



Towards Tuning-Free Minimum-Volume Nonnegative Matrix Factorization

Duc Toan Nguyen, *Texas Christian University, Fort Worth, Texas*

Advisor: Dr. Eric C. Chi, *Department of Statistics, Rice University, Houston, Texas*



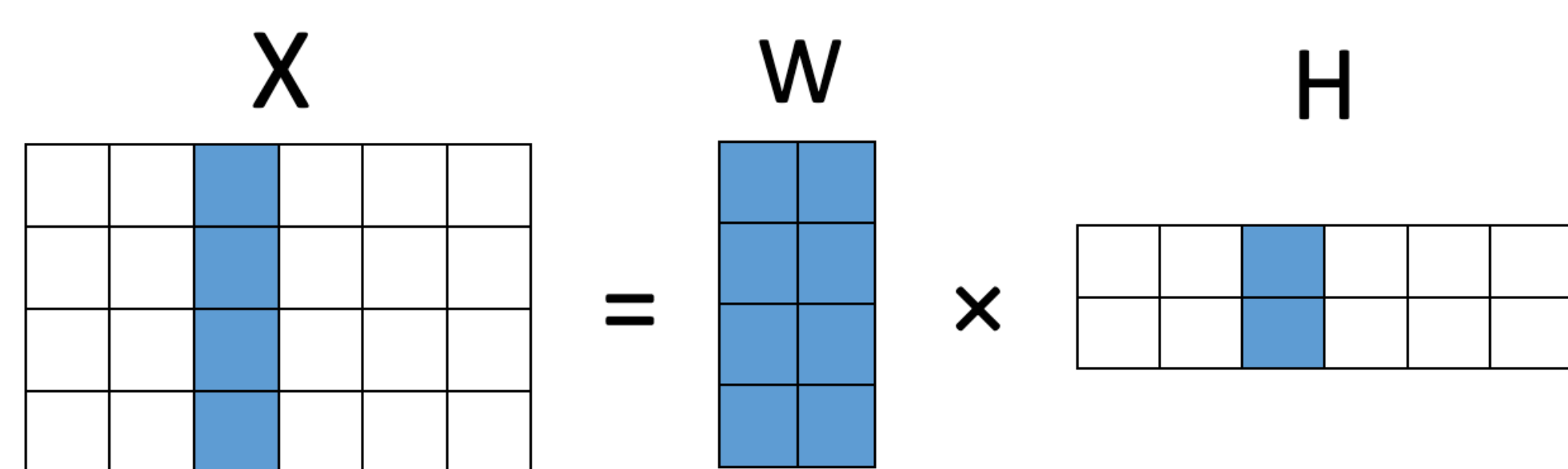
RICE

Introduction

Nonnegative matrix factorization (NMF): Given a **nonnegative** matrix $\mathbf{X} \in \mathbb{R}^{m \times n}$ (all entries are **nonnegative**) and a **target rank** r that $0 < r \ll m$. We want to find a **matrix factorization** model

$$\mathbf{X} \approx \mathbf{W}\mathbf{H},$$

where $\mathbf{W} \in \mathbb{R}^{m \times r}$ and $\mathbf{H} \in \mathbb{R}^{r \times n}$ are **nonnegative**.



Min-Vol Rank-Deficient NMF

Minimum-Volume Rank-Deficient NMF: In [3], Leplat et al. proposed the optimization problem for NMF

$$\begin{aligned} \min_{\mathbf{W}, \mathbf{H}} \quad & \|\mathbf{X} - \mathbf{W}\mathbf{H}\|_F^2 + \lambda \log \det(\mathbf{W}^T \mathbf{W} + \delta \mathbf{I}) \\ \text{s.t.} \quad & (\mathbf{W}, \mathbf{H}) \in \mathcal{S} \end{aligned}$$

Notation:

- $\mathcal{S} = \{\boldsymbol{\theta} = (\mathbf{W}, \mathbf{H}) | \mathbf{W}, \mathbf{H} \geq 0 \text{ and } \mathbf{1}^T \mathbf{H}(:, j) \leq 1, \forall j\}$ (**Constraint set**)
- $\mathbf{H}(:, j)$ is the j th column of matrix $\mathbf{H}$

Motivation

In the paper [3], Leplat et al. briefly choose the **tuning parameter** λ as

$$\lambda = \tilde{\lambda} \frac{\|\mathbf{X} - \mathbf{W}_0 \mathbf{H}_0\|_F^2}{\log \det(\mathbf{W}_0^T \mathbf{W}_0 + \delta \mathbf{I})}$$

such that

- $(\mathbf{W}_0, \mathbf{H}_0)$ is the initialization for $(\mathbf{W}, \mathbf{H})$
- $\tilde{\lambda}$ is between 1 and 10^{-3} depends on the noise level

Question: What is the **best** $\tilde{\lambda}$ for each noise level?

Experiment with λ

- True $\mathbf{W}^* = \begin{pmatrix} 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \\ 0 & 1 & 1 & 0 \\ 1 & 0 & 0 & 1 \end{pmatrix}$
- True $\mathbf{H}^*$: A stochastic 4×500 matrix generated by *Dirichlet distribution*
- True $\mathbf{X}^* = \mathbf{W}^* \mathbf{H}^*$
- Noise level $\sigma = 10^{-i}$ for $i \in \{1, \dots, 14\}$, Simulated matrix $\mathbf{X} = \mathbf{X}^* + \sigma \mathbf{N}$ ($\mathbf{N}$ is random, $n_{ij} \in [0, 1]$)
- $\text{rel-RMSE}(\mathbf{X}) = \|\mathbf{X}^* - \hat{\mathbf{X}}\|_F / \|\mathbf{X}^*\|_F$
- $\text{rel-RMSE}(\mathbf{W}) = \|\mathbf{W}^* - \hat{\mathbf{W}}\|_F / \|\mathbf{W}^*\|_F$

For each noise level σ , use multiple $\tilde{\lambda}$ between 1.5 and 10^{-11} , choose the **best results**

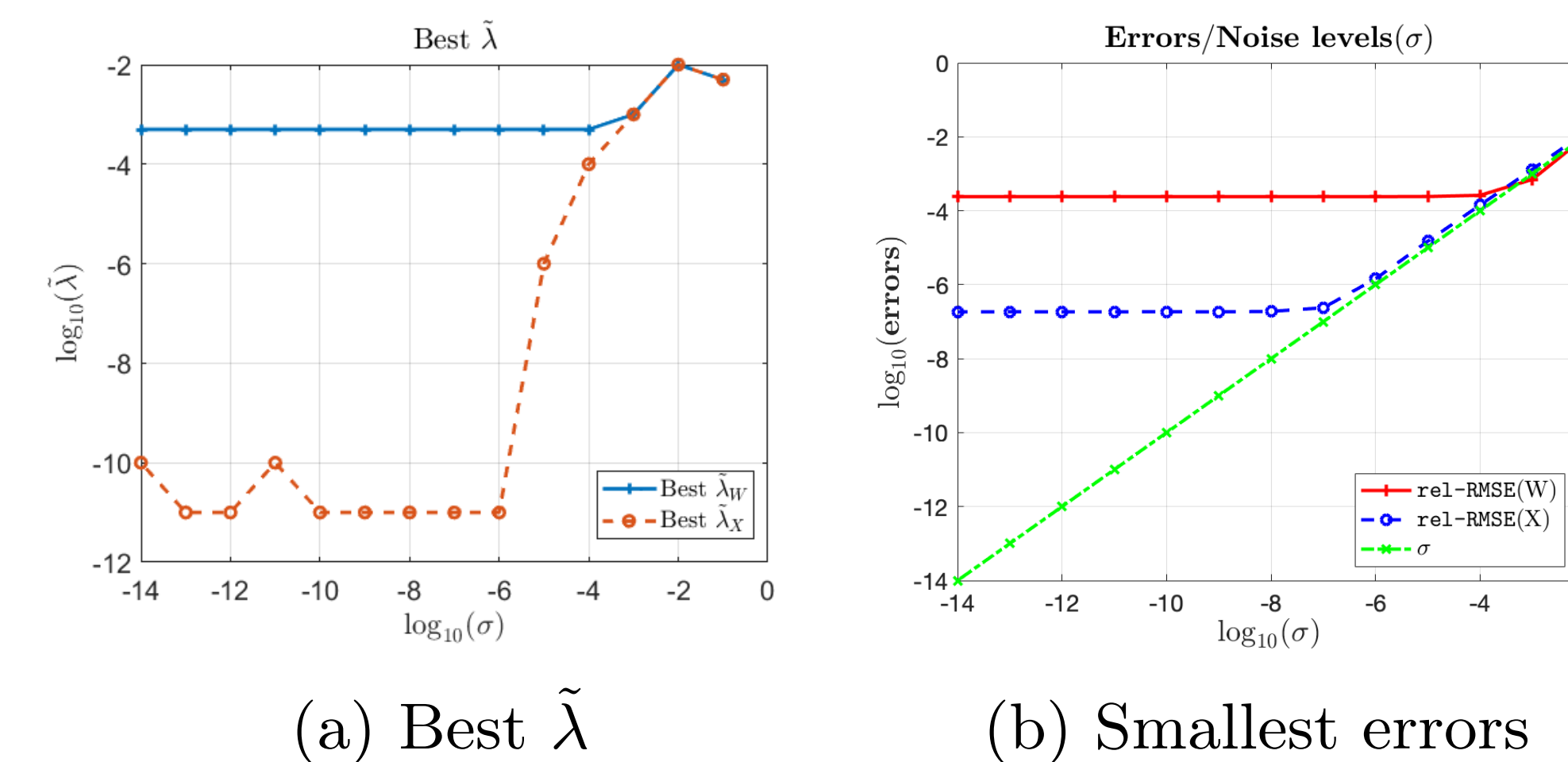


Figure 1: Best results with different noise levels

With the **best** $\tilde{\lambda}$, the errors **cannot get under** 10^{-8}

Majorization-Minimization Variant for Min-Vol

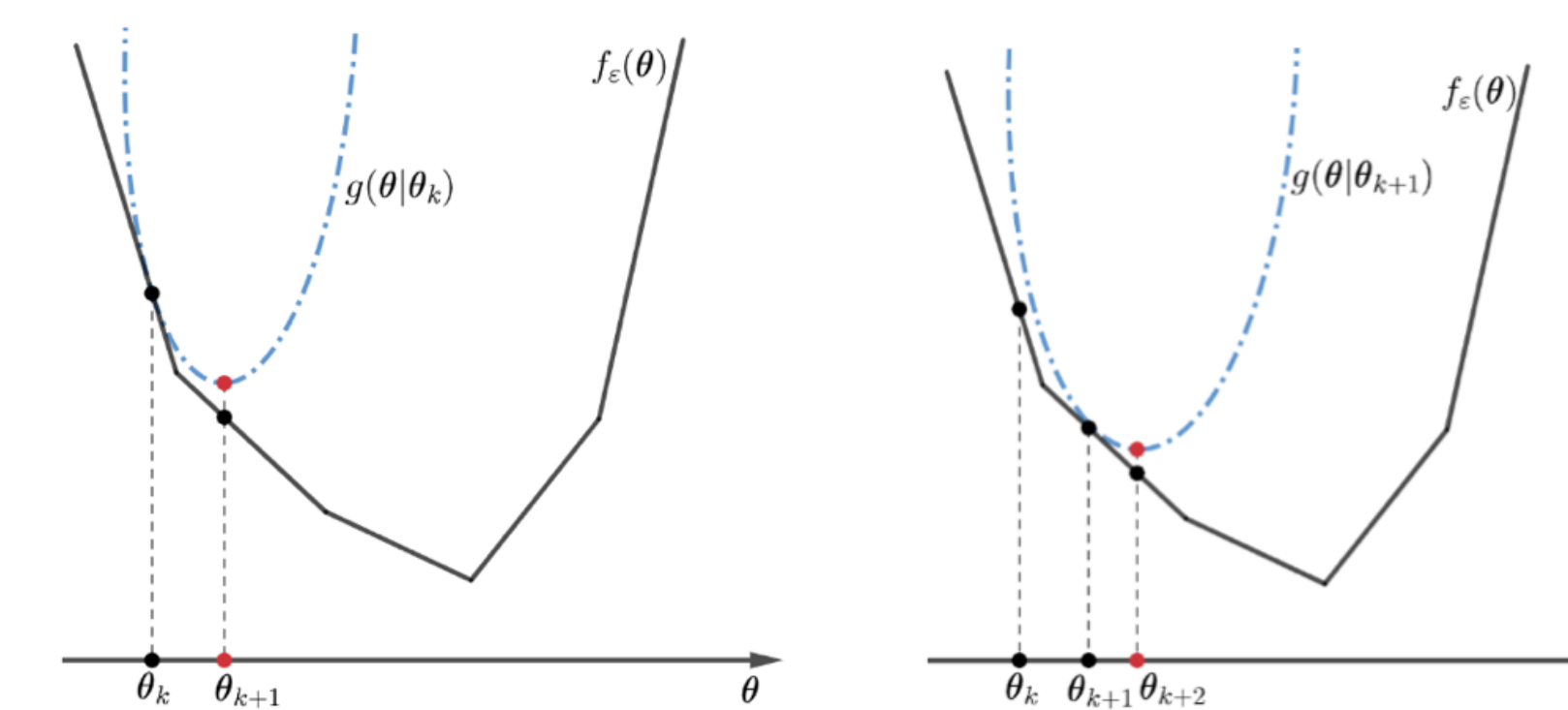
New problem (inspired by Square-Root Lasso)

$$\begin{aligned} \min_{\mathbf{W}, \mathbf{H}} \quad & f(\mathbf{W}, \mathbf{H}) := \sqrt{\|\mathbf{X} - \mathbf{W}\mathbf{H}\|_F^2} + \lambda \text{vol}(\mathbf{W}) \\ \text{s.t.} \quad & (\mathbf{W}, \mathbf{H}) \in \mathcal{S} \end{aligned}$$

Modified problem ("smoothen" square-root term)

$$\begin{aligned} \min_{\mathbf{W}, \mathbf{H}} \quad & f_\varepsilon(\mathbf{W}, \mathbf{H}) := \sqrt{\|\mathbf{X} - \mathbf{W}\mathbf{H}\|_F^2} + \varepsilon + \lambda \text{vol}(\mathbf{W}) \\ \text{s.t.} \quad & (\mathbf{W}, \mathbf{H}) \in \mathcal{S} \end{aligned}$$

where $\text{vol}(\mathbf{W}) = \log \det(\mathbf{W}^T \mathbf{W} + \delta \mathbf{I})$



Choose a surrogate function $g(\boldsymbol{\theta}|\boldsymbol{\theta}_k)$ for $f_\varepsilon(\boldsymbol{\theta})$ s.t.

$$\begin{aligned} f_\varepsilon(\boldsymbol{\theta}) &\leq g(\boldsymbol{\theta}|\boldsymbol{\theta}_k), \text{ for all } \boldsymbol{\theta} \in \mathcal{S}. \\ f_\varepsilon(\boldsymbol{\theta}_k) &= g(\boldsymbol{\theta}_k|\boldsymbol{\theta}_k) \end{aligned}$$

Results

With the same experiment set-up

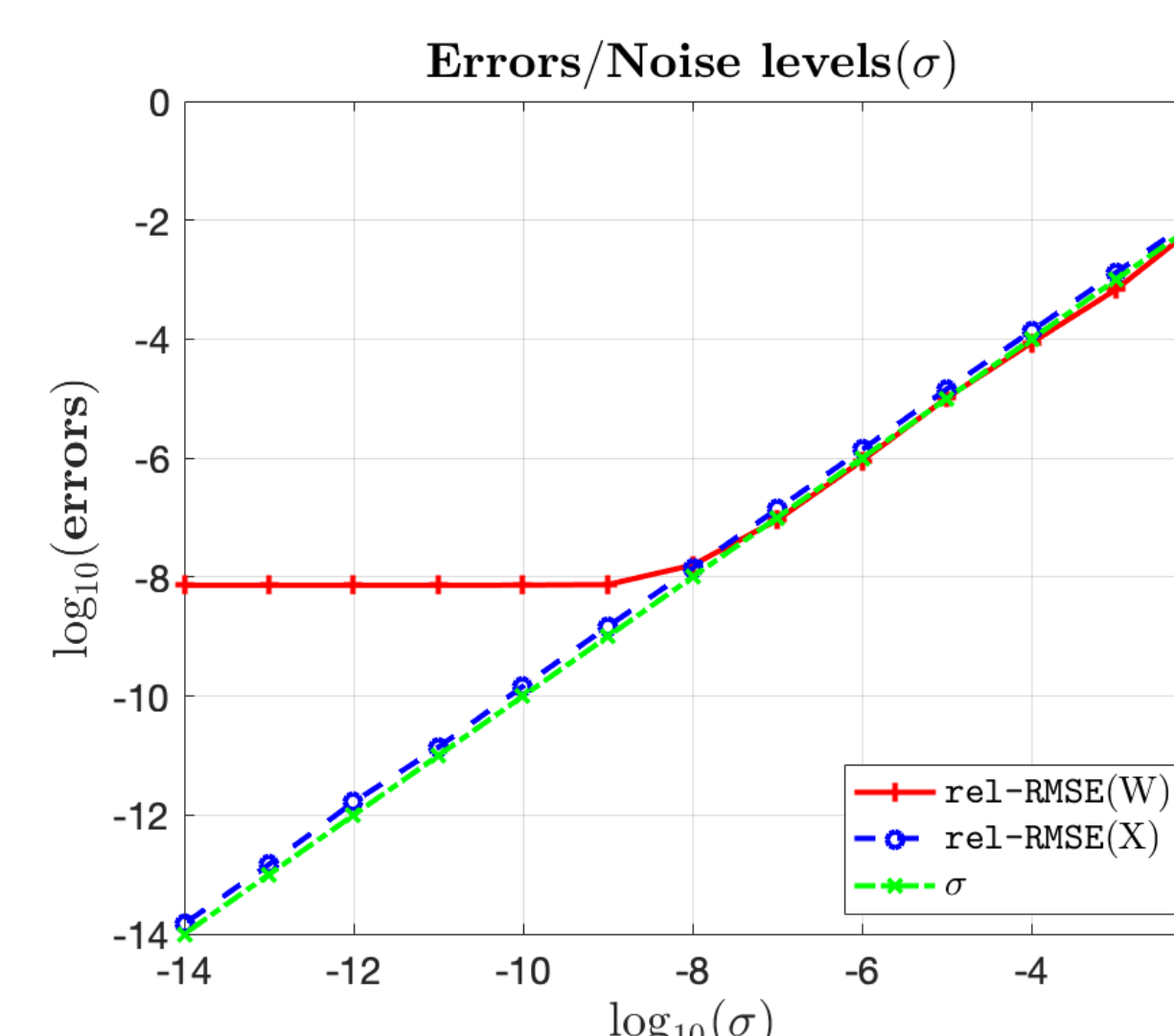


Figure 2: For **Square-Root Min-Vol**, both errors are much smaller, and $\text{rel-RMSE}(\mathbf{X}) = 10^{-14}$!

Applications

Applications: Hyperspectral unmixing (HU), topic modeling, representation learning.

Face Representation Learning:

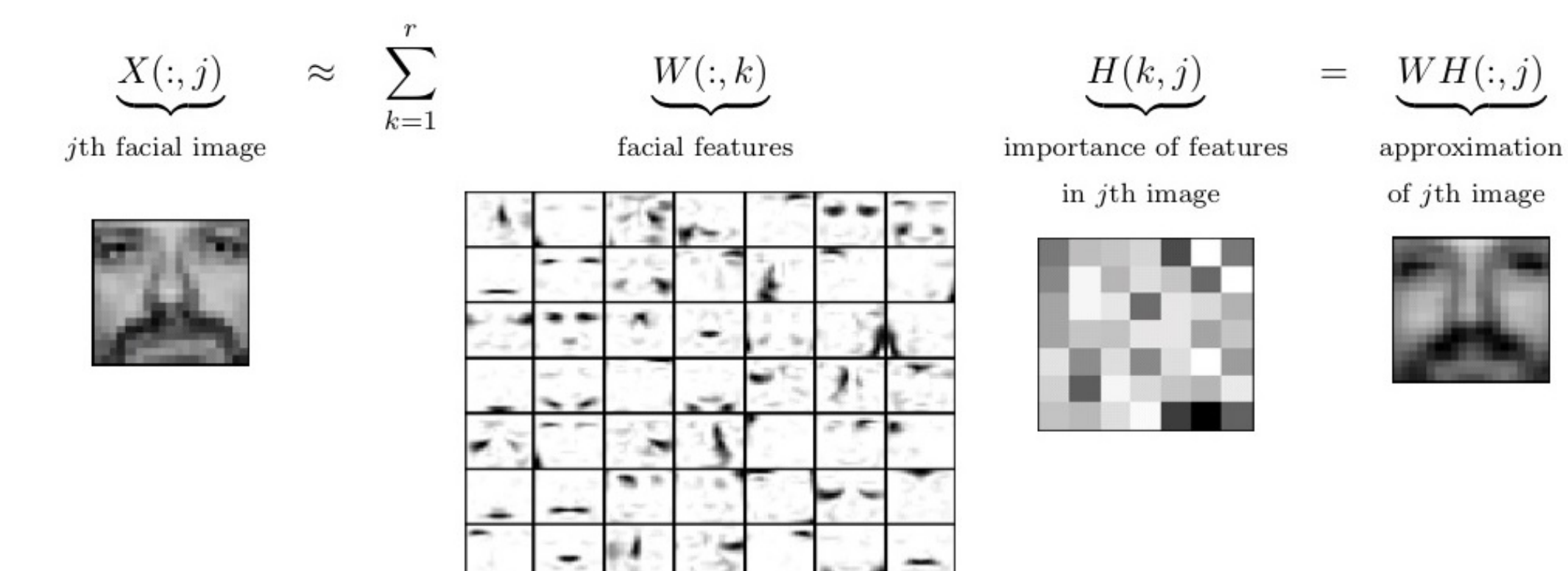


Figure 3: NMF for Representation Learning [2]. In matrix $\mathbf{X}$, each column represents a face image, while in matrix $\mathbf{W}$, each column represents a **facial feature** (learned basis).

Conclusion & Acknowledgement

Conclusion:

- The efficiency of **MinVol** algorithm depends on the initial λ choice
- The **MinVol** algorithm can be improved by the **Square-Root Min-Vol NMF**

Acknowledgement: This work was conducted in the 2023 REU STAT-DATASCI program at the Department of Statistics at Rice University and partially funded by a grant from the National Institute of General Medical Sciences (R01GM135928: EC).

References

- Nicolas Gillis. "Successive nonnegative projection algorithm for robust nonnegative blind source separation". In: *SIAM Journal on Imaging Sciences* 7.2 (2014), pp. 1420–1450.
- Nicolas Gillis. "The why and how of nonnegative matrix factorization". In: *Regularization, optimization, kernels, and support vector machines* 12.257 (2014), pp. 257–291.
- Valentin Leplat, Andersen M.S. Ang, and Nicolas Gillis. "Minimum-volume Rank-deficient Nonnegative Matrix Factorizations". In: *ICASSP 2019 - 2019 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. 2019, pp. 3402–3406.