\|\mathcal{X}\|$ is the Frobenius norm of the single remaining core.

Scorecard: Tensor Train

$$x_{ijklm} \approx \mathcal{T}_{x,1}(i, :)\mathcal{T}_{x,2}(:, j, :)\mathcal{T}_{x,3}(:, k, :)\mathcal{T}_{x,4}(:, l, :)\mathcal{T}_{x,5}(:, m)$$

unique / interpretable factors	✗
solves the approximation problem	✓
compresses (storage linear in dimensions)	✓
stable, predictable algorithms	✓

- ▶ TT is not unique — you cannot interpret the cores
- ▶ quasi-optimal for given ranks, and you can hit a prescribed ϵ
- ▶ **storage is linear in the order**: $(I_1 + \dots + I_N)R^2$ — this is the box Tucker could not tick, and the reason TT can handle an 11-way tensor
- ▶ stable, predictable algorithms

Scorecard: all four, together

	unique	optimal	compresses	stable
SVD	✓	✓	✓	✓
CP	✓	✗	✓	✗
Tucker	✗	✓	✓	✓
TT	✗	✓	✓	✓

- ▶ want to interpret the factors? CP is the only option, and you pay for it in stability
- ▶ want a guaranteed error at a given cost? Tucker or TT
- ▶ order 3 or 4, and want the modes treated symmetrically? Tucker
- ▶ order 10 or more? TT, because it is the only one whose storage does not blow up with the order

No format wins. Pick the column you cannot do without.

Three formats, three bottlenecks — and one shape

format	algorithm	bottleneck	cost
CP	ALS	MTTKRP / the LLS solve	$O(I^N r)$ per mode, per iteration
Tucker	ST-HOSVD	SVD of the modal unfoldings	$O(I^{N+1})$
TT	rounding	orthogonalization sweep	$O(NIR^3)$

Three formats, three bottlenecks — and one shape

format	algorithm	bottleneck	cost
CP	ALS	MTTKRP / the LLS solve	$O(I^N r)$ per mode, per iteration
Tucker	ST-HOSVD	SVD of the modal unfoldings	$O(I^{N+1})$
TT	rounding	orthogonalization sweep	$O(NIR^3)$

Every one of them has the same shape:

a huge, *implicitly defined* matrix, and we only want its range,
or a least-squares solution against it

- ▶ CP: the coefficient matrix is a Khatri-Rao product, $I^{N-1} \times r$
- ▶ Tucker: the unfolding $X_{(n)}$ is $I \times I^{N-1}$
- ▶ TT: the unfolding $X_{(1:n)}$ exists only through the cores

That is precisely the situation randomization was invented for. *After the break: why the obvious way to randomize does not work, and what does.*

Break

10 minutes

Randomization, and why it needs structure

Randomized rank- k decomposition

Randomized range finder (Halko, Martinsson & Tropp 2011). To find an approximate basis for the range of $A \in \mathbb{R}^{m \times n}$:

$$Y = A\Omega, \quad \Omega \in \mathbb{R}^{n \times \ell} \text{ random}, \quad Q = \text{orth}(Y)$$

with $\ell = r + \alpha$ a small oversampling parameter. Then $A \approx QQ^T A$.

Work out how you apply it to $A = XZ^T$ and $A = X_1Z_1^T + X_2Z_2^T$.

Why we cannot just use a Gaussian

Take an N -way tensor with all modes I . The mode- n unfolding is $I \times I^{N-1}$, so a dense Ω is $I^{N-1} \times \ell$.

	tensor I^N	dense Ω $I^{N-1}\ell$
$I=100, N=3, \ell=20$	10^6	2×10^5
$I=100, N=4, \ell=20$	10^8	2×10^7
$I=100, N=5, \ell=20$	10^{10}	2×10^9
$I=100, N=8, \ell=20$	10^{16}	2×10^{15}

The sketch is a constant fraction of the tensor — at every order.

And that is the *optimistic* reading. In the cases we actually care about:

- ▶ the tensor is **never stored in full** anyway — it lives in TT or Tucker form, and is far smaller than I^N
- ▶ so the random matrix would be **orders of magnitude larger than the data it is compressing**
- ▶ and forming $A\Omega$ would cost more than the deterministic algorithm we set out to replace

Why we cannot just use a Gaussian

Take an N -way tensor with all modes I . The mode- n unfolding is $I \times I^{N-1}$, so a dense Ω is $I^{N-1} \times \ell$.

	tensor I^N	dense Ω $I^{N-1}\ell$
$I=100, N=3, \ell=20$	10^6	2×10^5
$I=100, N=4, \ell=20$	10^8	2×10^7
$I=100, N=5, \ell=20$	10^{10}	2×10^9
$I=100, N=8, \ell=20$	10^{16}	2×10^{15}

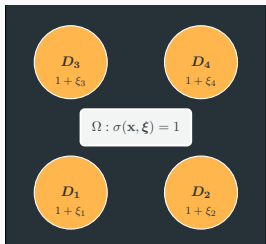
The sketch is a constant fraction of the tensor — at every order.

And that is the *optimistic* reading. In the cases we actually care about:

- ▶ the tensor is **never stored in full** anyway — it lives in TT or Tucker form, and is far smaller than I^N
- ▶ so the random matrix would be **orders of magnitude larger than the data it is compressing**
- ▶ and forming $A\Omega$ would cost more than the deterministic algorithm we set out to replace

We need an Ω that is never formed and never stored — only *applied* in a way exploiting the structure the data already has.

Where rounding actually comes from



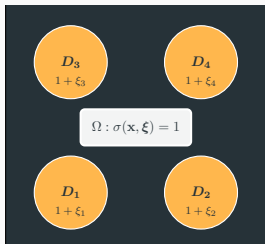
the cookies problem

Discretize the parameter-dependent PDE in space *and* in all p parameters at once and the linear system inherits the tensor structure:

$$\left(A_0 \otimes I \otimes \cdots \otimes I + \sum_{i=1}^p A_i \otimes I \otimes \cdots \otimes \underbrace{D_i}_{\text{mode } i+1} \otimes \cdots \otimes I \right) \text{vec}(\mathcal{U}) = f \otimes 1 \otimes \cdots \otimes 1$$

with $D_i = \text{diag}(\rho_i^{(1)}, \dots, \rho_i^{(I)})$, I samples per parameter: a **rank- $(p+1)$ operator**, a rank-one right-hand side, and an unknown $\mathcal{U}$ of order $N = p+1$ that is **never formed**.

Where rounding actually comes from



the cookies problem

Discretize the parameter-dependent PDE in space *and* in all p parameters at once and the linear system inherits the tensor structure:

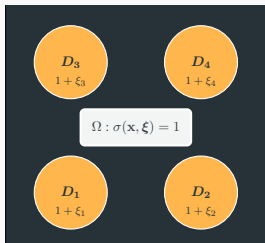
$$\left(A_0 \otimes I \otimes \cdots \otimes I + \sum_{i=1}^p A_i \otimes I \otimes \cdots \otimes \underbrace{D_i}_{\text{mode } i+1} \otimes \cdots \otimes I \right) \text{vec}(\mathcal{U}) = f \otimes 1 \otimes \cdots \otimes 1$$

with $D_i = \text{diag}(\rho_i^{(1)}, \dots, \rho_i^{(I)})$, I samples per parameter: a **rank- $(p+1)$ operator**, a rank-one right-hand side, and an unknown $\mathcal{U}$ of order $N = p+1$ that is **never formed**.

Solve it with low-rank GMRES (TT-GMRES, Dolgov 2013): every Krylov vector is a low-rank tensor, and every step **inflates the ranks**:

- ▶ the matvec $\mathcal{W} = \mathcal{A}\mathcal{V}_j$ multiplies the ranks by $p+1$
- ▶ Arnoldi orthogonalization $\mathcal{W} \leftarrow \mathcal{W} - \sum_{i \leq j} h_{ij} \mathcal{V}_i$ is a **sum of $j+1$ low-rank tensors** — ranks *add*; so is the update $\mathcal{X} = \sum_k y_k \mathcal{V}_k$

Where rounding actually comes from



the cookies problem

Discretize the parameter-dependent PDE in space *and* in all p parameters at once and the linear system inherits the tensor structure:

$$\left(A_0 \otimes I \otimes \cdots \otimes I + \sum_{i=1}^p A_i \otimes I \otimes \cdots \otimes \underbrace{D_i}_{\text{mode } i+1} \otimes \cdots \otimes I \right) \text{vec}(\mathcal{U}) = f \otimes 1 \otimes \cdots \otimes 1$$

with $D_i = \text{diag}(\rho_i^{(1)}, \dots, \rho_i^{(I)})$, I samples per parameter: a **rank- $(p+1)$ operator**, a rank-one right-hand side, and an unknown $\mathcal{U}$ of order $N = p+1$ that is **never formed**.

Solve it with low-rank GMRES (TT-GMRES, Dolgov 2013): every Krylov vector is a low-rank tensor, and every step **inflates the ranks**:

- ▶ the matvec $\mathcal{W} = \mathcal{A}\mathcal{V}_j$ multiplies the ranks by $p+1$
- ▶ Arnoldi orthogonalization $\mathcal{W} \leftarrow \mathcal{W} - \sum_{i \leq j} h_{ij} \mathcal{V}_i$ is a **sum of $j+1$ low-rank tensors** — ranks *add*; so is the update $\mathcal{X} = \sum_k y_k \mathcal{V}_k$

Additions were free; **this is where the bill arrives**. Each must be truncated back before the next step — and the kernel is always the same: **round a sum of low-rank tensors**.

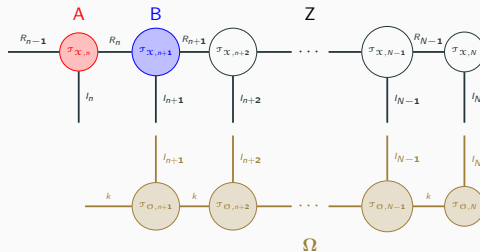
Family 1: tensor-train

Family 1: random matrices whose columns are tensor trains

For TT rounding we must sketch $X_{(1:n)}$, which is $(l_1 \cdots l_n) \times (l_{n+1} \cdots l_N)$ and **exists only implicitly**. So give Ω the same structure the data has: build it from its own cores $\{\mathcal{T}_{\Omega,n+1}, \dots, \mathcal{T}_{\Omega,N}\}$, each of rank k .

$$\Omega \in \mathbb{R}^{(l_{n+1} \cdots l_N) \times k}$$

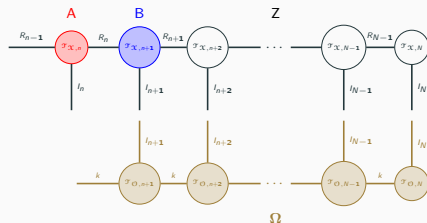
- ▶ stored in $O(\sum_j l_j k^2)$
- ▶ each *column* of Ω is itself a tensor in TT format



the data $\mathcal{X}$ (top) and the random Ω (bottom), as one network

Why this is cheap: the sketch is a contraction

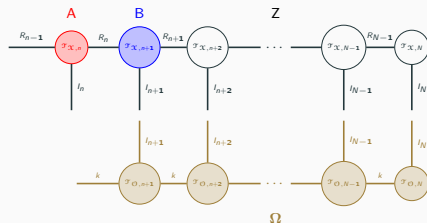
$X_{(1:n)}\Omega$ is not a matrix product we evaluate. It is a **partial contraction of two tensor networks** — read the picture from the right:



- ▶ contract $\mathcal{T}_{X,N}$ with $\mathcal{T}_{\Omega,N}$ over mode I_N : a small $R_{N-1} \times k$ matrix
- ▶ absorb it into $\mathcal{T}_{X,N-1}$, contract again, and sweep leftwards
- ▶ total cost $O(NIR^2k)$ — **linear in everything**, and the unfolding is never formed

Why this is cheap: the sketch is a contraction

$X_{(1:n)}\Omega$ is not a matrix product we evaluate. It is a **partial contraction of two tensor networks** — read the picture from the right:



- ▶ contract $\mathcal{T}_{X,N}$ with $\mathcal{T}_{O,N}$ over mode I_N : a small $R_{N-1} \times k$ matrix
- ▶ absorb it into $\mathcal{T}_{X,N-1}$, contract again, and sweep leftwards
- ▶ total cost $O(NIR^2k)$ — **linear in everything**, and the unfolding is never formed

And the partial contractions are reused. The sweep that produces the sketch for mode n also produces the sketches for every other mode along the way: *sketch once, sweep once*.

Randomized TT rounding

TT rounding is bottlenecked by the orthogonalization sweep.

Matrix idea: randomized SVD (Halko, Martinsson & Tropp 2011)

- ▶ randomized range finder with an overestimated rank
- ▶ deterministic SVD on the projection to get the estimated rank

Tensor idea: TT-Rounding-RandOrth (Al Daas, Ballard, Saibaba et al. 2023)

- ▶ apply the randomized SVD to the tensor unfoldings
- ▶ use **TT-structured** random matrices, so each sketch is a contraction
- ▶ then deterministically round the TT with overestimated ranks, as in RandSVD

Replaces the $O(NIR^3)$ orthogonalization with $O(NIR^2k)$ contractions.

Randomized TT rounding

TT rounding is bottlenecked by the orthogonalization sweep.

Matrix idea: randomized SVD (Halko, Martinsson & Tropp 2011)

- ▶ randomized range finder with an overestimated rank
- ▶ deterministic SVD on the projection to get the estimated rank

Tensor idea: TT-Rounding-RandOrth (Al Daas, Ballard, Saibaba et al. 2023)

- ▶ apply the randomized SVD to the tensor unfoldings
- ▶ use **TT-structured** random matrices, so each sketch is a contraction
- ▶ then deterministically round the TT with overestimated ranks, as in RandSVD

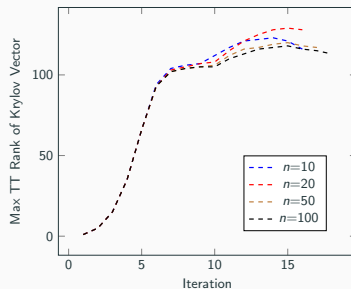
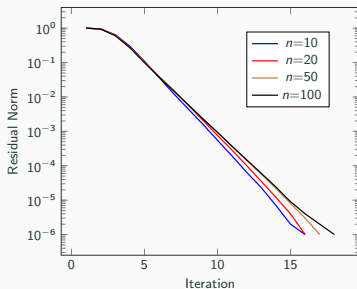
Replaces the $O(NIR^3)$ orthogonalization with $O(NIR^2k)$ contractions.

Same caveat as the Kronecker case, and more so: the entries of a TT-structured Ω are dependent, and the embedding constants degrade with both the sketch rank k and the order N . Oversample; the error is checkable afterwards.

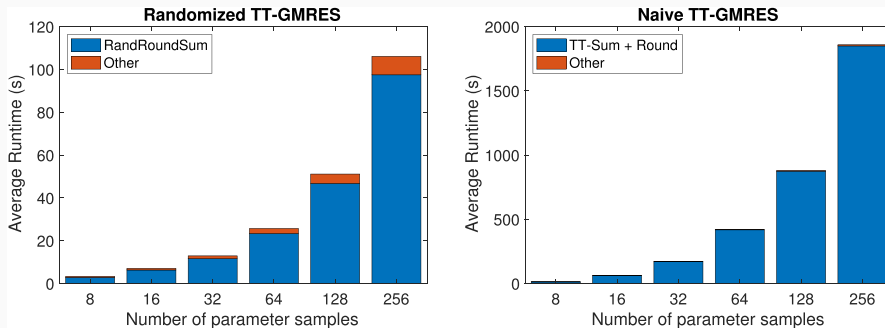
The cookies problem

Solve the cookies problem with 4 cookies, on a tensor of dimension $1781 \times n \times n \times n \times n$, with TT-GMRES (Dolgov 2013).

- ▶ a Krylov method in which **every vector is a TT tensor**
- ▶ full tensors are never explicitly formed
- ▶ the bottleneck is *rounding* the TT tensors after each operation



Randomized TT rounding: timing on the cookies problem



randomized TT-rounding, for varying numbers of parameters

Family 2: Khatri-Rao

Family 2: Khatri-Rao structure

Low-rank GMRES keeps handing us $\mathcal{Y} = \sum_{k=1}^K \alpha_k \mathcal{X}_k$, each summand a TT tensor. Sketch it with a **Khatri-Rao product of small Gaussians**, $\Omega = \Omega_N \odot \cdots \odot \Omega_2$ with $\Omega_n \in \mathbb{R}^{I_n \times \ell}$: storage $\ell \sum_n I_n$, every column a random rank-one tensor.

Family 2: Khatri-Rao structure

Low-rank GMRES keeps handing us $\mathcal{Y} = \sum_{k=1}^K \alpha_k \mathcal{X}_k$, each summand a TT tensor. Sketch it with a **Khatri-Rao product of small Gaussians**, $\Omega = \Omega_N \odot \cdots \odot \Omega_2$ with $\Omega_n \in \mathbb{R}^{I_n \times \ell}$: storage $\ell \sum_n I_n$, every column a random rank-one tensor.

Applying it is a right-to-left sweep of **partial contractions**, one MTTKRP per core:

$$W_N = \mathcal{H}(\mathcal{T}_{\mathcal{X},N}) \Omega_N, \quad W_n = \mathcal{H}(\mathcal{T}_{\mathcal{X},n}) [W_{n+1} \odot \Omega_n], \quad n = N-1, \dots, 2$$

- ▶ randomize-then-orthogonalize exactly as in Family 1: $S_n = \mathcal{V}(\mathcal{T}_{\mathcal{Y},n}) W_{n+1}(:, 1:\ell_n)$, thin QR, new core; $O(NIR^2\ell)$ in total
- ▶ **the sketch of the sum is the sum of the sketches** — contract each $\mathcal{X}_k$ and accumulate, so the block-diagonal formal sum is never built

Family 2: Khatri-Rao structure

Low-rank GMRES keeps handing us $\mathcal{Y} = \sum_{k=1}^K \alpha_k \mathcal{X}_k$, each summand a TT tensor. Sketch it with a **Khatri-Rao product of small Gaussians**, $\Omega = \Omega_N \odot \cdots \odot \Omega_2$ with $\Omega_n \in \mathbb{R}^{I_n \times \ell}$: storage $\ell \sum_n I_n$, every column a random rank-one tensor.

Applying it is a right-to-left sweep of **partial contractions**, one MTTKRP per core:

$$W_N = \mathcal{H}(\mathcal{X}_{\mathcal{X},N}) \Omega_N, \quad W_n = \mathcal{H}(\mathcal{X}_{\mathcal{X},n}) [W_{n+1} \odot \Omega_n], \quad n = N-1, \dots, 2$$

- ▶ randomize-then-orthogonalize exactly as in Family 1: $S_n = \mathcal{V}(\mathcal{Y}_{\mathcal{Y},n}) W_{n+1}(:, 1:\ell_n)$, thin QR, new core; $O(NIR^2\ell)$ in total
- ▶ **the sketch of the sum is the sum of the sketches** — contract each $\mathcal{X}_k$ and accumulate, so the block-diagonal formal sum is never built

Same leading order as the TT-structured sketch, with $O(NI\ell)$ random numbers instead of $O(NI\ell^2)$ — and **this sketch can be widened after the fact**.

Paying for a weaker embedding: adapt the sketch size

A KRP of N Gaussians is **not** a Gaussian: $\frac{1}{\ell} \|A\Omega\|_F^2$ is still unbiased for $\|A\|_F^2$, but the guarantee costs more:

$$\ell \gtrsim \max \left\{ \varepsilon^{-1} C''_{\alpha} \ln^N(2/\delta), \varepsilon^{-2} C'_{\alpha} \ln^2(2/\delta) \right\} \quad \text{with probability } 1 - \delta$$

That $\ln^N$ is the price of the structure — *prescribing ℓ up front is guesswork.*

Paying for a weaker embedding: adapt the sketch size

A KRP of N Gaussians is **not** a Gaussian: $\frac{1}{\ell} \|A\Omega\|_F^2$ is still unbiased for $\|A\|_F^2$, but the guarantee costs more:

$$\ell \gtrsim \max \left\{ \varepsilon^{-1} C''_{\alpha} \ln^N(2/\delta), \varepsilon^{-2} C'_{\alpha} \ln^2(2/\delta) \right\} \quad \text{with probability } 1 - \delta$$

That $\ln^N$ is the price of the structure — *prescribing ℓ up front is guesswork.*

Grow it instead. Per mode, start from b_{init} columns and enrich in blocks

- ▶ **the stopping test is that same estimator:** use the enriching block to estimate the residual $\|X - QQ^T X\|_F$
- ▶ enriching is cheap *because* the sketch is a KRP: draw new Ω_n and append $W_n \leftarrow [W_n \ W_n^{\text{new}}]$, keeping every contraction already done

Paying for a weaker embedding: adapt the sketch size

A KRP of N Gaussians is **not** a Gaussian: $\frac{1}{\ell} \|A\Omega\|_F^2$ is still unbiased for $\|A\|_F^2$, but the guarantee costs more:

$$\ell \gtrsim \max \left\{ \varepsilon^{-1} C''_{\alpha} \ln^N(2/\delta), \varepsilon^{-2} C'_{\alpha} \ln^2(2/\delta) \right\} \quad \text{with probability } 1 - \delta$$

That $\ln^N$ is the price of the structure — *prescribing ℓ up front is guesswork.*

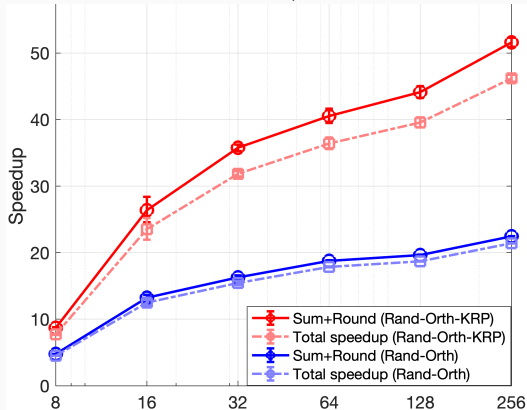
Grow it instead. Per mode, start from b_{init} columns and enrich in blocks

- ▶ **the stopping test is that same estimator:** use the enriching block to estimate the residual $\|X - QQ^T X\|_F$
- ▶ enriching is cheap *because* the sketch is a KRP: draw new Ω_n and append $W_n \leftarrow [W_n \ W_n^{\text{new}}]$, keeping every contraction already done

Adaptivity is what the Khatri-Rao structure buys: very cheap to append columns compared to a TT-format.

Adaptive KRP rounding: back to the cookies problem

TT-GMRES on 4 cookies, with *sum-then-round* done three ways:



red: KRP, blue: TT sketch; axis: mode size l

- ▶ **51×** on the sum-and-round kernel and **46×** on the whole solve
- ▶ **2×** faster than the Family 1 TT-structured sketch — at the same TT ranks per iteration
- ▶ Using 7 cookies, $N = 8$: the fixed-rank randomized solver stalls at residual 10^{-2} ; the adaptive one reaches 10^{-7} against a 10^{-8} target, 25× faster than naive — which runs out of memory at $l = 128$

Al Daas, Ballard, Grigori, Martinez Aguilar, Saibaba & Verma, *Adaptive randomized tensor train rounding using Khatri-Rao products*.

Family 3: Kronecker

Family 3: Kronecker-structured random matrices

Build Ω from one small random matrix per mode, $\Omega_j \in \mathbb{R}^{\ell_j \times I_j}$:

$$\Omega = (\Omega_N \otimes \cdots \otimes \Omega_{j+1} \otimes \Omega_{j-1} \otimes \cdots \otimes \Omega_1)$$

Storage

$$\sum_{i \neq j} I_i \ell_i \quad \text{instead of} \quad \prod_{i \neq j} I_i \ell_i$$

For $I = 100$, $N = 5$, $\ell = 20$: 8×10^3 instead of 2×10^9 .

Application — the key identity

$$\mathbf{Y}_j = \mathbf{X}_{(j)} \Omega^T$$

is exactly a **multi-TTM**:

$$\mathbf{y} = \mathbf{x} \times_1 \Omega_1 \cdots \times_{j-1} \Omega_{j-1} \times_{j+1} \Omega_{j+1} \cdots \times_N \Omega_N$$

Family 3: Kronecker-structured random matrices

Build Ω from one small random matrix per mode, $\Omega_j \in \mathbb{R}^{\ell_j \times l_j}$:

$$\Omega = (\Omega_N \otimes \cdots \otimes \Omega_{j+1} \otimes \Omega_{j-1} \otimes \cdots \otimes \Omega_1)$$

Storage

$$\sum_{i \neq j} l_i \ell_i \quad \text{instead of} \quad \prod_{i \neq j} l_i \ell_i$$

For $l = 100$, $N = 5$, $\ell = 20$: 8×10^3 instead of 2×10^9 .

Application — the key identity

$$\mathbf{Y}_j = \mathbf{X}_{(j)} \Omega^T$$

is exactly a **multi-TTM**:

$$\mathbf{Y} = \mathbf{X} \times_1 \Omega_1 \cdots \times_{j-1} \Omega_{j-1} \times_{j+1} \Omega_{j+1} \cdots \times_N \Omega_N$$

So sketching costs **TTM flops**, using the operation we discussed earlier — and we never form Ω , never form $\mathbf{X}_{(j)}$.

Better still: the partial products $\mathbf{X} \times_1 \Omega_1 \times_2 \Omega_2 \cdots$ are *shared* between the sketches for different modes j .

Randomized T-HOSVD

T-HOSVD is bottlenecked by the SVDs of the unfoldings.

Matrix idea: randomized range finder (Halko, Martinsson & Tropp 2011)

- ▶ post-multiply by a random matrix to approximate the range
- ▶ need an estimate of the rank, plus oversampling

Tensor idea: RandTucker (Minster, Li & Ballard 2022; Minster, Saibaba & Kilmer 2020)

1. Compute U = orthonormalization of $\mathcal{X} \times_2 \Omega_2 \times_3 \Omega_3$
2. Compute V = orthonormalization of $\mathcal{X} \times_1 \Omega_1 \times_3 \Omega_3$
3. Compute W = orthonormalization of $\mathcal{X} \times_1 \Omega_1 \times_2 \Omega_2$
4. $\mathcal{G} = \mathcal{X} \times_1 U^T \times_2 V^T \times_3 W^T$

Every line is a TTM. Parallelizable, cache-friendly, and the sequentially truncated variant randomizes the same way.

The catch: a Kronecker sketch is a weaker embedding

A Kronecker Ω has $\prod_i l_i \ell_i$ entries but only $\sum_i l_i \ell_i$ *degrees of freedom*. Its entries are heavily dependent, and the embedding guarantee degrades accordingly.

What changes

- ▶ for a Gaussian, $\ell = O(r/\epsilon^2)$
- ▶ for a Kronecker product of N factors, the required ℓ picks up a **dependence on N** — the known bounds degrade roughly exponentially in the number of factors
- ▶ the tail/concentration behaviour is heavier: rare bad draws are less rare

What to do about it

- ▶ **oversample per mode**, not globally: take $\ell_j = R_j + p$ with p a genuine margin, not $p = 5$
- ▶ the bounds are pessimistic — in practice these sketches behave far better than the theory promises
- ▶ and you can always *verify*: the Tucker error is computable after the fact

This is the honest trade. We buy a factor of 10^5 in storage and pay for it in sketch size and in theory. It is nearly always worth it — but know that you are paying.

Wrapping up

The three families, side by side

structure of Ω	where it arises	storage	cost to apply	embedding
dense Gaussian	— (the ideal)	$I^{N-1}\ell$	$O(I^N\ell)$	$\ell = O(r/\epsilon^2)$
TT columns $\{\mathcal{T}_{\odot,n}\}$, rank k	TT rounding: sketching $X_{(1:n)}$	$O(\sum_j l_j k^2)$	$O(NIR^2k)$ contraction; one sweep serves all n	degrades with k and N ; k is the dial
Khatri-Rao $\Omega_N \odot \cdots \odot \Omega_2$	rounding a sum of TT tensors, to a <i>tolerance</i>	$\ell \sum_n I_n$	$O(NIR^2\ell)$ partial contractions; $O(NI\ell)$ RNGs	weakest of the three ($\ln^N$), but the sketch can grow
Kronecker $\Omega_N \otimes \cdots \otimes \Omega_1$	Tucker: sketching modal unfoldings $X_{(j)}$	$\sum_i l_i \ell_i$	multi-TTM; partial products shared across modes	degrades with the number of factors

One idea throughout: give the random matrix the same structure as the data, so that “multiply by Ω ” becomes an operation the format already supports.