

Globally Aligned KDSUM: A Positive Semidefinite Mixed-Type Kernel for Low-Dimensional Embedding

Kejun Fang
McGill University

May 2026

Abstract

Mixed-type data combine continuous, unordered categorical, and ordered categorical variables, making low-dimensional representation depend on a pairwise geometry that can compare heterogeneous observations and remain valid for spectral or latent-variable embedding. Many mixed-type distances are useful for comparison or clustering, but their double-centered matrices need not be positive semidefinite, limiting their use in principal coordinate analysis, kernel principal component analysis, and fixed-kernel probabilistic principal coordinate analysis. We show that KDSUM, built from Gaussian, Aitken, and Wang–van Ryzin component kernels, defines a positive semidefinite mixed-type sample kernel. Its centered form therefore induces squared Hilbert-space distances and gives an admissible input to these embedding methods. We further propose globally aligned KDSUM (GA-KDSUM), a trace-balanced weighting scheme that favors component kernels aligned with the consensus geometry of the remaining variables while preserving positive semidefiniteness through nonnegative kernel combinations. In latent-factor simulations, GA-KDSUM gives its strongest gains in the setting with many irrelevant variables, outperforming EW-KDSUM and Gower-type PCO baselines, including PSD-corrected Gower PCO. Diagnostics confirm that KDSUM-based kernels are numerically PSD whereas Gower-type double-centered matrices are often indefinite. On two Student Performance prediction tasks, GA-KDSUM achieves the highest mean balanced accuracy and AUC among the compared embeddings based on mixed-type pairwise geometry. The method is therefore best viewed as a PSD kernel interface and embedding strategy for heterogeneous data with shared low-dimensional structure.

Keywords: mixed-type data, positive semidefinite kernels, KDSUM, principal coordinate analysis, kernel alignment, probabilistic embedding

1 Introduction

Modern applications often begin with heterogeneous records rather than purely numerical feature matrices. Electronic health records may contain laboratory measurements, diagnosis and procedure codes, medication indicators, and ordered clinical scores. Educational datasets often combine grades, demographic variables, school attributes, and ordinal behavioral responses. Credit and customer analytics similarly combine continuous financial measurements with unordered and ordered categorical descriptors. In these settings, a low-dimensional representation is useful not only for visualization but also for downstream classification, regression, cohort discovery, segmentation, and model diagnostics (Miotto et al., 2016, Cortez & Silva, 2008, Yeh & Lien, 2009).

The central obstacle is that an embedding method needs a geometry before it can be applied. PCA works with a numerical matrix and is tied to best low-rank approximation (Eckart & Young, 1936, Jolliffe, 2002). Principal coordinate analysis and classical multidimensional scaling start from pairwise dissimilarities, but the dissimilarities must be Euclidean for the centered inner product matrix to be positive semidefinite (Gower, 1966, Borg & Groenen, 1997). Kernel PCA replaces explicit coordinates by a positive semidefinite kernel matrix (Schölkopf et al., 1998, Williams, 2002). Probabilistic principal coordinate analysis gives PCO a latent-variable and maximum-likelihood interpretation, but it is also conditional on a fixed positive semidefinite matrix $\mathbf{Q}$ (Zhang & Jordan, 2009). Thus, for mixed-type data, the problem is not only how to compare heterogeneous observations, but how to build a comparison that is admissible as a kernel geometry.

Several approaches are commonly used to compare or model mixed-type observations. Dummy coding embeds categorical variables as indicator vectors, but it treats ordered categories as nominal unless extra structure is imposed and it can let category cardinality affect the geometry (Foss et al., 2019). Gower’s coefficient averages range-scaled numerical similarities and matching-based categorical similarities, making it a standard mixed-type baseline (Gower, 1971). Ordinal extensions of Gower replace exact matching on ordered variables by rank-aware comparisons (Podani, 1999). The k -prototypes objective combines squared Euclidean loss for numerical variables with matching loss for categorical variables (Huang, 1998). Semiparametric and latent-variable mixture methods, such as KAMILA and clustMD, model mixed-type clusters through distributional assumptions rather than through a single embedding kernel (Foss et al., 2016, McParland & Gormley, 2016). KDSUM takes a kernel-based route by combining Gaussian, Aitken, and Wang–van Ryzin component kernels into a mixed-type similarity and using bandwidth selection to control the contribution of each variable (Ghashti & Thompson, 2025).

These methods address important parts of the mixed-type problem, but they leave a gap for probabilistic and spectral embedding. A useful mixed-type distance or clustering objective is not automatically a valid positive semidefinite kernel. An indefinite similarity or double-centered distance matrix breaks the Hilbert-space interpretation required by kernel PCA, PCO, and fixed- $\mathbf{Q}$ PPCO. Conversely, a positive semidefinite mixed-type kernel immediately induces squared Hilbert-space distances, a centered PCO matrix, and a probabilistic embedding model. This paper focuses on that interface.

We show that the standard KDSUM similarity constructed from Gaussian kernels for continuous variables, Aitken kernels for unordered categorical variables, and Wang–van Ryzin kernels for ordered categorical variables is positive semidefinite. After centering, the resulting matrix

$$\mathbf{Q} = \mathbf{H}\Psi_{\text{KDSUM}}\mathbf{H}$$

is therefore a valid input to PCO, KPCA, and fixed- $\mathbf{Q}$ PPCO. The method is conceptually simple. It builds type-specific positive semidefinite component kernels, combines them with nonnegative weights, centers the resulting KDSUM kernel, and applies the PPCO embedding. Since PPCO treats $\mathbf{Q}$ as fixed, bandwidths and variable weights are chosen before PPCO is applied rather than by optimizing the PPCO likelihood.

We also introduce GA-KDSUM, a globally aligned weighted KDSUM construction for settings where several heterogeneous variables share a low-dimensional geometry but many additional variables are weakly related or irrelevant. Each component kernel is trace-balanced after centering, and its weight is determined by how strongly its centered geometry aligns with the average centered geometry of the other components. The idea is related to centered kernel alignment, which measures agreement between kernel matrices and has been used for learning convex combinations of kernels (Cortes et al.,

2012). The resulting GA-KDSUM kernel remains positive semidefinite because it is a nonnegative combination of positive semidefinite component kernels.

The proposed embedding is intended for settings where mixed-type observations need to be represented by low-dimensional coordinates. These settings include patient or student representation learning, survey and behavioral data visualization, credit approval and risk screening, customer segmentation, and preprocessing for downstream supervised models. Sections 5 and 6 evaluate this role using synthetic latent factor recovery and held-out downstream prediction. This paper makes three specific contributions:

1. It proves that the standard KDSUM similarity is a positive semidefinite mixed-type kernel under the Gaussian, Aitken, and Wang–van Ryzin construction.
2. It connects KDSUM to PCO, KPCA, and fixed- $\mathbf{Q}$ PPCO through the centered kernel matrix $\mathbf{Q} = \mathbf{H}\Psi\mathbf{H}$.
3. It proposes GA-KDSUM, a trace-balanced globally aligned KDSUM embedding that preserves positive semidefiniteness and improves mixed-type pairwise geometry in the setting with many irrelevant variables and in real downstream prediction.

Organization of the paper. Section 2 introduces the mixed-type KDSUM geometry and reviews its connection to PCO, KPCA, and PPCO. Section 3 presents the PSD result, the trace-balanced globally aligned kernel construction, and the embedding algorithm. Section 4 defines the evaluation criteria used in the experiments. Section 5 reports the simulation experiment. Section 6 presents downstream prediction experiments on mixed-type Student Performance datasets. Section 7 discusses interpretation, limitations, and future work, and Section 8 concludes. The appendix gives proof details and supplementary diagnostics.

2 Mixed-Type Kernel Geometry

2.1 Mixed-type data and KDSUM

Consider an $n \times p$ mixed-type data matrix $\mathbf{X}$ consisting of n observations and p variables. Assume that the variables are arranged so that the first p_c variables are continuous, the next p_u variables are unordered categorical, and the last p_o variables are ordered categorical, with

$$p = p_c + p_u + p_o.$$

The rows of $\mathbf{X}$ are observation vectors $\mathbf{x}_i$. Let $\mathbf{1}_n$ denote the n -vector of ones, and let

$$\boldsymbol{\lambda} = \{\boldsymbol{\lambda}^c, \boldsymbol{\lambda}^u, \boldsymbol{\lambda}^o\}$$

denote the collection of variable-specific bandwidths.

For a continuous variable, let

$$k(z) = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{z^2}{2}\right),$$

so that the continuous KDSUM contribution for variable k is

$$\frac{1}{\lambda_k} k \left(\frac{x_{ik} - x_{jk}}{\lambda_k} \right), \quad \lambda_k > 0.$$

For an unordered categorical variable, we use the Aitken kernel ([Aitchison & Aitken, 1976](#)),

$$L(a, b, \lambda_k) = \begin{cases} 1, & a = b, \\ \lambda_k, & a \neq b, \end{cases} \quad \lambda_k \in [0, 1].$$

For ordered categorical variables, levels are assumed to be encoded by their integer ranks. We use the Wang and van Ryzin kernel ([Wang & Van Ryzin, 1981](#)),

$$l(a, b, \lambda_k) = \begin{cases} 1 - \lambda_k, & a = b, \\ \frac{1}{2}(1 - \lambda_k)\lambda_k^{|a-b|}, & a \neq b, \end{cases} \quad \lambda_k \in [0, 1].$$

The pairwise KDSUM similarity is

$$\psi(\mathbf{x}_i, \mathbf{x}_j \mid \boldsymbol{\lambda}) = \prod_{k=1}^{p_c} \frac{1}{\lambda_k} k \left(\frac{x_{ik} - x_{jk}}{\lambda_k} \right) + \sum_{k=p_c+1}^{p_c+p_u} L(x_{ik}, x_{jk}, \lambda_k) + \sum_{k=p_c+p_u+1}^p l(x_{ik}, x_{jk}, \lambda_k).$$

When $p_c = 0$, the continuous product is interpreted as the empty product, namely the constant kernel equal to one. At the sample level this contributes the rank-one term $\mathbf{1}_n \mathbf{1}_n^\top$ to the raw similarity. Centering removes this term and it contributes zero to the induced squared distances.

The corresponding KDSUM similarity matrix is denoted by

$$\boldsymbol{\Psi}_{\text{KDSUM}}(\boldsymbol{\lambda}) = [\psi(\mathbf{x}_i, \mathbf{x}_j \mid \boldsymbol{\lambda})]_{i,j=1}^n.$$

The induced KDSUM dissimilarity is

$$d(\mathbf{x}_i, \mathbf{x}_j \mid \boldsymbol{\lambda}) = \psi(\mathbf{x}_i, \mathbf{x}_i \mid \boldsymbol{\lambda}) + \psi(\mathbf{x}_j, \mathbf{x}_j \mid \boldsymbol{\lambda}) - 2\psi(\mathbf{x}_i, \mathbf{x}_j \mid \boldsymbol{\lambda}).$$

Although we keep the original KDSUM notation d , in the PCO construction this quantity is used as a squared dissimilarity. Thus we write

$$\delta_{ij}^2 = d(\mathbf{x}_i, \mathbf{x}_j \mid \boldsymbol{\lambda}).$$

When $\boldsymbol{\Psi}_{\text{KDSUM}}$ is PSD, δ_{ij}^2 is a squared Hilbert-space distance, and δ_{ij} is the associated Hilbert-space distance.

The original KDSUM paper uses this distance mainly for clustering and uses maximum similarity cross-validation (MSCV) to select $\boldsymbol{\lambda}$ ([Ghashti & Thompson, 2025](#)). In this paper, MSCV is treated as an external KDSUM bandwidth selection rule. The experiments do not use the PPCO likelihood to tune $\boldsymbol{\lambda}$. The empirical implementation uses a trace-balanced weighted component version of the same PSD construction so that individual variables can be weighted. Accordingly, EW-KDSUM and GA-KDSUM are componentized PSD KDSUM extensions rather than algebraically identical copies of the raw KDSUM aggregation.

2.2 PCO, kernels, and PPCO

Given a squared dissimilarity matrix

$$\mathbf{\Delta} = [\delta_{ij}^2],$$

classical PCO forms

$$\mathbf{Q} = -\frac{1}{2}\mathbf{H}\mathbf{\Delta}\mathbf{H}, \quad \mathbf{H} = \mathbf{I}_n - \frac{1}{n}\mathbf{1}_n\mathbf{1}_n^\top.$$

If $\mathbf{Q} \succeq 0$, the dissimilarities are Euclidean. If a positive semidefinite kernel matrix $\mathbf{K}$ is available directly, the corresponding centered inner product matrix is

$$\mathbf{Q} = \mathbf{H}\mathbf{K}\mathbf{H}.$$

This connects PCO, kernel PCA, and metric multidimensional scaling (Gower, 1966, Williams, 2002, Borg & Groenen, 1997).

Zhang and Jordan’s probabilistic PCO (PPCO) model gives PCO a maximum-likelihood interpretation (Zhang & Jordan, 2009). Their model treats the high-dimensional feature vectors as virtual observations and estimates the low-dimensional coordinates conditional on a fixed centered PSD matrix $\mathbf{Q}$. Below, the PPCO noise variance is denoted by σ_{PPCO}^2 to distinguish it from the KDSUM bandwidth vector $\boldsymbol{\lambda}$.

Let

$$\mathbf{Q} = \mathbf{U}\mathbf{T}\mathbf{U}^\top, \quad \gamma_1 \geq \dots \geq \gamma_{n-1} \geq 0,$$

where $\gamma_1, \dots, \gamma_{n-1}$ are the eigenvalues of the centered matrix $\mathbf{Q}$ excluding the trivial eigenvalue associated with $\mathbf{1}_n$, with zero eigenvalues included after the nonzero rank. Let $d_Q = \text{rank}(\mathbf{Q})$. For a chosen embedding dimension $1 \leq q < n - 1$ with $q \leq d_Q$, the fixed- $\mathbf{Q}$ PPCO estimate is

$$\widehat{\mathbf{Y}} = \mathbf{U}_q(\mathbf{T}_q - \hat{\sigma}_{\text{PPCO}}^2 \mathbf{I}_q)^{1/2} \mathbf{S},$$

where $\mathbf{S}$ is an arbitrary $q \times q$ orthonormal matrix and

$$\hat{\sigma}_{\text{PPCO}}^2 = \frac{1}{n - q - 1} \sum_{j=q+1}^{n-1} \gamma_j.$$

Taking $\mathbf{S} = \mathbf{I}_q$ fixes the rotation. When $\hat{\sigma}_{\text{PPCO}}^2$ is small, PPCO is nearly the same as classical PCO on the same matrix $\mathbf{Q}$. Consequently, any downstream advantage should be interpreted as coming from the mixed-type kernel geometry rather than from PPCO alone.

3 KDSUM to PPCO Embedding

All PSD statements in this section are finite-sample kernel-matrix statements. The bandwidths, trace normalizers, and global-alignment weights are fitted from the training sample and then held fixed for the training kernel matrix and for out-of-sample projection.

3.1 PSD result for the KDSUM similarity

Proposition 1. Let Ψ_{KDSUM} be the sample KDSUM similarity matrix defined in Section 2.1, using Gaussian kernels for continuous variables, Aitken kernels for unordered categorical variables with bandwidths in $[0, 1]$, and Wang and van Ryzin kernels for ordered categorical variables with bandwidths in $[0, 1]$. Then

$$\Psi_{\text{KDSUM}} \succeq 0.$$

Consequently,

$$\mathbf{Q} = \mathbf{H}\Psi_{\text{KDSUM}}\mathbf{H} \succeq 0.$$

Proof. Each continuous Gaussian component matrix is PSD. The continuous KDSUM block is the Hadamard product of these component matrices, so it is PSD by the Schur product theorem (Horn & Johnson, 2012). If $p_c = 0$, the empty product convention gives the constant kernel $\mathbf{1}_n\mathbf{1}_n^\top$, which is PSD. Omitting this block changes the raw similarity by only a constant rank-one term. This term is removed by centering and contributes zero to the induced squared distances.

For an unordered categorical variable with one-hot category matrix $\mathbf{C}$, the Aitken component matrix can be written as

$$\mathbf{K}^{(u)} = \lambda \mathbf{1}_n \mathbf{1}_n^\top + (1 - \lambda) \mathbf{C} \mathbf{C}^\top, \quad \lambda \in [0, 1].$$

This is a nonnegative linear combination of PSD matrices, so $\mathbf{K}^{(u)} \succeq 0$.

For an ordered categorical variable with ordered levels $1, \dots, G$, let $\mathbf{C}$ be the one-hot level matrix and define

$$(\mathbf{T}_\lambda)_{rs} = \lambda^{|r-s|}, \quad r, s \in \{1, \dots, G\}.$$

For $\lambda \in [0, 1]$, $\mathbf{T}_\lambda \succeq 0$. Hence $\mathbf{A} = \mathbf{C} \mathbf{T}_\lambda \mathbf{C}^\top \succeq 0$. The Wang and van Ryzin component matrix can be written as

$$\mathbf{K}^{(o)} = \frac{1 - \lambda}{2} \mathbf{A} + \frac{1 + \lambda}{2} \mathbf{C} \mathbf{C}^\top.$$

It follows that $\mathbf{K}^{(o)} \succeq 0$, since it is a nonnegative linear combination of PSD matrices.

The KDSUM similarity matrix is the sum of the continuous block, the unordered categorical components, and the ordered categorical components. Since each summand is PSD, $\Psi_{\text{KDSUM}} \succeq 0$. Finally, for any $\mathbf{v} \in \mathbb{R}^n$,

$$\mathbf{v}^\top \mathbf{H} \Psi_{\text{KDSUM}} \mathbf{H} \mathbf{v} = (\mathbf{H} \mathbf{v})^\top \Psi_{\text{KDSUM}} (\mathbf{H} \mathbf{v}) \geq 0.$$

Thus, $\mathbf{Q} = \mathbf{H} \Psi_{\text{KDSUM}} \mathbf{H} \succeq 0$.

The standard PSD facts used above are recalled in Appendix A.1. □

Corollary 1. Define

$$\delta_{ij}^2 = (\Psi_{\text{KDSUM}})_{ii} + (\Psi_{\text{KDSUM}})_{jj} - 2(\Psi_{\text{KDSUM}})_{ij}, \quad \Delta = [\delta_{ij}^2].$$

Then δ_{ij}^2 is a squared Hilbert space distance and

$$-\frac{1}{2} \mathbf{H} \Delta \mathbf{H} = \mathbf{H} \Psi_{\text{KDSUM}} \mathbf{H}.$$

The proof is given in Appendix A.2.

3.2 Trace-balanced component KDSUM kernel

Raw component kernels may have different numerical scales. To prevent scale from determining the embedding, we normalize each PSD component kernel by its centered trace before combining the components. Let $\mathbf{K}_1, \dots, \mathbf{K}_M$ denote the PSD component kernels used in the embedding construction. In the implementation below, EW-KDSUM and GA-KDSUM use these variable-level components so that weights can be assigned while preserving PSD closure. For each component, define

$$\widetilde{\mathbf{K}}_m = \frac{\mathbf{K}_m}{\text{tr}(\mathbf{H}\mathbf{K}_m\mathbf{H}) + \epsilon_{\text{tr}}}, \quad \epsilon_{\text{tr}} > 0.$$

The equal-weight trace-balanced KDSUM baseline is denoted by EW-KDSUM and uses the kernel

$$\Psi_{\text{eq}} = \frac{1}{M} \sum_{m=1}^M \widetilde{\mathbf{K}}_m.$$

Since each $\widetilde{\mathbf{K}}_m$ is a positive scalar multiple of a PSD matrix, it is PSD. Hence

$$\Psi_{\text{eq}} \succeq 0.$$

Remark (Variable-level component construction). In the variable-level implementation, the components $\mathbf{K}_1, \dots, \mathbf{K}_M$ are the continuous Gaussian component kernels, the unordered Aitken component kernels, and the ordered Wang and van Ryzin component kernels. The raw KDSUM kernel uses a single continuous product block, whereas the variable-level construction keeps continuous components separate so that each variable can receive its own weight. Both constructions are PSD by the same closure properties.

3.3 Globally aligned weighted extension

Equal weighting can obscure informative components when many irrelevant variables are present. A component that reflects the shared mixed-type geometry should have a centered kernel that agrees with the geometry induced by other components. GA-KDSUM uses this idea directly: it gives a larger weight to a component when its centered trace-balanced kernel aligns with the average centered geometry of the remaining components.

For each trace-balanced component kernel, define

$$\mathbf{Q}_m = \mathbf{H}\widetilde{\mathbf{K}}_m\mathbf{H}.$$

For $M > 1$, define the leave-one-component-out global reference geometry

$$\mathbf{Q}_{-m} = \frac{1}{M-1} \sum_{r \neq m} \mathbf{Q}_r.$$

The global alignment score is

$$a_m = \max \left\{ 0, \frac{\langle \mathbf{Q}_m, \mathbf{Q}_{-m} \rangle_F}{\|\mathbf{Q}_m\|_F \|\mathbf{Q}_{-m}\|_F + \epsilon_{\text{align}}} \right\}, \quad \epsilon_{\text{align}} > 0.$$

The truncation at zero is used only for numerical stability. The final weight is

$$w_m = (1 - \alpha_w) \frac{1}{M} + \alpha_w \frac{\exp(a_m / \tau_w)}{\sum_{\ell=1}^M \exp(a_\ell / \tau_w)}, \quad \alpha_w \in [0, 1], \quad \tau_w > 0.$$

The experiments use the fixed defaults $\alpha_w = 0.70$ and $\tau_w = 0.05$. These values are chosen a priori, kept fixed across all reported simulation and real-data tasks, and not tuned using outcome labels. When $M = 1$, the global reference is undefined, and we set $w_1 = 1$.

Remark (Interpretation of globally aligned weights). The weighting rule emphasizes components that agree with the shared geometry represented across the mixed-type variables. It is not a supervised feature selector and it does not use the outcome labels or the true latent factors. If an irrelevant group of variables shares a strong nuisance geometry, a purely unsupervised global alignment rule can also emphasize that structure. The method is therefore best interpreted as an unsupervised geometry-alignment rule rather than as hard signal selection.

The GA-KDSUM kernel is

$$\Psi_{\text{GA}} = \sum_{m=1}^M w_m \widetilde{\mathbf{K}}_m.$$

Here $w_m \geq 0$ and $\sum_{m=1}^M w_m = 1$.

Corollary 2. The trace-balanced EW-KDSUM kernel Ψ_{eq} and the GA-KDSUM kernel Ψ_{GA} are PSD. Consequently, $\mathbf{H}\Psi_{\text{eq}}\mathbf{H} \succeq 0$ and $\mathbf{H}\Psi_{\text{GA}}\mathbf{H} \succeq 0$.

Proof. For each component m , $\mathbf{K}_m \succeq 0$. Since $\mathbf{H}$ is symmetric,

$$\mathbf{v}^\top \mathbf{H} \mathbf{K}_m \mathbf{H} \mathbf{v} = (\mathbf{H} \mathbf{v})^\top \mathbf{K}_m (\mathbf{H} \mathbf{v}) \geq 0$$

for every $\mathbf{v} \in \mathbb{R}^n$. Hence $\mathbf{H} \mathbf{K}_m \mathbf{H} \succeq 0$, and in particular

$$\text{tr}(\mathbf{H} \mathbf{K}_m \mathbf{H}) \geq 0.$$

Because $\epsilon_{\text{tr}} > 0$, the trace-normalization constant

$$c_m = \{\text{tr}(\mathbf{H} \mathbf{K}_m \mathbf{H}) + \epsilon_{\text{tr}}\}^{-1}$$

is strictly positive. Therefore

$$\widetilde{\mathbf{K}}_m = c_m \mathbf{K}_m \succeq 0.$$

The EW-KDSUM kernel Ψ_{eq} is a nonnegative linear combination of the matrices $\widetilde{\mathbf{K}}_m$, and is therefore PSD. Similarly, the globally aligned weights are nonnegative by construction and, conditional on the observed data, are fixed scalars summing to one. Thus

$$\Psi_{\text{GA}} = \sum_{m=1}^M w_m \widetilde{\mathbf{K}}_m$$

is a nonnegative linear combination of PSD matrices, and hence $\Psi_{\text{GA}} \succeq 0$. Finally, centering preserves positive semidefiniteness in the sense that, for any PSD matrix $\mathbf{A}$,

$$\mathbf{v}^\top \mathbf{H} \mathbf{A} \mathbf{H} \mathbf{v} = (\mathbf{H} \mathbf{v})^\top \mathbf{A} (\mathbf{H} \mathbf{v}) \geq 0$$

for all $\mathbf{v} \in \mathbb{R}^n$. Applying this with $\mathbf{A} = \Psi_{\text{eq}}$ and $\mathbf{A} = \Psi_{\text{GA}}$ gives

$$\mathbf{H} \Psi_{\text{eq}} \mathbf{H} \succeq 0, \quad \mathbf{H} \Psi_{\text{GA}} \mathbf{H} \succeq 0.$$

□

3.4 Algorithm

Algorithm 1 KDSUM-based embedding with equal or global-alignment weighting

Require: Mixed-type data $\mathbf{X} = (\mathbf{X}_c, \mathbf{X}_u, \mathbf{X}_o)$, target dimension $1 \leq q < n - 1$, component bandwidths, and weighting mode

Ensure: Embedding $\widehat{\mathbf{Y}}$, centered kernel $\mathbf{Q}$, and PSD diagnostics

- 1: Build PSD component kernels $\mathbf{K}_1, \dots, \mathbf{K}_M$ using Gaussian kernels for continuous variables, Aitken kernels for unordered categorical variables, and Wang and van Ryzin kernels for ordered categorical variables.
- 2: For each component, compute $\text{tr}(\mathbf{H}\mathbf{K}_m\mathbf{H})$ and set $\widetilde{\mathbf{K}}_m = \mathbf{K}_m / (\text{tr}(\mathbf{H}\mathbf{K}_m\mathbf{H}) + \epsilon_{\text{tr}})$.
- 3: **if** $M = 1$ **then**
- 4: Set $w_1 = 1$.
- 5: **else if** the weighting mode is equal **then**
- 6: Set $w_m = 1/M$ for all m .
- 7: **else**
- 8: Compute $\mathbf{Q}_m = \mathbf{H}\widetilde{\mathbf{K}}_m\mathbf{H}$, the global reference matrices $\mathbf{Q}_{-m}$, the global alignment scores a_m , and weights w_m .
- 9: **end if**
- 10: Form $\mathbf{\Psi} = \sum_{m=1}^M w_m \widetilde{\mathbf{K}}_m$.
- 11: Center the kernel as $\mathbf{Q} = \mathbf{H}\mathbf{\Psi}\mathbf{H}$.
- 12: Record PSD diagnostics including symmetry error, centering error, minimum eigenvalue, and the number of eigenvalues below -10^{-8} .
- 13: Compute the eigendecomposition $\mathbf{Q} = \mathbf{U}\mathbf{\Gamma}\mathbf{U}^\top$, with eigenvalues sorted in decreasing order.
- 14: Let $\tilde{\gamma}_j = \max\{\gamma_j, 0\}$ and estimate the PPCO noise variance

$$\hat{\sigma}_{\text{PPCO}}^2 = \frac{1}{n - q - 1} \sum_{j=q+1}^{n-1} \tilde{\gamma}_j.$$

- 15: Return $\widehat{\mathbf{Y}} = \mathbf{U}_q(\mathbf{\Gamma}_q - \hat{\sigma}_{\text{PPCO}}^2 \mathbf{I}_q)_+^{1/2}$, where $(\cdot)_+$ truncates negative diagonal entries to zero.
-

The component bandwidths are fixed before PPCO is applied. They may be chosen by an external KDSUM rule, such as MSCV, or by a prespecified experimental rule. They are not optimized by the PPCO likelihood. The eigenvalue truncation used to compute $\tilde{\gamma}_j$ is a numerical safeguard for roundoff. Under the PSD result above, all nontrivial eigenvalues γ_j are nonnegative in exact arithmetic. The final $(\cdot)_+$ operation similarly prevents numerical negative entries after subtracting the PPCO noise estimate. Algorithm 1 describes the training embedding. For held-out observations, the bandwidths, trace normalizers, and weights are fitted on the training set and then kept fixed.

For an out-of-sample point $\mathbf{x}_*$, let $\mathbf{k}_* \in \mathbb{R}^n$ contain its weighted KDSUM kernel similarities to the n training observations. Let $\mathbf{\Psi}$ be the training kernel before centering. The centered train–test kernel vector is

$$\mathbf{q}_* = \mathbf{k}_* - \frac{1}{n} \mathbf{\Psi} \mathbf{1}_n - \frac{1}{n} \mathbf{1}_n \mathbf{1}_n^\top \mathbf{k}_* + \frac{1}{n^2} \mathbf{1}_n \mathbf{1}_n^\top \mathbf{\Psi} \mathbf{1}_n.$$

Given the training eigendecomposition $\mathbf{Q} = \mathbf{U}\mathbf{\Gamma}\mathbf{U}^\top$, the out-of-sample PPCO coordinate is

$$\hat{\mathbf{y}}_*^\top = \mathbf{q}_*^\top \mathbf{U}_q \mathbf{\Gamma}_q^+ (\mathbf{\Gamma}_q - \hat{\sigma}_{\text{PPCO}}^2 \mathbf{I}_q)_+^{1/2},$$

where $\mathbf{\Gamma}_q^+$ is the Moore–Penrose inverse on the retained eigenspace.

4 Evaluation Criteria

This section defines the evaluation criteria used in Sections 5 and 6. These metrics are not part of the kernel construction. They are collected here to make the comparisons unambiguous.

Latent factor R^2 . For synthetic latent factor recovery, let $\mathbf{Z} \in \mathbb{R}^{n \times q_z}$ be the true latent factor matrix and let $\mathbf{Y} \in \mathbb{R}^{n \times r}$ be an estimated embedding. We first fit the least-squares regression

$$\widehat{\mathbf{Z}} = \mathbf{1}_n \widehat{\boldsymbol{\alpha}}^\top + \mathbf{Y} \widehat{\mathbf{B}}.$$

The latent recovery coefficient is

$$R_{\text{latent}}^2 = 1 - \frac{\|\mathbf{Z} - \widehat{\mathbf{Z}}\|_F^2}{\|\mathbf{Z} - \mathbf{1}_n \bar{\mathbf{z}}^\top\|_F^2}, \quad \bar{\mathbf{z}} = \frac{1}{n} \mathbf{Z}^\top \mathbf{1}_n.$$

Larger values indicate better recovery of the latent factors up to a linear transformation.

PSD diagnostics. For PSD diagnostics, the experiments record

$$\lambda_{\min}(\mathbf{Q}), \quad N_{-}(\mathbf{Q}) = \sum_{j=1}^n \mathbf{1}\{\lambda_j(\mathbf{Q}) < -10^{-8}\},$$

along with the centering and symmetry errors

$$E_{\text{center}} = \|\mathbf{Q} \mathbf{1}_n\|_{\infty}, \quad E_{\text{sym}} = \|\mathbf{Q} - \mathbf{Q}^\top\|_{\infty}.$$

These diagnostics separate the kernel validity claim from downstream predictive performance.

Paired improvement and win rate. For a score s evaluated over Monte Carlo replicates or cross-validation folds indexed by $b = 1, \dots, B$, paired improvement over a baseline is

$$\Delta_m = \frac{1}{B} \sum_{b=1}^B (s_{mb} - s_{0b}),$$

where s_{mb} is the score of method m and s_{0b} is the score of the baseline on the same split or replicate. In the real-data tables, the default baseline s_{0b} is standard Gower PCO unless another baseline is named explicitly. The win rate is

$$\text{WinRate}_m = \frac{1}{B} \sum_{b=1}^B \mathbf{1}\{s_{mb} > s_{0b}\}.$$

The paired comparison is the primary comparative summary because it removes replicate-to-replicate or fold-to-fold variation shared by the methods.

Signal-weight diagnostic. In simulations where the informative variables are known, let $\mathcal{S}$ denote the signal components and let $\mathcal{N}$ denote the noise components. For component weights w_m , define

$$W_{\mathcal{S}} = \sum_{m \in \mathcal{S}} w_m, \quad W_{\mathcal{N}} = \sum_{m \in \mathcal{N}} w_m.$$

This diagnostic is relevant for the weighted KDSUM construction because the method is designed to emphasize components sharing a common latent geometry. In the simulation experiment, signal

components include all continuous, unordered, and ordered variables generated from the latent factors.

Balanced accuracy. For classification tasks, the primary metric is balanced accuracy. Let $\mathcal{C}_{\mathcal{T}}$ be the set of classes represented in the test fold, and let $C_{\mathcal{T}} = |\mathcal{C}_{\mathcal{T}}|$. If TP_c and FN_c denote the one-vs-rest true positives and false negatives for class c , then

$$\text{BalAcc} = \frac{1}{C_{\mathcal{T}}} \sum_{c \in \mathcal{C}_{\mathcal{T}}} \frac{\text{TP}_c}{\text{TP}_c + \text{FN}_c}.$$

Macro one-vs-rest AUC. For macro one-vs-rest AUC, let $\mathcal{C}_{\text{AUC}}$ be the classes for which the test fold contains at least one positive and one negative observation. Let $n_c^+ = |\{i : y_i = c\}|$ and $n_c^- = |\{j : y_j \neq c\}|$ within the test fold. For the downstream experiments, the class scores s_{ic} are the logistic-regression predicted probabilities for class c . Then

$$\text{MacroAUC} = \frac{1}{|\mathcal{C}_{\text{AUC}}|} \sum_{c \in \mathcal{C}_{\text{AUC}}} \text{AUC}_c,$$

where

$$\text{AUC}_c = \frac{1}{n_c^+ n_c^-} \sum_{i: y_i = c} \sum_{j: y_j \neq c} \left[\mathbf{1}\{s_{ic} > s_{jc}\} + \frac{1}{2} \mathbf{1}\{s_{ic} = s_{jc}\} \right].$$

Undefined class-wise AUC values are omitted from the macro average. If $\mathcal{C}_{\text{AUC}}$ is empty, macro AUC is not reported for that fold.

5 Simulation Experiment

The simulation experiment is designed to answer three questions. First, does the proposed KDSUM construction produce numerically PSD centered kernels? Second, in a setting with many irrelevant mixed-type variables, does the globally aligned KDSUM geometry recover latent factors better than Gower-type pairwise geometries and EW-KDSUM? Third, does the global-alignment weighting rule increase the relative weight assigned to variables generated from the shared latent geometry?

5.1 Data generating process and compared methods

Each replicate used $n = 180$ observations and two latent factors

$$\mathbf{z}_i \sim N(0, \mathbf{I}_2).$$

The simulation used a mixed-type shared-signal design. The latent factors generated signal variables of all three types. Four continuous signal variables were noisy linear projections of z_1 , z_2 , $z_1 + z_2$, and $z_1 - z_2$. Six unordered categorical signal variables were generated from latent quadrants, latent directions, nearest latent prototypes, and nonlinear partitions of the latent plane. Four ordered categorical signal variables were generated by discretizing noisy latent projections into seven ordered levels. Additional irrelevant variables were split across continuous, unordered categorical, and ordered categorical types.

The four scenarios were `mixed_clean`, `mixed_noise`, `mixed_high_noise`, and `mixed_many_noise`. The clean scenario contains only the mixed-type signal variables. The noisy scenarios add 24, 48,

and 120 irrelevant mixed-type variables, respectively, with higher observation noise in the latter two settings. Each scenario used $B = 100$ Monte Carlo repetitions. The embedding dimension was fixed at $q = 2$, matching the true latent factor dimension.

The main comparisons are restricted to unsupervised mixed-type pairwise geometries, since the contribution of the paper is a PSD kernel interface for PCO, KPCA, and fixed- Q PPCO. The compared methods were GA-KDSUM, the equal-weight trace-balanced baseline EW-KDSUM, standard Gower PCO, ordinal-aware Gower PCO, and PSD-corrected Gower PCO. Standard Gower treats categorical variables by exact matching. The ordinal-aware Gower baseline treats ordered variables by scaled rank distance. The PSD-corrected Gower baseline repairs the Gower centered matrix by spectral clipping. Specifically, after forming $Q_G = -\frac{1}{2}H\Delta_G H$ from the squared Gower dissimilarity matrix Δ_G , negative eigenvalues are set to zero and PCO is applied to the resulting PSD matrix. This baseline tests whether generic PSD repair is enough to replace a native PSD mixed-type kernel.

5.2 Mixed-type simulations with irrelevant variables

The most informative simulation setting for GA-KDSUM is the setting with many irrelevant mixed-type variables. Table 1 reports latent factor recovery in `mixed_many_noise`. GA-KDSUM achieved the highest mean R_{latent}^2 , with a mean of 0.821. The next best method was ordinal Gower PCO with 0.781, followed by EW-KDSUM with 0.769. Standard Gower PCO and PSD-corrected Gower PCO both had mean R_{latent}^2 of 0.736.

Table 1: Latent factor recovery in `mixed_many_noise`. Larger R_{latent}^2 is better.

Method	Mean R_{latent}^2	SE	Rank
GA-KDSUM	0.821	0.002	1
ordinal Gower PCO	0.781	0.002	2
EW-KDSUM	0.769	0.002	3
Gower PCO	0.736	0.003	4
PSD-corrected Gower PCO	0.736	0.003	4

The paired comparison in Table 2 gives the strongest summary of the result for the setting with many irrelevant variables. On the same Monte Carlo replicates, GA-KDSUM improved over EW-KDSUM by 0.051, over standard Gower PCO by 0.085, over ordinal Gower PCO by 0.039, and over PSD-corrected Gower PCO by 0.085. The win rate was 1.00 against every baseline.

Table 2: Paired advantage of GA-KDSUM in `mixed_many_noise`. Positive values favor GA-KDSUM.

Baseline	GA-KDSUM – baseline	SE	Win rate
EW-KDSUM	0.051	0.001	1.00
Gower PCO	0.085	0.002	1.00
ordinal Gower PCO	0.039	0.001	1.00
PSD-corrected Gower PCO	0.085	0.002	1.00

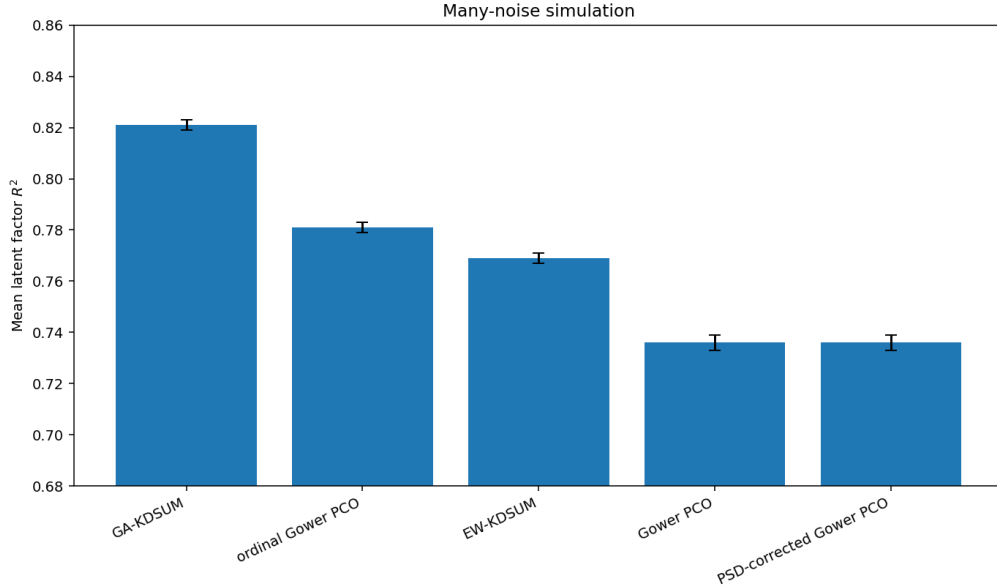


Figure 1: Latent factor recovery in the setting with many irrelevant variables. Error bars show Monte Carlo standard errors.

The full simulation results, reported in Appendix B, show the boundary of the method. In the clean and mild-noise settings, ordinal Gower PCO remains competitive and often gives the highest mean latent recovery. These results suggest that GA-KDSUM is most useful when the data contain a shared latent structure together with many irrelevant mixed-type variables. In this setting, equal weighting can allow unrelated variables to obscure the latent geometry, while the alignment weights emphasize variables that agree with the common pairwise structure and reduce the influence of unrelated variables.

5.3 PSD diagnostics and weighting mechanism

Table 3 summarizes the PSD diagnostics across all simulation replicates. The KDSUM-based methods had no centered-kernel eigenvalues below -10^{-8} . Their minimum eigenvalues were at numerical roundoff scale. In contrast, the double-centered matrices induced by standard Gower PCO and ordinal Gower PCO were often indefinite. Standard Gower PCO had a minimum eigenvalue of -8.22×10^{-1} and 32,833 total negative eigenvalues across the simulation panel. Ordinal Gower PCO had a minimum eigenvalue of -7.52×10^{-1} and 36,080 total negative eigenvalues. The PSD-corrected Gower baseline removed negative eigenvalues but did not improve latent recovery enough to match GA-KDSUM in the setting with many irrelevant variables.

Table 3: PSD diagnostics across all simulation replicates. N_- counts eigenvalues below -10^{-8} .

Method	Min $\lambda_{\min}(\mathbf{Q})$	Max N_-	Total N_-	Max E_{center}
GA-KDSUM	-4.09e-16	0	0	1.26e-15
EW-KDSUM	-3.11e-16	0	0	1.14e-15
Gower PCO	-8.22e-01	130	32833	3.43e-14
ordinal Gower PCO	-7.52e-01	123	36080	2.09e-14
PSD-corrected Gower PCO	-3.11e-14	0	0	1.60e-13

The weight diagnostics in Appendix B show that global alignment also changes the contribution of signal components. In `mixed_noise`, GA-KDSUM assigned mean signal weight 0.805, compared with 0.368 under EW-KDSUM. In `mixed_high_noise`, the corresponding weights were 0.615 and 0.226. In `mixed_many_noise`, they were 0.256 and 0.104. These values should be interpreted as evidence of geometry alignment rather than hard feature selection: GA-KDSUM increases the relative contribution of signal components, but its main objective is to align component geometries, not to recover a sparse set of true variables.

6 Downstream Prediction on Student Performance Data

The real-data experiments evaluate whether the proposed PSD mixed-type geometry improves held-out downstream prediction, rather than only recovering latent factors in simulation. The benchmark uses the UCI Student Performance data, which contains two course files with the same variable design: Portuguese and Mathematics (Cortez, 2008, Cortez & Silva, 2008). Both tasks combine continuous, unordered categorical, and ordered categorical variables describing school, demographic, family, and behavioral attributes. The target is pass or fail, defined as $1(G3 \geq 10)$. The previous grade variables $G1$ and $G2$ are excluded to avoid a near-deterministic grade proxy.

6.1 Dataset panel and inductive protocol

Table 4 summarizes the real-data panel. Because all compared embedding methods require pairwise matrices, each course file is capped at a deterministic benchmark subsample of at most 300 observations.

Table 4: Real-data benchmark panel. Both datasets are mixed-type educational classification tasks.

Dataset	n	p_c	p_u	p_o	Role
Student-Portuguese	300	2	17	11	mixed educational classification
Student-Math	300	2	17	11	same-source validation task

The evaluation is inductive. All preprocessing steps, trace normalizers, KDSUM weights, and embeddings are fitted within each training fold. Test observations are embedded only by out-of-sample projection, using the kernel out-of-sample formula for KDSUM methods and the classical MDS out-of-sample formula for Gower PCO. The same PSD repair for PSD-corrected Gower PCO is applied within each training fold. A logistic regression classifier is then fit on the training embedding

and evaluated on the held-out test fold. The reported results use five-fold cross-validation repeated five times, giving 25 paired train-test splits. The embedding dimension was fixed in advance at $q = 5$ for all methods and all splits.

6.2 Downstream performance

Table 5 reports downstream performance. Balanced accuracy is the primary score because pass or fail labels may be imbalanced. AUC is reported as a secondary threshold-free measure. The paired difference and win rate are computed relative to standard Gower PCO on the same cross-validation splits. GA-KDSUM achieved the highest mean balanced accuracy on both tasks. On Student-Portuguese, GA-KDSUM achieved balanced accuracy 0.703, compared with 0.677 for EW-KDSUM, 0.645 for PSD-corrected Gower PCO, 0.639 for standard Gower PCO, and 0.621 for ordinal Gower PCO. On Student-Math, GA-KDSUM achieved balanced accuracy 0.661, compared with 0.605 for EW-KDSUM, 0.578 for ordinal Gower PCO, 0.555 for PSD-corrected Gower PCO, and 0.553 for standard Gower PCO.

Table 5: Real-data downstream performance under repeated five-fold inductive cross-validation. Positive Δ values favor the listed method over standard Gower PCO.

Dataset	Method	Bal. acc.	Δ vs. Gower	Win rate	AUC	Rank
Student-Portuguese	GA-KDSUM	0.703 (0.014)	0.064 (0.015)	0.84	0.735 (0.017)	1
Student-Portuguese	EW-KDSUM	0.677 (0.015)	0.038 (0.014)	0.68	0.734 (0.018)	2
Student-Portuguese	Gower PCO	0.639 (0.015)	—	—	0.679 (0.019)	4
Student-Portuguese	ordinal Gower PCO	0.621 (0.016)	-0.019 (0.012)	0.32	0.657 (0.017)	5
Student-Portuguese	PSD-corrected Gower PCO	0.645 (0.015)	0.006 (0.001)	0.48	0.678 (0.019)	3
Student-Math	GA-KDSUM	0.661 (0.012)	0.108 (0.010)	0.92	0.699 (0.013)	1
Student-Math	EW-KDSUM	0.605 (0.011)	0.052 (0.008)	0.84	0.657 (0.013)	2
Student-Math	Gower PCO	0.553 (0.014)	—	—	0.600 (0.014)	5
Student-Math	ordinal Gower PCO	0.578 (0.013)	0.025 (0.007)	0.76	0.607 (0.014)	3
Student-Math	PSD-corrected Gower PCO	0.555 (0.014)	0.001 (0.001)	0.12	0.600 (0.014)	4

The paired advantages of GA-KDSUM are summarized in Table 6. On Student-Portuguese, GA-KDSUM improved over standard Gower PCO by 0.064 in balanced accuracy, with win rate 0.84. On Student-Math, the paired improvement over Gower PCO was 0.108, with win rate 0.92. The improvement over EW-KDSUM was also positive on both tasks, with 0.026 on Student-Portuguese and 0.056 on Student-Math. These results show that the benefits of GA-KDSUM are not limited to synthetic latent-factor recovery. Under a fully inductive evaluation protocol, the proposed PSD mixed-type geometry produces embeddings that improve held-out prediction relative to both Gower-type PCO baselines and the equal-weight KDSUM baseline.

Table 6: Paired balanced-accuracy advantage of GA-KDSUM over each real-data baseline. Positive values favor GA-KDSUM.

Dataset	Baseline	GA-KDSUM – baseline	SE	Win rate
Student-Math	EW-KDSUM	0.056	0.010	0.88
Student-Math	Gower PCO	0.108	0.010	0.92
Student-Math	ordinal Gower PCO	0.083	0.013	0.88
Student-Math	PSD-corrected Gower PCO	0.107	0.011	0.92
Student-Portuguese	EW-KDSUM	0.026	0.010	0.60
Student-Portuguese	Gower PCO	0.064	0.015	0.84
Student-Portuguese	ordinal Gower PCO	0.082	0.017	0.80
Student-Portuguese	PSD-corrected Gower PCO	0.058	0.015	0.84

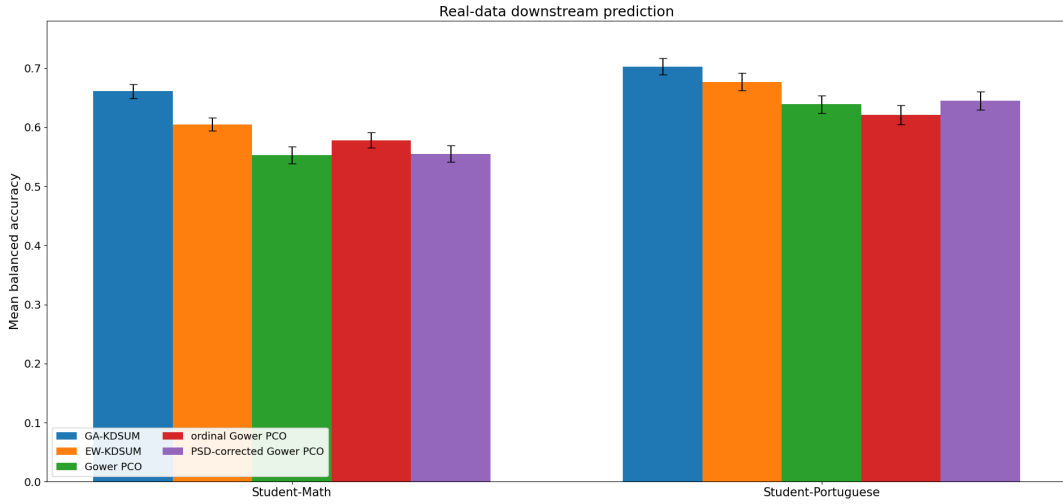


Figure 2: Mean balanced accuracy on the two Student Performance tasks. Error bars show standard errors across repeated cross-validation splits.

As a secondary metric, AUC gives the same broad ordering. GA-KDSUM had the highest mean AUC on Student-Math (0.699) and the highest mean AUC on Student-Portuguese (0.735), although on Student-Portuguese it was nearly tied with EW-KDSUM (0.734). The AUC advantage table is reported in Appendix C.

7 Discussion

The central message of this paper is that KDSUM can be used as a native PSD kernel for mixed-type embedding. The original KDSUM construction was developed as a similarity and dissimilarity for clustering. Here the same component kernels, once bandwidths are fixed, define a sample kernel that is PSD before centering and remains PSD after centering by $\mathbf{H}$. This moves the construction from a distance-based clustering tool to an admissible kernel geometry for PCO, KPCA, and fixed- $\mathbf{Q}$ PPCO. The distinction is important because a mixed-type dissimilarity can be useful for comparison while still producing an indefinite double-centered matrix. In that case the embedding requires truncation or spectral repair. KDSUM avoids this ambiguity by supplying the inner-product object directly.

The role of GA-KDSUM is to make this kernel geometry adaptive without losing PSD validity. Equal weighting treats all trace-balanced components as equally informative, which can be unstable when many variables are weakly related to the shared low-dimensional structure. The globally aligned weights use agreement with the consensus mixed-type geometry to increase the contribution of aligned components and reduce the influence of weakly aligned components. Because the final kernel is still a nonnegative combination of PSD component kernels, the weighting step changes the geometry while preserving the kernel guarantees. This explains why the clearest simulation gains appear in the setting with many irrelevant variables, where GA-KDSUM achieves the strongest latent factor recovery among the pairwise geometry baselines. It also explains why the method is not expected to dominate in every setting. When ordered variables directly encode the latent factors and few nuisance variables are present, ordinal Gower PCO can remain highly competitive. An unsupervised alignment rule can also emphasize a mutually aligned nuisance block, underweight isolated but useful variables, or miss task-specific supervised signal.

The Gower comparisons sharpen the interpretation of the results. Gower-type distances remain strong and useful baselines for heterogeneous data, but they do not necessarily produce PSD centered matrices. PSD-corrected Gower removes negative eigenvalues and makes PCO feasible, yet the simulations show that this repair does not by itself improve latent recovery. The same pattern appears in the downstream experiments, where GA-KDSUM improves over both standard Gower PCO and PSD-corrected Gower PCO. These results suggest that the advantage comes from constructing a type-aware PSD kernel and learning an aligned variable-weighted geometry, not merely from forcing a matrix to be PSD after the fact.

The real-data experiments support the same conclusion beyond synthetic latent factors. The Student Performance tasks contain continuous, unordered categorical, and ordered categorical variables, and the evaluation embeds test observations only through out-of-sample projection under repeated inductive cross-validation. On both tasks, GA-KDSUM gives the highest mean balanced accuracy and AUC. This shows that the proposed geometry can improve held-out predictive embeddings in real mixed-type data. The evidence is still limited to two related educational datasets, so broader benchmarks are needed before claiming broad empirical dominance across domains.

Several limitations remain. The current implementation forms dense pairwise matrices, so larger datasets require scalable kernel approximations, landmark or blockwise construction, or scalable fixed- Q PPCO solvers. The experiments also use fixed bandwidth choices and a fixed embedding dimension q , leaving systematic comparison of MSCV, task-tuned bandwidths, and data-driven choices of q for future work. Broader evaluations should compare GA-KDSUM with raw-feature factorization, supervised representation learning, and task-specific clustering or prediction methods, while also addressing missingness, high-cardinality categorical variables, and mixed local and global structure.

This discussion frames GA-KDSUM as a principled PSD mixed-type embedding method. Its key value is the combination of type-aware kernel construction, unsupervised geometry alignment, and positive semidefinite validity in one pairwise representation.

8 Conclusion

This paper establishes GA-KDSUM as a positive semidefinite kernel framework for mixed-type dimensionality reduction. Starting from the KDSUM similarity, we show that Gaussian, Aitken, and Wang–van Ryzin component kernels define valid PSD sample kernels after the bandwidths are

fixed, and that their nonnegative combinations preserve PSD structure. This result moves KDSUM beyond its original role as a mixed-type dissimilarity for clustering and turns it into a native kernel geometry for PCO, KPCA, and fixed- $\mathbf{Q}$ PPCO.

The globally aligned extension strengthens this kernel construction by making the geometry adaptive to shared structure. Through trace normalization and nonnegative alignment weights, GA-KDSUM combines type-aware treatment of continuous, unordered categorical, and ordered categorical variables with automatic downweighting of weakly aligned variables. It preserves positive semidefiniteness without post-hoc spectral correction. The method therefore does not merely repair a mixed-type distance for embedding. It constructs a valid and adaptive mixed-type kernel from the outset.

The empirical results support this interpretation. In simulations, GA-KDSUM is most effective when informative mixed-type variables are embedded among many irrelevant variables, precisely the regime where equal weighting and generic mixed-type distances are expected to struggle. In the Student Performance benchmarks, the same kernel embedding pipeline achieves the best held-out balanced accuracy and AUC among the compared pairwise geometry methods. These results indicate that the gains come from learning a more informative mixed-type geometry while maintaining PSD validity.

In general, GA-KDSUM provides a practical and principled route from heterogeneous tabular data to kernel-based low-dimensional representations. It avoids homogenizing variable types, avoids relying on indefinite dissimilarities, and avoids separating PSD correction from variable relevance weighting. This makes it a strong foundation for mixed-type spectral embedding, kernel PCA, and probabilistic principal coordinate analysis in modern tabular settings where heterogeneous variables, irrelevant features, and downstream embedding requirements appear together.

Acknowledgements

I would like to thank Daniel Krasnov for his mentorship, guidance, and helpful feedback throughout this project.

References

- Aitchison, J. and Aitken, C. G. G. (1976). Multivariate binary discrimination by the kernel method. *Biometrika*, 63(3), 413–420.
- Borg, I. and Groenen, P. J. F. (1997). *Modern Multidimensional Scaling: Theory and Applications*. Springer.
- Cortes, C., Mohri, M., and Rostamizadeh, A. (2012). Algorithms for learning kernels based on centered alignment. *Journal of Machine Learning Research*, 13, 795–828.
- Cortez, P. (2008). Student Performance [Dataset]. UCI Machine Learning Repository. <https://doi.org/10.24432/C5TG7T>.
- Cortez, P. and Silva, A. M. G. (2008). Using data mining to predict secondary school student performance. In *Proceedings of the 5th Annual Future Business Technology Conference*, 5–12.
- Eckart, C. and Young, G. (1936). The approximation of one matrix by another of lower rank. *Psychometrika*, 1(3), 211–218.

- Foss, A. H., Markatou, M., Ray, B., and Heching, A. (2016). A semiparametric method for clustering mixed data. *Machine Learning*, 105(3), 419–458.
- Foss, A. H., Markatou, M., and Ray, B. (2019). Distance metrics and clustering methods for mixed-type data. *International Statistical Review*, 87(1), 80–109.
- Ghashti, J. S. and Thompson, J. R. J. (2025). Mixed-type distance shrinkage and selection for clustering via kernel metric learning. *Journal of Classification*, 42(2), 311–334. <https://doi.org/10.1007/s00357-024-09493-z>.
- Gower, J. C. (1966). Some distance properties of latent root and vector methods used in multivariate analysis. *Biometrika*, 53(3/4), 325–338. <https://doi.org/10.1093/biomet/53.3-4.325>.
- Gower, J. C. (1971). A general coefficient of similarity and some of its properties. *Biometrics*, 27(4), 857–871.
- Huang, Z. (1998). Extensions to the k-means algorithm for clustering large data sets with categorical values. *Data Mining and Knowledge Discovery*, 2(3), 283–304.
- Horn, R. A. and Johnson, C. R. (2012). *Matrix Analysis*. Second edition. Cambridge University Press.
- Jolliffe, I. T. (2002). *Principal Component Analysis*. Second edition. Springer.
- McParland, D. and Gormley, I. C. (2016). Model based clustering for mixed data: clustMD. *Advances in Data Analysis and Classification*, 10(2), 155–169.
- Miotto, R., Li, L., Kidd, B. A., and Dudley, J. T. (2016). Deep Patient: An unsupervised representation to predict the future of patients from the electronic health records. *Scientific Reports*, 6, 26094.
- Podani, J. (1999). Extending Gower’s general coefficient of similarity to ordinal characters. *Taxon*, 48(2), 331–340.
- Schölkopf, B., Smola, A., and Müller, K.-R. (1998). Nonlinear component analysis as a kernel eigenvalue problem. *Neural Computation*, 10(5), 1299–1319.
- Wang, M.-C. and Van Ryzin, J. (1981). A class of smooth estimators for discrete distributions. *Biometrika*, 68(1), 301–309.
- Williams, C. K. I. (2002). On a connection between kernel PCA and metric multidimensional scaling. *Machine Learning*, 46, 11–19.
- Yeh, I.-C. and Lien, C.-H. (2009). The comparisons of data mining techniques for the predictive accuracy of probability of default of credit card clients. *Expert Systems with Applications*, 36(2), 2473–2480.
- Zhang, Z. and Jordan, M. I. (2009). Latent variable models for dimensionality reduction. In *Proceedings of the Twelfth International Conference on Artificial Intelligence and Statistics*, 655–662.

A Proofs

A.1 PSD details used in Proposition 1

This subsection records the two standard PSD facts used in Proposition 1.

First, the Gaussian component is PSD. Let

$$\mathbf{G}_{ij} = \exp \left\{ -\frac{(x_i - x_j)^2}{2\lambda^2} \right\}, \quad \lambda > 0.$$

For any $\mathbf{a} = (a_1, \dots, a_n)^\top \in \mathbb{R}^n$,

$$\begin{aligned} \mathbf{a}^\top \mathbf{G} \mathbf{a} &= \sum_{i,j=1}^n a_i a_j \exp \left\{ -\frac{x_i^2}{2\lambda^2} \right\} \exp \left\{ -\frac{x_j^2}{2\lambda^2} \right\} \exp \left\{ \frac{x_i x_j}{\lambda^2} \right\} \\ &= \sum_{m=0}^{\infty} \frac{1}{m! \lambda^{2m}} \left(\sum_{i=1}^n a_i \exp \left\{ -\frac{x_i^2}{2\lambda^2} \right\} x_i^m \right)^2 \geq 0. \end{aligned}$$

Thus $\mathbf{G} \succeq 0$.

Since $(\lambda\sqrt{2\pi})^{-1} > 0$, the normalized Gaussian component matrix is also PSD.

Second, the matrix $\mathbf{T}_\lambda$ with entries $(\mathbf{T}_\lambda)_{rs} = \lambda^{|r-s|}$ is PSD for every $\lambda \in [0, 1]$. For $0 \leq \lambda < 1$, let $\varepsilon_1, \dots, \varepsilon_G$ be independent standard normal random variables and define

$$Z_1 = \varepsilon_1, \quad Z_t = \lambda Z_{t-1} + \sqrt{1 - \lambda^2} \varepsilon_t, \quad t = 2, \dots, G.$$

Then $\text{Var}(Z_t) = 1$ and, for $r \leq s$,

$$\text{Cov}(Z_r, Z_s) = \lambda^{s-r}.$$

Hence $\mathbf{T}_\lambda$ is a covariance matrix and is PSD. For $\lambda = 1$, $\mathbf{T}_\lambda = \mathbf{1}_G \mathbf{1}_G^\top$, which is PSD. $\square$

A.2 Proof of Corollary 1

Proof. By Proposition 1, $\Psi_{\text{KDSUM}} \succeq 0$. Hence there exists a matrix $\mathbf{B}$ such that

$$\Psi_{\text{KDSUM}} = \mathbf{B} \mathbf{B}^\top.$$

Let $\phi_i^\top$ be the i th row of $\mathbf{B}$. Then

$$(\Psi_{\text{KDSUM}})_{ij} = \langle \phi_i, \phi_j \rangle.$$

Therefore

$$\begin{aligned} \delta_{ij}^2 &= (\Psi_{\text{KDSUM}})_{ii} + (\Psi_{\text{KDSUM}})_{jj} - 2(\Psi_{\text{KDSUM}})_{ij} \\ &= \|\phi_i\|^2 + \|\phi_j\|^2 - 2\langle \phi_i, \phi_j \rangle \\ &= \|\phi_i - \phi_j\|^2. \end{aligned}$$

Thus δ_{ij}^2 is a squared Hilbert space distance.

Let $\mathbf{d} = \text{diag}(\Psi_{\text{KDSUM}})$ be the vector of diagonal entries of Ψ_{KDSUM} . Then

$$\Delta = \mathbf{d} \mathbf{1}_n^\top + \mathbf{1}_n \mathbf{d}^\top - 2\Psi_{\text{KDSUM}}.$$

Using $\mathbf{H}\mathbf{1}_n = 0$ and $\mathbf{1}_n^\top \mathbf{H} = 0^\top$, we obtain

$$\begin{aligned} -\frac{1}{2}\mathbf{H}\Delta\mathbf{H} &= -\frac{1}{2}\mathbf{H}\left(d\mathbf{1}_n^\top + \mathbf{1}_n d^\top - 2\Psi_{\text{KDSUM}}\right)\mathbf{H} \\ &= \mathbf{H}\Psi_{\text{KDSUM}}\mathbf{H}. \end{aligned}$$

This proves the claim. $\square$

B Additional Simulation Results

This appendix contains secondary simulation diagnostics and illustrative figures omitted from the main text. The main simulation is the mixed-type shared-signal design described in Section 5. All final simulation tables use $B = 100$, $n = 180$, and embedding dimension $q = 2$.

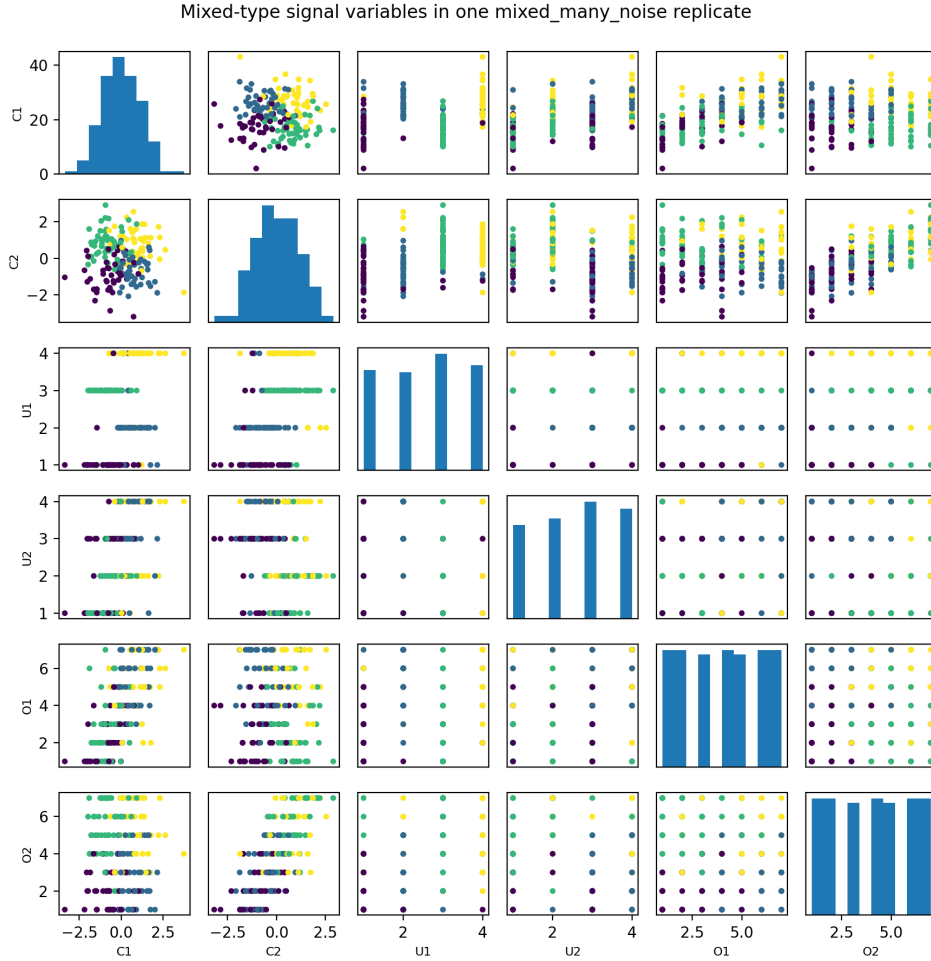


Figure 3: Pair plot of representative observed variables in one mixed-signal replicate. Points are colored by latent quadrant only for visualization.

Table 7: Full latent factor recovery across simulation scenarios. Larger R^2_{latent} is better.

Scenario	Method	Mean R^2_{latent}	SE
mixed clean	GA-KDSUM	0.758	0.004
mixed clean	EW-KDSUM	0.831	0.002
mixed clean	Gower PCO	0.830	0.002
mixed clean	ordinal Gower PCO	0.868	0.001
mixed clean	PSD-corrected Gower PCO	0.830	0.002
mixed noise	GA-KDSUM	0.794	0.002
mixed noise	EW-KDSUM	0.823	0.002
mixed noise	Gower PCO	0.811	0.002
mixed noise	ordinal Gower PCO	0.844	0.002
mixed noise	PSD-corrected Gower PCO	0.811	0.002
mixed high noise	GA-KDSUM	0.796	0.002
mixed high noise	EW-KDSUM	0.793	0.002
mixed high noise	Gower PCO	0.767	0.002
mixed high noise	ordinal Gower PCO	0.803	0.002
mixed high noise	PSD-corrected Gower PCO	0.767	0.002
mixed many noise	GA-KDSUM	0.821	0.002
mixed many noise	EW-KDSUM	0.769	0.002
mixed many noise	Gower PCO	0.736	0.003
mixed many noise	ordinal Gower PCO	0.781	0.002
mixed many noise	PSD-corrected Gower PCO	0.736	0.003

Table 8: Best method by simulation scenario. This table highlights that GA-KDSUM is strongest in the many-noise setting but not uniformly best in clean or mild-noise settings.

Scenario	Best method	Best R^2_{latent}	GA-KDSUM R^2_{latent}
mixed clean	ordinal Gower PCO	0.868	0.758
mixed noise	ordinal Gower PCO	0.844	0.794
mixed high noise	ordinal Gower PCO	0.803	0.796
mixed many noise	GA-KDSUM	0.821	0.821

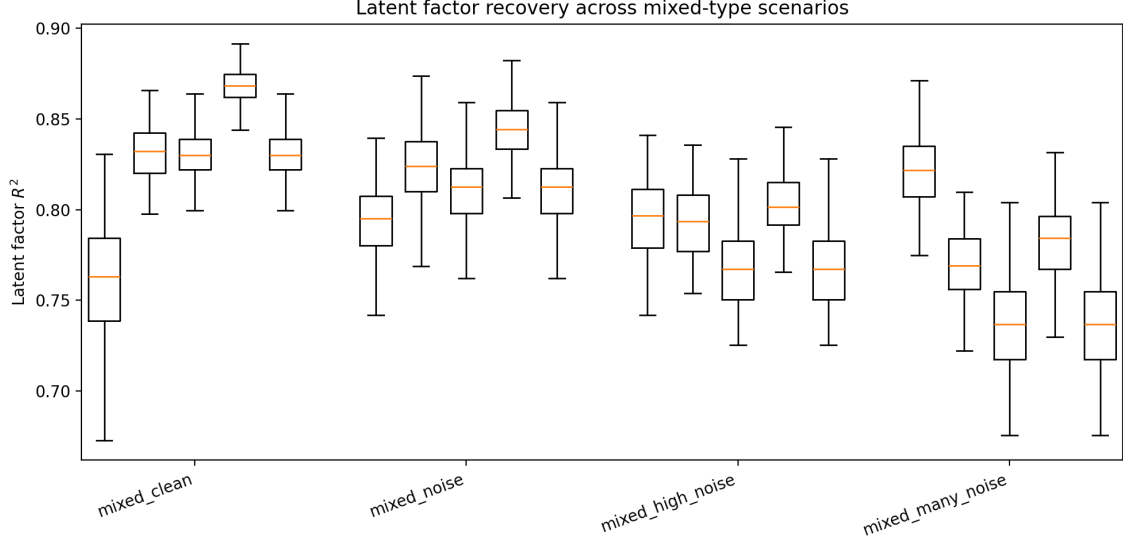


Figure 4: Latent factor R^2 boxplots by scenario and method.

Table 9: Full paired latent factor R^2 difference relative to Gower PCO. Positive values favor the listed method.

Scenario	Method	ΔR^2 vs. Gower	SE	Win rate
mixed clean	GA-KDSUM	-0.072	0.004	0.00
mixed clean	EW-KDSUM	0.001	0.001	0.59
mixed clean	ordinal Gower PCO	0.038	0.000	1.00
mixed clean	PSD-corrected Gower PCO	0.000	0.000	0.16
mixed noise	GA-KDSUM	-0.017	0.002	0.23
mixed noise	EW-KDSUM	0.013	0.001	0.94
mixed noise	ordinal Gower PCO	0.033	0.001	1.00
mixed noise	PSD-corrected Gower PCO	0.000	0.000	0.18
mixed high noise	GA-KDSUM	0.028	0.002	0.92
mixed high noise	EW-KDSUM	0.026	0.001	0.98
mixed high noise	ordinal Gower PCO	0.035	0.001	1.00
mixed high noise	PSD-corrected Gower PCO	0.000	0.000	0.14
mixed many noise	GA-KDSUM	0.085	0.002	1.00
mixed many noise	EW-KDSUM	0.034	0.002	0.99
mixed many noise	ordinal Gower PCO	0.046	0.001	1.00
mixed many noise	PSD-corrected Gower PCO	0.000	0.000	0.20

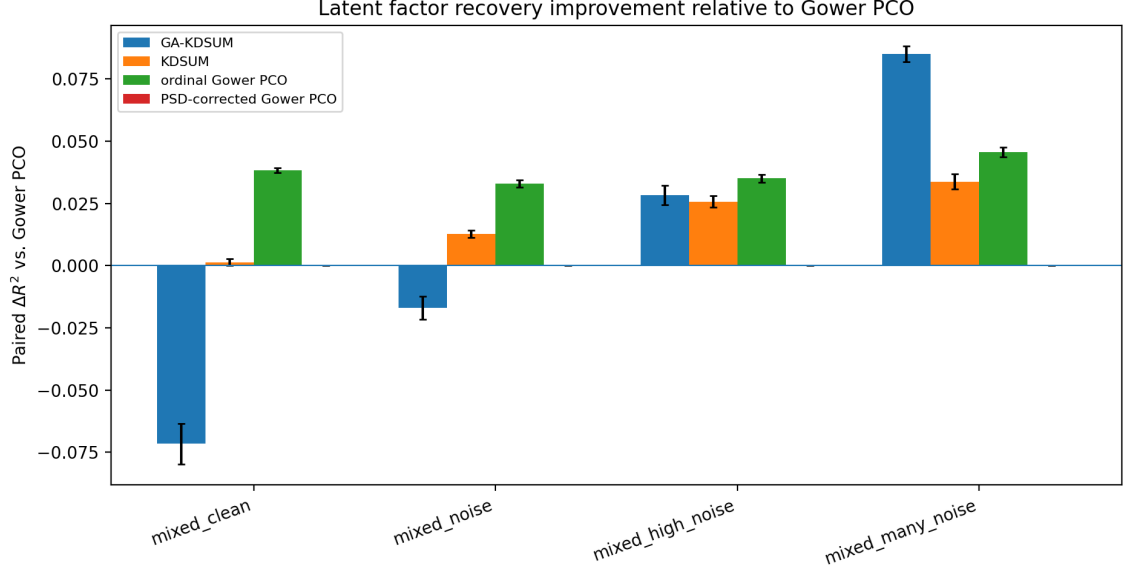


Figure 5: Paired ΔR^2 relative to Gower PCO across all simulation scenarios. Error bars show ± 1.96 Monte Carlo standard errors.

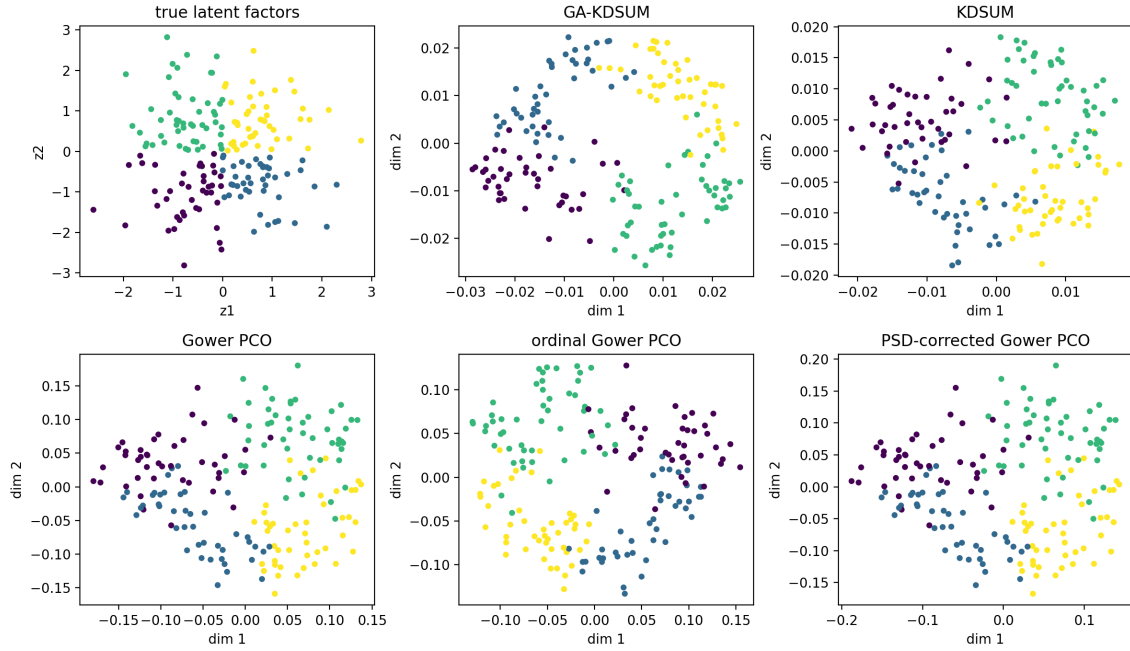


Figure 6: Embedding comparison for one mixed-signal replicate. Colors indicate latent quadrants and are used only for visualization.

Table 10: Mean total kernel weight by variable type and by signal/noise status. In this simulation, signal components include continuous, unordered, and ordered variables generated from the latent factors.

Scenario	Method	Continuous	Unordered	Ordered	Signal weight	Noise weight
mixed clean	GA-KDSUM	0.219	0.664	0.117	1.000	0.000
mixed clean	EW-KDSUM	0.286	0.429	0.286	1.000	0.000
mixed noise	GA-KDSUM	0.266	0.535	0.199	0.805	0.195
mixed noise	EW-KDSUM	0.316	0.368	0.316	0.368	0.632
mixed high noise	GA-KDSUM	0.243	0.464	0.293	0.615	0.385
mixed high noise	EW-KDSUM	0.323	0.355	0.323	0.226	0.774
mixed many noise	GA-KDSUM	0.224	0.391	0.385	0.256	0.744
mixed many noise	EW-KDSUM	0.328	0.343	0.328	0.104	0.896

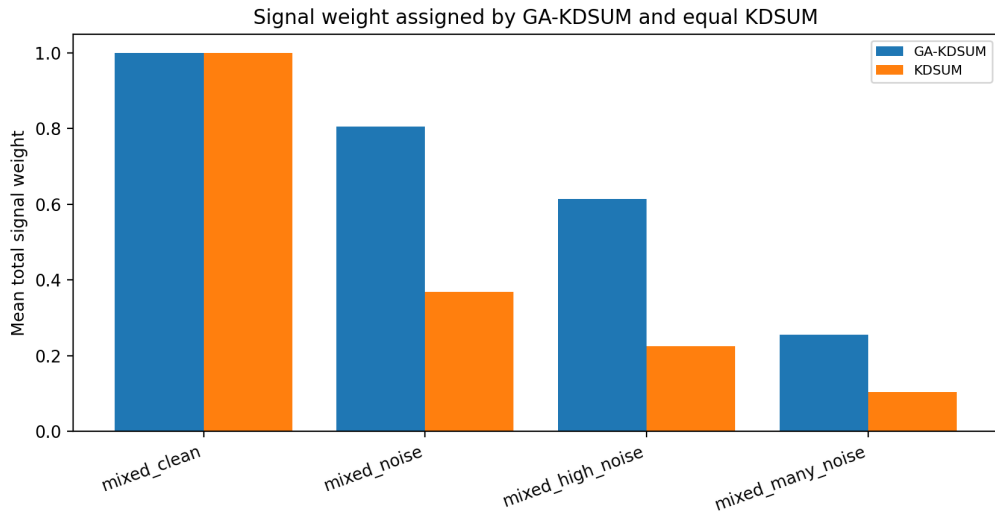


Figure 7: Mean signal weight assigned by GA-KDSUM and the equal-weight trace-balanced baseline EW-KDSUM across simulation scenarios.

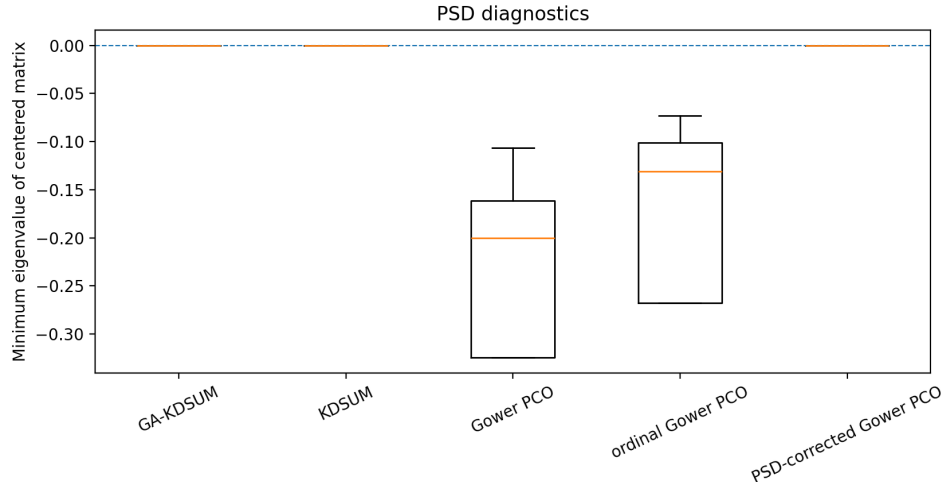


Figure 8: Minimum eigenvalue of the centered KDSUM kernel matrix $\mathbf{Q}$ for the equal-weight trace-balanced baseline EW-KDSUM and GA-KDSUM. Values are at numerical precision.

C Additional Real-Data Results

This appendix reports supplementary real-data metrics for the two Student Performance tasks used in the main text. The appendix is limited to results needed to support the discussion of the Student benchmark and does not include unrelated off-target datasets.

Table 11: Full real-data downstream performance under repeated five-fold inductive cross-validation. The primary score is balanced accuracy.

Dataset	Method	Bal. acc.	AUC	Rank
Student-Math	GA-KDSUM	0.661 (0.012)	0.699 (0.013)	1
Student-Math	EW-KDSUM	0.605 (0.011)	0.657 (0.013)	2
Student-Math	ordinal Gower PCO	0.578 (0.013)	0.607 (0.014)	3
Student-Math	PSD-corrected Gower PCO	0.555 (0.014)	0.600 (0.014)	4
Student-Math	Gower PCO	0.553 (0.014)	0.600 (0.014)	5
Student-Portuguese	GA-KDSUM	0.703 (0.014)	0.735 (0.017)	1
Student-Portuguese	EW-KDSUM	0.677 (0.015)	0.734 (0.018)	2
Student-Portuguese	PSD-corrected Gower PCO	0.645 (0.015)	0.678 (0.019)	3
Student-Portuguese	Gower PCO	0.639 (0.015)	0.679 (0.019)	4
Student-Portuguese	ordinal Gower PCO	0.621 (0.016)	0.657 (0.017)	5

Table 12: Paired balanced-accuracy difference relative to Gower PCO on the same cross-validation splits. Positive values favor the listed method.

Dataset	Method	Δ bal. acc. vs. Gower	SE	Win rate
Student-Math	GA-KDSUM	0.108	0.010	0.92
Student-Math	EW-KDSUM	0.052	0.008	0.84
Student-Math	ordinal Gower PCO	0.025	0.007	0.76
Student-Math	PSD-corrected Gower PCO	0.001	0.001	0.12
Student-Portuguese	GA-KDSUM	0.064	0.015	0.84
Student-Portuguese	EW-KDSUM	0.038	0.014	0.68
Student-Portuguese	ordinal Gower PCO	-0.019	0.012	0.32
Student-Portuguese	PSD-corrected Gower PCO	0.006	0.001	0.48

Table 13: Paired AUC advantage of GA-KDSUM over each real-data baseline. Positive values favor GA-KDSUM.

Dataset	Baseline	GA-KDSUM AUC – baseline	SE	Win rate
Student-Math	EW-KDSUM	0.043	0.007	0.96
Student-Math	Gower PCO	0.099	0.009	0.96
Student-Math	ordinal Gower PCO	0.092	0.012	0.92
Student-Math	PSD-corrected Gower PCO	0.099	0.009	0.96
Student-Portuguese	EW-KDSUM	0.002	0.008	0.44
Student-Portuguese	Gower PCO	0.056	0.011	0.92
Student-Portuguese	ordinal Gower PCO	0.078	0.012	0.92
Student-Portuguese	PSD-corrected Gower PCO	0.057	0.012	0.92

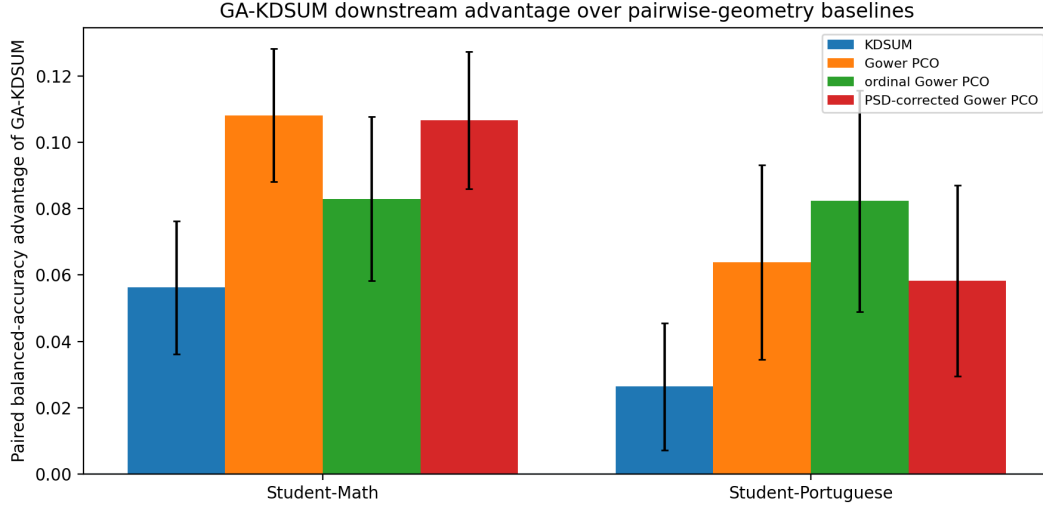


Figure 9: Paired balanced-accuracy advantage of GA-KDSUM over real-data baselines. Error bars show ± 1.96 standard errors across repeated cross-validation splits.

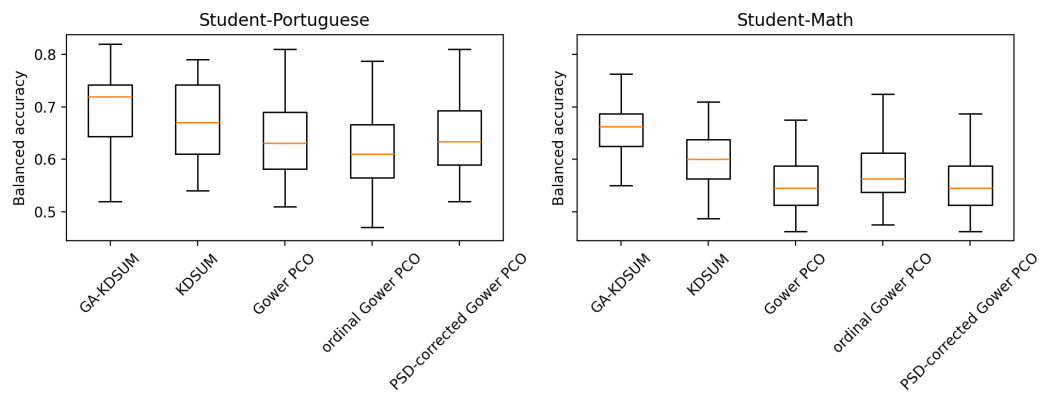


Figure 10: Balanced-accuracy distributions across repeated cross-validation splits for the two Student Performance tasks.