

Optimizing Class Separability via Projection-Based Discriminative Basis Selection

İsmail Özçil

*Department of Mechanical Engineering
Middle East Technical University
Ankara, Türkiye
iozcil@metu.edu.tr*

A. Bugra Koku

*Department of Mechanical Engineering
Middle East Technical University
Ankara, Türkiye
kbugra@metu.edu.tr*

Ali Sekmen

*Department of Computer Science
Tennessee State University
Nashville, TN 37209
asekmen@tnstate.edu*

Bahadır Bilgin

*Department of Computer Science
Tennessee State University
Nashville, TN 37209
bbilgin@tnstate.edu*

Abstract—In high-dimensional classification tasks, data from different classes often lie in a union of lower-dimensional subspaces. Identifying the basis vectors for each subspace that effectively differentiates between classes can enhance the explainability and accuracy of classification methods. This study proposes a novel approach that uses singular value decomposition to identify class-specific basis vectors that maximize the separability of classes. Instead of selecting the most significant n number of basis vectors using traditional heuristics for basis selection, the mean average precision for each basis vector is calculated, and the top-performing n basis vectors are selected. Furthermore, this study extends the methodology by integrating feature vector outputs from two different pre-trained deep learning models, as input for classification evaluation in two different cases. The proposed methodology is validated through simulations, demonstrating its potential for improving classification in high-dimensional spaces.

Index Terms—Subspace clustering, subspace segmentation, multidimensional classification, dimensionality reduction.

I. INTRODUCTION

According to the Manifold Hypothesis, data in high-dimensional classification tasks often come from a union of lower-dimensional subspaces. Identifying these subspaces is fundamentally equivalent to solving the classification problem, as it enables a more structured and interpretable separation of classes. For example, in facial expression recognition, images of the same person under different expressions form a smooth low-dimensional manifold within the high-dimensional image space [1]. Similarly, in motion segmentation, the trajectories of different moving objects in a video typically reside in distinct subspaces [2].

Singular Value Decomposition (SVD) is a widely used mathematical tool for determining the dimension of lower-dimensional subspaces and extracting a corresponding set of basis vectors. Given a data matrix $X \in \mathbb{R}^{m \times n}$, SVD decomposes it as $X = U\Sigma V^T$, where $U \in \mathbb{R}^{m \times m}$ and $V \in \mathbb{R}^{n \times n}$ contain orthonormal basis vectors for the column

and row spaces of X , respectively, and $\Sigma \in \mathbb{R}^{m \times n}$ is a diagonal matrix containing the singular values. The rank of X , determined by the number of nonzero singular values in Σ , provides an estimate of the intrinsic dimension of the subspace, while the corresponding left and right singular vectors in U and V span the basis of that subspace. However, selecting the most informative basis vectors for optimal class differentiation remains an open research problem.

This paper introduces a method for selecting differentiating bases by using projections and mean average precision (mAP) analysis. Additionally, it explores how feature vector outputs from pre-trained ResNet-18 [3] and ViT-L/16 [4] can be used to evaluate the classification performance of the proposed method.

A. Related Work

SVD is a foundational matrix factorization technique that is applicable in dimensionality reduction and data compression [5]. One common application of SVD is Principal Component Analysis (PCA), which identifies the directions of maximum variance in high-dimensional data [6]. PCA is mathematically equivalent to SVD applied to the covariance matrix of the data and is extensively used in feature selection and noise reduction [7], [8]. Several extensions of SVD have been introduced to address computational challenges. Skinny SVD reduces complexity by selecting the top k singular values [9], while Incremental SVD updates the decomposition dynamically for streaming data [10]. Robust PCA (RPCA) handles outliers by decomposing data into low-rank and sparse components [11], and Randomized SVD efficiently approximates SVD for large-scale datasets [12].

SVD and PCA are applied in image compression [13], recommender systems [14], natural language processing (NLP) through Latent Semantic Analysis (LSA) [15], and bioinformatics for genomic data clustering [16]. While SVD remains computationally intensive for large datasets, modern adaptations such as Randomized SVD and Incremental PCA improve

efficiency, making these methods highly relevant for large-scale and real-time data processing.

Subspace analysis techniques have been extensively studied in machine learning and signal processing [17]–[19]. SVD has been employed for dimensionality reduction, with applications ranging from PCA to subspace clustering [5], [20]. Previous studies have explored the use of subspaces for classification, such as in [21], which analyzes geometric relationships between subspaces. Feature selection methods, including filter, wrapper, and embedded approaches, aim to identify discriminative features but often lack a direct connection to the underlying subspace structure [20].

Recent advancements in deep learning, particularly with convolutional neural networks (CNNs) and transformer-based architectures, have demonstrated the effectiveness of feature extraction using pre-trained models such as ResNet-18 [3] and ViT-L/16 [4]. Feature vectors obtained from these networks have been successfully used for transfer learning and classification tasks [22]. This study extends these findings by evaluating how SVD-based basis selection interacts with feature vectors extracted from both CNN-based (ResNet-18) and transformer-based (ViT-L/16) architectures.

Linear Discriminant Analysis (LDA), originally introduced by Fisher [23], is a well-established method for dimensionality reduction and classification. It optimizes class separability by projecting high-dimensional data onto a lower-dimensional space while maximizing the ratio of between-class to within-class scatter. Traditional LDA assumes normally distributed classes with equal covariance, which limits its applicability. To address these constraints, several extensions have been developed, including Regularized LDA (RLDA) [24] for handling singular scatter matrices, Kernel LDA (KLDA) [25] for non-linear separability, and Incremental LDA [26] for real-time data updates.

Recent studies have focused on selecting discriminative bases for multi-class classification problems. Sun et al. [27] proposed the Sparse Softmax Feature Selection (S²FS) method, which enhances multi-class classification by combining $\ell_{2,0}$ -norm regularization with the Softmax model to enforce sparsity while preserving discriminative power. Zhu et al. [28] introduced a feature selection approach integrating Fisher’s Linear Discriminant Analysis (FLDA) and Locality Preserving Projection (LPP) to select class-discriminative and noise-resistant features, notably applied in Alzheimer’s disease classification. Aguilar-Ruiz [29] presented a class-specific feature selection technique using a deep one-versus-each strategy, enhancing classification explainability by selecting features most relevant to each individual class.

B. Paper Contributions

- A novel methodology for selecting class-specific basis vectors that maximize class separability, using projection-based discriminative basis selection.
- The integration SVD and mean average precision (mAP) analysis to evaluate and rank basis vectors based on their discriminative power.

- A systematic approach to selecting the most informative basis vectors, improving classification performance over traditional heuristic-based selection methods.
- Validation of the proposed methodology through simulations and feature vectors extracted from pre-trained deep learning models (ResNet-18 and ViT-L/16).
- Experimental comparisons showing significant improvements in true positive rates (TPR) and false positive rates (FPR), particularly in cases where class subspaces exhibit overlap.

II. METHOD

Selecting the most representative bases from the subspace of a given class is a crucial step in classification tasks. This section outlines the approach for obtaining basis vectors using SVD and further refining the selection using mAP analysis. The methodology involves decomposing class-specific data matrices, selecting discriminative basis vectors, and utilizing projection-based classification.

A. Singular Value Decomposition (SVD)

Singular Value Decomposition (SVD) is a fundamental technique used to extract basis vectors representing the subspace in which a set of data points reside. By decomposing a matrix into orthonormal components, SVD reveals the underlying structure of the data in a ranked order, where the most significant basis vectors capture the greatest variance.

For a given class i , feature vectors belonging to that class, C_i , are stacked into a matrix M_i , where each column corresponds to an individual data point:

$$M_i = \text{horizontalStack}(C_i) \quad (1)$$

Applying SVD to this matrix results in:

$$U_i \Sigma_i V_i^T = M_i \quad (2)$$

where:

- U_i contains the left singular vectors, representing an orthonormal basis for the subspace.
- Σ_i is a diagonal matrix with singular values, ranking the significance of each basis vector.
- V_i^T provides the right singular vectors, which represent the data points in the transformed space.

The columns of U_i serve as basis vectors of the subspace in which class i resides. If the feature vectors are arranged as rows instead of columns, the rows of V_i^T provide an equivalent basis.

This decomposition facilitates the selection of a subset of basis vectors to represent the class while maintaining essential information. Traditional selection methods choose the first n most significant basis vectors based on singular values. However, this approach does not necessarily optimize class separability.

B. Selecting the Most Differentiating Bases

SVD provides an ordered set of basis vectors, but these vectors do not necessarily correspond to the most discriminative features for class separation. While traditional methods select the first n singular vectors, this approach may include bases that represent intra-class variance rather than inter-class separability. To address this limitation, a selection strategy based on mAP is introduced.

The mAP-based selection process ensures that only the most class-distinguishing bases are chosen. Instead of relying on singular value magnitude, each basis vector is evaluated based on its ability to differentiate between data points belonging to the target class and those from other classes. This process starts by grouping vectors corresponding to the target class as shown previously in Equation (1). Moreover, vectors of all other non-target classes, $\tilde{\mathbf{C}}_i$, are similarly stacked into $\tilde{\mathbf{M}}_i$ matrix as in Equation (3):

$$\tilde{\mathbf{M}}_i = \text{horizontalStack}(\tilde{\mathbf{C}}_i) \quad (3)$$

As a result of the SVD analysis shown in Equation (2), each basis vector $u_{i,k}$ obtained from the columns of $\mathbf{U}_i$ is used to project data points from both the target class and non-target classes. However, instead of directly using the projections, we normalize them by the ℓ_2 norm of each feature vector. The modified projections are given in Equations (4) and (5) for r number of target class vectors and s number of non-target class vectors:

$$X_{i,k} = \left[\frac{u_{i,k}^T c_1}{\|c_1\|_2}, \frac{u_{i,k}^T c_2}{\|c_2\|_2}, \dots, \frac{u_{i,k}^T c_r}{\|c_r\|_2} \right] \quad (4)$$

$$\tilde{X}_{i,k} = \left[\frac{u_{i,k}^T \tilde{c}_1}{\|\tilde{c}_1\|_2}, \frac{u_{i,k}^T \tilde{c}_2}{\|\tilde{c}_2\|_2}, \dots, \frac{u_{i,k}^T \tilde{c}_s}{\|\tilde{c}_s\|_2} \right] \quad (5)$$

The mAP score of each basis vector is then computed by treating the normalized projections of the target class as positive instances and those of the non-target classes as negative instances as shown in Equation (6):

$$mAP_{i,k} = mAP(\mathbf{X}_{i,k} : \tilde{\mathbf{X}}_{i,k}) \quad (6)$$

Then, the basis vectors are sorted in descending order based on their mAP scores. Now, instead of selecting the first n singular vectors, the n basis vectors with the highest mAP scores are chosen.

This method prioritizes bases that maximize inter-class separability while preserving computational efficiency. By selecting bases based on discriminative power rather than variance alone, classification performance is improved significantly.

C. Calculation of Projection Matrix

As the projection matrix of the base method, the most significant n number of columns of $\mathbf{U}_i$ are selected, forming the matrix shown in Equation (7)

$$\mathbf{U}_{i,base} = [\mathbf{u}_1 \quad \mathbf{u}_2 \quad \dots \quad \mathbf{u}_n] \quad (7)$$

Using these basis vectors, the projection matrix $\mathbf{P}_i$ for the n -dimensional subspace of class i is computed as shown in Equation (8):

$$\mathbf{P}_{i,base} = \mathbf{U}_{i,base} \mathbf{U}_{i,base}^T \quad (8)$$

Similarly, basis vectors sorted by their mAP scores are also used to form a basis matrix as shown in Equation (9)

$$\mathbf{U}_{i,mAP} = [\mathbf{u}_{1,mAP} \quad \mathbf{u}_{2,mAP} \quad \dots \quad \mathbf{u}_{n,mAP}] \quad (9)$$

As shown previously, using the matrix composed by the selected base vector, the projection matrix for the selected subspace based on mAP scores is calculated as shown in Equation (10):

$$\mathbf{P}_{i,mAP} = \mathbf{U}_{i,mAP} \mathbf{U}_{i,mAP}^T \quad (10)$$

D. Classification of a Given Test Point

After computing the projection matrix for a class, a given test point is analyzed to determine its class membership. The classification process involves projecting the test vector $\mathbf{z}$ onto each class subspace and measuring its alignment with the class representation. The projection ratio, defined in Equation (11), is used as a similarity measure:

$$d_{i,z} = \frac{\|\mathbf{P}_i \mathbf{z}\|_2}{\|\mathbf{z}\|_2} \quad (11)$$

A test vector is assigned to a class if its projection ratio exceeds a predefined threshold. The determination of this threshold relies on analyzing the projection ratios of both in-class and out-of-class data points. The projection ratios for a given projection matrix $\mathbf{P}_i$ are computed as shown in Equations (12) and (13):

$$X_i = \left[\frac{\|\mathbf{P}_i^T c_1\|_2}{\|c_1\|_2}, \frac{\|\mathbf{P}_i^T c_2\|_2}{\|c_2\|_2}, \dots, \frac{\|\mathbf{P}_i^T c_r\|_2}{\|c_r\|_2} \right] \quad (12)$$

$$\tilde{X}_i = \left[\frac{\|\mathbf{P}_i^T \tilde{c}_1\|_2}{\|\tilde{c}_1\|_2}, \frac{\|\mathbf{P}_i^T \tilde{c}_2\|_2}{\|\tilde{c}_2\|_2}, \dots, \frac{\|\mathbf{P}_i^T \tilde{c}_s\|_2}{\|\tilde{c}_s\|_2} \right] \quad (13)$$

The optimal threshold value should ideally be lower than the smallest value in $\mathbf{X}_i$ and greater than the largest value in $\tilde{\mathbf{X}}_i$. However, this condition is not always met. To find a suitable threshold value, using True Positive (TP), False Negative (FN), False Positive (FP), and True Negative (TN) values, True Positive Rate (TPR) and False Positive Rate (FPR) are calculated for each threshold value ranging from 0 to 1. These metrics are defined in Equations (14) and (15):

$$TPR = \frac{TP}{TP + FN} \quad (14)$$

$$FPR = \frac{FP}{FP + TN} \quad (15)$$

To ensure an optimal balance between classification sensitivity and specificity, a threshold score ts is computed using Equation (16):

$$ts = \sqrt{(1 - TPR)^2 + FPR^2} \quad (16)$$

The value of ts measures the proximity of a given threshold to the ideal classification point, where $TPR = 1$ and $FPR = 0$. The threshold value that minimizes ts is selected as the optimal threshold for each class projection matrix.

III. EXPERIMENTAL SETUP

The proposed classification approach is evaluated using both synthetic data and feature vectors extracted from pre-trained ResNet-18 and ViT/L-16 networks.

A. Synthetic Data Experiments

The synthetic data experiments are structured into two stages:

1) *Stage 1: Distinct Subspaces*: In the first stage, data points are generated from two distinct subspaces of a high-dimensional ambient space. The classification accuracy is evaluated using both the proposed basis selection method and a standard approach that selects the top n singular vectors using SVD.

2) *Stage 2: Overlapping Bases*: In the second stage, synthetic data is generated such that the two subspaces share a subset of principal components. This setup is designed to highlight the advantages of the proposed method, which prioritizes discriminative basis selection over purely ranking singular vectors by magnitude. The classification performance in this scenario is expected to demonstrate the superiority of the proposed approach in handling overlapping feature distributions.

B. Feature Vector Experiment Using a Pre-Trained Network

The feature vectors extracted from pre-trained ResNet-18 [3] and ViT/L-16 [4] trained on ImageNet [30] are further analyzed to assess classification performance. The dataset used for this evaluation is the "A Large Scale Multi-View RGBD Visual Affordance Learning Dataset" [31], which consists of 23,605 images across 37 object categories with 15 affordance labels.

Only RGB images are processed through the pre-trained networks, with the fully connected layers removed. The feature vectors extracted from the final convolutional layers before classification are used as input for the proposed basis selection method. Due to the multi-label nature of affordance classification, overlapping class regions exist, making the proposed method particularly advantageous in distinguishing between multiple affordances.

IV. RESULTS AND DISCUSSION

In Stage 1 of the synthetic data experiments, the parameter n varies from 1 to 8 in a two-class classification scenario. The ambient space has a dimensionality of 175, with each class represented by a 28-dimensional subspace. Each class contains 10,000 data points, with 5,000 used for training (projection matrix and threshold computation) and 5,000 used

for validation. The same data split is maintained for all subsequent experiments.

Synthetic data results are averaged over 100 independent trials, and the mean performance is reported in Table I. In Stage 2, all parameters remain the same, except that the two classes share a 10-dimensional subspace over their 28-dimensional subspaces. Results for ResNet-18 and ViT/L-16 feature vectors are also included.

TABLE I
RESULTS OF THE BASE METHOD AND PROPOSED METHOD

Dataset	Base Method		Proposed Method	
	TPR(%)	FPR(%)	TPR(%)	FPR(%)
Synthetic Stage 1	83.71	9.33	88.65	6.80
Synthetic Stage 2	62.26	30.28	87.73	7.82
ResNet-18 Feature Output	80.36	4.16	86.47	4.26
ViT/L-16 Feature Output	88.77	2.15	92.17	1.67

As expected, the proposed method demonstrates a significant improvement in Stage 2, where the classes share overlapping subspaces. In Stage 1, the proposed method achieves a 4.94% increase in TPR and a 2.53% reduction in FPR. However, in Stage 2, where overlapping bases exist, the improvements are significantly more pronounced: a 25.47% increase in TPR and a 22.46% reduction in FPR.

For real-world data, the proposed method improves TPR by 6.11% for ResNet-18 feature vectors and 3.40% for ViT/L-16 feature vectors. These results, observed consistently across both synthetic and real-world datasets, confirm the robustness and effectiveness of the proposed basis selection method.

V. CONCLUSION

This study proposed a novel basis selection method for subspace classification, addressing challenges in overlapping subspaces where traditional SVD-based selection fails to capture discriminative features effectively. The classification approach involved projecting test vectors onto class-specific subspaces and determining class membership based on an optimized projection threshold.

Experimental validation was performed using synthetic data and feature vectors from pre-trained deep-learning models. In synthetic experiments, the proposed method achieved a 25.47% increase in TPR and a 22.46% reduction in FPR when class subspaces exhibited overlap. On real-world data, using ResNet-18 and ViT/L-16 feature vectors, the method improved TPR by 6.11% and 3.40%, respectively, over the base method.

These results confirm the robustness of the proposed approach in handling complex class structures. Unlike conventional methods that prioritize global variance, the proposed technique selects bases that maximize class separability. Future work may explore its extension to non-linear subspaces and evaluate its applicability to other multi-label classification tasks.

ACKNOWLEDGMENT

The authors thank the Center for Robotics and Artificial Intelligence (ROMER) of Middle East Technical University for its continuous support to this research.

REFERENCES

- [1] A. S. Georghiades, P. N. Belhumeur, and D. J. Kriegman, "From few to many: Illumination cone models for face recognition under variable lighting and pose," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 23, no. 6, pp. 643–660, 2001.
- [2] K. Kanatani, "Motion segmentation by subspace separation and model selection," in *8th International Conference on Computer Vision*, vol. 2, pp. 301–306, 7 2001.
- [3] K. He, X. Zhang, S. Ren, and J. Sun, "Deep residual learning for image recognition," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 770–778, 2016.
- [4] A. Dosovitskiy, L. Beyer, A. Kolesnikov, D. Weissenborn, X. Zhai, T. Unterthiner, M. Dehghani, M. Minderer, G. Heigold, S. Gelly, *et al.*, "An image is worth 16x16 words: Transformers for image recognition at scale," in *International Conference on Learning Representations*, 2021.
- [5] G. H. Golub and C. F. Van Loan, *Matrix Computations*. Johns Hopkins University Press, 4th ed., 2013.
- [6] I. T. Jolliffe, *Principal Component Analysis*. Springer, 2nd ed., 2002.
- [7] K. Pearson, "On lines and planes of closest fit to systems of points in space," *Philosophical Magazine*, vol. 2, no. 11, pp. 559–572, 1901.
- [8] H. Hotelling, "Analysis of a complex of statistical variables into principal components," *Journal of Educational Psychology*, vol. 24, no. 6, pp. 417–441, 1933.
- [9] P. C. Hansen, "The truncated svd as a method for regularization," *BIT Numerical Mathematics*, vol. 27, no. 4, pp. 534–553, 1987.
- [10] M. Brand, "Incremental singular value decomposition of uncertain data with missing values," in *European Conference on Computer Vision*, pp. 707–720, Springer, 2002.
- [11] E. J. Candès, X. Li, Y. Ma, and J. Wright, "Robust principal component analysis?," *Journal of the ACM*, vol. 58, no. 3, pp. 1–37, 2011.
- [12] N. Halko, P.-G. Martinsson, and J. A. Tropp, "Finding structure with randomness: Probabilistic algorithms for constructing approximate matrix decompositions," *SIAM Review*, vol. 53, no. 2, pp. 217–288, 2011.
- [13] H. C. Andrews and C. L. Patterson, "Singular value decomposition (svd) image coding," *IEEE Transactions on Communications*, vol. 24, no. 4, pp. 425–432, 1976.
- [14] Y. Koren, R. Bell, and C. Volinsky, "Matrix factorization techniques for recommender systems," *Computer*, vol. 42, no. 8, pp. 30–37, 2009.
- [15] S. Deerwester, S. T. Dumais, G. W. Furnas, T. K. Landauer, and R. Harshman, "Indexing by latent semantic analysis," *Journal of the American Society for Information Science*, vol. 41, no. 6, pp. 391–407, 1990.
- [16] M. Ringnér, "What is principal component analysis?," *Nature Biotechnology*, vol. 26, pp. 303–304, 2008.
- [17] A. Aldroubi, A. Sekmen, A. B. Koku, and A. F. Cakmak, "Similarity matrix framework for data from union of subspaces," *Applied and Computational Harmonic Analysis*, 2017.
- [18] A. Aldroubi and A. Sekmen, "Reduced row echelon form and non-linear approximation for subspace segmentation and high-dimensional data clustering," *Applied and Computational Harmonic Analysis*, vol. 37, no. 2, pp. 271–287, 2014.
- [19] E. Elhamifar and R. Vidal, "Sparse subspace clustering: Algorithm, theory, and applications," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2013.
- [20] I. Guyon and A. Elisseeff, "An introduction to variable and feature selection," *Journal of Machine Learning Research*, vol. 3, pp. 1157–1182, 2003.
- [21] S. Watanabe, "Karhunen–loève expansion and factor analysis theorems," in *Proc. Fourth Symposium on System Theory*, pp. 283–318, 1965.
- [22] R. Novak, Y. Bahri, D. A. Abolafia, J. Pennington, and J. Sohl-Dickstein, "Sensitivity and generalization in neural networks: An empirical study," in *International Conference on Learning Representations*, 2018.
- [23] R. A. Fisher, "The use of multiple measurements in taxonomic problems," *Annals of Eugenics*, vol. 7, pp. 179–188, 1936.
- [24] J. H. Friedman, "Regularized discriminant analysis," *Journal of the American Statistical Association*, vol. 84, no. 405, pp. 165–175, 1989.
- [25] S. Mika, G. Ratsch, J. Weston, B. Schölkopf, and K.-R. Müller, "Fisher discriminant analysis with kernels," in *Neural Networks for Signal Processing IX*, pp. 41–48, IEEE, 1999.
- [26] S. Pang, S. Ozawa, and N. Kasabov, "Incremental linear discriminant analysis for classification of data streams," *IEEE Transactions on Systems, Man, and Cybernetics, Part B (Cybernetics)*, vol. 35, no. 5, pp. 905–914, 2005.
- [27] Z. Sun, Z. Chen, J. Liu, and Y. Yu, "Multi-class feature selection via sparse softmax with a discriminative regularization," *[Journal Name]*, 2021.
- [28] X. Zhu, H.-I. Suk, and D. Shen, "Sparse discriminative feature selection for multi-class alzheimer's disease classification," *[Journal Name]*, 2020.
- [29] J. S. Aguilar-Ruiz, "Class-specific feature selection for classification explainability," *[Journal Name]*, 2022.
- [30] O. Russakovsky, J. Deng, H. Su, J. Krause, S. Satheesh, S. Ma, Z. Huang, A. Karpathy, A. Khosla, M. Bernstein, A. C. Berg, and L. Fei-Fei, "ImageNet Large Scale Visual Recognition Challenge," *International Journal of Computer Vision (IJCV)*, vol. 115, no. 3, pp. 211–252, 2015.
- [31] Z. Khalifa and S. A. A. Shah, "A large scale multi-view rgbd visual affordance learning dataset," pp. 1325–1329, 2023.