

# From Bits to Images: Inversion of Local Binary Descriptors

Emmanuel d'Angelo, Laurent Jacques, Alexandre Alahi and Pierre Vanderghenst

**Abstract**—Local Binary Descriptors are becoming more and more popular for image matching tasks, especially when going mobile. While they are extensively studied in this context, their ability to carry enough information in order to infer the original image is seldom addressed. In this work, we leverage an inverse problem approach to show that it is possible to directly reconstruct the image content from Local Binary Descriptors. This process relies on very broad assumptions besides the knowledge of the pattern of the descriptor at hand. This generalizes previous results that required either a prior learning database or non-binarized features. Furthermore, our reconstruction scheme reveals differences in the way different Local Binary Descriptors capture and encode image information. Hence, the potential applications of our work are multiple, ranging from privacy issues caused by eavesdropping image keypoints streamed by mobile devices to the design of better descriptors through the visualization and the analysis of their geometric content.

**Index Terms**—Computer Vision, Inverse problems, Image reconstruction, BRIEF, FREAK, Privacy



## 1 INTRODUCTION

How much, and what type of information is encoded in a keypoint descriptor? Surprisingly, the answer to this question has seldom been addressed directly. Instead, the performance of keypoint descriptors is studied extensively through several image-based benchmarks following the seminal work of Mikolajczyk and Schmid [1] using Computer Vision and Pattern Recognition task-oriented metrics. These stress tests aim at measuring the stability of a given descriptor under geometric and radiometric changes, which is a key to success in matching templates and real world observations. While precision/recall scores are of primary interest when building object recognition systems, they do not tell much about the intrinsic quality and quantity of information that are embedded in the descriptor. Indeed, these benchmarks are informative about the context in which a descriptor performs well or poorly, but not why. As a consequence, descriptors were mostly developed empirically by benchmarking new ideas against some image matching datasets.

Furthermore, there is a growing trend towards the use of image recognition technologies from mobile handheld devices such as the smartphones combining high quality imaging parts and a powerful computing platform. Application examples include image search and landmark recognition [2] or augmented media and advertisement [3]. To reduce the amount of data exchanged between the mobile and the online knowledge database,

it is tempting to use the terminal to extract image features and send only these features over the network. This data is obviously sensitive since it encodes what the user is viewing. Hence it is legitimate to wonder if its interception could lead to a privacy breach.

Recently, an inspirational paper [4] showed that ubiquitous interest points such as SIFT [5] and SURF [6] suffice to reconstruct plausible source images. This method is based on an image patch database indexed by their SIFT descriptors and then proceeds by successive queries, replacing each input descriptor by the corresponding patch retrieved in the learning set. Although it produces good image reconstruction results, it eventually tells us little about the information embedded in the descriptor: retrieving an image patch from a query descriptor leverages the matching capabilities of SIFT which are now well established by numerous benchmarks and were actually key for its wide adoption.

In this paper, we propose instead two algorithms that aim at reconstructing image patches from local descriptors *without* any external information and with very little additional constraints. We consider descriptors made of local image intensity differences, which are increasingly popular in the Computer Vision community, for they are not very demanding in computational power and hence well suited for embedded applications. The first algorithm that we describe works on *real-valued* difference descriptors, and addresses the reconstruction process as a regularized deconvolution problem. The second algorithm leverages some recent results from 1-bit Compressive Sensing (CS) [7] to reconstruct image parts from *binarized* difference descriptors, and hence is of great practical interest because these descriptors are usually available as bitstrings rather than as real-valued vectors.

- E. d'Angelo and P. Vanderghenst are with the Signal Processing Labs (LTS2), Ecole Polytechnique Fédérale de Lausanne, 1015 Lausanne, Switzerland. E-mail: {emmanuel.dangelo,pierre.vanderghenst}@epfl.ch
- Alexandre Alahi is with the Stanford Vision Lab, Stanford University, CA 94305-9020, USA. E-mail: alahi@stanford.edu
- Laurent Jacques is with the ICTeAM institute, ELEN Department, Université catholique de Louvain (UCL), 1348 Louvain-la-Neuve, Belgium. E-mail: laurent.jacques@uclouvain.be

## Contributions

The contributions of this paper are twofold. First, we extend the seminal work of [4] by showing that an inverse problem approach suffices to invert a local image patch descriptor provided that the descriptor is a local difference operator, thus avoiding the need to build an external database beforehand. Second, we present the first 1-bit reconstruction algorithm with practical applications: even Compressive Sensing (CS) related developments are still focused on theoretical issues tested on toy signals.

An earlier version of this work appeared in [8]. It was however limited to real-valued descriptors, hence we greatly extend it by proposing an algorithm for 1-bit measurements. We also replace the Total Variation constraint of [8] by the sparsity of wavelet analysis coefficients in order to have two reconstruction algorithms (for real and binarized descriptors) that optimize over similar quantities in the current paper, indeed easing the reading. Furthermore, we had to drop the detailed derivation of the real-valued algorithm therein for brevity concerns, and take advantage of the current paper to make the technical steps more explicit.

## Notations

In this paper, we make extensive use of the following notations. Matrices and vectors are denoted by bold letters or symbols (e.g.,  $\Phi$ ,  $\mathbf{x}$ ) while light letters are associated to scalar values (e.g., scalar functions, vector components or dimensions). The scalar product between two  $N$ -length vectors  $\mathbf{x}$  and  $\mathbf{y}$  is written  $\langle \mathbf{x}, \mathbf{y} \rangle = \sum_{i=1}^N x_i y_i$ , while their Hadamard product  $\mathbf{x} \odot \mathbf{y}$  is such that  $(\mathbf{x} \odot \mathbf{y})_i = x_i y_i$  for  $1 \leq i \leq N$ . Since we work only with real matrices, the adjoint of a matrix  $\mathbf{A}$  is  $\mathbf{A}^* = \mathbf{A}^T$ . The vector of ones is written  $\mathbf{1} = (1, \dots, 1)^T$  and the identity matrix is denoted  $\text{Id}$ .

Most of the time, we will “vectorize” 2-D images, i.e., an image or a patch image  $\mathbf{x}$  of dimension  $N_1 \times N_2$  is represented as a  $N$ -dimensional vector  $\mathbf{x} \in \mathbb{R}^N$  with  $N = N_1 N_2$ . This allows us to represent any linear operation on  $\mathbf{x}$  as a simple matrix-vector multiplication. One important linear operator is the 2-D wavelet analysis operator  $\mathbf{W}$  with  $\mathbf{W}^T$  the corresponding synthesis operator. For  $\mathbf{x}, \mathbf{y} \in \mathbb{R}^N$ ,  $\mathbf{W}\mathbf{x}$  is then a vector of wavelet coefficients and  $\mathbf{W}^T \mathbf{y}$  a patch with the same size as  $\mathbf{x}$ .

We denote by  $\|\mathbf{x}\|_p = (\sum_i |x_i|^p)^{1/p}$  with  $p \geq 1$  the  $\ell_p$ -norm of  $\mathbf{x} \in \mathbb{R}^N$ , reserving the notation  $\|\cdot\|$  for  $p = 2$  and with  $\|\mathbf{x}\|_\infty = \max_i |x_i|$ . The  $\ell_0$  “norm” of  $\mathbf{x}$  is  $\|\mathbf{x}\|_0 = \#\{i : x_i \neq 0\}$ . Correspondingly, for  $1 \leq p \leq +\infty$ , a  $\ell_p$ -ball of radius  $\lambda$  is the set  $B_p(\lambda) = \{\mathbf{x} \in \mathbb{R}^N : \|\mathbf{x}\|_p \leq \lambda\}$ .

We use also the following functions. We denote by  $(\mathbf{x})_+$  the non-negativity thresholding function, which is defined componentwise as  $(\lambda)_+ = (\lambda + |\lambda|)/2$ , and  $(\mathbf{x})_- = -(\mathbf{x})_+$ . The sign function  $\text{sign } \lambda$  is equal to 1 if  $\lambda > 0$  and  $-1$  otherwise.

In the context of convex optimization, we denote by  $\Gamma^0(\mathbb{R}^N)$  the class of proper, convex and lower-

semicontinuous functions of the finite dimensional vector space  $\mathbb{R}^N$  to  $(-\infty, +\infty]$  [9]. The indicator function  $\iota_S \in \Gamma^0(\mathbb{R}^N)$  of a set  $S$  maps  $\iota_S(\mathbf{x})$  to 0 if  $\mathbf{x} \in S$  and to  $+\infty$  otherwise. For any  $F \in \Gamma^0(\mathbb{R}^N)$  and  $\mathbf{z} \in \mathbb{R}^N$ , the Fenchel-Legendre conjugate function  $F^*$  is

$$F^*(\mathbf{z}) = \max_{\mathbf{x} \in \mathbb{R}^N} \langle \mathbf{z}, \mathbf{x} \rangle - F(\mathbf{x}), \quad (1)$$

while, for any  $\lambda > 0$ , its proximal operator reads

$$\text{prox}_{\lambda F} \mathbf{z} = \arg \min_{\mathbf{x} \in \mathbb{R}^N} \lambda F(\mathbf{x}) + \frac{1}{2} \|\mathbf{x} - \mathbf{z}\|^2. \quad (2)$$

For  $F = \iota_S$  for some convex set  $S \subset \mathbb{R}^N$ , the proximal operator of  $\text{prox}_{\lambda F}$  simply reduces to the orthogonal projection operator on  $S$  denoted by  $\text{proj}_S$ .

## 2 LOCAL BINARY DESCRIPTORS

In this paper, we are interested in reconstructing image patches from binary descriptors obtained by quantization of local image differences, such as BRIEF [10] or FREAK [11]. Hence, we will refer to these descriptors as Local Binary Descriptors (LBDs) in the sequel. In a standard Computer Vision and Pattern Recognition application, such as object recognition or image retrieval, an interest point detector such as Harris corners [12], SIFT [5] or FAST [13] is first applied on the images to locate interest points. The regions surrounding these keypoints are then described by a feature vector, thus replacing the raw light intensity values by more meaningful information such as histograms of gradient orientation or Haar-like analysis coefficients. In the case of LBDs, the feature vectors are made of local binarized differences computed according to the generic process described below.

### 2.1 Generic Local Binary Descriptor model

A LBD of length  $M$  describing a given image patch of  $\sqrt{N} \times \sqrt{N} = N$  pixels can be computed by iterating  $M$  times the following three-step process:

- 1) compute the Gaussian average of the patch at two locations  $\mathbf{x}_i$  and  $\mathbf{x}'_i$  with variance  $\sigma_i$  and  $\sigma'_i$  respectively;
- 2) form the difference between these two measurements;
- 3) binarize the result by retaining only its sign.

Reshaping the input patch as a column vector  $\mathbf{p} \in \mathbb{R}^N$ , the first two steps in the above procedure can be merged into the application of a single linear operator  $\mathcal{L}$ :

$$\mathcal{L} : \mathbb{R}^N \rightarrow \mathbb{R}^M$$

$$\mathbf{p} \mapsto (\langle \mathcal{G}_{\mathbf{q}_i, \sigma_i}, \mathbf{p} \rangle - \langle \mathcal{G}_{\mathbf{q}'_i, \sigma'_i}, \mathbf{p} \rangle)_{1 \leq i \leq M}, \quad (3)$$

where  $\mathcal{G}_{\mathbf{q}, \sigma} \in \mathbb{R}^N$  denotes a (vectorized) two-dimensional Gaussian of width  $\sigma$  centered in  $\mathbf{q} \in \mathbb{R}^2$  (Fig. 1, top) with  $\|\mathcal{G}_{\mathbf{q}, \sigma}\|_1 = 1$ . As illustrated in Fig. 1-bottom, since  $\mathcal{L}$  is a linear operator, it can be represented

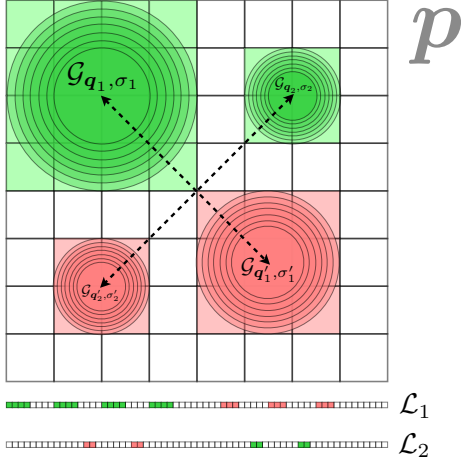


Fig. 1. Example of a local descriptor for an  $8 \times 8$  pixels patch and the corresponding sensing matrix. Only two measurements of the descriptor are depicted; each one is produced by subtracting the Gaussian mean in the lower (red) area from the corresponding upper (green) one. All the integrals are normalized by their area to have values in  $[0, 1]$ . In the bottom, the corresponding vectors.

by a matrix  $\mathcal{L} \in \mathbb{R}^{M \times N}$  multiplying  $p$  and whose each row  $\mathcal{L}_i$  is given by

$$\mathcal{L}_i = \mathcal{G}_{q_i, \sigma_i} - \mathcal{G}_{q'_i, \sigma'_i}, \quad 1 \leq i \leq M. \quad (4)$$

We will take advantage of this decomposition interpretation to avoid explicitly writing  $\mathcal{L}$  later on.

The final binary descriptor is obtained by the composition of this sensing matrix with a component-wise quantization operator  $\mathcal{B}$  defined by  $\mathcal{B}(x)_i = \text{sign } x_i$ , so that, given a patch  $p$ , the corresponding LBD reads

$$\bar{p} := \mathcal{B}(\mathcal{L}p) \in \{-1, +1\}^M. \quad (5)$$

Note that we have chosen this definition of  $\mathcal{B}$  to be consistent with the notations of [7]. Implementations of LBDs will of course use the binary space  $\{0, 1\}^M$  instead, since it fits naturally with the digital representation found in computers.

From the description of LBDs, it is clear that they involve only simple arithmetic operations. Furthermore, the distance between two LBDs is measured using the Hamming distance, which is a simple bitwise exclusive-or (XOR) instruction [10, 11]. Hence, computation and matching of LBDs can be implemented efficiently, sometimes even using hardware instructions (XOR), allowing their use on mobile platforms where computational power and electric consumption are strong limiting constraints. Since they also provide good matching performances, LBDs are getting more and more popular over SIFT and SURF: combined with FAST for the key-point detection, they provide a fast and efficient feature extraction and matching pipe-line, producing compact descriptors that can be streamed over networks.

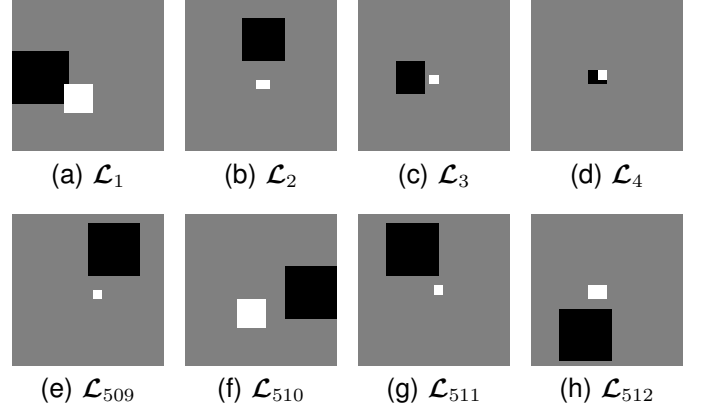


Fig. 2. The first and last frame vectors for a FREAK descriptor of 512 measurements (2-D representations). The white square depicts the positive lobe and the dark square the negative one. Each approximated Gaussian is normalized to be of unit  $\ell_1$ -norm.

Typically, a 32-by-32 pixels image patch (1024 bytes in 8 bit grayscale format) can be reduced to a vector of only 256 measurements [10] coded with 256 *bits*. A typical floating-point descriptor such as SIFT or SURF would require instead 64 float values, *i.e.*, 256 *bytes* for the same patch, eight times the LBD size, and the distances would be measured with the  $\ell_2$ -norm using slower floating-point instructions.

## 2.2 LBDs, LBPs, and other integral descriptors

Unlike [4], we use LBDs in this work instead of SIFT descriptors. As we will see in Section 3, it is actually the knowledge of the spatial measurement pattern used by an LBD that allows us to properly define the matrix of the operator  $\mathcal{L}$  in (3) as a convolution matrix. SIFT and SURF use histograms of gradient orientation instead, thus losing the precise localization information through an integration step. Hence, it seems very unlikely that our approach could be extended to these descriptors. On the other hand however it is possible to reproduce most of the algorithm described in [4] by replacing SIFT with a correctly chosen LBD to index the reference patch database, but this would bring only minor novelty.

Note also that we have coined the descriptors used here as LBDs, which are not the same as the Local Binary Patterns (LBPs) popularized by [14] for face detection. Although both LBDs and LBPs produce bit string descriptors, LBPs are obtained after binarization of image direction histograms. As such, LBPs are integral descriptors and suffer from the same lack of spatial awareness as SIFT and SURF.

## 2.3 The BRIEF and FREAK descriptors

Given two LBDs, the differences reside in the pattern used to select the size and the location of the measurement pairs  $(\mathcal{G}_{q_i, \sigma_i}, \mathcal{G}_{q'_i, \sigma'_i})_{i=1}^M$ . The authors of the

pioneering BRIEF [10] chose small Gaussians of fixed width to bring some robustness against image noise, and tested different spatial layouts. Among these, two random patterns outperformed the others: the first one corresponds to a normal distribution of the measurement points centered in the image patch, and the second one to a uniform distribution.

Working on improving BRIEF, the authors of ORB [15] introduced a measurement selection process based on their matching performance and retained pairs with the highest selectivity. On the other hand, BRISK [16] introduced a concentric pattern to distribute the measurements inside the patch but retained only the innermost points for the descriptor, keeping the peripheral ones to estimate the orientation of the keypoint.

Eventually, the FREAK descriptor was proposed in [11] to leverage the advantages of both approaches: the learning procedure introduced with ORB and the concentric measurement layout of BRISK. The pattern was modified to resemble the retinal sampling and can be seen in Figure 3. Note that it allows for a wider overlap than the BRISK pattern. All the rings were allowed to contribute in the training phase. Consequently, the FREAK descriptor implicitly captures the image details at a coarser scale when going away from the center of the patch.

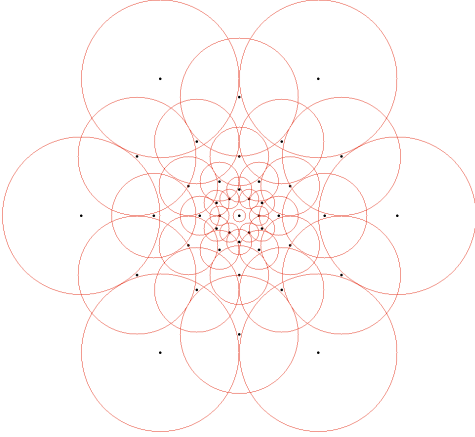


Fig. 3. The retinal pattern used by FREAK. The further a point from the center, the wider the averaging area. Hence, FREAK captures the image variations at a coarser scale on the border of the patch than in its center.

### 3 RECONSTRUCTION AS AN INVERSE PROBLEM

In this work, our goal is to demonstrate that the knowledge of the particular measurement layout of an LBD is sufficient to infer the original image patch *without any external information*, using only an inverse problem approach. Typically, a  $32 \times 32$  pixels patch (1024 values) will be represented by a descriptor with 512 components. Hence, the reconstruction task is ill-posed:

even without binarization of the features, there are half less measurements than unknowns. Assuming that this feature vector is represented with floating-point values, the binarization will then divide by an additional factor of 32 (the standard size of a float in bits) the amount of available information! Classically, to make this problem tractable we introduce a regularization constraint that should be highly generic since we do not know a priori the type of image that we need to reconstruct. Thus, the sparsity of the reconstructed patch in some wavelet frame appeared as a natural choice: it only imposes that a patch should have few nonzero coefficients when analyzed in this wavelet frame, which is quite general.

#### 3.1 Real-valued descriptor reconstruction with convex optimization

Ignoring first the quantization operator by replacing  $\mathcal{B}$  by the identity function, we choose the  $\ell_1$ -norm to penalize the error in the data term and the  $\ell_1$ -norm of the wavelet coefficients as a sparsity promoting regularizer. The  $\ell_1$ -norm is more robust than the usual  $\ell_2$ -norm to the actual value of the error and it is more connected with its sign. Hence, it is hopefully a better choice when dealing with binarized descriptors. The problem of reconstructing an image patch  $\hat{p} \in \mathbb{R}^N$  given an observed binary descriptor  $\bar{p} \in \mathbb{R}^M$  then reads:

$$\hat{p} = \arg \min_{x \in \mathbb{R}^N} \lambda \|\mathcal{L}x - \bar{p}\|_1 + \|Wx\|_1 + \iota_S(x), \quad (6)$$

which is a sparse  $\ell_1$  deconvolution problem [17]. In Eq. (6),  $\|\mathcal{L}x - \bar{p}\|_1$  is the *data term* that ties the solution to the observation  $\bar{p}$ ,  $\|Wx\|_1$  is the regularizer that constrains the patch candidate  $x$  to have a sparse representation, and  $\iota_S(\cdot)$  is the indicator function of the *validity domain* of  $x$  that we will make explicit later.

While the objective function in (6) is convex, it is not differentiable since the  $\ell_1$ -norm has singular points on the axes of  $\mathbb{R}^N$ . Hence, we chose the *primal-dual* algorithm presented in [18] to solve this minimization problem. Instead of using derivatives of the objective which may not exist, it relies on *proximal calculus* and proceeds by alternate minimizations on the primal and dual unknowns.

The generic version of this algorithm aims at solving minimization problems of the form:

$$\hat{x} = \arg \min_{x \in \mathbb{R}^N} F(Kx) + G(x), \quad (7)$$

where  $x \in \mathbb{R}^N$  is the *primal* variable,  $K \in \mathbb{R}^{D \times N}$  is a linear operator, and  $F \in \Gamma^0(\mathbb{R}^D)$  and  $G \in \Gamma^0(\mathbb{R}^N)$  are convex (possibly non-smooth) functions. The algorithm proceeds by restating (7) as a *primal-dual* saddle-point problem on both the *primal*  $x$  and its *dual* variable  $u$ :

$$\min_x \max_u \langle Kx, u \rangle + G(x) - F^*(u). \quad (8)$$

For the problem at hand, we start by decoupling the data term and the sparsity constraint in Eq. (6) by



introducing the auxiliary unknowns  $\mathbf{y}, \mathbf{z} \in \mathbb{R}^N$  such that:

$$\hat{\mathbf{p}} = \arg \min_{\mathbf{y}, \mathbf{z} \in \mathbb{R}^N} \lambda \|\mathcal{L}\mathbf{y} - \bar{\mathbf{p}}\|_1 + \|\mathbf{W}\mathbf{z}\|_1 + \iota_S(\mathbf{y}) + \iota_{\{0\}}(\mathbf{y} - \mathbf{z}). \quad (9)$$

The term  $\iota_{\{0\}}(\mathbf{y} - \mathbf{z})$  guarantees that the optimization occurs on the bisector plane  $\mathbf{y} = \mathbf{z}$ . Then, defining  $\mathbf{K} = \begin{pmatrix} \mathcal{L} & 0 \\ 0 & \mathbf{W} \end{pmatrix} \in \mathbb{R}^{(M+N) \times 2N}$ , we can perform our optimization (6) in the product space  $\mathbf{x} = (\mathbf{y}^T, \mathbf{z}^T)^T \in \mathbb{R}^{2N}$  [19, 20]. In this case, (6) can be written as (7) by setting  $F(\mathbf{K}\mathbf{x}) = F_1(\mathcal{L}\mathbf{y}) + F_2(\mathbf{W}\mathbf{z})$ , where:

$$F_1(\cdot) = \lambda \|\cdot - \bar{\mathbf{p}}\|_1, \quad (10)$$

$$F_2(\cdot) = \|\cdot\|_1, \quad (11)$$

$$G(\frac{\mathbf{y}}{\mathbf{z}}) = \iota_S(\mathbf{y}) + \iota_{\{0\}}(\mathbf{y} - \mathbf{z}). \quad (12)$$

Introducing  $\mathbf{r} \in \mathbb{R}^M$  and  $\mathbf{s} \in \mathbb{R}^N$  the dual counterparts of  $\mathbf{y}$  and  $\mathbf{z}$  respectively, and taking the Fenchel-Legendre transform of  $F$  yields the desired primal-dual formulation of (6):

$$\min_{\mathbf{y}, \mathbf{z}} \max_{\mathbf{r}, \mathbf{s}} \langle \mathcal{L}\mathbf{y}, \mathbf{r} \rangle + \langle \mathbf{W}\mathbf{z}, \mathbf{s} \rangle + G(\frac{\mathbf{y}}{\mathbf{z}}) - F_1^*(\mathbf{r}) - F_2^*(\mathbf{s}). \quad (13)$$

An explicit formulation of  $F_1^*(\mathbf{r})$  can be obtained by noting that, for  $\varphi(\mathbf{r}) = \varphi'(\mathbf{r} - \mathbf{u})$ ,  $\varphi^*(\mathbf{r}) = \varphi'^*(\mathbf{r}) + \langle \mathbf{r}, \mathbf{u} \rangle$ , and that the conjugate of the  $\ell_1$ -norm is  $\iota_{B_\infty(1)}$  [9].  $F_2^*$  is derived in [18], eventually yielding:

$$F_1^*(\mathbf{r}) = \iota_{B_\infty(\lambda)}(\mathbf{r}) + \langle \mathbf{r}, \bar{\mathbf{p}} \rangle, \quad (14)$$

$$F_2^*(\mathbf{s}) = \iota_{B_\infty(1)}(\mathbf{s}). \quad (15)$$

The algorithm presented in [18] requires explicit solutions for the proximal mappings of  $F_1^*$ ,  $F_2^*$  and  $G$ . The first two are easily computed pointwise [9, 18] as:

$$(\text{prox}_{\sigma F_1^*} \mathbf{r})_i = \text{sign}(r_i - \sigma \bar{p}_i) \cdot \min(\lambda, |r_i - \sigma \bar{p}_i|), \quad (16)$$

$$(\text{prox}_{\sigma F_2^*} \mathbf{s})_i = \text{sign}(s_i) \cdot \min(1, |s_i|). \quad (17)$$

The function  $G$  is formed by the indicator of the set  $S$  and the indicator of the bisector plane  $\{(\frac{\mathbf{y}}{\mathbf{z}}) \in \mathbb{R}^{2N} : \mathbf{y} = \mathbf{z}\}$ . An easy computation provides [20]:

$$\text{prox}_{\sigma G}(\frac{\mathbf{y}}{\mathbf{z}}) = \begin{pmatrix} \text{proj}_S \frac{1}{2}(\mathbf{y} + \mathbf{z}) \\ \text{proj}_S \frac{1}{2}(\mathbf{y} + \mathbf{z}) \end{pmatrix}. \quad (18)$$

Thus, its proximal mapping does not depend on any parameter.

Let us now precise our *validity domain*  $S$ . It is defined in order to remove ambiguities in the definition of the program (6) that could lead to a non-uniqueness of the solution. They are due to the differential nature of  $\mathcal{L}$ , i.e., the descriptor of any constant patch is zero. This involves both that  $\bar{\mathbf{p}}$  does not include any information about the average of the initial patch  $\mathbf{p}$ , and the average of  $\mathbf{x}$  in (6) cannot be determined by the optimization.

This problem is removed by defining  $S$  as the intersection of two convex sets  $S_1$  and  $S_2$ . The first set  $S_1$  arbitrarily constrains the minimization domain to stay in the set of patches whose pixel dynamic lies in  $[0, h_{\text{pix}}]$ , i.e.,  $S_1 = \{\mathbf{x} \in \mathbb{R}^N : 0 \leq x_i \leq h_{\text{pix}}\}$ . In our experiments, we simply consider pixels with real values in  $[0, 1]$  and consequently fix  $h_{\text{pix}} = 1$ . The second domain  $S_2$  is

associated to the space of patches whose pixel mean is equal to 0.5, i.e.,  $S_2 = \{\mathbf{x} \in \mathbb{R}^N : \frac{1}{N} \sum_i x_i = 0.5\}$ .

This gives us a first set  $S_1$  whose proximal mapping  $\text{proj}_{S_1}$  is a simple clipping of the values in  $[0, 1]$ , while  $S_2$  is an hyperplane in  $\mathbb{R}^N$  whose corresponding proximal mapping  $\text{proj}_{S_2}$  is the projection onto the simplex of  $\mathbb{R}^N$  of vectors with mean 0.5. This projection can be solved efficiently using [21]. While the desired constraint set  $S$  is the intersection of  $S_1$  and  $S_2$ , we approximate the projection on  $S$  by  $\text{proj}_S \simeq \text{proj}_{S_1} \circ \text{proj}_{S_2}$ . The correct treatment of  $\text{proj}_S$  would normally require to iteratively combine  $\text{proj}_{S_1}$  and  $\text{proj}_{S_2}$  (e.g., running Generalized Forward-Backward splitting [19] until convergence). In our experiments, this approximation did not lead to differences in the estimated patches.

Alg. 1 summarizes the different steps involved in the resolution scheme. It requires a bound  $\Gamma$  on the operator norm  $\mathbf{K} \in \mathbb{R}^{(M+N) \times 2N}$ , i.e., on  $\|\mathbf{K}\| = \max_{\mathbf{x}: \|\mathbf{x}\|=1} \|\mathbf{K}\mathbf{x}\|$ . This is obtained by observing that  $\|\mathbf{W}\|^2 = 1$  with a proper rescaling due to energy conservation constraints, leaving:

$$\|\mathbf{K}\|^2 = \|\begin{pmatrix} \mathcal{L} & 0 \\ 0 & \mathbf{W} \end{pmatrix}\|^2 \leq \|\mathcal{L}\|^2 + 1, \quad (19)$$

where  $\|\mathcal{L}\|$  can be efficiently estimated without any spectral decomposition of  $\mathcal{L}$  by using the power method [17].

While (13) may seem unnecessarily complicated at first because it involves both a minimization and a maximization subproblems, the resolution scheme is actually very efficient: it is a first-order method that involves mostly pointwise normalization and thresholding operations.

---

#### Algorithm 1 Primal-dual $\ell_1$ sparse patch reconstruction.

---

- 1: Take  $\Gamma \geq \|\mathbf{K}\|_2$ , choose  $\tau, \sigma, \theta$  such that  $\Gamma^2 \sigma \tau \leq 1$ ,  $\theta \in [0, 1]$  and  $n$  the number of iterations
  - 2: Initialize:  $\mathbf{x}^{(0)}, \tilde{\mathbf{x}}^{(0)} \leftarrow \mathbf{0}$ , and  $\mathbf{r}^{(0)}, \mathbf{s}^{(0)} \leftarrow \mathbf{0}$
  - 3: **for**  $i = 0$  to  $n - 1$  **do**
  - 4:    $\mathbf{r}^{(i+1)} \leftarrow \text{prox}_{\sigma F_1^*}(\mathbf{r}^{(i)} + \sigma \mathcal{L} \tilde{\mathbf{x}}^{(i)})$
  - 5:    $\mathbf{s}^{(i+1)} \leftarrow \text{prox}_{\sigma F_2^*}(\mathbf{s}^{(i)} + \sigma \mathbf{W} \tilde{\mathbf{x}}^{(i)})$
  - 6:    $\mathbf{x}^{(i+1)} \leftarrow \text{proj}_S(\mathbf{x}^{(i)} - \frac{\tau}{2} \mathcal{L}^T \mathbf{r}^{(i+1)} - \frac{\tau}{2} \mathbf{W}^T \mathbf{s}^{(i+1)})$
  - 7:    $\tilde{\mathbf{x}}^{(i+1)} \leftarrow \mathbf{x}^{(i+1)} + \theta (\mathbf{x}^{(i+1)} - \mathbf{x}^{(i)})$
  - 8: **end for**
  - 9: **return**  $\hat{\mathbf{p}} \leftarrow \mathbf{x}^{(n)}$ .
- 

### 3.2 Iterative binary descriptor reconstruction

To our surprise, when implementing and testing Alg. 1 it turned out that it was able to reconstruct not only real-valued descriptors but also binarized ones, i.e., it still worked without modifications for some  $\bar{\mathbf{p}} \in \{-1, 1\}^M$  instead of  $\mathbb{R}^M$ . This is probably due to the choice of the  $\ell_1$ -norm in the data term, which tends to attach more importance to the sign of the error than to its actual value. However, the behavior of our solver in the binarized descriptor case was unstable and it consistently failed to reconstruct some image patches, yielding a null solution. Hence, we chose to leverage some recent results from

1-bit Compressive Sensing [7] to work out a dedicated binary reconstruction scheme.

Keeping the same inverse-problem approach, we substantially modified the functional of (6) in two ways:

- 1) the data term was changed to enforce bitwise consistency between the LBD computed from the reconstructed patch and the input binary descriptor;
- 2) to apply the same Binary Iterative Hard Thresholding (BIHT) algorithm as [7], we take as sparsity measure the  $\ell_0$ -norm of the wavelet coefficients instead of the relaxed version obtained with the  $\ell_1$ -norm.

We are interested by the solution of this new Lasso-type program [22]:

$$\hat{\mathbf{p}} = \arg \min_{\mathbf{x} \in \mathbb{R}^N} \mathcal{J}(\mathbf{x}) \text{ s.t. } \|\mathbf{W}\mathbf{x}\|_0 \leq k \text{ and } \mathbf{x} \in \mathcal{S}, \quad (20)$$

where the constraints enforces both the validity and the  $k$ -sparsity of  $\mathbf{x}$  in the wavelet domain.

Inspired by [7], we set our data term as

$$\mathcal{J}(\mathbf{x}) = \|\lceil \bar{\mathbf{p}} \odot \mathcal{B}(\mathcal{L}\mathbf{x}) \rceil - \mathbf{1}\|_1. \quad (21)$$

Qualitatively,  $\mathcal{J}$  measures the LBD consistency of  $\mathbf{x}$  with  $\bar{\mathbf{p}}$ , with  $\mathcal{J}(\mathbf{x}) = 0$  iff  $\mathcal{B}(\mathcal{L}\mathbf{x}) = \bar{\mathbf{p}}$ . Each component of the Hadamard product in the definition of  $\mathcal{J}$  is either positive (both signs are the same) or negative. Since the negative function sets to 0 the consistent components, the  $\ell_1$ -norm finally adds the contribution of each inconsistent entry. Note that at the time of writing we do not know a solution for the proximal mapping associated to this data term  $\mathcal{J}$ , which explains our choice for BIHT over the primal-dual solver used in the previous section.

Similarly to the way Iterative Hard Thresholding aims at solving an  $\ell_0$ -Lasso problem [23], BIHT finds one solution of (20) by repeating the three following steps until convergence:

- 1) computing a step of gradient descent of the data term;
- 2) enforcing sparsity by projecting the intermediate estimate to the set of patches with at most  $K$  non-zero coefficients;
- 3) enforcing the mean-value constraint on the result.

This last operation was already studied in the previous section for the real-valued case: it is the projection onto the set  $\mathcal{S}$ . The  $\ell_0$ -norm constraint is applied by Hard Thresholding and amounts to keeping the  $K$  biggest coefficients of the wavelet transform of the estimate and discarding the others. We write this operation  $\mathcal{H}_K$ . Finally, unlike in the primal-dual algorithm, the gradient descent of the data term has to be computed. The result of Lemma 5 in [7] applies in our case and a subgradient of the data term in (20) is:

$$\partial \mathcal{J}(\mathbf{x}) \ni \frac{1}{2} \mathcal{L}^T (\mathcal{B}(\mathcal{L}\mathbf{x}) - \bar{\mathbf{p}}), \quad (22)$$

*i.e.*, the back-projection of the binary error.

Putting everything together, we obtain Alg. 2 that is the adaptation of BIHT to the reconstruction of image

patches from their LBD representation. Again, this algorithm is made of simple elementary steps. The parameter  $\tau = 1/M$  guarantees that the current solution  $\mathbf{x}^{(i)}$  and the gradient step  $\frac{\tau}{2} \mathcal{L}^T (\bar{\mathbf{p}} - \mathcal{B}(\mathcal{L}\mathbf{x}^{(i)}))$  have comparable amplitudes [7]. Since  $M$  is determined from the LBD size, only the patch sparsity level  $K$  in the wavelet basis must be tuned (see Sec. 4.1). In our experiments, however, the algorithm was not very sensitive to the value of  $K$ .

---

#### Algorithm 2 BIHT patch reconstruction.

---

- 1: Take  $\tau = 1/M$ , choose  $K$  the number of non-zero coefficients and  $n$  the number of iterations
  - 2: Initialize:  $\mathbf{x}^{(0)} \leftarrow 0$  and  $\mathbf{a}_0 \leftarrow 0$ .
  - 3: **for**  $i = 0$  to  $n - 1$  **do**
  - 4:    $\mathbf{a}^{(i+1)} \leftarrow \mathbf{x}^{(i)} + \frac{\tau}{2} \mathcal{L}^T (\bar{\mathbf{p}} - \text{sign}(\mathcal{L}\mathbf{x}^{(i)}))$
  - 5:    $\mathbf{b}^{(i+1)} \leftarrow \mathcal{H}_K(\mathbf{W}\mathbf{a}^{(i+1)})$
  - 6:    $\mathbf{x}^{(i+1)} \leftarrow \text{proj}_{\mathcal{S}}(\mathbf{W}^T \mathbf{b}^{(i+1)})$
  - 7: **end for**
  - 8: **return**  $\hat{\mathbf{p}} \leftarrow \mathbf{x}^{(n)}$ .
- 

## 4 RESULTS AND DISCUSSION

### 4.1 Implementation details

For the reconstruction tests presented in this Section, we re-implemented two of the different LBDs: BRIEF and FREAK. For BRIEF, we chose a uniform distribution for the location of the Gaussian measurements, whose support was fixed to  $3 \times 3$  pixels, following the original paper [10]. For FREAK, we did not take into account the orientation of the image patches (see [11], Sec. 4.4) but we also implemented two variants:

- EX-FREAK, for EXhaustive-Freak, computes all the possible pairs from the retinal pattern;
- RA-FREAK, for RAndomized-FREAK, randomly selects its pairs from the retinal pattern.

All the operators were implemented in C++ with the OpenCV library<sup>1</sup> and used the same codebase for fair comparisons, varying only in the measurement pair selection. The code used to generate the examples in this paper is available online and can be retrieved from the page <http://lts2www.epfl.ch/software/>.

The sensing operator was implemented in the following way:

- the forward operator  $\mathcal{L}$  is obtained through the use of integral images for a faster computation of the Gaussian weighted integrals  $\langle \mathcal{G}_{q_i, \sigma_i}, \mathcal{G}_{q'_i, \sigma'_i} \rangle$ . This approximation has become standard in feature descriptor implementations since it allows a huge acceleration of the computations, see for example [6];
- the backward operator  $\mathcal{L}^T$  is the combination of the frame vectors of the considered LBD weighted by the input vector of coefficients. Hence, we avoid explicitly forming  $\mathcal{L}^T$  by computing on the fly the image representation of each vector  $\mathcal{L}_i$  of (4).

1. Freely available at <http://opencv.org>

In the previous sections, we have proposed two algorithms that aimed at reconstructing an image patch from the corresponding local descriptor. In order to assess their relevance and the quality of the reconstructions we have applied them on whole images according to the following protocol:

- 1) an image is divided into patches of size  $\sqrt{N} \times \sqrt{N}$  pixels, with an horizontal and vertical offset of  $N_{\text{off}}$  pixels between each patch;
- 2) each patch is reconstructed independently from its LBD representation using the additional constraint that its mean should be 0.5, *i.e.*, the mean of the input dynamic range;
- 3) reconstructed patches are back-projected to their original image position. Wherever patches overlap, the final result is simply the average of the reconstructions.

Hence, the experiments introduce an additional parameter which is the offset between the selected patches.

Note that in contrast with [4] our methods do not require the use of a seamless patch blending algorithm. Simple averaging does not introduce artifacts. Also, we do not require the knowledge of the scale and orientation of the keypoint to reconstruct: we assume it is a patch of fixed size aligned with the image axes. This is in line with the 1-bit feature detection and extraction pipeline: the genuine FAST detector does not consider scale and orientation, and the BRIEF descriptor is not rotation or scale-invariant. Later descriptors such as FREAK were trained in an affine-invariant context; hence by considering only fixed width and orientation our algorithm is suboptimal.

We used patches of  $32 \times 32$  pixels, LBDs of 512 measurements and run Alg. 1 and Alg. 2 for 1000 and 200 iterations respectively. In Alg. 1, the trade-off parameter  $\lambda$  was set to 0.1. We tried different values for the sparsity  $K$  in Alg. 2 (retaining between 10% and 40% of the wavelet coefficients) but the results did not vary in a meaningful way. Thus we fixed  $K$  throughout all the experiments to keep the 40% greater coefficients of  $Wp$ , choosing the Haar wavelet as analysis operator.

The original *Lena*, *Barbara* and *Kata* images can be seen in Fig. 4.

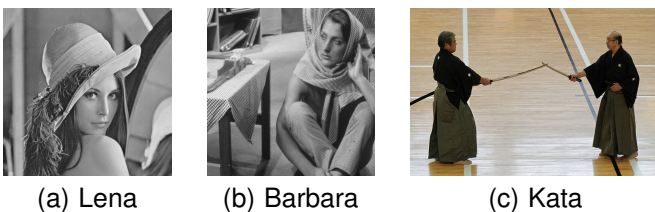


Fig. 4. Original images and designated name in the text.

## 4.2 Reconstruction results

At first glance, the reconstruction results for non-overlapping patches visible in Fig. 5 seem very weird and have sometimes very little in common with the original image. However, if we overlay the original edges on top of the reconstructed images, one can see that each estimated patch contains a correct version of the original gradient direction. This shows that all the four LBDs that we have experimented capture the local gradient and that this information is enough for Alg. 2 to infer the original value. Even curved lines and cluttered area are encoded by the binary descriptors: see the shoulder of *Lena* and the feathers of her hat (Fig. 5). While there is a significant difference between the reconstruction from BRIEF and FREAK, the variants RA-FREAK and EX-FREAK lead to results almost identical with the original FREAK.

Keeping the patch size constant at  $32 \times 32$  pixels, some results for various offsets between the patches can be seen in Fig. 6. An increased number of overlapping patches does dramatically improve the quality of the reconstruction using FREAK. This can be understood easily by considering the peculiar shape of the reconstruction without overlap: each estimated patch contains the correct gradient information at its center only. By introducing more overlap between the patches these small parts of contour sum up to recreate the original objects.

Instead of computing patches at fixed positions and offsets, an experiment more relevant with respect to privacy concerns consists in reconstructing only the patches associated with an interest point detector. For the results shown in Fig. 7, we have first applied the FAST feature detector of OpenCV with its default parameters and discarded the remaining part of the image, hence the black areas, and used real-valued descriptors. Fig. 7 shows the results of the same experiment with binary descriptors. Since FAST keypoints tend to aggregate near angular points and corners, this lead to a relatively dense reconstruction. In each of the three images, the original content can clearly be recognized and a great part of the background clutter has been removed. Thus, one can add as a side note that FAST keypoints are a good indicator of image content saliency. The results in Fig. 8 and Fig. 12 extend to binary descriptors an important privacy issue that was raised before by [4] for SIFT: if one can intercept keypoint data sent over a network (*e.g.*, to an image search engine), then it is possible to find out what the legitimate user was seeing.

Reconstruction results from BRIEF and FREAK are strikingly different (Fig. 5). While BRIEF leads to large, blurred edge estimates that occupy almost entirely the original patch, FREAK produces small accurate edges almost confined in the center of the patch. This allows us to point at a fundamental difference between BRIEF and FREAK. While the former does randomly sample a

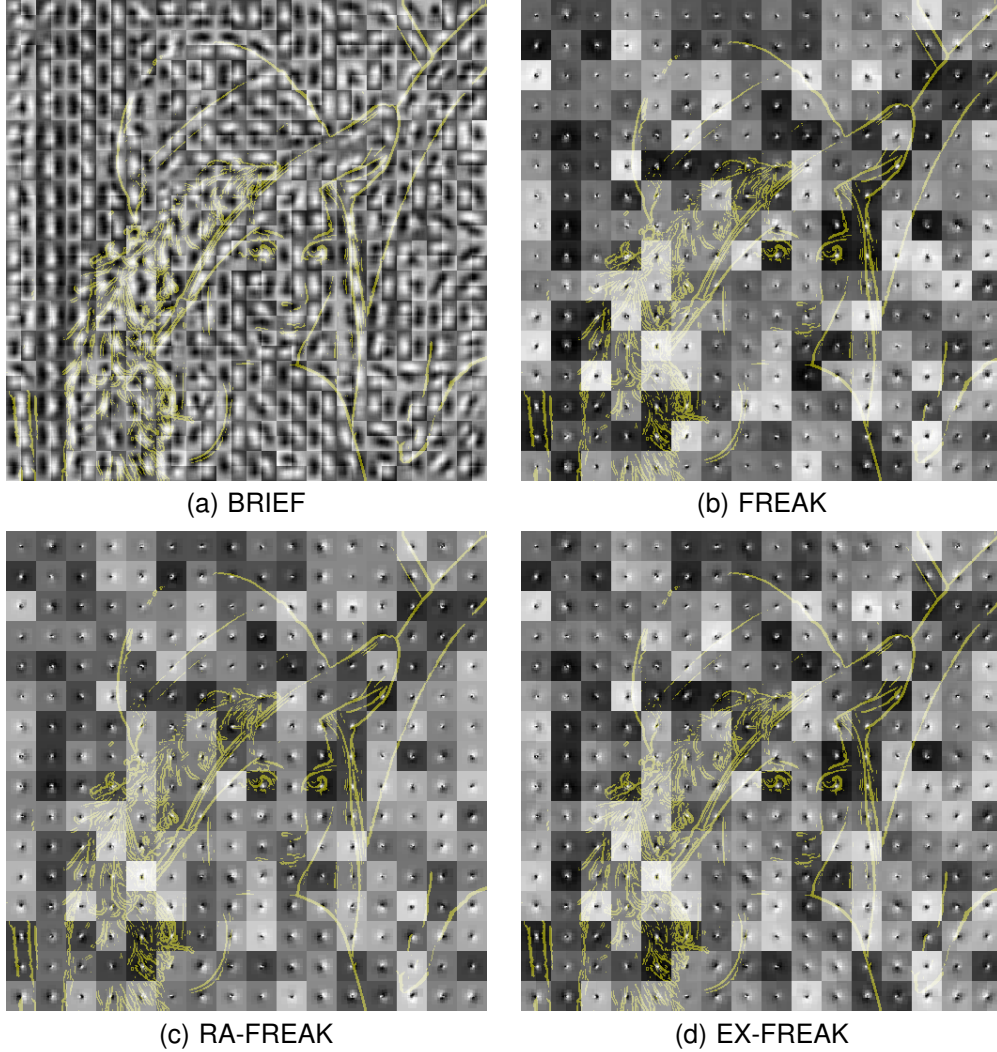


Fig. 5. Reconstruction of Lena from binary LBDs using Alg. 2. There is no overlap between the patches used in the experiment, thus giving a blockwise aspect. We have overlaid some edges from the original image. In each case, the orientation selected for the output patch is consistent with the original main gradient direction. Note also the difference between BRIEF (random measurements spread over an image patch) and FREAK and its derivatives (fine measurements with higher density near the center of the patch): the former gives large blurred edges covering the whole patch, while the latter affect the dominant gradient direction to the central pixel, leaving the periphery almost untouched.

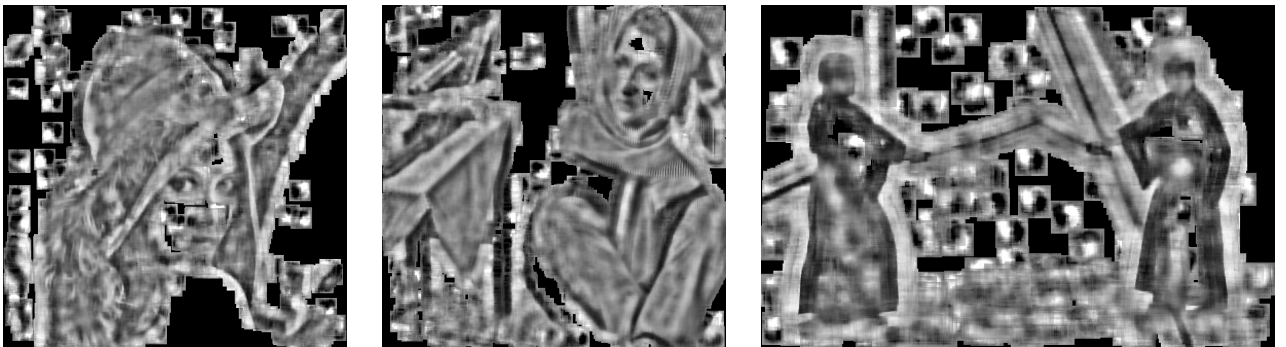


Fig. 7. Reconstruction of floating-point (non binarized) BRIEFs centered on FAST keypoints.

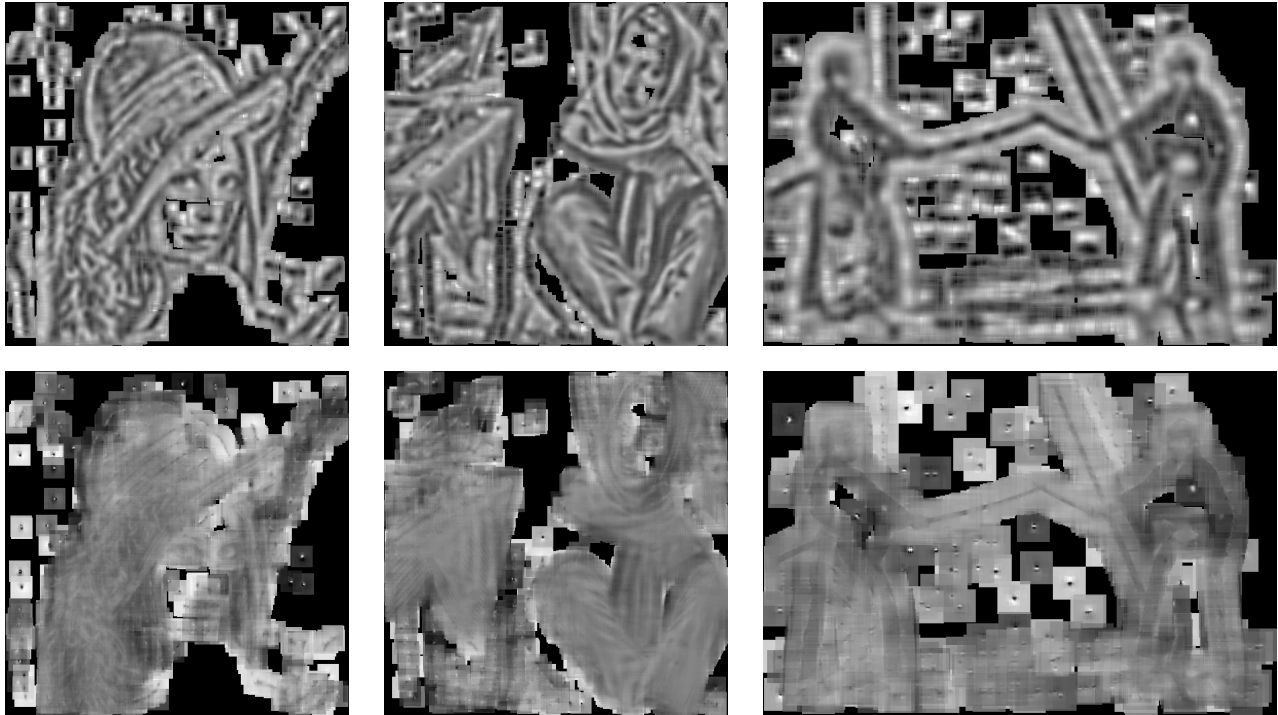


Fig. 8. Reconstruction of LBDs centered on FAST keypoints only. *Top row*: using BRIEF. *Bottom row*: using FREAK. Since the detected points are usually very clustered there is a dense overlap between patches, yielding a visually plausible reconstruction. The original image content has been correctly recovered by Alg. 2 from binary descriptors, and eavesdropping the communications of a mobile camera (e.g., embedded in a smartphone) could reveal private data.

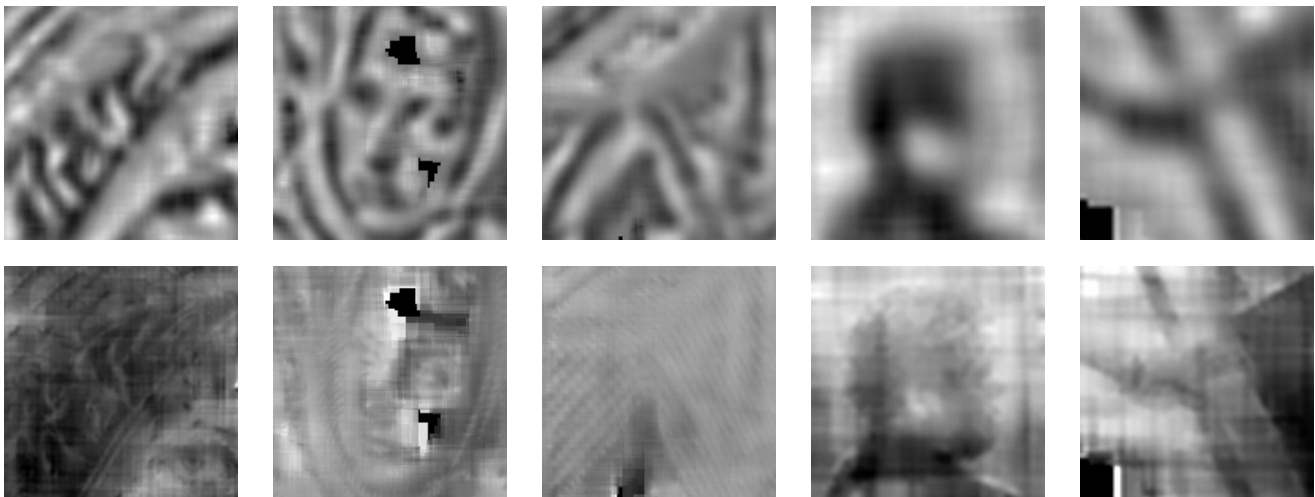


Fig. 9. Details of the reconstructions from Fig. 8. *Top row*: using BRIEF as LBD. *Bottom row*: using FREAK. The reconstructed patches were selected by the application of the FAST detector with identical parameters. While BRIEF is successful at capturing large gradient orientations, hence giving pleasant results when the image is seen from a distance, FREAK captures more accurate orientations in the center of the patches. Thus finer details are recovered: notice for example the eyes in the pictures of Lena and Barbara, the textures from Barbara or the face and the fingers in the kata image. For this Figure, some additional contrast enhancement post-processing was applied to emphasize the point.



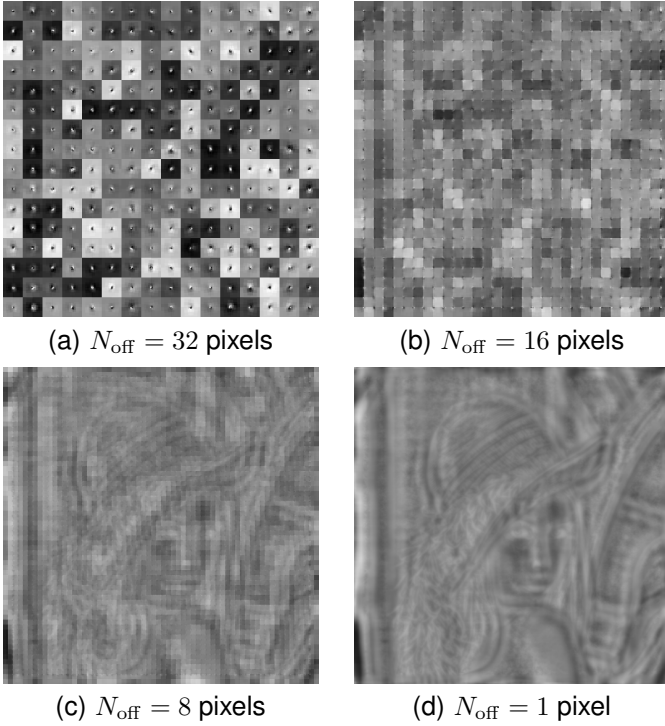


Fig. 6. Reconstruction of Lena from binary FREAKs. The size of the patches is kept fixed at  $32 \times 32$  pixels, while their spacing is gradually reduced. We start with an offset of 32 pixels, *i.e.*, no overlap, until a dense reconstruction. In the limit when each pixel is reconstructed from its neighborhood the individual edge bits chain up and one can clearly distinguish the original image contours, like after the application of a Laplacian filter.

rough estimate of the dominant gradient in the neighborhood, the latter concentrates its finest measurements and allows more bits (Fig. 10) to the innermost part. Thus, the inversion of BRIEF leads to a fuzzy blurred edge dividing two areas since the information is spread spatially over the whole patch, while the reconstruction of FREAK produces a small but accurate edge surrounded by a large low-resolution area. This is confirmed by the experiments shown in Fig. 9. One can especially remark the eyes of Lena and Barbara and the crossed pattern of the table blanket from Barbara which exhibit fine details using FREAK that are missing with BRIEF. In the Kata image, one can almost recognize the face of the characters with FREAK, while the fingers holding the sword are clearly distinguishable.

Figure 10 compares the measurement strategies of BRIEF and FREAK for 512 measurements. The top row displays the sum of the absolute values of the weights applied to a pixel when computing the descriptor, *i.e.*,  $\sum_{i=1}^M |(\mathcal{G}_{q_i, \sigma_i})_j| + |(\mathcal{G}_{q'_i, \sigma'_i})_j|$  in (3) for the  $N$  pixels  $1 \leq j \leq N$ . We clearly see that BRIEF measures patch intensity almost uniformly over its domain, while FREAK focuses its patch observation on the patch domain center. Yet, when plotting in how many LBD measurements a

pixel contributed (Fig. 10, bottom row) one can see that FREAK also uses peripheral pixels. Since both the weight and occurrence patterns are similar with BRIEF, it means that this LBD is democratic and gives all the pixels a similar importance.

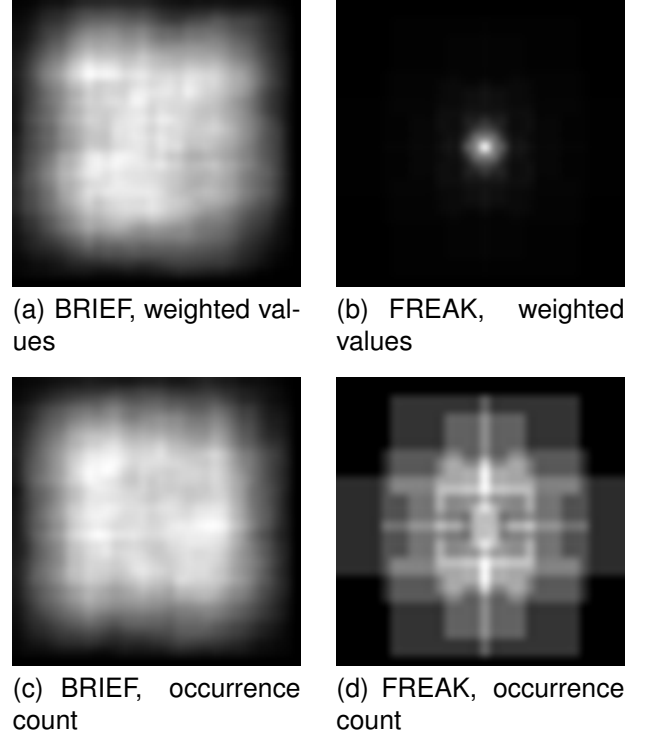


Fig. 10. Comparison of the spatial weights in BRIEF and FREAK basis functions. In the top row, we display the sum of the absolute values of the weight of each pixel when computing a descriptor. Brighter means higher importance. One can see that BRIEF considers almost equally pixels all over the patch, while FREAK gives a very high weight to the centre. The bottom row shows how many times a pixel value was read to generate the description vector. Here, brighter means often retained. This shows that FREAK uses peripheral values, but with a low ponderation. BRIEF is clearly more democratic since the weight pattern is similar to the occurrence pattern.

### 4.3 Quality and stability of the reconstruction

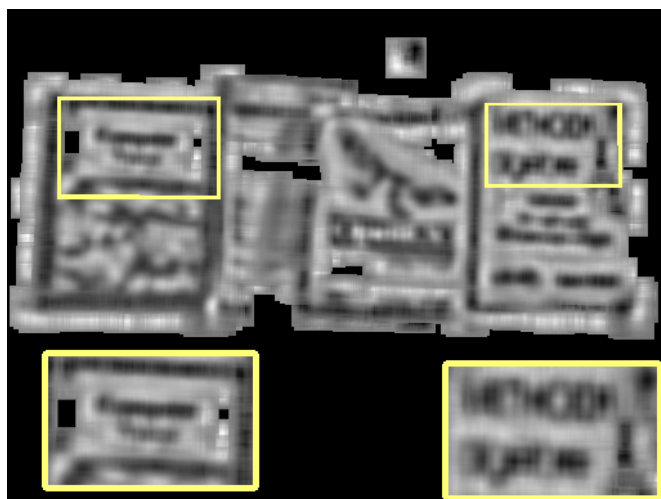
Because of the very peculiar structure of the LBD operators, establishing strong mathematical properties on these matrices is a very arduous task, especially in the 1-bit case. As a consequence, finding indubitable theoretical grounds to the success of our BIHT reconstruction algorithm still remains to be investigated. Intuitively, one can however remark that the conditions used to ensure the existence of a reconstruction in Compressive Sensing, like the famous Restrictive Isometry Property (RIP), are only *sufficient* conditions and are by no means necessary conditions. Since LBDs were designed to accurately



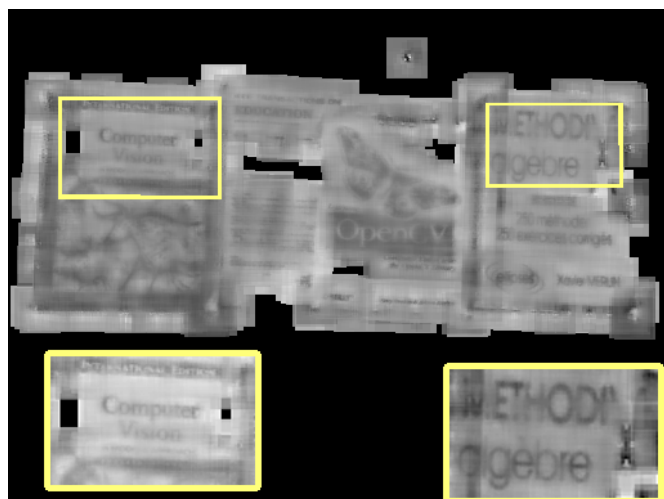
(a) Original image



(b) FAST + Floating-point BRIEF



(c) FAST + binarized BRIEF



(d) FAST + binarized FREAK

Fig. 12. Reconstruction of book covers. The bottom part of each image shows in inset a close-up view of two book titles. This experiment confirms the difference between BRIEF and FREAK: while the former extracts salient shapes such as the auroch and the butterfly, the latter is more successful at reconstructing the text. Note that FREAK allows to read 3 titles out of 4, hence demonstrating the potential existing privacy breach in case of mobile communications eavesdropping.

describe some image content, they are probably more efficient than random sensing matrices. Hence, they can capture more information from an input patch with very few measurements and with a more brutal quantization at the cost of a loss in genericity: they are specialized sensors tuned to image keypoints.

An important parameter with respect to the expected quality of the reconstructions is of course the length  $M$  of the LBDs. Since we lack a reliable quality metric to assess the reconstructions, we have proceeded to visual comparisons between the original image contours and the reconstructed gradient directions on a synthetic image. As can be seen in Fig. 11, dominant orientations are reconstructed correctly until  $M = 128$  measures. Smaller sizes yield blocky estimated patches where it is hard to infer any edge direction.

## 5 CONCLUSION AND FUTURE WORK

In this work, we have presented two algorithms that can successfully reconstruct small image parts from a subset of local differences without requiring external data or prior training. Both algorithms leverage an inverse problem approach to tackle this task and use as regularization constraint the sparsity of the reconstructed image patches in some wavelet frame. They rely however on different frameworks to solve the corresponding problem.

The first method relies on proximal calculus to minimize a convex non-smooth objective function, adopting a deconvolution-like approach. While this functional was not specifically designed for 1-bit LBDs, it has proved to be robust enough to provide some 1-bit reconstructions, but it does lack stability in this case. On the other

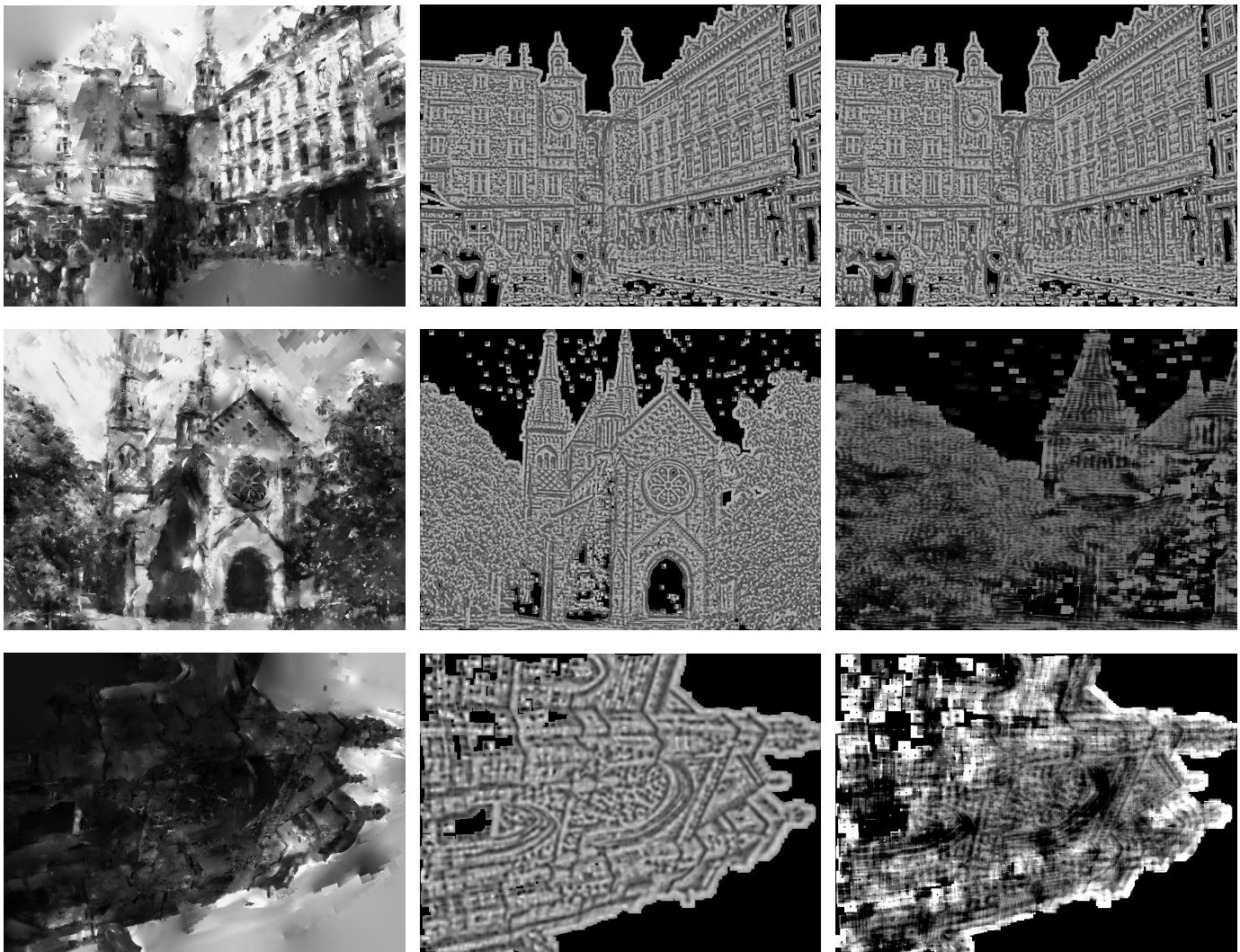


Fig. 13. Comparison with [4]. From left to right: image reconstructed with the method of [4], our binary reconstruction algorithm using FAST+BRIEF (middle) and FAST+FREAK (right). We used patches of  $32 \times 32$  pixels in our algorithm. The contrast of the FREAK results was enhanced for readability, but this Figure is best viewed online in electronic version.

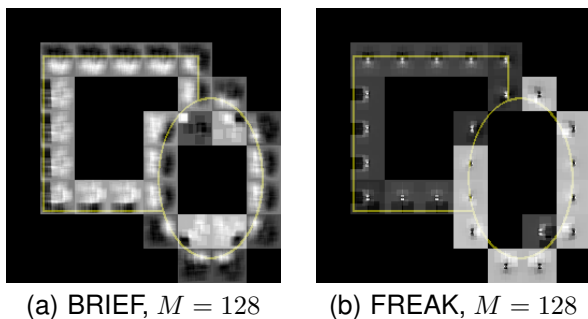


Fig. 11. Zero-overlap reconstruction of a synthetic image using LBDs of 128 measurements instead of 512 as in the other experiments (and 256 in most image matching softwares). Note that in spite of the huge information loss (compression ratio of 256:1 for each patch) the directions of the edges are correctly estimated.

hand, the second method was built from the ground up to handle 1-bit LBDs, and thus provides stable results. The reconstruction process is guided by a hard sparsity constraint in the wavelet domain.

There are several levels to exploit and interpret our results. First, they can have an important industrial impact. Since it is possible to easily invert LBDs without additional information, mobile application developers cannot simply move from SIFT to LBDs in order to avoid the privacy issues raised by [4]. Hence, they need to add an additional encryption tier to their feature point transmission process if the conveyed data can be sensitive or private.

Second, the differences in the reconstruction from different LBDs can help researchers to design their own LBDs. For example, our experiments have pointed out that BRIEF encodes information at a coarser scale than



SIFT, and maybe both descriptors could be combined in some way to create a scale-aware descriptor taking advantage of both patterns. Hence, our work can be used as a tool to study and compare binary descriptors providing a different information than standard matching benchmarks. Furthermore, the fact that real-value differences yield comparable results as binarized descriptors legitimates *a posteriori* the performance of LBDs in matching benchmarks: they encode most of the originally available information.

Finally, our framework for 1-bit contour reconstruction could be combined with the previously proposed Gradient Camera concept [24], leading to the development of a 1-bit Compressive Sensing Gradient Camera. This disruptive device would ally the qualities of both worlds with an extended dynamic range and low power consumption. Exploiting the retinal pattern of FREAK and our reconstruction framework could also yield neuromorphic cameras mimicking the human visual system that could be useful for medical and physiological studies.

Of course, the reconstruction algorithms still need to be improved before reaching this application level. Among the possible improvements, we believe that an interesting extension would be to make our framework *scale-sensitive*. While some feature point detectors provide a scale space location of the detected feature, we discarded the scale coordinate and used patches of fixed width instead. This does not depreciate our experiments with FAST points because we used an implementation that is not scale aware, but reconstructions of better quality can probably be achieved by mixing smooth coarse scale patches with finer details. Additionally, this work did not investigate the issues linked to the geometric transformation invariance enabled by most descriptors. Our model can be interpreted in terms of reconstruction of canonical image patches that correspond to a reference orientation and scale. As far as we have seen, this omission did not create artifacts in our results. This absence by itself is worth of investigating.

## ACKNOWLEDGMENT

The *Kata* image is taken from the iCoseg dataset<sup>2</sup> [25]. The results obtained with the algorithm proposed in [4] were kindly provided by its authors, and the corresponding original images are courtesy of INRIA through the INRIA Copydays dataset. The authors would also like to thank Jérôme Gilles from UCLA who contributed greatly to the implementation of the wavelet transforms. Laurent Jacques is funded by the Belgian Fund for Scientific Research – F.R.S-FNRS.

## REFERENCES

- [1] K. Mikolajczyk and C. Schmid, "A performance evaluation of local descriptors," *Pattern Analysis and Machine Intelligence, IEEE Transactions on*, vol. 27, no. 10, pp. 1615–1630, 2005.
- [2] "Googles Goggles." [Online]. Available: {<http://www.google.com/mobile/goggles/>}
- [3] "Kooaba Image Recognition." [Online]. Available: {<http://www.kooaba.com/en/home>}
- [4] P. Weinzaepfel, H. Jegou, and P. Pérez, "Reconstructing an image from its local descriptors," in *Computer Vision and Pattern Recognition (CVPR), 2011 IEEE Conference on*, 2011, pp. 337–344.
- [5] D. Lowe, "Distinctive image features from scale-invariant keypoints," *International Journal of Computer Vision*, vol. 60, no. 2, pp. 91–110, 2004.
- [6] H. Bay, T. Tuytelaars, and L. van Gool, "Surf: Speeded up robust features," *Computer Vision–ECCV 2006*, pp. 404–417, 2006.
- [7] L. Jacques, J. N. Laska, P. T. Boufounos, and R. G. Baraniuk, "Robust 1-Bit Compressive Sensing via Binary Stable Embeddings of Sparse Vectors," *arXiv.org*, vol. cs.IT, Apr. 2011.
- [8] E. d'Angelo, A. Alahi, and P. Vanderghenst, "Beyond Bits: Reconstructing Images from Local Binary Descriptors," *International Conference On Pattern Recognition (ICPR)*, pp. 1–4, Sep. 2012.
- [9] P. L. Combettes and J. C. Pesquet, "Proximal splitting methods in signal processing," *Fixed-Point Algorithms for Inverse Problems in Science and Engineering*, pp. 185–212, 2011.
- [10] M. Calonder, V. Lepetit, C. Strecha, and P. Fua, "BRIEF: Binary Robust Independent Elementary Features," *Computer Vision–ECCV 2010*, pp. 778–792, 2010.
- [11] A. Alahi, R. Ortiz, and P. Vanderghenst, "FREAK: Fast Retina Keypoint," in *Computer Vision and Pattern Recognition (CVPR), 2012 IEEE Conference on*, 2012, pp. 510–517.
- [12] C. Harris and M. Stephens, "A combined corner and edge detector," *Alvey vision conference*, vol. 15, p. 50, 1988.
- [13] E. Rosten and T. Drummond, "Machine learning for high-speed corner detection," *Computer Vision–ECCV 2006*, pp. 430–443, 2006.
- [14] T. Ahonen, A. Hadid, and M. Pietikainen, "Face Description with Local Binary Patterns: Application to Face Recognition," *Pattern Analysis and Machine Intelligence, IEEE Transactions on*, vol. 28, no. 12, pp. 2037–2041, 2006.
- [15] E. Rublee, V. Rabaud, K. Konolige, and G. Bradski, "ORB: An efficient alternative to SIFT or SURF," in *Computer Vision (ICCV), 2011 IEEE International Conference on*, 2011, pp. 2564–2571.
- [16] S. Leutenegger, M. Chli, and R. Siegwart, "BRISK: Binary Robust invariant scalable keypoints," in *Computer Vision (ICCV), 2011 IEEE International Conference on*, 2011, pp. 2548–2555.
- [17] E. Sidky and J. Jørgensen, "Convex optimization problem prototyping with the Chambolle-Pock algorithm for image reconstruction in computed tomography," *arXiv.org*, 2011.
- [18] A. Chambolle and T. Pock, "A First-Order Primal-Dual Algorithm for Convex Problems with Applications to Imaging," *Journal of Mathematical Imaging and Vision*, vol. 40, no. 1, pp. 120–145, Dec. 2010.
- [19] H. Raguét, J. Fadili, and G. Peyré, "Generalized forward-backward splitting," Preprint Hal-00613637, Tech. Rep., 2011. [Online]. Available: <http://hal.archives-ouvertes.fr/hal-00613637/>
- [20] A. Gonzalez, L. Jacques, C. De Vleeschouwer, and P. Antoine, "Compressive Optical Deflectometric Tomography: A Constrained Total-Variation Minimization Approach," *Submitted, preprint, arXiv:1209.0654*, 2012.
- [21] C. Michelot, "A finite algorithm for finding the projection of a point onto the canonical simplex of  $\mathbb{R}^n$ ," *Journal of Optimization Theory and Applications*, vol. 50, no. 1, pp. 195–200, 1986.
- [22] R. Tibshirani, "Regression shrinkage and selection via the lasso," *Journal of the Royal Statistical Society. Series B (Methodological)*, pp. 267–288, 1996.
- [23] T. Blumensath and M. Davies, "Iterative hard thresholding for compressed sensing," *Applied and Computational Harmonic Analysis*, vol. 27, no. 3, pp. 265–274, 2009.
- [24] J. Tumblin, A. Agrawal, and R. Raskar, "Why I want a gradient camera," *Computer Vision and Pattern Recognition, 2005. CVPR 2005. IEEE Computer Society Conference on*, vol. 1, pp. 103–110 vol. 1, 2005.
- [25] D. Batra, A. Kowdle, D. Parikh, J. Luo, and T. Chen, "Interactively Co-segmenting Topically Related Images with Intelligent Scribble Guidance," *International Journal of Computer Vision*, vol. 93, no. 3, pp. 273–292, Jan. 2011.

2. Available online at <http://chenlab.ece.cornell.edu/projects/touch-coseg/>