

A085 · Stacked SVD or SVD stacked? A Random Matrix Theory perspective on data integration

Presenter: Tavor Baharav — Broad Institute

Authors: Tavor Z. Baharav, Phillip B. Nicol, Rafael A. Irizarry, Rong Ma

Session: Day 1 · Thu Oct 1, 2026 · 2:15–4:15 PM

Modern data analysis increasingly requires identifying shared latent structure across multiple high-dimensional datasets. A commonly used model assumes that the data matrices are noisy observations of low-rank matrices with a shared singular subspace. In this case, two primary methods have emerged for estimating this shared structure, which vary in how they integrate information across datasets. The first approach, termed , concatenates all the datasets, and then performs a singular value decomposition (SVD). The second approach, termed , first performs an SVD separately for each dataset, then aggregates the top singular vectors across these datasets, and finally computes a consensus amongst them. While these methods are widely used, they have not been rigorously studied in the proportional asymptotic regime, which is of great practical relevance in today's world of increasing data size and dimensionality. Consequently, it remains unclear when one method should be preferred over another. In this work, we derive exact expressions for the asymptotic performance and phase transitions of these two methods and develop optimal weighting schemes to further improve both methods. Our analysis reveals that while neither method uniformly dominates the other in the unweighted case, optimally weighted dominates optimally weighted when the low rank signal is fully shared across the datasets. We then extend our analysis to handle multiple, partially shared components per dataset and demonstrate that can yield improved performance without requiring estimation of subspace alignment. Finally, we provide practical algorithms for estimating optimal weights from data, offering theoretical guidance for method selection in practical data integration problems. Extensive numerical simulations and semi-synthetic experiments on genomic data corroborate our theoretical findings.

STACKED SVD OR SVD STACKED? A RANDOM MATRIX THEORY PERSPECTIVE ON DATA INTEGRATION

BY TAVOR Z. BAHARAV^{1,2,*[a](#)}, PHILLIP B. NICOL^{2,3,*[c](#)}, RAFAEL A. IRIZARRY^{2,3,[b](#)} AND RONG MA^{1,2,3,[†](#)[d](#)}

¹ERIC AND WENDY SCHMIDT CENTER, BROAD INSTITUTE, CAMBRIDGE, MA, 02142, USA,
^aBAHARAV@BROADINSTITUTE.ORG

²DEPARTMENT OF DATA SCIENCE, DANA FARBER CANCER INSTITUTE, BOSTON, MA, 02115, USA,
^bRAFAEL_IRIZARRY@DFCI.HARVARD.EDU

³DEPARTMENT OF BIostatistics, HARVARD T.H. CHAN SCHOOL OF PUBLIC HEALTH, BOSTON, MA, 02115, USA, ^cPHILLIPNICOL@G.HARVARD.EDU; ^dRONGMA@HSPH.HARVARD.EDU

Modern data analysis increasingly requires identifying shared latent structure across multiple high-dimensional datasets. A commonly used model assumes that the data matrices are noisy observations of low-rank matrices with a shared singular subspace. In this case, two primary methods have emerged for estimating this shared structure, which vary in how they integrate information across datasets. The first approach, termed **Stack-SVD**, concatenates all the datasets, and then performs a singular value decomposition (SVD). The second approach, termed **SVD-Stack**, first performs an SVD separately for each dataset, then aggregates the top singular vectors across these datasets, and finally computes a consensus amongst them. While these methods are widely used, they have not been rigorously studied in the proportional asymptotic regime, which is of great practical relevance in today's world of increasing data size and dimensionality. Consequently, it remains unclear when one method should be preferred over another. In this work, we derive exact expressions for the asymptotic performance and phase transitions of these two methods and develop optimal weighting schemes to further improve both methods. Our analysis reveals that while neither method uniformly dominates the other in the unweighted case, optimally weighted **Stack-SVD** dominates optimally weighted **SVD-Stack** when the low rank signal is fully shared across the datasets. We then extend our analysis to handle multiple, partially shared components per dataset and demonstrate that **SVD-Stack** can yield improved performance without requiring estimation of subspace alignment. Finally, we provide practical algorithms for estimating optimal weights from data, offering theoretical guidance for method selection in practical data integration problems. Extensive numerical simulations and semi-synthetic experiments on genomic data corroborate our theoretical findings.

*Equal contributions, listed alphabetically.

[†]Corresponding author.

MSC2020 subject classifications: Primary 15A18, 62H25.

Keywords and phrases: Data integration, Random matrix theory, Spectral methods.