

SPARSE REPRESENTATION BASED INPAINTING FOR THE RESTORATION OF DOCUMENT IMAGES AFFECTED BY BLEED-THROUGH

Muhammad Hanif, Anna Tonazzini, Pasquale Savino, and Emanuele Salerno

Istituto di Scienza e Tecnologie dell'Informazione
Consiglio Nazionale delle Ricerche
Via G. Moruzzi, 1, 56124 Pisa, Italy

ABSTRACT

Bleed-through is a commonly encountered degradation in ancient printed documents and manuscripts, which severely impair their readability. Digital image restoration techniques can be effective to remove or significantly reduce this degradation. In bleed-through document image restoration the main issue is to identify the bleed-through pixels and replace them with appropriate values, in accordance to their surroundings. In this paper, we propose a two stage method, where a pair of properly registered images of the document recto and verso is first used to locate the bleed-through pixels in each side, and then a sparse representation based image inpainting technique is used to fill-in the bleed-through areas according to the neighbourhood, in such a way to preserve the original appearance of the document. The advantages of the proposed inpainting technique over state-of-the-art methods are illustrated by the improvement in the visual results.

Index Terms— Bleed-through image restoration, dictionary learning, sparse representation, ancient documents restoration.

1. INTRODUCTION

Ancient documents are valuable sources of history, culture, and heritage information. Nevertheless, when the ancient manuscripts are written on both pages (recto and verso) of the sheet, often the ink penetrates through the page due to aging, humidity, ink chemical property or paper porosity, and becomes visible as degradation on the other side. This type of degradation is termed as bleed-through, and impairs legibility and interpretation of the document contents. In the past, physical restoration techniques were applied to remove the bleed-through degradation, but unfortunately they are costly, invasive, and can cause permanent damage to the document. Nowadays, ancient documents are usually archived in the form of digital images for preservation, distribution, retrieval, or analysis. Thus, digital image processing techniques have become a standard for their ability to perform any number

of alterations to the document appearance, while leaving the original intact, and have been used with impressive results to remove bleed-through. Indeed, besides being necessary to improve readability, the removal of such degradation is a critical preprocessing step in many tasks such as optical character recognition (OCR), feature extraction, segmentation and automatic transcription.

Generally, bleed-through restoration is handled as a classification problem, where image pixels are labeled as either background, bleed-through, or foreground (original text) [1]. Broadly speaking, the approaches to bleed-through restoration can be divided into two categories: blind or single-sided and non-blind or double-sided. In blind methods only one side of the document is used, whereas non-blind methods exploit accurately pre-registered recto and verso images of the document. Most of the earlier blind methods involve an intensity based thresholding step [2]. However, thresholding is not suitable when the aim is to preserve the original appearance of the document. In order to remove bleed-through selectively, in [3] an independent component analysis (ICA) based method is applied to an RGB image to separate the foreground, background, and bleed-through layers. A dual-layer Markov random field (MRF) is suggested in [4], whereas in [5] a conditional random field (CRF) based method is proposed. A non-blind ICA based method is outlined in [6]. Other methods of this category are proposed in [7][8][9].

A common drawback of all the methods above is the search for an optimal way to replace the bleed-through pattern to be removed. In some of those papers, as a preliminary step, a “clean” background for the entire image is estimated (e.g. [4][9]), but this is usually a very laborious task.

In this paper, we propose a two-step method to address bleed-through document restoration from pre-registered recto and verso images. According to the point of view of document appearance preservation, the two steps consist in the selective identification of the bleed-through pattern, followed by its inpainting with the texture of the close background. Bleed-through identification can be performed using, in principle, any of the methods proposed in the literature for this purpose. Here we adopt the bleed-through cancellation algorithm de-

scribed in [10], which is simple and very fast. Although efficient in locating bleed-through, also this algorithm suffers from the lack of a proper strategy to replace the unwanted pixels. Indeed, these are assigned with the predominant graylevel of the background. This causes imprints of the bleed-through pattern that are visible in the restored image. An interpolation based inpainting technique for such imprints is presented in [11], but the filled-in areas are mostly smooth. Here, in the second step of our method, we use a sparse image representation based inpainting to find a befitting fill-in for the bleed-through imprints, according to their neighborhood. This inpainting step, which constitutes the main contribution of the paper, enhances the quality of the bleed-through restored image and preserves the natural paper texture. The rest of this paper is organized as follows. Next section briefly introduces sparse image representation and dictionary learning. The non-blind bleed-through identification is presented in Section 3. Section 4 describes the proposed inpainting technique and the comparative results are discussed in Section 5. Section 6 presents some future prospects.

2. SPARSE DICTIONARY LEARNING

Sparse representation based modeling emerges as a powerful tool in the image processing field [12], with applications in compression, restoration, inpainting, clustering, recognition, and more. The underlying assumption in this direction is that images contain highly redundant information, therefore they can be approximated by sparse representations of prototype signal-atoms such that very few samples are sufficient to reconstruct a whole image [13]. The objective in this case is to find a basis, called dictionary, which can be used to represent the given image data by a sparse combination of a few elements of the learned dictionary. Formally, for a given set of signals $\mathbf{Y} = [\mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_N]$ where $\mathbf{y}_i \in \mathbb{R}^n$, the sparse representation consists in learning a suitable dictionary, $\mathbf{D} \in \mathbb{R}^{n \times K}$ (with $n < K \ll N$), and the sparse vectors with minimal number of non-zero coefficients, $\mathbf{X} \in \mathbb{R}^{K \times N}$, such that $\mathbf{y}_i \approx \mathbf{D}\mathbf{x}_i$, where $\mathbf{x}_i \in \mathbf{X}$. In this setting, the sparse decomposition can be formulated as an optimization problem of the following form:

$$\{\mathbf{D}, \hat{\mathbf{X}}\} = \arg \min_{\mathbf{D}, \mathbf{x}_i} \|\mathbf{Y} - \mathbf{D}\mathbf{X}\|_F^2; \text{ s.t. } \forall i \|\mathbf{x}_i\|_p \leq z \quad (1)$$

where z is the desired sparsity level in the coefficients $\mathbf{X}$ and $\|\cdot\|_p$ is the l_p norm, with $0 \leq p \leq 1$. An iterative-alternate approach is a natural way to find a solution, by minimizing with respect to one set of variables while keeping the other fixed until some convergence condition is fulfilled [14][13]. The sparse vectors indicate the dictionary atoms that are used to generate each signal. Finding the exact sparse representation for signals, given a dictionary, is a NP-hard problem [15], however using approximate solutions has proven to be a good choice. Commonly used sparse approximation algorithms are

Matching Pursuit (MP) [16], Basis Pursuits (BP)[17], Focal Undetermined System Solver (FOCUSS)[18] and Orthogonal Matching Pursuit (OMP)[19].

The dictionary that can lead to sparse decomposition can either be selected from a set of pre-defined transforms like Wavelets, Fourier, Cosine etc., or can be adaptively learned from the set of training signals [13][20]. Although working with pre-defined dictionaries may be simple and fast, their performance might be not good for every case due to their global-adaptivity nature. Instead, learned dictionaries are adaptive to both the signals and the processing task at hand, thus resulting in far better performances [21][22].

For a given set of signals $\mathbf{Y}$ and dictionary $\mathbf{D}$, the dictionary learning algorithms generate a representation of signal $\mathbf{y}_i$ as a sparse linear combination of the atoms $\mathbf{d}_k$ for $k = 1, \dots, K$,

$$\hat{\mathbf{y}}_i = \mathbf{D}\mathbf{x}_i \quad (2)$$

Dictionary learning algorithms distinguish themselves from traditional model-based methods by the fact that, in addition to $\mathbf{x}_i$, they also train the dictionary $\mathbf{D}$ to better fit the data set $\mathbf{Y}$. The solution is generated by iteratively alternating between the sparse coding stage

$$\hat{\mathbf{x}}_i = \arg \min_{\mathbf{x}_i} \|\mathbf{y}_i - \mathbf{D}\mathbf{x}_i\|^2; \text{ subject to } \|\mathbf{x}_i\|_0 \leq z \quad (3)$$

for $i = 1, \dots, N$ and the dictionary update stage for the $\mathbf{X}$ obtained from the sparse coding stage

$$\mathbf{D} = \arg \min_{\mathbf{D}} \|\mathbf{Y} - \mathbf{D}\mathbf{X}\|_F^2. \quad (4)$$

Dictionary learning algorithms are often sensitive to the choice of z . The update step can either be sequential as in [14] or parallel as in [23]. In sequential dictionary learning, the dictionary update minimization problem (4) is split into K sequential minimizations, optimizing the cost function (4) for each individual atom while keeping fixed the remaining atoms. In the method proposed in [14], which has become a benchmark in dictionary learning, each column $\mathbf{d}_k$ of $\mathbf{D}$ and its corresponding row of coefficients $\mathbf{x}_k^{row}$ are updated based on a rank-1 matrix approximation of the error for all the signals when $\mathbf{d}_k$ is removed

$$\begin{aligned} \{\mathbf{d}_k, \mathbf{x}_k^{row}\} &= \arg \min_{\mathbf{d}_k, \mathbf{x}_k^{row}} \|\mathbf{Y} - \mathbf{D}\mathbf{X}\|_F^2 \\ &= \arg \min_{\mathbf{d}_k, \mathbf{x}_k^{row}} \|\mathbf{E}_k - \mathbf{d}_k \mathbf{x}_k^{row}\|_F^2. \end{aligned} \quad (5)$$

where $\mathbf{E}_k = \mathbf{Y} - \sum_{i=1, i \neq k}^K \mathbf{d}_i \mathbf{x}_i^{row}$. The singular value decomposition (SVD) of $\mathbf{E}_k = \mathbf{U}\Delta\mathbf{V}^\top$ is used to find the closest rank-1 matrix approximation of $\mathbf{E}_k$. The $\mathbf{d}_k$ update is taken as the first column of $\mathbf{U}$ and the $\mathbf{x}_k^{row}$ update is taken as the first column of $\mathbf{V}$ multiplied by the first element of Δ . To avoid the loss of sparsity in $\mathbf{x}_k^{row}$ that will be created by the direct application of the SVD on $\mathbf{E}_k$, in [14] it was proposed to modify only the non-zero entries of $\mathbf{x}_k^{row}$ resulting from the sparse

coding stage. This is achieved by taking into account only the signals $\mathbf{y}_i$ that use the atom $\mathbf{d}_k$ in (5) or by taking the SVD of $\mathbf{E}_k^R = \mathbf{E}_k \mathbf{I}_{w_k}$, where $w_k = \{i | 1 \leq i \leq N; \mathbf{x}_k^{row}(i) \neq 0\}$ and $\mathbf{I}_{w_k}$ is the $N \times |w_k|$ submatrix of the $N \times N$ identity matrix obtained by retaining only those columns whose index numbers are in w_k , instead of the SVD of $\mathbf{E}_k$.

3. BLEED-THROUGH IDENTIFICATION

Provided that the two sides of the manuscript have been written with the same ink and at two close moments, in the majority of the cases the bleed-through text is almost everywhere lighter than the foreground text. Nonetheless, the intensity of bleed-through is usually very variable, and this makes hard to find suitable thresholds to selectively identify and remove it. Indeed, once again we stress the fact that thresholding would likely destroy also the background texture and other slight document marks that genuinely belong to the side at hand. Thus, in [10] we propose to selectively identify bleed-through and discriminate it from other desired features by looking for those pixels that are lighter than the corresponding pixels in the opposite side.

Three main exceptions to this basic rule are however accounted for. The first one concerns the occlusions, i.e. those areas where the two texts overlap, and that should not be identified as bleed-through in none of the two sides. In these pixels the intensity (graylevel) is the lowest and almost equal in the two corresponding pixels of the recto and verso sides. Therefore, at present we identify the occlusion areas with the aid of suitable thresholds. A specular situation occurs for pixels of background in both sides. Here the intensity is the highest and almost equal in the two sides. Again, proper thresholds can help to discriminate this situation. A third case generates from the typical phenomenon of diffusion of the ink that penetrates through the paper. It may thus happen that, especially at the border of the bleed-through characters, the pixel intensity is lower than that of the (background) pixel in the opposite side, since the text causing the degradation is sharper. To cope with this effect, rather than simply compare the two recto and verso intensities, we compare the intensity in one side with the “smeared” version of the intensity in the opposite side.

Therefore, at each pixel except those identified as occlusions or background, we apply the following formulas to compute the two “seeping levels” q_r and q_v :

$$\begin{aligned} q_r(i, j) &= \frac{1 - \frac{s_v(i, j)}{b_v}}{h_r(i, j) \otimes \left(1 - \frac{s_r(i, j)}{b_r}\right) + \epsilon} \\ q_v(i, j) &= \frac{1 - \frac{s_r(i, j)}{b_r}}{h_v(i, j) \otimes \left(1 - \frac{s_v(i, j)}{b_v}\right) + \epsilon} \end{aligned} \quad (6)$$

where the subscripts r, v indicate recto and verso, respectively, $s(i, j)$ is the intensity at pixel (i, j) , b is an estimate of the predominant graylevel in the background, and the small

constant ϵ is included to avoid indeterminacies or infinity. These formulas comprise two Point Spread Functions (PSF) of unit volume, h_r and h_v , which describe the smearing of the seeping ink, thus allowing a pattern in a side to better match the corresponding one in the opposite side. We assume h_r and h_v as two Gaussian functions, with size and standard deviation approximately estimated from the extent of the character smearing. At each pixel, we retain the smallest between the two “seeping levels” computed with eq. 6, to indicate a bleed-through pixel in the related side, and set to zero the other, to indicate a foreground pixel in the opposite side.

In [10], the strategy described above is applied, with similar performance, using the optical density rather than the intensity, where the optical density is interpreted as a sort of “quantity” of ink, and with reference to a linear additive, non-stationary model to describe the two observed recto and verso densities. The restoration algorithm is completed by substituting the pixels having positive “seeping level” with the background predominant graylevel b , and leaving unchanged the areas where $q = 0$. The experimental results proved the efficiency of the algorithm in removing bleed-through while leaving unaltered other salient marks, such as stamps, pencil annotations, paper textures, etc. These marks can be of paramount importance for the scholar studying the document contents and its origin.

4. BLEED-THROUGH IMPRINT INPAINTING

To each side separately, a dictionary based inpainting method is applied to find a suitable fill-in to replace the identified bleed-through pixels. Our aim is to preserve the background texture to ensure the original look of the document. A patch based sparse representation model is used for image inpainting, with a learned dictionary.

Mathematically, the image inpainting problem is to reconstruct the underlying complete image (lexicographically ordered) $\mathbf{C} \in \mathbb{R}^N$ from its observed incomplete version $\mathbf{I} \in \mathbb{R}^L$, where $L < N$. We assume a sparse representation of $\mathbf{C}$ over a dictionary $\mathbf{D} \in \mathbb{R}^{n \times K}$: $\mathbf{C} \approx \mathbf{D}\mathbf{X}$. The incomplete image $\mathbf{I}$ and the complete image $\mathbf{C}$ are related through $\mathbf{I} = \mathbf{M}\mathbf{C}$, where $\mathbf{M} \in \mathbb{R}^{L \times N}$ represents the layout of the missing pixels. The image inpainting problem can be formulated as:

$$\begin{aligned} \mathbf{I} &= \mathbf{M}\mathbf{C} \\ &\approx \mathbf{M}\mathbf{D}\mathbf{X} \approx \mathbf{\Phi}\mathbf{X} \end{aligned} \quad (7)$$

Assuming that a well trained dictionary $\mathbf{D}$ is available, the problem boils down to the estimation of sparse coefficients $\hat{\mathbf{X}}$ such that the underlying complete image $\hat{\mathbf{C}}$ is given by $\hat{\mathbf{C}} = \mathbf{D}\hat{\mathbf{X}}$. To learn the dictionary $\mathbf{D}$ a training set $\mathbf{Y}$ is created by extracting overlapping patches of size $\sqrt{p_s} \times \sqrt{p_s}$ from the image at location $j = 1, 2, \dots, P$, where P is the total

number of patches. Then we have $\mathbf{y}_j = \mathbf{R}_j \mathbf{C}$, where $\mathbf{R}_j(\cdot)$ is an operator that extracts the patch $\mathbf{y}_j$ from the image $\mathbf{y}$, and its transpose, denoted by $\mathbf{R}_j^T(\cdot)$, is able to put back a patch into the j^{th} position in the reconstructed image. Considering that patches are overlapped, the recovery of $\mathbf{C}$ from $\{\mathbf{y}_j\}$ can be obtained by averaging all the overlapped patches.

We learned a dictionary $\mathbf{D} \in \mathbb{R}^{p_s \times K}$ from the training set $\mathbf{Y}$, using the method described in Section 2. For optimization we used only complete patches from $\mathbf{Y}$, i.e. the patches with no bleed-through imprint pixels. This choice speeds up the training process and excludes the 'non-informative' imprint pixels. For each patch $\mathbf{y}_j$ to be inpainted, we first search for its mutual similar patches in a $W \times W$ bounded neighborhood window. Similar to [24] we used block matching technique with Euclidean distance metric as similarity criterion. The similar patches are grouped together in a matrix $\Omega \in \mathbb{R}^{p_s \times c}$, where c is the total number of similar patches. Using the learned dictionary $\mathbf{D}$ and Ω , a grouped-based sparse representation method similar to [25] is used to find the befitting fill-in for the bleed-through imprints. The use of the similar patch group Ω incorporates the local information in the sparse estimation of the patch and preserves the natural look of the document.

5. EXPERIMENTAL RESULTS

We compared the proposed method with other state-of-the-art methods for images of the database of 25 manuscript recto-verso pairs, presented in [26] [27]. The input image with bleed-through is first processed for bleed-through detection and removal as discussed in Section 3.

For imprints inpainting, the training data set $\mathbf{Y}$ is constructed by selecting all the overlapping patches of size 8×8 from the input image. We learned a dictionary $\mathbf{D}$ of size 64×256 from $\mathbf{Y}$, with the sparsity level $z = 3$ and overcomplete DCT as initial dictionary. For each patch to be inpainted, the sparse coefficients are estimated under the learned dictionaries and the similar patches in its 25×25 neighbourhood window, as discussed in Section 4. The estimated sparse coefficient vector of each patch is denoted by $\mathbf{x}_j$, where j indicates the number of the patch. The reconstructed patch is then obtained as $\hat{\mathbf{y}}_j = \mathbf{D} \mathbf{x}_j$.

In bleed-through restoration the efficacy is generally compared qualitatively, as in most cases the original clean image is not available. A visual comparison of the proposed method with other state-of-the-art methods is presented in Fig. 1. As can be seen, the proposed method (Fig. 1 (e)) produces a better result compared with the other methods, especially in the area bounded by the red rectangle. The dual-layer MRF approach with user trained likelihood of [7] (Fig. 1 (b)) flattens the background and removes the original text in the overlapping cases. The wavelet based method in [8] (Fig. 1 (c)) produces a binary image and fails to handle darker bleed-through. The non-parametric method of [9] (Fig. 1 (d)) produces better results but again fails in some dark bleed-through regions.

The proposed method copes very well with bleed-through removal and the dictionary based inpainting preserves the original look of the document.

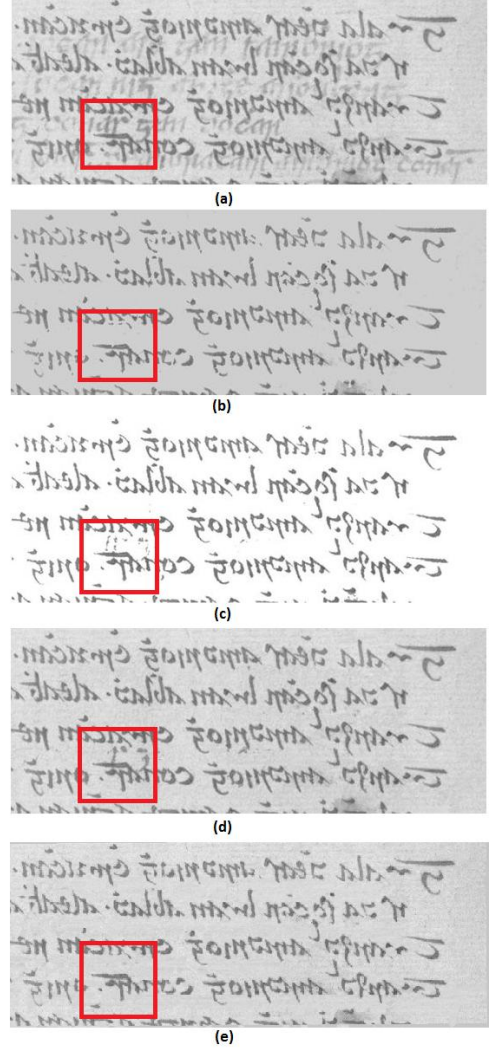


Fig. 1: Visual Comparison: (a) original; (b) result of [7]; (c) result of [8]; (d) result of [9]; (e) proposed.

6. CONCLUSIONS

We plan to extend the sparse representation based inpainting method to manage the cases where some occlusions are wrongly identified as bleed-through, and where some bleed-through strokes are as dark as the foreground text. As per the first case, a proper dictionary of foreground patches must be additionally learned. In both cases, a finer criterion, e.g. based on a measure of proximity, needs to be found to discriminate dark strokes that belong to the bleed-through text from those that belong to the foreground text.

7. REFERENCES

- [1] F. Drira, F. L. Bourgeois, and H. Emptoz, "Restoring ink bleed-through degraded document images using a recursive unsupervised classification technique," *Proc. DAS*, pp. 38–49, 2006.
- [2] R. Estrada and C. Tomasi, "Manuscript bleed-through removal via hysteresis thresholding," *Proc. ICDAR*, pp. 753–757, 2009.
- [3] A. Tonazzini, L. Bedini, and E. Salerno, "Independent component analysis for document restoration," *Int. J. Doc. Anal. Recogn.*, vol. 7, pp. 17–27, 2004.
- [4] C. Wolf, "Document ink bleed-through removal with two hidden markov random fields and a single observation field," *IEEE Trans. Pattern Anal. Mach. Intell.*, pp. 431–447, 2010.
- [5] B. Sun, S. Li, X. P. Zhang, and J. Sun, "Blind bleed-through removal for scanned historical document image with conditional random fields," *IEEE Trans. Image Process.*, pp. 5702–5712, 2016.
- [6] A. Tonazzini, E. Salerno, and L. Bedini, "Fast correction of bleed-through distortion in grayscale documents by a blind source separation technique," *Int. J. Doc. Anal. Recogn.*, vol. 10, pp. 17–27, 2007.
- [7] H. Yi, M. S. Brown, and X. Dong, "User-assisted ink-bleed reduction," *IEEE Trans. Image Process.*, vol. 19, pp. 2646–2658, 2010.
- [8] R. F. Moghaddam and M. Cheriet, "A variational approach to degraded document enhancement," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 32, pp. 1347–1361, 2010.
- [9] R. Rowley-Brooke, F. Pitié, and A. C. Kokaram, "A non-parametric framework for document bleed-through removal," *Proc. CVPR*, pp. 2954–2960, 2013.
- [10] A. Tonazzini, P. Savino, and E. Salerno, "A non-stationary density model to separate overlapped texts in degraded documents," *Signal, Image and Video Processing*, vol. 9, pp. 155–164, 2015.
- [11] I. Gerace, C. Palomba, and A. Tonazzini, "An inpainting technique based on regularization to remove bleed-through from ancient documents," *Proc. IWCIM*, pp. 1–5, 2016.
- [12] M. Elad, M. Figueiredo, and Y. Ma, "On the role of sparse and redundant representations in image processing," *Proc. IEEE*, vol. 98, pp. 972–982, 2010.
- [13] I. Tosic and P. Frossard, "Dictionary learning," *IEEE Signal Proc. Mag.*, vol. 28, pp. 27–38, 2011.
- [14] M. Aharon, M. Elad, and A. Bruckstein, "K-SVD: An algorithm for designing of overcomplete dictionaries for sparse representation," *IEEE Trans. Signal Process.*, vol. 54, pp. 4311–4322, 2006.
- [15] J. A. Tropp and S. J. Wright, "Computational methods for sparse solution of linear inverse problems.," *Proc. IEEE*, vol. 98, pp. 948–958, 2010.
- [16] S. G. Mallat and Z. Zhang, "Matching pursuits with time-frequency dictionaries.," *IEEE Trans. Signal Process.*, vol. 41, pp. 3397–3415, 1993.
- [17] S. Chen, D. Donoho, and M. Saunders, "Atomic decomposition by basis pursuit.," *SIAM J. Sci. Comput.*, vol. 20, pp. 33–61, 1999.
- [18] I. Gorodnitsky and B. Rao, "Sparse signal reconstruction from limited data using FOCUSS: A re-weighted minimum norm algorithm," *IEEE Trans. Signal Process.*, vol. 45, pp. 600–616, 1997.
- [19] J. Tropp, "Greed is Good: Algorithmic results for sparse approximation.," *IEEE Trans. Information Theory*, vol. 50, pp. 2231–2242, 2004.
- [20] K. Kreutz-Delgado, J. Murray, B. Rao, K. Engan, T. Lee, and T. Sejnowski, "Dictionary learning algorithms for sparse representation," *Neural Comput.*, vol. 15(2), pp. 349396, 2003.
- [21] M. Lewicki and T. Sejnowski, "Learning overcomplete representation.," *Neural Comput.*, vol. 12, pp. 337–365, 2000.
- [22] B.A. Olshausen and D.J. Field, "Sparse coding with an overcomplete basis set: A strategy employed by V1?," *J. Vision Research*, vol. 37, pp. 3311–3325, 1997.
- [23] K. Engan, S. O. Aase, and J. Hakon-Husoy, "Method of optimal directions for frame design," *Proc. ICASSP*, pp. 2443–2446, 1999.
- [24] Y. He, T. Gan, W. Chen, and H. Wang, "Adaptive denoising by singular value decomposition," *IEEE Signal Process. Lett.*, vol. 18, pp. 215–218, 2011.
- [25] J. Zhang and D. Zhaocand W. Gao, "Group-based sparse representation for image restoration," *IEEE Trans. Image Process.*, vol. 32, pp. 1307–1314, 2016.
- [26] Irish Script On Screen Project, "http://www.isos.dias.ie," 2012.
- [27] R. Rowley-Brooke, F. Pitié, and A. C. Kokaram, "A ground truth bleed-through document image database," in *Theory and Practice of Digital Libraries*, P. Zaphiris and G. Buchanan, E. Rasmussen, and F. Loizides, Eds. 2012, vol. 7489 of *LNCS*, pp. 185–196, Springer.