

Efficient Structure-preserving Support Tensor Train Machine

Kirandeep Kour¹, Sergey Dolgov², Martin Stoll³, Peter Benner^{1,3}

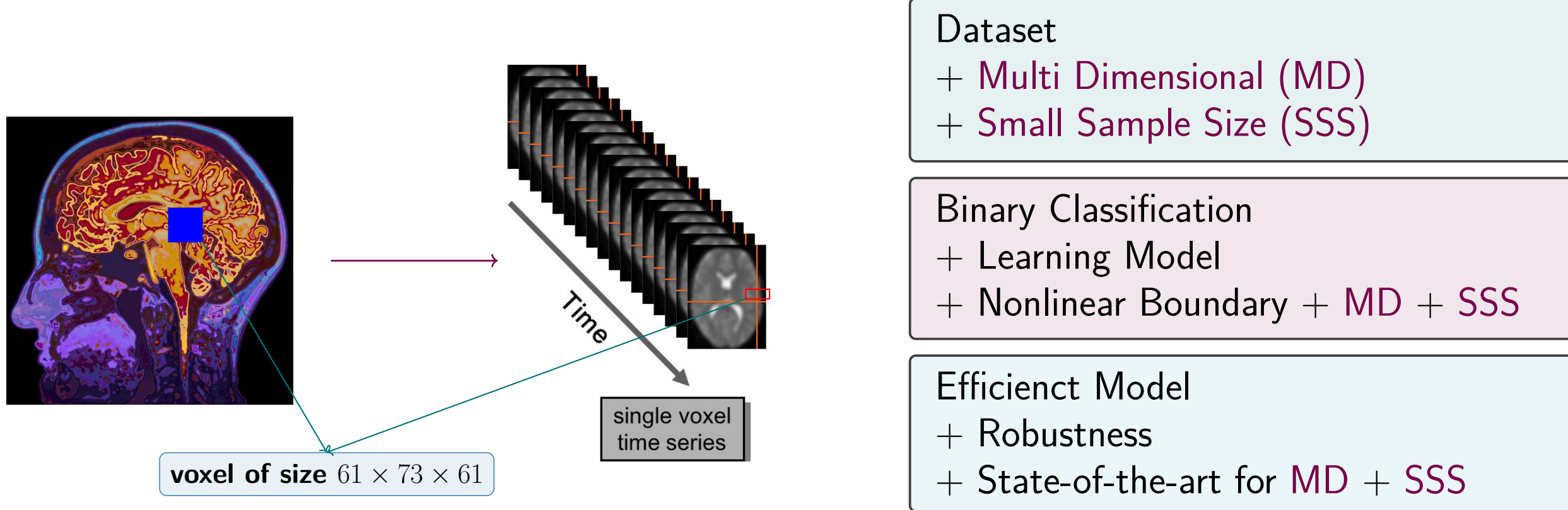
¹Max Planck Institute for Dynamics of Complex Technical Systems, Magdeburg, Germany

²University of Bath, Bath, UK

³Technische Universität, Chemnitz, Germany

Motivation

An example: classification of **ADHD** infected brain image and normal brain image:



Tensor-Train (TT) Decomposition

- A tensor is a multidimensional array.
- Matricization is unfolding of a tensor along the M^{th} -mode for an M -dimensional tensor.
- Manipulating a tensor is often prone to the curse of dimensionality $\mathcal{O}(\mathcal{I}^M)$.

$$\mathbf{x}_{i_1 i_2 \dots i_M} \cong \sum_{r_0, \dots, r_M} \mathbf{g}_{i_1, i_1, r_1}^{(1)} \mathbf{g}_{i_2, i_2, r_2}^{(2)} \dots \mathbf{g}_{i_M, i_M, r_M}^{(M)}$$

$$\mathbf{x} \cong \langle \langle \mathbf{g}^{(1)}, \mathbf{g}^{(2)}, \dots, \mathbf{g}^{(M)} \rangle \rangle,$$

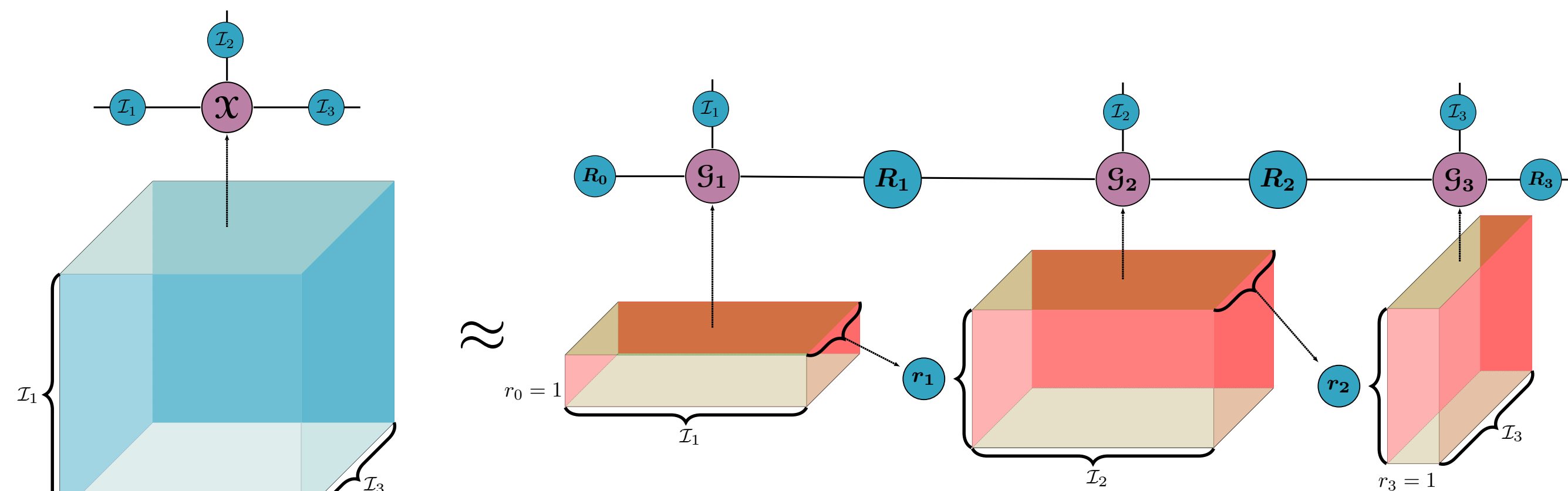


Figure 1: TT decomposition of a 3D tensor

- A tensor decomposition works efficiently avoiding the curse of dimensionality.
- Specially, TT decomposition provides a compact form with storage demands in $\mathcal{O}(MIR^2)$.
- TT decomposition adequately exploits low-rank structures.

Support Tensor Machine: Binary Classification

Model for **Small Sample Size**:

- Binary Classification $\rightarrow \mathbf{x}$ - input (all the data points), y - output $\{-1, 1\}$
- Learning Model $\rightarrow$ **Support Vector Machine (SVM)**
- Nonlinear Boundary $\rightarrow$ **Projection into higher-dimensional space**

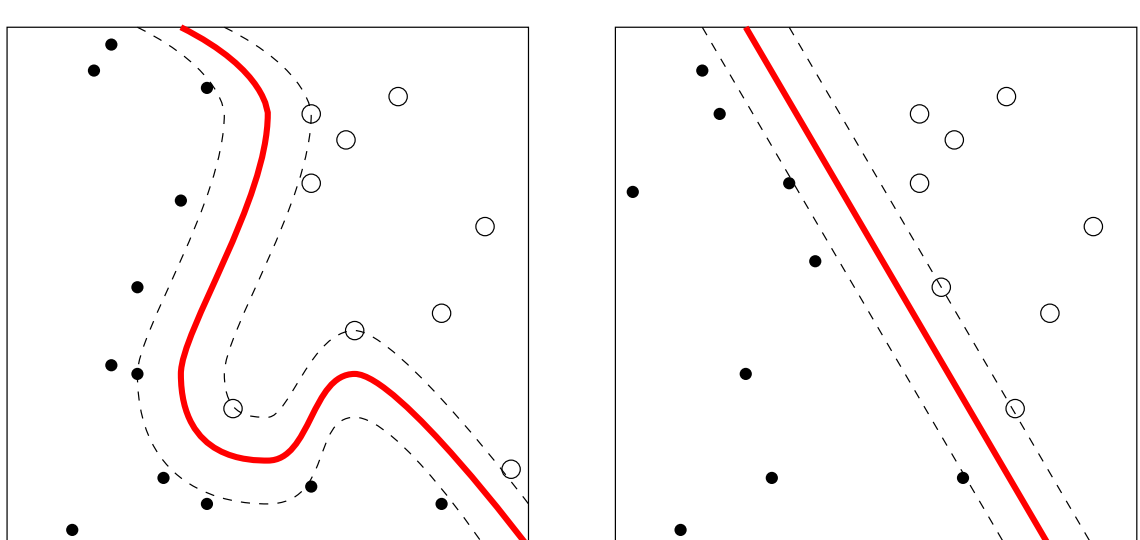


Figure 2: Non-linear(kernel) and linear SVM

Linear Boundary Optimization for SVM (Dual):

$$\max_{\alpha_1, \dots, \alpha_N} \left(\sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i,j=1}^N \alpha_i \alpha_j y_i y_j \langle \mathbf{x}_i, \mathbf{x}_j \rangle \right)$$

subject to $0 \leq \alpha_i \leq C, \sum_{i=1}^N \alpha_i y_i = 0.$

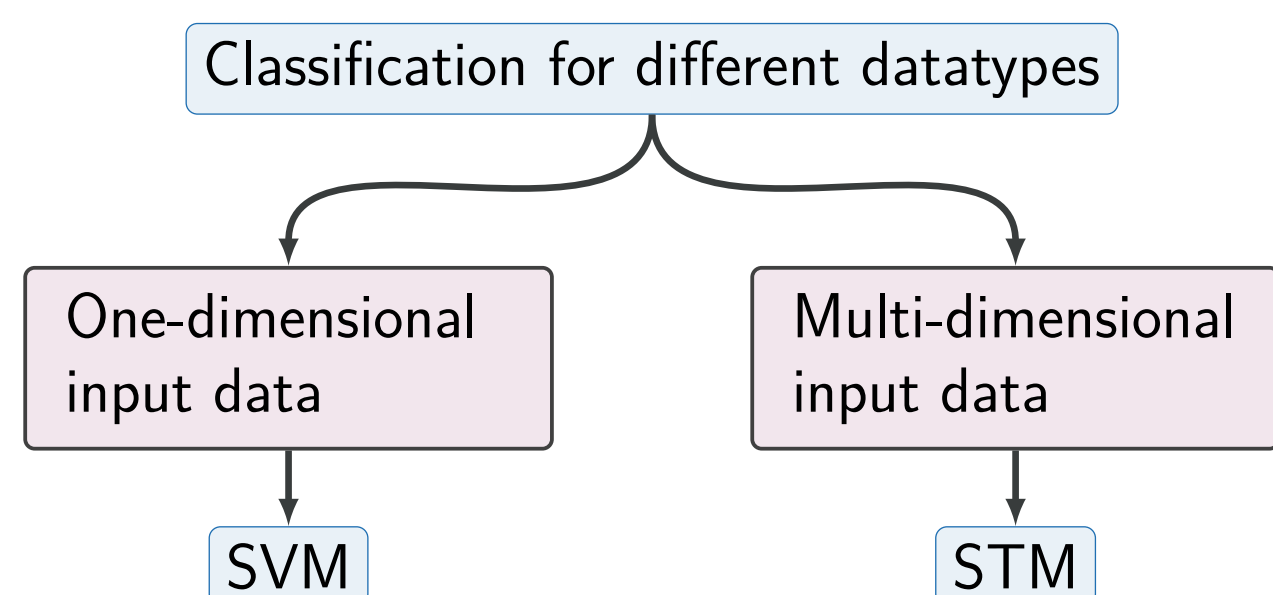


Figure 3: SVM vs. STM

Dual optimization problem corresponding to input data tensor for **Support Tensor Machine (STM)** is as follows:

$$\max_{\alpha_1, \dots, \alpha_N} \left(\sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i,j=1}^N \alpha_i \alpha_j y_i y_j \langle \mathbf{x}_i, \mathbf{x}_j \rangle \right)$$

subject to $0 \leq \alpha_i \leq C, \sum_{i=1}^N \alpha_i y_i = 0.$

Nonlinear Boundary

Projection into higher-dimensional space

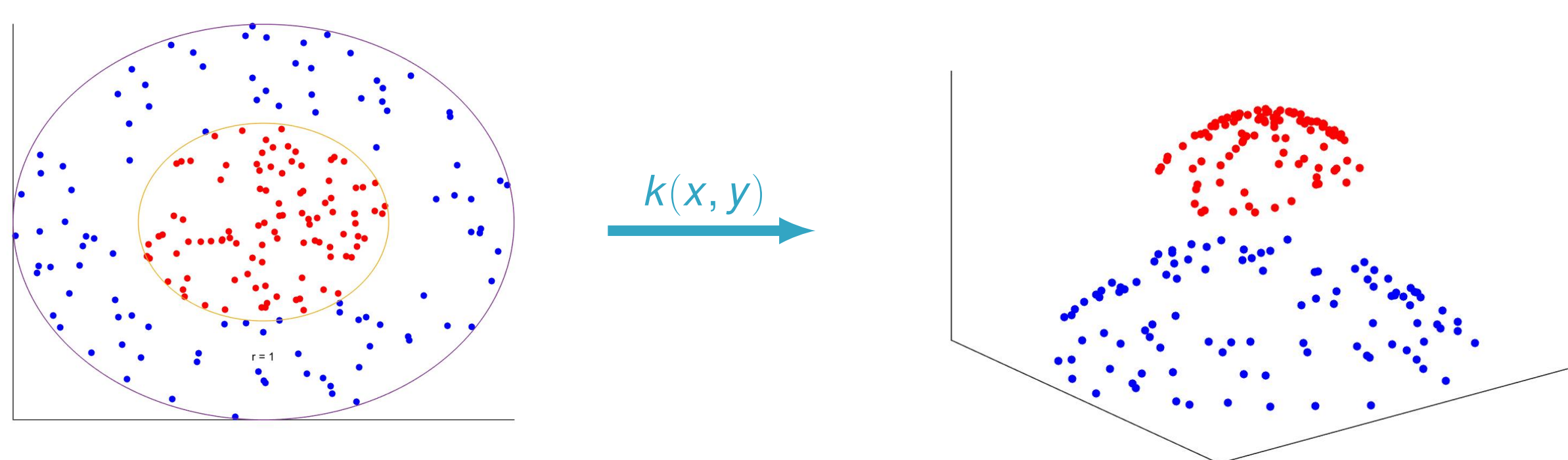


Figure: Nonlinear classification of data in ($\mathbb{R}^2$)

Figure: Linear classification in higher-dimension ($\mathbb{R}^3$)

The Feature map from lower dimensional nonlinear boundary value problem to linear high dimensional problem,

$$\phi: \mathcal{X}(\text{input}) \rightarrow \mathcal{V}(\text{feature space})$$

Using **Representation theorem**

$$\langle \phi(\mathbf{x}), \phi(\mathbf{x}') \rangle = k(\mathbf{x}, \mathbf{x}')$$

Here, k is called **kernel function**, $k: \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$. We use the **Gaussian kernel**

$$k(\mathbf{x}, \mathbf{x}') = \exp \left(- \frac{\|\mathbf{x}_i - \mathbf{x}'_i\|^2}{2\sigma^2} \right)$$

For STM $k(\mathbf{x}, \mathbf{x}') \rightarrow k(\mathbf{x}, \mathbf{x}')$, Computationally expensive term.

Multi-Way Nonlinear Mapping

A pictorial representation of a nonlinear mapping:

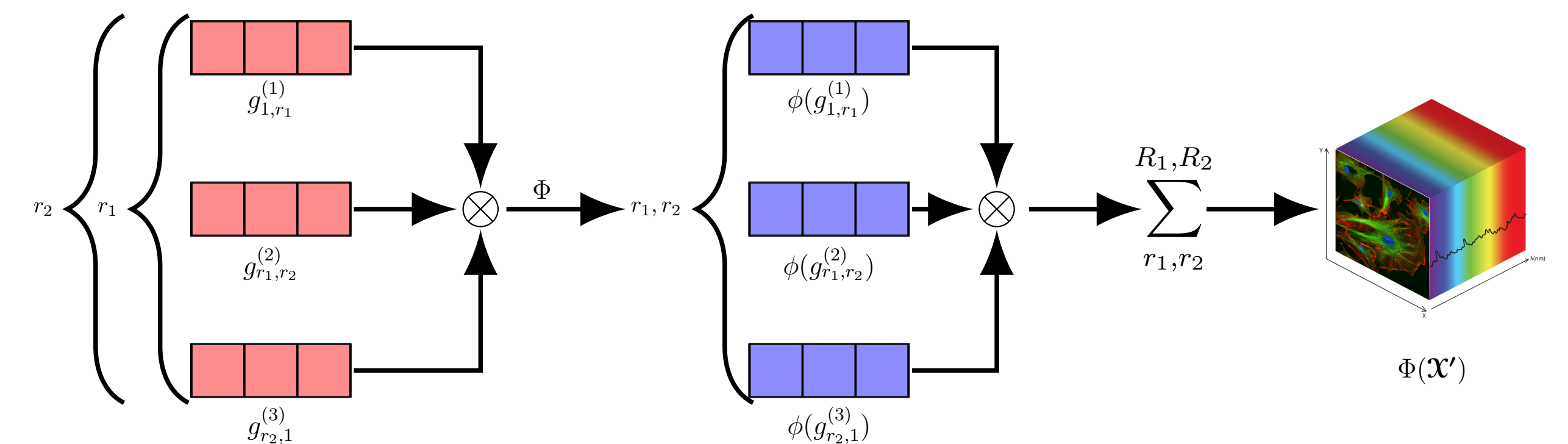


Figure 4: Multiway nonlinear mapping of TT vectors

$$\langle \Psi(\mathbf{x}), \Psi(\mathbf{y}) \rangle = k(\mathbf{x}, \mathbf{y}) = \sum_{r_1, r_2, q_1, q_2}^{R_1, R_2, Q_1, Q_2} k(g_{i_1, r_1}^{(1)}, p_{i_1, q_1}^{(1)}) k(g_{i_2, r_2}^{(2)}, p_{i_2, q_2}^{(2)}) k(g_{i_3, r_3}^{(3)}, p_{i_3, q_3}^{(3)})$$

Building Efficient Model

- Using TT decomposition
 - Stable algorithm
 - Depends on matrix SVD
 - Leads to over-fitting
 - Blocks might not be unique
- Using CP decomposition
 - Simplicity
 - Finding best rank: NP-hard

TT-CP Expansion

Given $\mathbf{x} \in \mathbb{R}^{l_1 \times l_2 \times l_3}$, $\mathbf{x}_{i_1, i_2, i_3} \approx \sum_{r_1, r_2=1}^{R_1, R_2} \mathbf{g}_{i_1, r_1}^{(1)} \mathbf{g}_{i_2, r_2}^{(2)} \mathbf{g}_{i_3, r_3}^{(3)}$

CP approximation $\mathbf{x}_{i_1, i_2, i_3} \approx \sum_{r=1}^{R_1 R_2} \mathbf{H}_{i_1, r}^{(1)} \mathbf{H}_{i_2, r}^{(2)} \mathbf{H}_{i_3, r}^{(3)}$

$r_1, r_2 \rightarrow r = r_1 + (r_2 - 1)R_1$

$\mathbf{H}_{i_1, r}^{(1)} = \mathbf{g}_{i_1, r_1}^{(1)} \mathbf{1}_{r_2}$, $\mathbf{H}_{i_3, r}^{(3)} = \mathbf{1}_{r_1} \mathbf{g}_{i_3, r_2}^{(3)}$

Stability: Uniqueness of SVD (UoSVD)

UoSVD ensures that "close" tensor produce "close" TT-cores. Also, makes model stable w.r.t. different rank truncation.

$$\mathbf{x}_{(1)} = \sigma_1 \mathbf{u}_1 \mathbf{v}_1^T + \dots + \sigma_l \mathbf{u}_l \mathbf{v}_l^T,$$

$$\bar{\mathbf{u}}_i := \mathbf{u}_i / \text{sign}(\mathbf{u}_i(i_i^*)),$$

$$\bar{\mathbf{v}}_i := \mathbf{v}_i / \text{sign}(\mathbf{v}_i(i_i^*)),$$

$$i_i^* = \arg \max_{j=1, \dots, l_i} |\mathbf{u}_i(j)|.$$

State-of-the-art: Norm Equilibration

It makes a significant increment in the classification accuracy. For a given rank- r TT-CP decomposition $\llbracket \mathbf{H}^{(1)}, \mathbf{H}^{(2)}, \mathbf{H}^{(3)} \rrbracket$,

$$\|\mathbf{N}_r\| = \|\mathbf{H}_r^{(1)}\| * \|\mathbf{H}_r^{(2)}\| * \|\mathbf{H}_r^{(3)}\|,$$

$$\mathbf{H}_r^{(k)} := \frac{\mathbf{H}_r^{(k)}}{\|\mathbf{H}_r^{(k)}\|} * \|\mathbf{N}_r\|^{1/3}, \quad k = 1, 2, 3.$$

$$\langle \Psi(\mathbf{x}), \Psi(\mathbf{y}) \rangle = k(\mathbf{x}, \mathbf{y}) = \sum_{i,j=1}^{R_1, R_2} k(h_{i_1}^{(1)}, q_j^{(1)}) k(h_{i_2}^{(2)}, q_j^{(2)}) k(h_{i_3}^{(3)}, q_j^{(3)})$$

Numerical Results

Experiment Setup:

- Using TT-Toolbox for TT decomposition,
- Using LIBSVM library for SVM.
- k-fold cross validation for splitting training and testing data

Parameter Tuning:

- TT-ranks $R \in \{1, 2, \dots, 10\}$
- Gaussian Kernel $\sigma \in \{2^{-8}, 2^{-7}, \dots, 2^7, 2^8\}$
- Trade-off parameter in SVM $C \in \{2^{-8}, 2^{-7}, \dots, 2^7, 2^8\}$

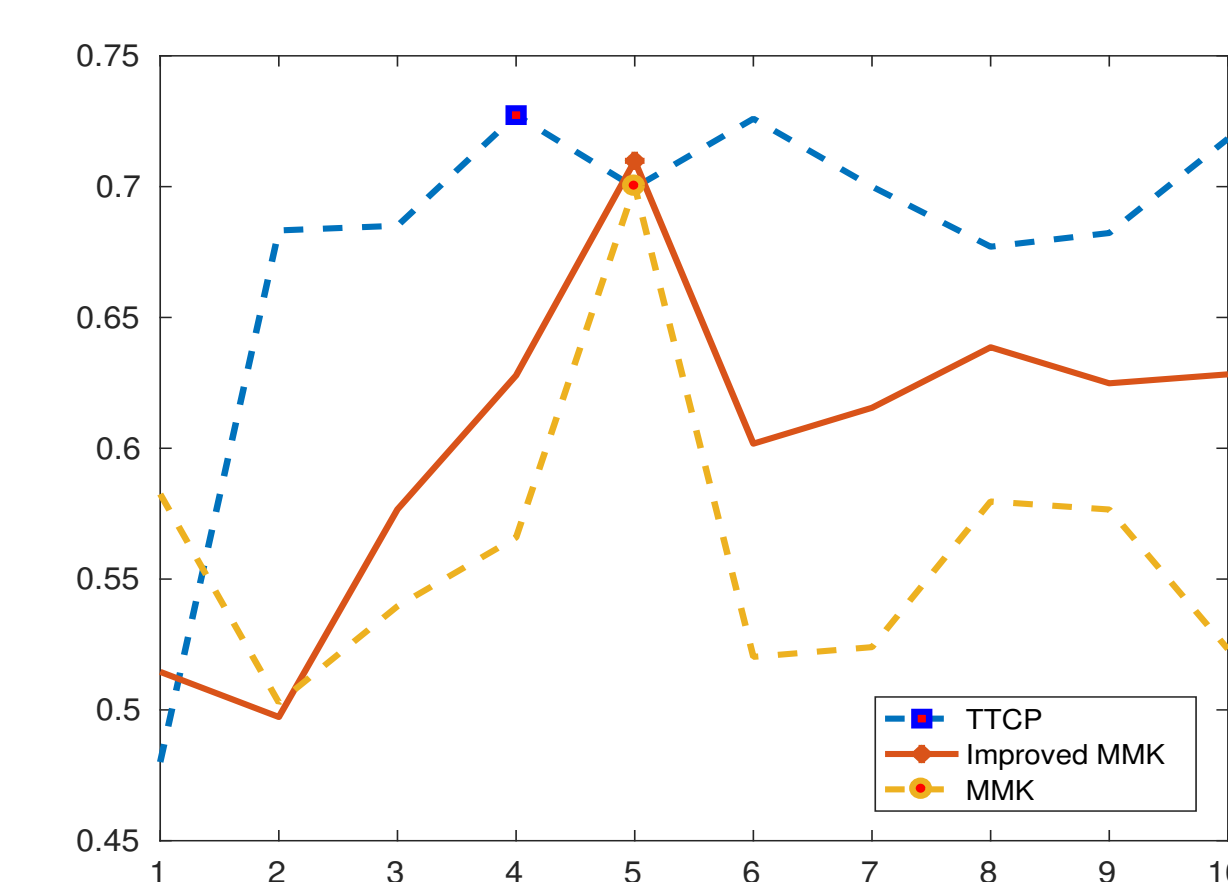


Figure 5: Accuracy for ADNI dataset w.r.t. truncated TT-CP rank

<http://adni.loni.usc.edu/>
<http://neurobureau.projects.nitrc.org/ADHD200/Data.html>

Datasets:

- Alzheimer Disease (ADNI¹)
- Attention Deficit Hyperactivity Disorder (ADHD²)

Table 1: Best average classification accuracy for different methods and datasets

Methods	ADNI	ADHD
SVM	49 %	52%
STuM	51 %	54%
DuSK	55 %	58%
MMK	69 %	60%
Improved MMK	71 %	61%
TT-MMK	73%	64%

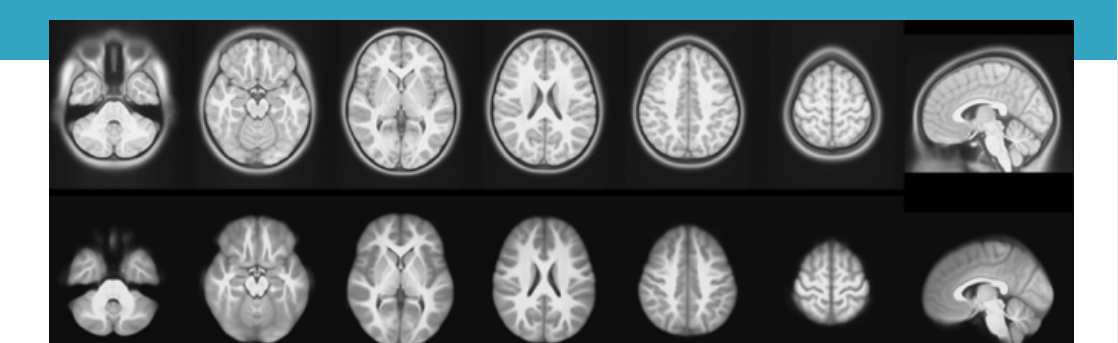
Conclusion and Outlook

Conclusion

- Low rank-method "Exact TT-CP Expansion"
- Kernel Model for less data with higher dimensions
- Approximating kernel not only reduced cost, also increases accuracy
- Stability and computational cost

Toolboxes and Codes:

- TT-Toolbox by Ivan Oseledets and Sergey Dolgov: <https://github.com/oseledets/TT-Toolbox>
- LIBSVM by Chih-J. Lin: <https://www.csie.ntu.edu.tw/~cjlin/libsvm/>
- <https://arxiv.org/abs/2002.05079>



Selected References

- [1] L. HE, X. KONG, P. S. YU, X. YANG, A. B. RAGIN, AND Z. HAO, *Dusk: A dual structure-preserving kernel for supervised tensor learning with applications to neuroimages*, in Proc. of the 2014 SIAM Internat. Conference on Data Mining, SIAM, 2014, pp. 127–135.
- [2] L. HE, C.-T. LU, H. DING, S. WANG, L. SHEN, P. S. YU, AND A. B. RAGIN, *Multi-way multi-level kernel modeling for neuroimaging classification*, in Proc. of the IEEE Conference Computer Vision and Pattern Recognition, 2017, pp. 356–364.
- [3] I. V. OSELEDETS, *Tensor-train decomposition*, SIAM J. Sci. Comput., 33 (2011), pp. 2295–2317.