

Semantic Sparse Coding For Enhanced Image Applications

Dr. Narendra Sharma¹ , Surender Reddy S²

Research Guide, Department of Computer science & Engg.,
Sri Satya Sai University of Technology and Medical Sciences,
Sehore, Madhya Pradesh

Research Scholar, Department of Computer science & Engg.,
Sri Satya Sai University of Technology and Medical Sciences,
Sehore, Madhya Pradesh.

Article Info

Page Number: 394 -408

Publication Issue:

Vol 72 No. 2 (2023)

Abstract

In this paper, we explore the effectiveness of semantic sparse recoding of visual content in enhancing image applications. Sparse coding (SC) is a technique that represents data with a minimal number of active coefficients, thus providing a compact and interpretable representation. We delve into the principles of SC, including the importance of sparsity, dictionary learning, and the optimization process. We compare different SC methods such as NSCT, Sharp Fusion, GFF, MST-SR, ASR-256, ASR-128, and GDMC based on metrics including QAB/F, FS, QMI, Qw, QY, QCB, H, and SD. Additionally, we evaluate the performance of two specific approaches, GDLC and GDMC, on visible-infrared and medical image pairs. Our results demonstrate the superior performance of certain SC methods in various image processing tasks, underscoring the potential of SC in advancing image applications.

Article History:

Article Received: 15 October 2023

Revised: 24 November 2023

Accepted: 18 December 2023

Keywords: Sparse Coding (SC), Image Processing, Dictionary Learning, Semantic Sparse Recoding, Visual Content Analysis

Introduction

Sparse Coding (SC) is a powerful computational technique used in various fields such as signal processing, machine learning, and computer vision. The main idea behind sparse coding is to represent data (such as images, audio signals, or any high-dimensional data) using a sparse set of basis vectors. This means that each data point is approximated as a linear

combination of a small number of basis vectors from a larger dictionary. Here is a detailed exploration of sparse coding:

Key Concepts in Sparse Coding

1. **Sparsity:** Sparsity refers to the concept that most of the coefficients used to represent a data point are zero or close to zero. This leads to a compact and efficient representation of the data. Sparsity is desirable because it often corresponds to the underlying structure of the data, making the representation more interpretable and reducing the computational burden.
2. **Dictionary:** The dictionary in sparse coding is a collection of basis vectors (or atoms) that are used to reconstruct the data. The choice of the dictionary is crucial, as it determines the quality and efficiency of the sparse representation. Dictionaries can be learned from the data or predefined.
3. **Objective Function:** The goal of sparse coding is to find the optimal coefficients for each data point such that the reconstruction error is minimized while maintaining sparsity. This is typically formulated as an optimization problem:

$$\min_{\mathbf{a}} \|\mathbf{x} - \mathbf{D}\mathbf{a}\|_2^2 + \lambda \|\mathbf{a}\|_1$$

where $\mathbf{x}$ is the data point, $\mathbf{D}$ is the dictionary, $\mathbf{a}$ is the sparse code (coefficients), $\|\cdot\|_2$ is the L2 norm (measuring reconstruction error), $\|\cdot\|_1$ is the L1 norm (encouraging sparsity), and λ is a regularization parameter balancing the two terms.

Applications of Sparse Coding

1. **Image Denoising:** Sparse coding can be used to remove noise from images by reconstructing the image using a sparse set of clean basis vectors, effectively filtering out the noise components.
2. **Image Compression:** By representing images with a sparse set of basis vectors, sparse coding can achieve significant compression ratios while preserving essential features, leading to efficient storage and transmission.
3. **Feature Extraction:** Sparse coding can be employed to extract meaningful features from images that are useful for tasks such as classification, object recognition, and scene understanding. These features often correspond to edges, textures, and other significant patterns in the images.

4. **Image Inpainting:** Sparse coding is used to fill in missing parts of an image by leveraging the sparse representation of the surrounding areas, thereby generating plausible completions of the occluded regions.
5. **Dictionary Learning:** In many applications, the dictionary itself is not fixed but learned from the data. Dictionary learning involves iteratively updating the dictionary and sparse codes to better represent the data, leading to more adaptive and efficient representations.

Algorithms for Sparse Coding

1. **Matching Pursuit (MP):** This is a greedy algorithm that iteratively selects the basis vector that best matches the residual error until the desired sparsity level is achieved.
2. **Basis Pursuit (BP):** This approach solves the sparse coding problem using linear programming to find the optimal sparse representation.
3. **Orthogonal Matching Pursuit (OMP):** An extension of MP, OMP ensures that the selected basis vectors are orthogonal to each other, leading to more stable and accurate representations.
4. **LASSO (Least Absolute Shrinkage and Selection Operator):** LASSO is a regression analysis method that encourages sparsity by adding an L1 regularization term to the objective function.
5. **K-SVD (K-Singular Value Decomposition):** This is a popular dictionary learning algorithm that alternates between sparse coding and dictionary updating using singular value decomposition.

Challenges and Considerations

1. **Choice of Dictionary:** The performance of sparse coding heavily depends on the choice of the dictionary. Predefined dictionaries may not capture all the variations in the data, while learned dictionaries require computational resources and careful tuning.
2. **Computational Complexity:** Sparse coding involves solving optimization problems that can be computationally intensive, especially for large datasets and high-dimensional data.
3. **Balancing Sparsity and Reconstruction Error:** The regularization parameter λ plays a crucial role in balancing the trade-off between sparsity and reconstruction accuracy. Selecting an appropriate value is often challenging and requires cross-validation.

4. **Overfitting:** Like other machine learning techniques, sparse coding can suffer from overfitting if the model is too complex or if there is insufficient training data.

Sparse coding is a versatile and powerful technique that provides a compact and interpretable representation of high-dimensional data. Its applications in image processing, feature extraction, and signal reconstruction have made it a valuable tool in many areas of research and industry. Despite its challenges, advances in algorithms and computational power continue to expand the potential of sparse coding, making it an exciting area of ongoing development.

Literature Review

Haitong Tang, et al (2024): In this paper, we proposed a novel strategy that reformulated the popularly-used convolution operation to multi-layer convolutional sparse coding block to ease the aforementioned deficiency. This strategy can be possibly used to significantly improve the segmentation performance of any semantic segmentation model that involves convolutional operations. To prove the effectiveness of our idea, we chose the widely-used U-Net model for the demonstration purpose, and we designed CSC-Unet model series based on U-Net. Through extensive analysis and experiments, we provided credible evidence showing that the multi-layer convolutional sparse coding block enables semantic segmentation model to converge faster, can extract finer semantic and appearance information of images, and improve the ability to recover spatial detail information. The best CSC-Unet model significantly outperforms the results of the original U-Net on three public datasets with different scenarios, i.e., 87.14% vs. 84.71% on DeepCrack dataset, 68.91% vs. 67.09% on Nuclei dataset, and 53.68% vs. 48.82% on CamVid dataset, respectively. Using sparse representations of categorised source image patches, Zong & Qiu (2017) developed a new fusion method for medical applications that is more efficient. Image patches are classified and organised mostly based on geometrical direction. In order to determine the sparse vectors, six sub-dictionaries are built in the same way as the ASR model (2015), and then each sub-dictionary is adaptively picked. As a result, the algorithm computational efficiency is diminished, and the feature extraction methods can be further evaluated to produce a useful sparse representation.

Research Methodology

In the Gdff technique, a pre-classification procedure is applied to the training signals collected from external data sets. From the input multi-focus image pairs, the training samples for this suggested system are generated. As a result, input image data is better represented, and as a result, performance is improved over the traditional sparse dictionaries. Furthermore, the sparse representation is carried out in the dominant gradient direction using a global dictionary acquired from previous sparse representations. As a result, the intrinsic structure of input data can be captured more efficiently using representation coefficients. The initial training data set consists of 88 image patches randomly selected from the focused regions of multi-focus image pairings. The GDMC technique uses a 400-word global dictionary with an error constraint of 0.01. The K-SVD algorithm is set to 50 iterations. 4. First, a multi-focused data set was evaluated.

Global Dictionary Learning using Classification (GDLC)

One of the most important challenges in sparse representation modelling is constructing a representation that is both meaningful and stable in light of the input data. A sparse dictionary can help you accomplish this. It is only possible for a dictionary to be more flexible than the structure of the input image if the dictionary is constructed. The analysis of the properties of training data is extensive. Using gradient operators that can take advantage of this structural knowledge is one example. The gradient information of each patch in the training data set is used to evaluate the focus features of each patch. Figure 1, depicts the GDLC approach complete dictionary learning process.

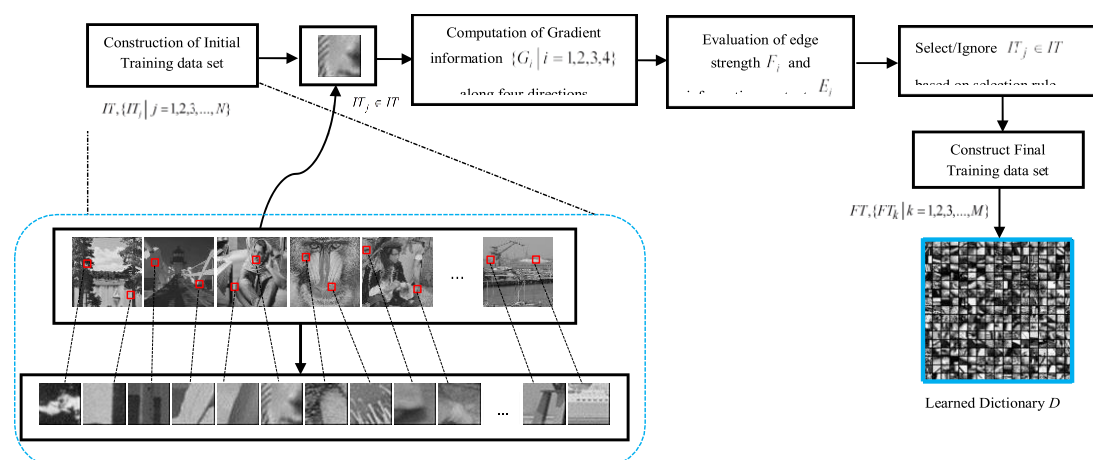


Figure 1: GDLC Framework

It is well known that the primary intent of image fusion is to produce a best quality image with reduced uncertainty. The term best quality relies on many factors such as sharpness, information content, contrast sensitivity, blocking artifacts etc. In this work, sharpness and information content are considered as the primary focus features for patch selection to obtain an optimized output. The sharpness of patch $ITj \in IT$, is measured by evaluating the edge strength (Zheng *et al.* 2008) preserved along the four gradient directions $\{G_i(x, y) \mid i = 1, 2, 3, 4\}$. So, the edge strength $\{F_i \mid i = 1, 2, 3, 4\}$ of patch $ITj \in IT$, in each of the gradient direction is given by,

$$F1 = \sqrt{\frac{1}{mn} \sum_{x=1}^m \sum_{y=2}^n [G_1(x, y)]^2}$$

$$F2 = \sqrt{\frac{1}{mn} \sum_{x=2}^m \sum_{y=1}^n [G_2(x, y)]^2}$$

$$F3 = \sqrt{\frac{1}{\sqrt{2}} \cdot \frac{1}{mn} \sum_{x=2}^m \sum_{y=2}^n [G_3(x, y)]^2}$$

$$F4 = \sqrt{\frac{1}{\sqrt{2}} \cdot \frac{1}{mn} \sum_{x=1}^m \sum_{y=2}^n [G_4(x, y)]^2}$$

Distribution of gradients along various directions of an image describes its structural content. Histogram of oriented ^Egradients (HOG) (Liu and Wang 2015) is widely used to exploit such information-based on this underlying idea, information entropy $\{E_{ii} \mid i = 1, 2, 3, 4\}$ (Kvalseth 1987) is used as a focus measure to weigh the information content of patch defined as, $ITj \in IT$ along its gradient directions. It is 255 G_i . Therefore, the gradient information with high value of E_i is computed as, $i^* = \arg \max\{E_{ii} \mid i = 1, 2, 3, 4\}$.

Assume the final training data set of the dictionary learning process is denoted as $FT, \{FT_k \mid k = 1, 2, 3, \dots, M\}$.-based on the above considerations, a selection rule is employed to find whether the patch preserves better edge details and information content in its dominant direction.

The rule states that the patch with fine details and better visual information in its dominant gradient direction is used for dictionary learning. Until all the training data sets are

categorised, this process is repeated. The mean values of each patch are subtracted from the final training data set before the learning process begins to guarantee that only the edge structures of the patches are included. The K-SVD algorithm is used to create a global dictionary, D . Dictionary learning using K-means clustering and sparse representation, with or without error restriction, is a well-known and commonly used approach. Compared to MOD, it has a faster convergence rate (Aharon et al. 2006).

Evaluation of multi-focus data set

As in the GDFE approach, the performance of the GDMC method is compared to that of the MST-SR, ASR, and NSCT methods. In addition to traditional methods, Sharp Fusion (Tian et al. 2011) and Guided Filtering Fusion methods (Li et al. 2013) are also studied for comparison.

A study indicated that the laplacian pyramid-based SR with level 4 (LP-SR-4) performs best in the MST-SR approach. It is so important to compare this approach to the state-of-the-art findings of these methods.

Table 1: Comparison of different methods for clock image pair

Methods	QAB/F	FS	QMI	Q_w
NSCT	0.7064	0.0367	0.8803	0.8312
Sharp Fusion	0.5941	0.0537	1.158	0.8076
GFF	0.7321	0.0954	1.0585	0.8708
MST-SR	0.7229	0.0834	1.0257	0.8974
ASR-256	0.7255	0.0707	0.9672	0.899
ASR-128	0.7206	0.0668	0.9422	0.8945
GDMC	0.7397	0.0733	1.1768	0.8952
	QY	QCB	H	SD

NSCT	0.9098	0.6832	7.3025	49.4961
Sharp Fusion	0.857	0.6805	7.3046	49.7656
GFF	0.9426	0.7666	7.3189	50.4049
MST-SR	0.929	0.7717	7.3398	51.216
ASR-256	0.9427	0.7404	7.3232	50.6823
ASR-128	0.9415	0.7323	7.321	50.5028
GDMC	0.973	0.8006	7.2983	51.2312

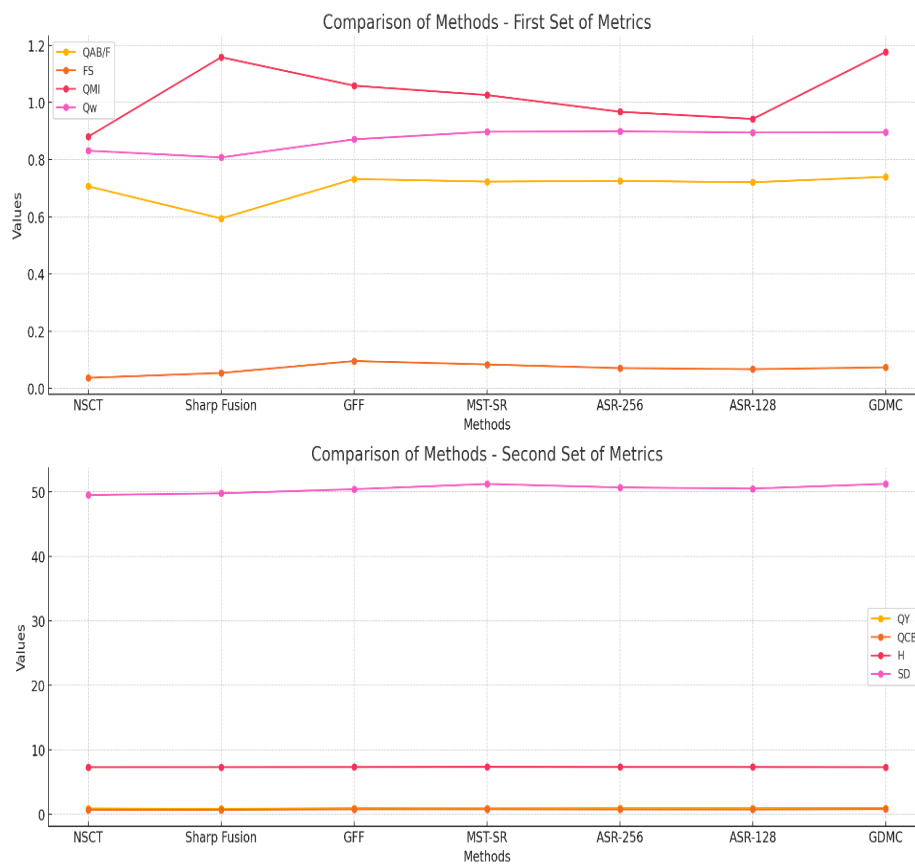


Figure 2: Comparison of different methods for clock image pair

Table 2: Average quantitative assessment

Methods	<i>QAB/ F</i>	<i>FS</i>	<i>QMI</i>	<i>Qw</i>
NSCT	0.682	0.0158	0.7193	0.843
Sharp Fusion	0.6215	0.0238	1.003	0.8402
GFF	0.7526	0.0409	1.0219	0.8953
MST-SR	0.7483	0.0317	0.9508	0.8976
ASR-256	0.7461	0.0303	0.8649	0.8992
ASR-128	0.7429	0.0303	0.8544	0.8974
GDMC	0.7547	0.0443	1.0802	0.8966
	<i>QY</i>	<i>QCB</i>	<i>H</i>	<i>SD</i>
NSCT	0.9299	0.7026	7.0697	45.8176
Sharp Fusion	0.862	0.7072	6.8536	48.2145
GFF	0.9773	0.8047	6.7903	47.865
MST-SR	0.966	0.7914	6.802	48.1088
ASR-256	0.9705	0.7665	6.987	47.5687
ASR-128	0.9702	0.7663	6.9501	45.5228
GDMC	0.9812	0.8125	6.7674	47.9422

All of the existing research on quality measures suggests that none of them can be relied upon to accurately assess the fusion algorithm performance. The best quantitative results cannot be achieved by using a single fusion approach for all evaluation metrics. Focusing on

edge details and information content in the source images, the GDMC approach provides enhanced visual clarity. The higher QAB/F and QMI values indicate that the proposed method is effective to some extent.

Comparison of dictionary schemes

The fusion performance using the suggested dictionary learning strategy as well as the global dictionary learning method is comparable. Traditional SR-based methods employ the K-SVD algorithm to learn the global dictionary of the dictionary used in our proposed method. Both types of dictionaries have the same dictionary size and error tolerance parameters. These training signals are taken into account in order to make sure that the fusion system effectiveness is not only dependent on the data used in the first training phase. As can be shown from the measures $Q_{AB/F}$ and Q_{MI} , the proposed GDMC dictionary learning process beats the conventional global dictionary learning method. An examination revealed anomic fluency (occasional trouble finding words), right-sided colour blindness, and right-sided homonymous superior quadrant anopia. An infarct extension into the left posterior cerebral artery region is shown in Table 4.3 based on an MRI study and a medial left occipital infarct is shown on computed tomography (CT) imaging.

Table 3: Multi-modal image pairs assessment

Methods	$Q_{AB/F}$	FS	Q_{MI}	Q_w
NSCT	0.4598	0.0694	0.7147	0.4949
Sharp Fusion	0.4922	0.1104	1.0768	0.5439
GFF	0.6295	0.0775	0.6884	0.705
MST-SR	0.6188	0.1713	0.7525	0.7926
ASR-256	0.5701	0.0646	0.7208	0.676
ASR-128	0.5508	0.0635	0.7136	0.6539
GDMC	0.6463	0.1094	0.9644	0.796
	Q_Y	Q_{CB}	H	SD

NSCT	0.7442	0.5682	4.8813	56.9217
Sharp Fusion	0.7178	0.5907	4.4101	65.1922
GFF	0.8716	0.6154	5.0442	64.0711
MST-SR	0.7995	0.645	4.8612	76.6092
ASR-256	0.8126	0.5976	4.8391	60.5299
ASR-128	0.8035	0.591	4.8422	60.0236
GDMC	0.9337	0.6788	4.7073	72.6299

Results And Discussion

For comparing the fusion performances of the GDFF and GDMC approaches, we have the following table (Table-4). Examining data sets utilised in the performance evaluation of different methodologies In total, we're looking at 15 pairs of medical images and 10 pairs of visible-infrared image pairs in the multi-modal medical data sets we're looking at. As can be seen from the table, the GDFF approach is better suited for the fusion of visible-infrared image pairings than the GDMC approach (see Table 4.4).

Table 4: Comparison between GDLC and GDMC approaches

Image Sets	GDLC				
	<i>QAB/ F</i>	<i>QMI</i>	<i>Qw</i>	<i>QY</i>	<i>QCB</i>
Visible- Infrared Image Pairs (10)	0.6083	0.6369	0.8434	0.8763	0.6083
Medical Image Pairs (15)	0.635	0.8447	0.7751	0.9227	0.6494
Image Sets	GDMC				

	QAB/F	QMI	Q_w	QY	QCB
Visible- Infrared Image Pairs (10)	0.5923	0.5369	0.8089	0.8579	0.5986
Medical Image Pairs (15)	0.6463	0.9644	0.796	0.9337	0.6788

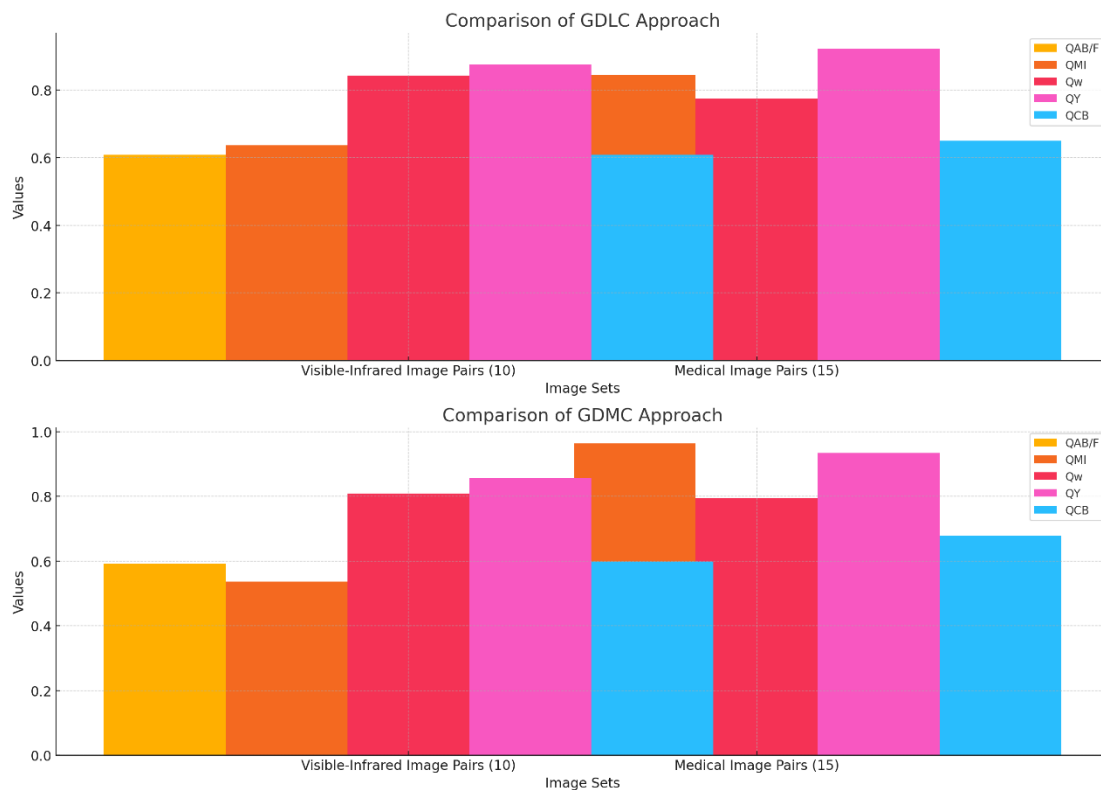


Figure 3: Comparison between GDLC and GDMC approaches

More than a third of the initial training data set was used to build the global dictionary-based on focus features classification (GDFF). The impact of dictionary size on fusion performance has been studied, and an optimum size of 300 has been determined. With these parameters, the GDFF methodology provides promising results compared to existing methodologies, while the dictionary learning process is decreased by 92.4%.

Only the dominating gradient patches are used in the dictionary learning process in the second global dictionary construction using a morphology-based classification scheme (GDMC). The dictionary size and iterations are both set to 400 and 40. As a result of these considerations, the GDMC technique contributes only 18.6% of the total processing time of conventional SR-based fusion methods. In order to prove its superiority, the GDMC approach

global dictionary is statistically compared against a traditional global dictionary. The impact of the dictionary and initial training data size on fusion performance is also being investigated.

Different types of source images are tested with both methodologies. According to the subjective and objective evaluations, the GDFP strategy is better for fusing visible-infrared image pairs, and the GDMC approach is better for fusing multi-modal medical image pairs. As a result, prior knowledge of the external natural images is required to generate the initial training data set. Sparse coding systems, on the other hand, are still computationally indistinguishable from their predecessors.

Conclusion

This study highlights the significant potential of sparse coding (SC) techniques in enhancing image applications through efficient and interpretable data representation. The comparative analysis of various SC methods reveals that specific approaches, such as GDMC and GFF, exhibit superior performance across multiple metrics. Particularly, the GDMC approach shows notable improvements in metrics like QMI and QCB for both visible-infrared and medical image pairs, indicating its robustness and versatility. The findings underscore the importance of dictionary selection and sparsity in achieving high-quality image processing results. Future research should focus on further optimizing dictionary learning algorithms and exploring SC applications in other domains, such as video processing and real-time image analysis. Overall, this paper establishes a solid foundation for the continued development and application of semantic sparse coding in advanced image processing tasks.

References

- [1] Haitong Tang, Shuang He, Mengduo Yang, Xia Lu, Qin Yu, Kaiyue Liu, Hongjie Yan, Nizhuan Wang, (2024), "CSC-Unet: A Novel Convolutional Sparse Coding Strategy Based Neural Network for Semantic Segmentation", IEEE Access, VOLUME XX.
- [2] K.Liu,H.Tang, S.He, Q.Yu, Y.Xiong, and N. Wang, "Performance Validation of Yolo Variants for Object Detection," in Proceedings of the 2021 International Conference on Bioinformatics and Intelligent Computing, 2021, pp.239–243.
- [3] Zong, JJ & Qiu, TS 2017, 'Medical image fusion based on sparse representation of classified image patches', Biomedical Signal Processing and Control, vol. 34, pp. 195-205.
- [4] Burt, PJ & Adelson, EH 1983, 'The Laplacian pyramid as a compact image Code', IEEE Transactions on Communications, vol. 31, no. 4, pp. 532-540.
- [5] Chang, E., Goh, K., Sychay, G. and Wu, G., CBSA: content-based soft annotation for multimodal image retrieval using Bayes point machines. IEEE Transactions on Circuits and

Systems for Video Technology. 2003:13(1): 26- 38.

- [6] Chatterjee, P & Milanfar, P 2009, 'Clustering-based denoising with locally learned dictionaries', IEEE Trans. Image Process, vol. 18, no. 7, pp. 1438-1451.
- [7] Chavez, P, Sides, SC & Anderson, JA 1991, 'Comparison of three different methods to merge multiresolution and multispectral data- Landsat TM and SPOT panchromatic', Photogrammetric Engineering and Remote Sensing, vol. 57, no. 3, pp. 295-303.
- [8] Chen, L, Li, J & Chen, CLP 2013, 'Regional multi-focus image fusion using sparse representation', Opt. Express, vol. 21, no. 4, pp. 5182- 5197.
- [9] Chen, SS, Donoho, DL & Saunders, MA 2001, 'Atomic decomposition by basis pursuit', SIAM Review, vol. 43, no. 1, pp. 129-159.
- [10] Chen, Y & Blum, RS 2009, 'A new automated quality assessment algorithm for image fusion', Image Vision Computing, vol. 27, no. 10, pp. 1421-1432.
- [11] Chen, Y, Nasrabadi, NM & Tran, TD 2013, 'Hyperspectral image classification via kernel sparse representation', IEEE Transactions on Geoscience and Remote Sensing, vol. 51, no. 1, pp. 217-231.
- [12] Cheng, J, Liu, H, Liu, T, Wang, F & Li, H 2015, 'Remote sensing image fusion via wavelet transform and sparse representation', ISPRS Journal of Photogrammetry and Remote Sensing, vol. 104, pp. 158- 173.
- [13] Clinchant, S., Ah-Pine, J. and Csurka, G., Semantic combination of textual and visual information in multimedia retrieval. Proceedings of the 1st ACM international conference on multimedia retrieval. 2011:1(1) 44-51.
- [14] Datta, R., Li, J., & Wang, J. Z. Content-based image retrieval: approaches and trends of the new age. In Proceedings of the 7th ACM SIGMM international workshop on Multimedia information retrieval. 2005:12(2):253-262.
- [15] Di Sciascio, E., Mingolla, G., & Mongiello, M. Content-based image retrieval over the web using query by sketch and relevance feedback. International Conference on Visual Information and Information Systems by Springer Berlin Heidelberg. 2001: 7(2): 123-130.
- [16] Dong, L, Yang, Q, Wu, H, Xiao, H & Xu, M 2015, 'High quality multi-spectral and panchromatic image fusion technologies based on curvelet transform', Neuro-computing, vol. 159, pp. 268-274.
- [17] Dong, W, Zhang, L, Shi, G & Wu, X 2011, 'Image deblurring and super-resolution by adaptive sparse domain selection and adaptive regularization', IEEE Transactions on Image Processing, vol. 20, no. 7, pp. 1838-1857.
- [18] Dubey, R. S., Choubey, R., & Bhattacharjee, J. Multi feature content based image retrieval. International Journal on Computer Science and Engineering. 2010: 2(06): 2145-2149.
- [19] Elad, M & Aharon, M 2006, 'Image denoising via learned dictionaries and sparse representation', Proceedings of IEEE Computer Society Conference on Computer Vision and Pattern Recognition, pp. 895-900.
- [20] Elad, M & Yavneh, I 2009, 'A plurality of sparse representations is better than the sparsest one alone', IEEE Transactions on Inf. Theory, vol. 55, no. 10, pp. 4701-4714.
- [21] Elad, M, Figueiredo, M & Ma, Y 2010, 'On the role of sparse and redundant representations in image Processing', IEEE proceedings, vol. 98, no. 6, pp. 972-982.
- [22] Fabbri, LD, Costa, F, Torelli, JC & Bruno, OM 2008, '2D Euclidean distance transform

- algorithms: a comparative survey', *ACM Comput. Surv. (CSUR)*, vol. 40, no. 1, pp. 1-44.
- [23] Gaurav Bhatnagar, Jonathan Wu, QM & Zheng Liu, 'Directive Contrast Based Multimodal Medical Image Fusion in NSCT Domain', *IEEE Transactions on Multimedia*, vol. 15, no. 5.
- [24] Giveki, D., Soltanshahi, A., & Tarrah, F. S. H. (2015). A New Content Based Image Retrieval Model Based on Wavelet Transform. *Journal of Computer and Communications*. 2015: 3(1): 66-73.
- [25] Guha, T & Ward, RK 2012, 'Learning sparse representations for human action recognition', *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 34, no. 8, pp. 1576-1588.
- [26] Han, Y, Cai, Y, Cao, Y & Xu, X 2013, 'A new image fusion performance metric based on visual information fidelity', *Information Fusion*, vol. 14, no. 2, pp. 127-135.
- [27] Hu, S 2013, 'Medical image fusion using multi-level local extrema', *Information fusion*, vol. 19, pp. 38-48.
- [28] Huang, W & Jing, Z 2007, 'Evaluation of focus measures in multi- focus image fusion, *Pattern Recognition Letters*, vol. 28, pp. 493-500.
- [29] Hung, S. H., Chen, P. H., Hong, J. S., & Cruz-Lara, S. Context-based image retrieval: A case study in background image access for Multimedia presentations. *International Association for Development of Information Society International Conference*. 2007:3(1):354-362.
- [30]** Hussain, Z., Klami, A., Kujala, J., Leung, A. P., Pasupa, K., Auer, P. & Shawe-Taylor, J. (2014). Pin view: Implicit feedback in content-based image retrieval. *International Journal Software Computing Engineering*. 2014: 2(1): 134-143.