



Classification of SD-OCT Volumes using Local Binary Patterns: Experimental Validation for DME Detection

Guillaume Lemaître, Mojdeh Rastgoo, Joan Massich, Carol Y. Cheung, Tien y Wong, Ecosse Lamoureux, Dan Milea, Fabrice Mériaudeau, Désiré Sidibé

► To cite this version:

Guillaume Lemaître, Mojdeh Rastgoo, Joan Massich, Carol Y. Cheung, Tien y Wong, et al.. Classification of SD-OCT Volumes using Local Binary Patterns: Experimental Validation for DME Detection. Journal of Ophthalmology, 2016, 2016. hal-01320791

HAL Id: hal-01320791

<https://u-bourgogne.hal.science/hal-01320791>

Submitted on 24 May 2016

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Classification of SD-OCT Volumes using Local Binary Patterns: Experimental Validation for DME Detection

Guillaume Lemaître^{a,*}, Mojdeh Rastgoo^{a,*}, Joan Massich^{a,*}, Carol Y. Cheung^c, Tien Y. Wong^c, Ecosse Lamoureux^c, Dan Milea^c, Fabrice Mériaudeau^{a,b}, Désiré Sidibé^a

^a*LE2I UMR6306, CNRS, Arts et Métiers, Univ. Bourgogne Franche-Comté, 12 rue de la Fonderie, 71200 Le Creusot, France*

^b*Center for Intelligent Signal and Imaging Research (CISIR), Electrical & Electronic Engineering Department, Universiti Teknologi Petronas, 32610 Seri Iskandar, Perak, Malaysia*

^c*Singapore Eye Research Institute, Singapore National Eye Center, Singapore*

Abstract

This paper addresses the problem of automatic classification of Spectral Domain OCT (SD-OCT) data for automatic identification of patients with Diabetic Macular Edema (DME) versus normal subjects. Optical Coherence Tomography (OCT) has been a valuable diagnostic tool for DME, which is among the most common causes of irreversible vision loss in individuals with diabetes. Here, a classification framework with five distinctive steps is proposed and we present an extensive study of each step. Our method considers combination of various pre-processings in conjunction with Local Binary Patterns (LBP) features and different mapping strategies. Using linear and non-linear classifiers, we tested the developed framework on a balanced cohort of 32 patients.

Experimental results show that the proposed method outperforms the previous studies by achieving a Sensitivity (SE) and Specificity (SP) of 81.2% and 93.7%, respectively. Our study concludes that the 3D features and high-level representation of 2D features using patches achieve the best results. However,

[☆]Document source available in GitHub [1]

^{*}Corresponding author

Email addresses: g.lemaitre58@gmail.com (Guillaume Lemaître),
mojdeh.rastgoo@gmail.com (Mojdeh Rastgoo), joan.massich@u-bourgogne.fr
(Joan Massich)

the effects of pre-processing is inconsistent with respect to different classifiers and feature configurations.

Keywords: Diabetic Macular Edema, Optical Coherence Tomography, DME, OCT, LBP

1. Introduction

Eye diseases such as Diabetic Retinopathy (DR) and Diabetic Macular Edema (DME) are the most common causes of irreversible vision loss in individuals with diabetes. Just in United States alone, health care and associated costs related to eye diseases are estimated at almost \$500 M [2]. Moreover, the prevalent cases of DR are expected to grow exponentially affecting over 300 M people worldwide by 2025 [3]. Given this scenario, early detection and treatment of DR and DME play a major role to prevent adverse effects such as blindness. DME is characterized as an increase in retinal thickness within 1 disk diameter of the fovea center with or without hard exudates and sometimes associated with cysts [4]. Fundus images which have proven to be very useful in revealing most of the eye pathologies [5, 6] are not as good as Optical Coherence Tomography (OCT) images which provide information about cross-sectional retinal morphology [7].

Many of the previous works on OCT image analysis have focused on the problem of retinal layers segmentation, which is a necessary step for retinal thickness measurements [8, 9]. However, few have addressed the specific problem of DME and its associated features detection from OCT images. Figure 1 shows one normal B-scan and two abnormal B-scans.

A summary of the existing work can be found in Table 1. Srinivasan *et al.* [10] proposed a classification method to distinguish DME, Age-related Macular Degeneration (AMD) and normal SD-OCT volumes. The OCT images are pre-processed by reducing the speckle noise by enhancing the sparsity in a transform-domain and flattening the retinal curvature to reduce the inter-patient variations. Then, Histogram of Oriented Gradients (HOG) are extracted for each slice of a volume and a linear Support Vector Machines (SVM) is used for clas-

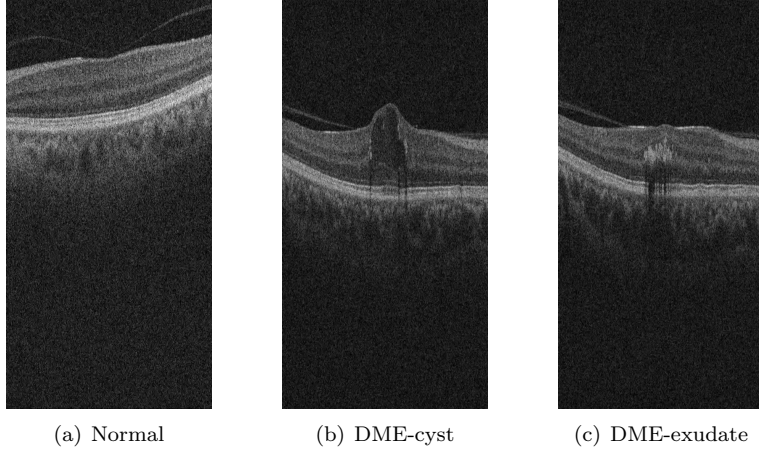


Figure 1: Example of SD-OCT images for normal (a) and DME patients (b)-(c) with cyst and exudate, respectively.

sification. On a dataset of 45 patients equally subdivided into the three aforementioned classes, this method leads to a correct classification rate of 100%, 100% and 86.67% for normal, DME and AMD patients, respectively. The images that have been used in their paper, are publicly available but are already
 30 preprocessed (i.e., denoised), have different sizes for the OCT volumes, do not offer a huge variability in term of DME lesions, and some of them, without specifying which, have been excluded for the training phase; all these reasons prevent us from using this dataset to benchmark our work.

Venhuizen *et al.* proposed a method for OCT images classification using the
 35 Bag-of-Words (BoW) models [11]. The method starts with the detection and selection of keypoints in each individual B-scan, by keeping the most salient points corresponding to the top 3% of the vertical gradient values. Then, a texton of size 9×9 pixels is extracted around each keypoint, and Principal Component Analysis (PCA) is applied to reduce the dimension of every texton
 40 to get a feature vector of size 9. All extracted feature vectors are used to create a codebook using k -means clustering. Then, each OCT volume is represented in terms of this codebook and is characterized as a histogram that captures the

codebook occurrences. These histograms are used as feature vector to train a Random Forest (RF) with a maximum of 100 trees. The method was used to
 45 classify OCT volumes between AMD and normal cases and achieved an Area Under the Curve (AUC) of 0.984 with a dataset of 384 OCT volumes.

Liu *et al.* proposed a methodology for detecting macular pathology in OCT images using Local Binary Patterns (LBP) and gradient information as attributes [12]. The method starts by aligning and flattening the images and
 50 creating a 3-level multi-scale spatial pyramid. The edge and LBP histograms are then extracted from each block of every level of the pyramid. All the obtained histograms are concatenated into a global descriptor whose dimensions are reduced using PCA. Finally a SVM with an Radial Basis Function (RBF) kernel is used as classifier. The method achieved good results in detection OCT
 55 scan containing different pathology such as DME or AMD, with an AUC of 0.93 using a dataset of 326 OCT scans.

Lemaître *et al.* [13] proposed to use 2D and 3D LBP features extracted from denoised volumes and dictionary learning using the BoW models [14]. In the proposed method all the dictionaries are learned with same size of “visual
 60 words” ($k = 32$) and final descriptors are classified using RF classifier.

The work described in this paper is an extension of our previous work [13]. In this research, beside the comparison of 2D and 3D features, we explore different possible representations of the features, and different pre-processing steps for OCT data (i.e. aligning, flattening, denoising). We also compare the per-
 65 formances of different classifiers.

This paper is organized as follows: the proposed framework is explained in Sect. 2, while the experiments and results are discussed through Sect. 3 and Sect. 4. Finally, the conclusion and avenue for future directions are drawn in Sect. 5.

Ref	Diseases			Data size	Pre-processing				Features	Representation	Classifier	Evaluation	Results
	AMD	DME	Normal		De-noise	Flatten	Aligning	Cropping					
[10]	✓	✓	✓	45	✓	✓		✓	HOG		linear-SVM	ACC	86.7%,100%,100%
[11]	✓		✓	384					Texton	BoW, PCA	RF	AUC	0.984
[12]	✓	✓	✓	326		✓	✓		Edge, LBP	PCA	SVM-RBF	AUC	0.93
[13]		✓	✓	62	✓				LBP-LBP-TOP	PCA, BoW, histogram	RF	SE,SP	87.5%, 75%

CT

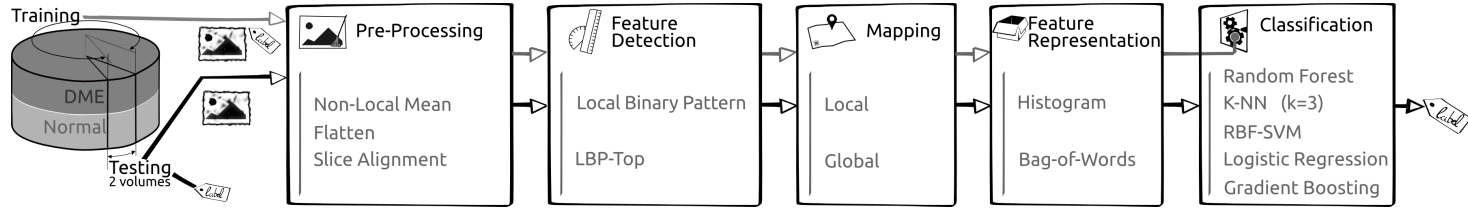


Figure 2: Our proposed classification pipeline.

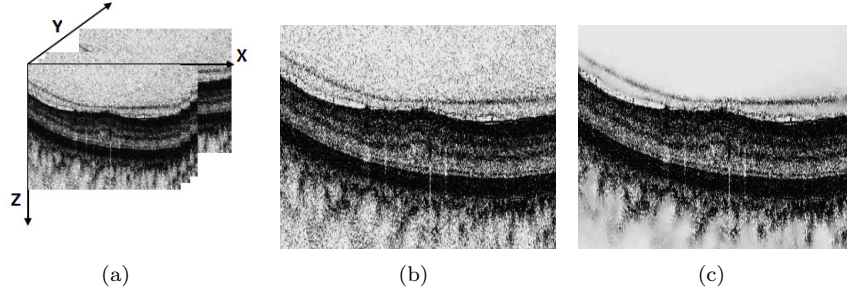


Figure 3: OCT: (a) Organization of the OCT data - (b) Original image - (c) NLM filtering. Note that the images have been negated for visualization purposes.

2. Materials and Methods

The proposed method, as well as, its experimental set-up for OCT volume classification are outlined in Fig. 2. The methodology is formulated as a standard classification procedure which consists of five steps. First, the OCT volumes are pre-processed as presented in details in Sect. 2.1. Then, LBP and LBP-
TOP features are detected, mapped and represented as discussed in depth in Sect. 2.2, Sect. 2.3, and Sect. 2.4, respectively. Finally, the classification step is presented in Sect. 2.5.

2.1. Image pre-processing

This section describes the set of pre-processing techniques which aim at
enhancing the OCT volume. The influence of these pre-processing methods and their possible combinations are extensively studied in Sect. 3.

2.1.1. Non-Local Means (NLM)

OCT images suffer from speckle noise, like other image modalities such as Ultra-Sound (US) [15]. The OCT volumes are enhanced by denoising each B-scan (i.e. each $(x-z)$ slice) using the NLM [16], as shown in Fig. 3. NLM has
been successfully applied to US images to reduce speckle noise and outperforms other common denoising methods [17]. NLM filtering preserves fine structures as well as flat zones, by using all the possible self-predictions that the image can provide rather than local or frequency filters such as Gaussian, anisotropic, or
Wiener filters [16].

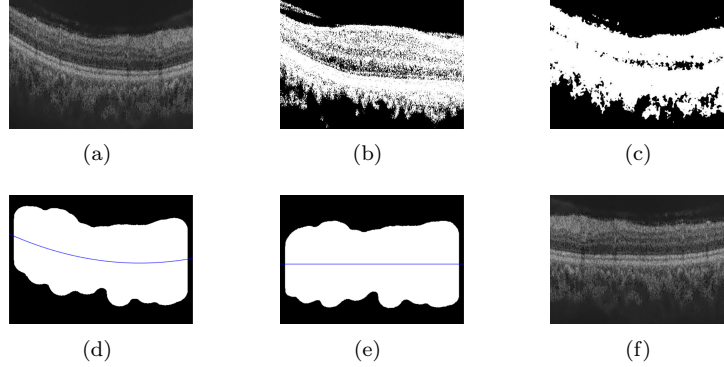


Figure 4: Flattening procedure: (a) original image, (b) thresholding, (c) median filtering, (d) curve fitting, (e) warping, (f) flatten image.

2.1.2. Flattening

Textural descriptors characterize spatial arrangement of intensities. However, the OCT scans suffer from large type of variations: inclination angles, positioning, and natural curvature of the retina [12]. Therefore, these variations have to be taken into account to ensure a consistent characterization of the tissue disposition, regardless of the location in the retina. This invariance can be achieved from different manners: (i) using a rotation invariant descriptor (cf. Sect. 2.2), or (ii) by unfolding the curvature of the retina. This latter correction is known as image flattening which theoretically consists of two distinct steps: (i) estimate and fit the curvature of the Retinal Pigment Epithelium (RPE) and (ii) warp the OCT volume such that the RPE becomes flat.

Our correction is similar to the one of Liu *et al.* [12]: each B-scan is thresholded using Otsu's method followed by a median filtering to detect the different retina layers (see Fig 4(c) and Fig 4(b)). Then, a morphological closing and opening is applied to fill the holes and the resulting area is fitted using a second-order polynomial (see Fig. 4(d)). Finally, the scan is warped such that the curve becomes a line as presented in Fig. 4(e) and Fig. 4(f).

2.1.3. Slice alignment

The flattening correction does not enforce an alignment through the OCT
 110 volume. Thus, in addition to the flattening correction, the warped curves of
 each B-scan are positioned at the same altitude in the z axis.

2.2. Feature detection

In this research, we choose to detect simple and efficient LBP texture fea-
 tures with regards to each OCT slice and volume. LBP is a texture descriptor
 115 based on the signs of the differences of a central pixel with respect to its neigh-
 boring pixels [18]. These differences are encoded in terms of binary patterns as
 in Eq. (1):

$$LBP_{P,R} = \sum_{p=0}^{P-1} s(g_p - g_c)2^p, \quad s(x) = \begin{cases} 1 & \text{if } x \geq 0 \\ 0 & \text{otherwise} \end{cases}, \quad (1)$$

where g_c , g_p are the intensities of the central pixel and a given neighbor pixel,
 respectively; P is the number of sampling points in the circle of radius R .

120 Ojala *et al.* further extended the original LBP formulation to achieve rota-
 tion invariance at the expense of limiting the texture description to the notion
 of circular “uniformity” [18]. Referring to the coordinate system defined in
 Fig. 3(a), the LBP codes are computed on each $(x-z)$ slice, leading to a set of
 LBP maps, a map for each $(x-z)$ slice.

125 Volume encoding is later proposed by Zhao *et al.* by computing LBP de-
 scriptors in three orthogonal planes, so called LBP-TOP [19]. More precisely,
 the LBP codes are computed considering the $(x-z)$ plane, $(x-y)$ plane, and $(y-z)$
 plane, independently. Thus, three sets of LBP maps are obtained, one for each
 orthogonal plane.

130 In this work, we consider rotation invariant and uniform LBP and LBP-TOP
 features with various sampling points (i.e., $\{8, 16, 24\}$) with respect to different
 radius, (i.e., $\{1, 2, 3\}$). The number of patterns ($LBP_{\#pat}$) in regards with each
 configuration is reported in Table 2.

Table 2: Number of patterns ($LBP_{\#pat}$) for different sampling points and radius ($\{P, R\}$) of the LBP descriptor.

	Sampling point for a radius ($\{P, R\}$)		
	$\{8, 1\}$	$\{16, 2\}$	$\{24, 3\}$
$LBP_{\#pat}$	10	18	26

Table 3: Size of a descriptor for an SD-OCT volume. d denotes the number of slices in the volume, N the number of 2D windows, and N' the number of 3D sub-volumes, respectively.

	Global mapping	Local mapping
LBP	$d \times LBP_{\#pat}$	$(N \times d) \times LBP_{\#pat}$
LBP-TOP	$1 \times (3 \times LBP_{\#pat})$	$N' \times (3 \times LBP_{\#pat})$

2.3. Mapping

135 The mapping stage is used to partition the previously computed LBP maps; for this work, two mapping strategies are defined: (i) *global* and (ii) *local* mapping. The size of the feature descriptor is summarized in Table 3.

Global mapping extracts the final descriptors from the 2D feature image for LBP and 3D volume for LBP-TOP. Therefore, for a volume with d slices, 140 the *global*-LBP mapping will lead to the extraction of d elements. While the *global*-LBP-TOP represents the whole volume as a single element. The *global* mapping for 2D images and 3D volume is shown in Fig. 5(a) and 5(b).

Local mapping extracts the final descriptors from a set of $(m \times m)$ 2D patches for LBP and a set of $(m \times m \times m)$ sub-volumes for LBP-TOP. Given N 145 and N' the total number of 2D patches and 3D sub-volumes respectively, the *local*-LBP approach provides $N \times d$ elements, while *local*-LBP-TOP provides N' elements. This mapping is illustrated in Fig. 5(c) and 5(d).

2.4. Feature representation

150 Two strategies are used to describe each OCT volume’s texture.

Low-level representation The texture descriptor of an OCT volume is defined as the concatenation of the LBP histograms with the *global*-mapping.

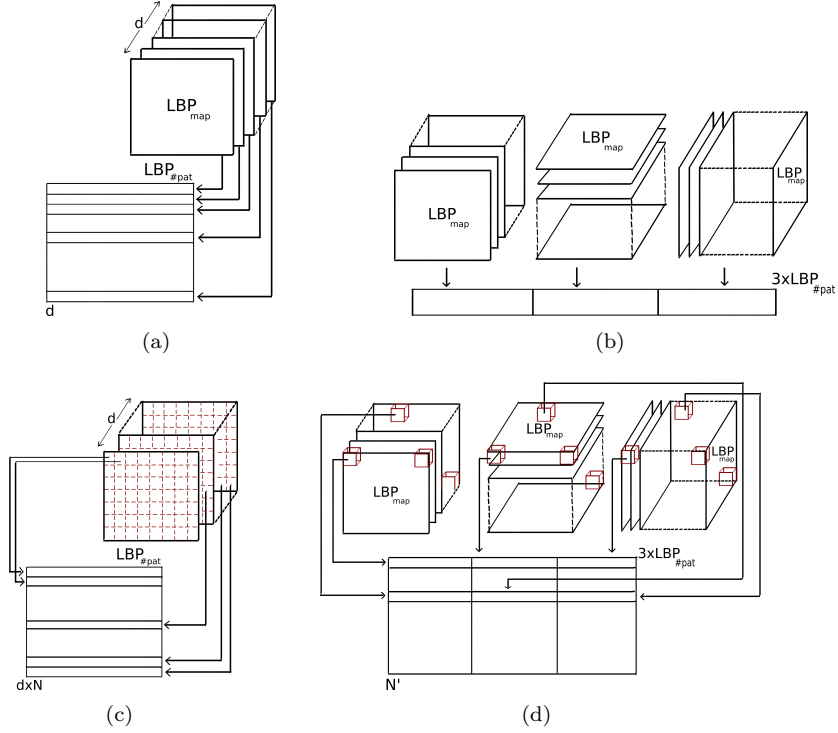


Figure 5: Graphical representation of the feature extraction: (a) extraction of LBP for global mapping - (b) extraction of LBP-TOP for global mapping - (c) extraction of LBP for local mapping - (d) extraction of LBP-TOP for local mapping.

The LBP histograms are extracted from the previously computed LBP maps (see Sect. 2.2). Therefore, the LBP-TOP final descriptor is computed through the concatenation of the LBP histograms of the three orthogonal planes with the final size of $3 \times LBP_{\#pat}$. More precisely, an LBP histogram is computed for each set of LBP maps ($x-z$) plane, ($x-y$) plane, and ($y-z$) plane, respectively. Similarly, the LBP descriptor is defined through concatenation of the LBP histograms per each ($x-z$) slice with the final size of $d \times LBP_{\#pat}$.

High-level representation The concatenation of histograms employed in the low-level representation in conjunction with either *global*- or *local*-mapping can lead to a high dimensional feature space. For instance, *local*-mapping

results in a size of $N \times d \times LBP_{\#pat}$ for the final LBP descriptor and
 165 $N' \times LBP_{\#pat}$ for the final LBP-TOP descriptor, where N and N' are the
 total number of 2D patches and 3D sub-volumes, respectively. High-level
 representation simplifies this high dimensional feature space into a more
 discriminant lower space. BoW approach is used for this purpose [14].
 This model represents the features by creating a codebook or visual dic-
 170 tionary, from the set of low-level features. The set of low-level features are
 clustered using k -means to create the codebook with k clusters or visual
 words. After creating the codebook from the training set, the low-level
 descriptors are replaced by their closest word within the codebook. The
 final descriptor is a histogram of size k which represents the codebook
 175 occurrences for a given mapping.

2.5. Classification

The last step of our framework consists in the classification of SD-OCT vol-
 umes as normal or DME. For that matter, five different classifiers are used:
 (i) k -Nearest Neighbor (NN), (ii) Logistic Regression (LR) [20], (iii) Random
 180 Forest (RF) [21], (iv) Gradient Boosting (GB) [22, 23], and (v) Support Vec-
 tor Machines (SVM) [24, 25]. Details regarding the parameters used in our
 experiments are provided in Sect. 3.

Table 4: The outline and summary of the performed experiments. $\sim$ indicate that common configuration applies.

	Dataset	Pre-processing	Features	Mapping	Representation	Classification	Evaluation
Common:	SERI	NLM	LBP,LBP-TOP $P = \{8, 16, 24\}$ $R = \{1, 2, 3\}$				Leave-One-Patient Out Cross-Validation (LOPO-CV) SE, SP
Baseline [13]: Goal: Evaluation of features, mapping and representation	+ Duke	$\sim$	$\sim$	<i>global</i> <i>local</i>	BoW Histogram	RF	+ comparison with [11]
Experiment#1: Goal: Finding the optimum number of words	$\sim$	+ F + F+A	$\sim$	<i>global</i> <i>local</i>	BoW $k \in K$	LR	+ACC, F1-score (F1)
Experiment#2: Goal: Evaluation of different pre-processing and classifiers for high-level features	$\sim$	+F +F+A	$\sim$	<i>global</i> <i>local</i>	BoW optimal k	3-NN RF SVM GB	$\sim$
Experiment#3: Goal: Evaluation of different pre-processing and classifier for low-level features	$\sim$	+F +F+A	$\sim$	<i>global</i>	Histogram	3-NN RF SVM GB	$\sim$

3. Experiments

A set of three experiments is designed to test the influence of the different
185 blocks of the proposed framework in comparison to our previous work [13].
These experiments are designed such as:

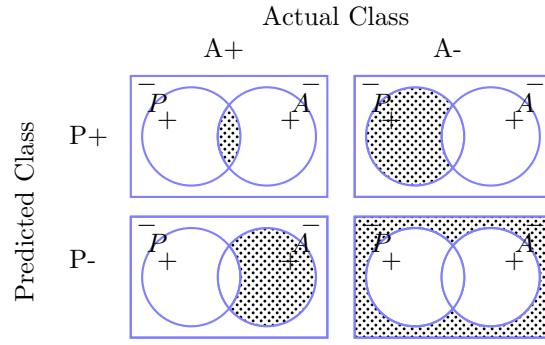
- (i) *Experiment #1* evaluates the effects of number of words used in BoW
(high-level representation).
- (ii) *Experiment #2* evaluates the effects of different pre-processing steps and
190 classifiers on high-level representation.
- (iii) *Experiment #3* evaluates the effects of different pre-processing steps and
classifiers on low-level representation.

Table 4 reports the experiments which have been carried out in [13] as a baseline
and outlines the complementary experimentation here proposed. The reminder
195 of this section details the common configuration parameters across the experi-
ments, while the detailed explanations are presented in the following subsections.

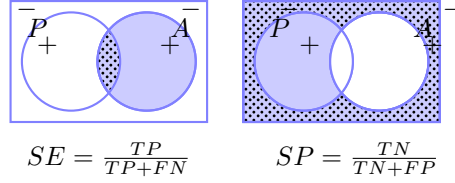
All the experiments are performed using a private dataset (see Sect. 3.1) and
are reported as presented in Sect. 3.2. In all the experiments, LBP and LBP-
TOP features are extracted using both *local* and *global*-mapping for different
200 sampling points of 8, 16, and 24 for radius of 1, 2, and 3 pixels, respectively.
The partitioning for *local*-mapping is set to (7×7) pixels patch for 2D LBP and
 $(7 \times 7 \times 7)$ pixels sub-volume for LBP-TOP.

3.1. SERI-Dataset

This dataset was acquired by the Singapore Eye Research Institute (SERI),
205 using CIRRUS TM (Carl Zeiss Meditec, Inc., Dublin, CA) SD-OCT device.
The dataset consists of 32 OCT volumes (16 DME and 16 normal cases). Each
volume contains 128 B-scan with resolution of 512×1024 pixels. All SD-OCT
images are read and assessed by trained graders and identified as normal or
DME cases based on evaluation of retinal thickening, hard exudates, intraretinal
210 cystoid space formation and subretinal fluid.



(a) Confusion matrix with truly and falsely positive samples detected (TP, FP) in the first row, from left to right and the falsely and truly negative samples detected (FN, TN) in the second row, from left to right.



(b) SE and SP evaluation, corresponding to the ratio of the dotted area over the blue area.

Figure 6: Evaluation metrics: (a) confusion matrix, (b) SE - SP

3.2. Validation

All the experiments are evaluated in terms of Sensitivity (SE) and Specificity (SP) using the LOPO-CV strategy, in line with [13]. SE and SP are statistics driven from the confusion matrix as depicted in Fig. 6. The SE evaluates the performance of the classifier with respect to the positive class, while the SP evaluates its performance with respect to negative class. The use of LOPO-CV implies that at each round, a pair DME-normal volume is selected for testing while the remaining volumes are used for training. Subsequently, no SE or SP variance can be reported. However, LOPO-CV strategy has been adopted despite this limitation due to the reduced size of the dataset.

3.3. Experiment #1

This experiment intends to find the optimal number of words and its effect on the different configurations (i.e., pre-processing and feature representation), on the contrary to [13], where the codebook size was arbitrarily set to $k = 32$.

Several pre-processing strategies are used: (i) NLM, (ii) a combination of NLM and flattening (NLM+F), and (iii) a combination of NLM, flattening, and aligning (NLM+F+A). LBP and LBP-TOP descriptors are detected using the default configuration. Volumes are represented using BoW, where the codebook size ranging for $k \in \{10, 20, 30, \dots, 100, 200, \dots, 500, 1000\}$. Finally, the volumes are classified using LR. The choice of this linear classifier avoids that the results get boosted by the classifier. In this manner, any improvement would be linked to the pre-processing and the size of the codebook.

The usual build of the codebook consists of clustering the samples in the feature space using k -means (see Sect. 2.4). However, this operation is rather computationally expensive and the convergence of the k -means algorithm for all codebook sizes is not granted. Nonetheless, Nowak *et al.* [26] pointed out that randomly generated codebooks can be used at the expenses of accuracy. Thus, the codebook are randomly generated since the final aim is to asses the influence of the codebook size and not the performance of the framework. For this experiment, the codebook building is carried out using random initialization us-

ing *k*-means++ algorithm [27], which is usually used as a *k*-means initialization algorithm.

For this experiment, SE and SP are complemented with ACC and F1 score (see Eq. (2)). ACC offers an overall sense of the classifier performance, and F1 illustrates the trade off between SE and precision. Precision or positive predictive value is a measure of algorithm exactness and is defined as a ratio of True Positive over the total predicted positive samples.

$$ACC = \frac{TP + TN}{TP + TN + FP + FN} \quad F1 = \frac{2TP}{2TP + FP + FN} \quad (2)$$

Appendix A - Table 6 shows the results obtained for the optimal dictionary size while the complete set of all ACC and F1 graphics can be found at [1].

245 According to the obtained results, it is observed that optimum number of words is smaller for *local*-LBP features in comparison to *local*-LBP-TOP and *global*-LBP, respectively. Using LR classifier, the best performances were achieved using *local*-LBP with 70 words (SE and SP of 75.0%) and *local*-LBP-TOP with 500 words (SE and SP of 75.0% as well). These results are highlighted in
250 Appendix A - Table 6.

3.4. Experiment #2

This experiments explores the improvement associated with: (i) different pre-processing methods and (ii) using larger range of classifiers (i.e., linear and non-linear) on the high-level representation.

255 All the pre-processing stages are evaluated (NLM, NLM+F, and NLM+F+A). In this experiment, the codebooks for the BoW representation of LBP and LBP-TOP features are computed using regular *k*-means algorithm which is initialized using *k*-means++, where *k* is chosen according to the findings of *Experiment #1*. Finally, the volumes are classified using *k*-NN, RF, GB, and SVM. The *k*-NN
260 classifier is used in conjunction with the 3 nearest-neighbors rule to classify the test set. The RF and GB classifier are trained using 100 un-pruned trees, while SVM classifier is trained using an RBF kernel and its parameters *C*, and γ are optimized through grid-search.

Complete list of the obtained results from this experiment are shown in
265 Appendix A - Table 7. Despite that highest performances are achieved when
NLM+F or NLM+F+A are used, most configurations decline when applied with
extra pre-processing stages. The best results are achieved using SVM followed
by RF.

3.5. Experiment #3

270 This experiment replicates the *Experiment #2* for the case of low-level rep-
resentation of LBP and LBP-TOP features extracted using *global*-mapping.

The obtained results from this experiment are listed in Appendix A - Ta-
ble 8. In this experiment, flattening the B-scan boosts the results of the best
performing configuration. However, its effects is not consistent across all the
275 configurations. RF has a better performance by achieving better SE (81.2%,
75.0%, 68.7%), while SVM achieve the highest SP (93.7%), see Appendix A -
Table 8.

In terms of classifier, RF has a better performance than the others despite
the fact that the highest SP is achieved using SVM.

Table 5: Summary of all the results. The best results for each experiment are denoted in bold.

[illegible]

280 4. Results and discussion

Table 5 combines the obtained results from Sect. 3 with those reported by Lemaître *et al.* [13], while detailing the frameworks configurations. This table shows the achieved performances with SE higher than 55%.

The obtained results indicate that expansion and tuning of our previous
285 framework improves the results. Tuning the codebook size, based on the finding of *Experiment #1*, leads to an improvement of 6% in terms of SE (see Table 5 at line 7 and 13). Furthermore, the fine tuning of our framework (see Sect. 2) also leads to an improvement of 6% in both SE and SP (see Table 5 at line 1 and 13). Our framework also outperforms the proposed method of [11] with an
290 improvement of 20% and 36% in terms of SE and SP, respectively.

Note that although the effects of pre-processing are not consistent through all the performance, the best results are achieved with NLM +F and NLM +F+A configurations as pre-processing stages. In general, the configurations presented in *Experiment #2* outperform the others, in particular the high-level
295 representation of locally mapped features with an SVM classifier. Focusing on the most desirable radius and sampling point configuration, smaller radius and sampling points are more effective in conjunction with local mapping, while global mapping benefit from larger radius and sampling points.

5. Conclusions

300 The work presented here addresses automatic classification of SD-OCT volumes as normal or DME. In this regard, an extensive study is carried out covering the (i) effects of different pre-processing steps, (ii) influence of different mapping and feature extraction strategies, (iii) impact of the codebook size in BoW, and (iv) comparison of different classification strategies.

305 While outperforming the previous studies [13, 11], the obtained results in this research showed the impact and importance of optimal codebook size, the potential of 3D features and high level representation of 2D features while extracted from local patches.

The strength of SVM while used along BoW approach and RF classifier
 310 while used with global mapping were shown. In terms of pre-processing steps,
 although the highest performances are achieved while alignment and flattening
 were used in the pre-processing, it was shown that the effects of these extra steps
 are not consistent for all the cases and do not guaranty a better performance.

Several avenues for future directions can be explored. The flattening method
 315 proposed by Liu *et al.* flattens roughly the RPE due to the fact that the RPE is
 not segmented. Thus, in order to have a more accurate flattening pre-processing,
 the RPE layer should be pre-segmented as proposed by Garvin *et al.* [28]. In
 this work, the LBP invariant to rotation was used and the number of pattern
 encoded is reduced. Once the data are flattened, the non-rotation invariant
 320 LBP could be studied since this descriptor encode more patterns. In addition
 to LBP, other feature descriptors can be included in the framework.

Acknowledgments

This project was supported by the Singapore French Institute (IFS) and
 the Singapore Eye Research Institute (SERI) through the PHC Merlion pro-
 325 gram (2015-2016) and the Regional Council of Burgundy (grant nb. 2015-
 9201AAO050S02760). Calculations were performed using HPC resources from
 DSI-CCUB (Université de Bourgogne).

Conflict of interest statement

The authors declare no conflict of interest.

330 **Appendix A Complementary results for *Experiment #1* & #2 & #3**

Table 6: Experiment #1 - Optimum number of words for each configuration as a result of LR Classification, for high-level feature extraction of *global* and *local*-LBP, and *local*-LBP-TOP features with different pre-processing. The pre-processing includes: NF, F, and F+A. The achieved performances are indicated in terms of ACC, F1, SE, and SP.

Features	Pre-processing	{8, 1}					{16, 2}					{24, 3}				
		ACC%	F1%	SE%	SP%	W#	ACC%	F1%	SE%	SP%	W#	ACC%	F1%	SE%	SP%	W#
<i>global</i> -LBP																
	NF	81.2	78.5	68.7	93.7	500	62.5	58.0	56.2	62.5	80	62.5	62.5	62.5	62.5	80
	F	71.9	71.0	68.7	75.0	400	68.7	66.7	62.5	75.0	300	68.7	66.7	62.5	75.0	300
	F+A	71.9	71.0	68.7	75.0	500	71.9	71.0	68.7	75.0	200	75.0	68.7	68.7	68.7	500

<i>local</i> -LBP																
	NF	75.0	75.0	75.0	75.0	70	65.6	64.5	62.5	68.7	90	62.5	60.0	56.2	68.7	30
	F	75.0	73.3	68.7	81.2	30	71.8	61.0	68.7	75.0	70	62.5	62.5	62.5	62.5	100
	F+A	75.0	69.0	62.5	81.2	40	71.9	71.0	68.7	75.0	200	68.7	66.7	68.7	62.5	10

<i>local</i> -LBP-TOP																
	NF	68.7	68.7	68.7	68.7	400	75.0	75.0	75.0	75.0	500	71.9	71.0	68.7	75.0	60
	F	68.7	68.7	68.7	68.7	300	68.7	66.7	62.5	75.0	50	75.0	76.5	81.2	68.7	80
	F+A	75.0	73.3	68.7	81.2	100	75.0	73.3	68.7	81.2	90	75.0	69.0	62.5	81.2	70

Table 7: Experiment #2 - k -NN, SVM, RF, and GB classification with BoW for the *global* and *local* LBP and *local* LBP-TOP features with different pre-processing. The optimum number of words were selected based on experiment #1. The most relevant configurations are shaded. The configurations which their performances declines with additional pre-processing are shaded in light gray while those with the opposite behavior are shaded with darker gray color. The highest results which are specified in Table 5 are highlighted in **bold**.

k-NN								SVM					
Features	Pre-processing	{8, 1}		{16, 2}		{24, 3}		{8, 1}		{16, 2}		{24, 3}	
		SE%	SP%	SE%	SP%	SE%	SP%	SE%	SP%	SE%	SP%	SE%	SP%
global-LBP													
	NF	43.7	93.7	43.7	87.5	43.7	62.5	68.7	87.5	62.5	62.5	50.0	56.2
	F	43.7	56.2	50.0	75.0	62.5	56.2	56.2	56.2	56.2	75.0	56.2	68.7
	FA	56.2	62.5	43.7	81.2	68.7	56.2	56.2	68.7	68.7	68.7	56.2	75.0
local-LBP													
	NF	75.0	87.5	50.0	68.7	43.7	43.7	75.0	93.7	50.0	75.0	56.2	56.2
	F	56.2	56.2	50.0	50.0	50.0	43.7	81.2	93.7	68.7	68.7	68.7	75.0
	FA	56.2	43.7	50.0	75.0	50.0	62.5	75.0	93.7	75.0	68.7	68.7	68.7
local-LBP-TOP													
	NF	56.2	75.0	56.2	75.0	62.5	56.2	81.2	87.5	75.0	100	56.2	75.0
	F	62.5	43.7	37.5	68.7	43.7	62.5	81.2	81.2	75.0	68.7	81.2	68.7
	F+A	56.2	56.2	68.7	50.0	43.7	62.5	62.5	75.0	68.7	75.0	62.5	81.2
RF								GB					
Features	Pre-processing	8 ^{riu2}		16 ^{riu2}		24 ^{riu2}		8 ^{riu2}		16 ^{riu2}		24 ^{riu2}	
		SE%	SP%	SE%	SP%	SE%	SP%	SE%	SP%	SE%	SP%	SE%	SP%
global-LBP													
	NF	68.7	93.7	43.7	62.5	50.0	68.7	56.2	50.0	37.5	31.2	50.0	43.7
	F	56.2	50.0	56.2	75.0	50.0	75.0	50.0	56.2	56.2	75.0	43.7	62.5
	FA	68.7	50.0	56.2	62.5	62.5	56.2	56.2	50.0	68.7	50.0	43.7	75.0
local-LBP													
	NF	81.2	81.2	62.5	56.2	56.2	56.2	75.0	62.5	68.7	87.5	50.0	75.0
	F	56.2	81.2	62.5	68.7	68.7	62.5	68.7	75.0	50.0	75.0	50.0	62.5
	FA	68.7	62.5	62.6	68.7	43.7	43.7	56.2	50.0	68.7	56.2	50.0	50.0
local-LBP-TOP													
	NF	68.7	62.5	68.7	81.2	68.7	68.7	37.5	68.7	62.5	81.2	62.5	50.0
	F	50.0	62.5	62.5	62.5	43.7	75.0	50.0	56.2	43.7	62.5	50.0	62.5
	F+A	50.0	62.5	81.2	87.5	50.0	68.7	56.2	62.5	81.2	68.7	75.0	68.7

Table 8: Experiment #3 - Classification results obtained from low-level representation of global LBP and LBP-TOP features with different pre-processing. Pre-processing steps include: NF, F, F+A. Different classifiers such as RF, GB, SVM, and k -NN are used. The most relevant configurations are shaded. The configurations which their performances declines with additional pre-processing are shaded in light gray while those with the opposite behavior are shaded with darker gray color. The highest results which are specified in Table 5 are highlighted in **bold**.

Features	Pre-processing	<i>k</i> -NN						SVM					
		{8, 1}		{16, 2}		{24, 3}		{8, 1}		{16, 2}		{24, 3}	
		SE%	SP%	SE%	SP%	SE%	SP%	SE%	SP%	SE%	SP%	SE%	SP%
<i>global</i> -LBP													
	NF	37.5	50.0	25.0	50.0	37.5	68.7	56.2	62.5	56.2	43.7	56.2	68.7
	F	62.5	50.0	56.2	75.0	62.5	68.7	75.0	68.7	62.5	62.5	62.5	68.7
	FA	56.2	50.0	56.2	75.0	62.5	68.7	75.0	68.7	62.5	62.5	62.5	68.7
<i>global</i> -LBP-TOP													
	NF	31.2	93.7	37.5	100.0	37.5	81.2	62.5	75.0	62.5	93.7	56.2	87.5
	F	50.0	56.2	56.2	75.0	56.2	62.5	68.7	75.0	43.7	68.7	68.7	56.2
	F+A	75.0	43.7	56.2	43.7	68.7	50.0	68.7	62.5	62.5	56.2	56.2	68.7
Features	Pre-processing	RF						GB					
		8^{riu2}		16^{riu2}		24^{riu2}		8^{riu2}		16^{riu2}		24^{riu2}	
		SE%	SP%	SE%	SP%	SE%	SP%	SE%	SP%	SE%	SP%	SE%	SP%
<i>global</i> -LBP													
	NF	43.7	62.5	43.7	62.5	56.2	75	43.7	43.7	43.7	37.5	37.5	31.25
	F	56.2	56.2	68.7	62.5	62.5	68.7	25	56.2	50.0	43.7	25.0	43.7
	F+A	65.2	56.2	50.0	50.0	56.2	68.7	43.75	62.5	62.5	50.0	31.2	31.2
<i>global</i> -LBP-TOP													
	NF	56.2	68.7	68.7	87.5	68.7	81.2	68.7	68.7	75.0	50.0	56.2	43.7
	F	56.2	62.5	81.2	68.7	81.2	81.2	56.2	62.5	62.5	68.7	68.7	81.2
	F+A	68.7	62.5	75.0	68.7	75.0	81.2	56.2	43.7	62.5	62.5	75.0	75.0

References

- [1] G. Lemaitre, M. Rastgoo, J. Massich, retinopathy: Jo-omia-2015 (Nov. 2015). doi:10.5281/zenodo.34277.
URL <http://dx.doi.org/10.5281/zenodo.34277>
- 335 [2] S. Sharma, A. Oliver-Hernandez, W. Liu, J. Walt, The impact of diabetic retinopathy on health-related quality of life, *Current Opinion in Ophthalmology* 16 (2005) 155–159.
- [3] S. Wild, G. Roglic, A. Green, R. Sicree, H. King, Global prevalence of diabetes estimates for the year 2000 and projections for 2030, *Diabetes*
340 *Care* 27 (5) (2004) 1047–1053.
- [4] Early Treatment Diabetic Retinopathy Study Group, Photocoagulation for diabetic macular edema: early treatment diabetic retinopathy study report no 1, *JAMA Ophthalmology* 103 (12) (1985) 1796–1806.
- [5] M. R. K. Mookiah, U. R. Acharya, C. K. Chua, C. M. Lim, E. Ng, A. Laude,
345 Computer-aided diagnosis of diabetic retinopathy: A review, *Computers in Biology and Medicine* 43 (12) (2013) 2136–2155.
- [6] E. Trucco, A. Ruggeri, T. Karnowski, L. Giancardo, E. Chaum, J. Hub-
schman, B. al Diri, C. Cheung, D. Wong, M. Abramoff, G. Lim, D. Ku-
mar, P. Burlina, N. M. Bressler, H. F. Jelinek, F. Meriaudeau, G. Quelled,
350 T. MacGillivray, B. Dhillon, Validation retinal fundus image analysis algo-
rithms: issues and proposal, *Investigative Ophthalmology & Visual Science* 54 (5) (2013) 3546–3569.
- [7] Y. T. Wang, M. Tadarati, Y. Wolfson, S. B. Bressler, N. M. Bressler, Com-
parison of Prevalence of Diabetic Macular Edema Based on Monocular
355 Fundus Photography vs Optical Coherence Tomography, *JAMA Ophthalmology* (2015) 1–7.

- [8] S. J. Chiu, X. T. Li, P. Nicholas, C. A. Toth, J. A. Izatt, S. Farsiu, Automatic segmentation of seven retinal layers in sd-oct images congruent with expert manual segmentation, *Optic Express* 18 (18) (2010) 19413–19428.
- 360 [9] R. Kafieh, H. Rabbani, M. D. Abramoff, M. Sonka, Intra-retinal layer segmentation of 3d optical coherence tomography using coarse grained diffusion map, *Medical Image Analysis* 17 (2013) 907–928.
- [10] P. P. Srinivasan, L. A. Kim, P. S. Mettu, S. W. Cousins, G. M. Comer, J. A. Izatt, S. Farsiu, Fully automated detection of diabetic macular edema and
365 dry age-related macular degeneration from optical coherence tomography images, *Biomedical Optical Express* 5 (10) (2014) 3568–3577.
- [11] F. G. Venhuizen, B. van Ginneken, B. Bloemen, M. J. P. P. van Grisen, R. Philipsen, C. Hoyng, T. Theelen, C. I. Sanchez, Automated age-related macular degeneration classification in OCT using unsupervised feature
370 learning, in: *SPIE Medical Imaging*, Vol. 9414, 2015, p. 941411.
- [12] Y.-Y. Liu, M. Chen, H. Ishikawa, G. Wollstein, J. S. Schuman, R. J. M., Automated macular pathology diagnosis in retinal oct images using multi-scale spatial pyramid and local binary patterns in texture and shape encoding, *Medical Image Analysis* 15 (2011) 748–759.
- 375 [13] G. Lemaitre, M. Rastgoo, J. Massich, S. Sankar, F. Meriaudeau, D. Sidibe, Classification of SD-OCT volumes with LBP: Application to dme detection, in: *Medical Image Computing and Computer-Assisted Intervention (MICCAI), Ophthalmic Medical Image Analysis Workshop (OMIA)*, 2015.
- 380 [14] J. Sivic, A. Zisserman, Video google: a text retrieval approach to object matching in videos, in: *IEEE ICCV*, 2003, pp. 1470–1477.
- [15] J. M. Schmitt, S. H. Xiang, K. M. Yung, Speckle in optical coherence tomography, *Journal of biomedical optics* 4 (1) (1999) 95–105.

- [16] A. Buades, B. Coll, J.-M. Morel, A non-local algorithm for image denoising,
 385 in: Computer Vision and Pattern Recognition, 2005. CVPR 2005. IEEE
 Computer Society Conference on, Vol. 2, IEEE, 2005, pp. 60–65.
- [17] P. Coupe, P. Hellier, C. Kervrann, C. Barillot, Nonlocal means-based
 speckle filtering for ultrasound images, IEEE TIP (2009) 2221–2229.
- [18] T. Ojala, M. Pietikäinen, T. Mäenpää, Multiresolution gray-scale and ro-
 390 tation invariant texture classification with local binary patterns, Pattern
 Analysis and Machine Intelligence, IEEE Transactions on 24 (7) (2002)
 971–987.
- [19] G. Zhao, T. Ahonen, J. Matas, M. Pietikäinen, Rotation-invariant image
 and video description with local binary pattern features, Image Processing,
 395 IEEE Transactions on 21 (4) (2012) 1465–1477.
- [20] D. R. Cox, The regression analysis of binary sequences, Journal of the Royal
 Statistical Society. Series B (Methodological) (1958) 215–242.
- [21] L. Breiman, Random forests, Machine learning 45 (1) (2001) 5–32.
- [22] J. H. Friedman, Stochastic gradient boosting, Computational Statistics &
 400 Data Analysis 38 (4) (2002) 367–378.
- [23] G. Lemaitre, J. Massich, R. Marti, J. Freixenet, J. C. Vilanova, P. M.
 Walker, D. Sidibe, F. Meriaudeau, A boosting approach for prostate can-
 cer detection using multi-parametric mri, in: International Conference on
 Quality Control and Artificial Vision (QCAV2015), SPIE, 2015.
- [24] V. Vapnik, A. J. Lerner, Generalized portrait method for pattern recogni-
 405 tion, Automation and Remote Control 24 (6) (1963) 774–780.
- [25] A. Aizerman, E. M. Braverman, L. I. Rozoner, Theoretical foundations of
 the potential function method in pattern recognition learning, Automation
 and Remote Control 25 (1964) 821–837.

- 410 [26] E. Nowak, F. Jurie, B. Triggs, Sampling strategies for bag-of-features image classification, in: Computer Vision–ECCV 2006, 2006, pp. 490–503.
- [27] D. Arthur, S. Vassilvitskii, k-means++: The advantages of careful seeding, in: Proceedings of the eighteenth annual ACM-SIAM symposium on Discrete algorithms, Society for Industrial and Applied Mathematics, 2007, 415 pp. 1027–1035.
- [28] M. Garvin, M. Abramoff, X. Wu, S. Russell, T. Burns, M. Sonka, Automated 3-d intraretinal layer segmentation of macular spectral-domain optical coherence tomography images, Medical Imaging, IEEE Transactions on 28 (9) (2009) 1436–1447.