

Segmentation of 4D Echocardiography Using Stochastic Online Dictionary Learning*

Xiaojie Huang¹, Donald P. Dione⁴, Ben A. Lin⁴, Alda Bregasi⁴,
Albert J. Sinusas^{3,4}, and James S. Duncan^{1,2,3}

¹ Departments of Electrical Engineering,

² Departments of Biomedical Engineering,

³ Departments of Diagnostic Radiology,

⁴ Departments of Internal Medicine,

Yale University, New Haven, CT, USA

xiaojie.huang@yale.edu

Abstract. Dictionary learning has been shown to be effective in exploiting spatiotemporal coherence for echocardiographic segmentation. To overcome the limitations of previous methods, we present a stochastic online dictionary learning approach for segmenting left ventricular borders from 4D echocardiography. It is based on stochastic approximations and processes a mini-batch of samples at a time, which results in lower memory consumption and lower computational cost than classical batch algorithms. In contrast to the previous methods, where dictionaries and their weights are optimized only on the most recently segmented frame, our stochastic online learning procedure optimizes the dictionaries and the corresponding weights by aggregating all the past information while adapting them to the dynamically changing data. The rate of updating the past information is controlled and varied according to the appearance scale to seek a balance between old and new information. Results on 26 4D echocardiographic images show the proposed method is more accurate, more robust, and faster than the previous batch algorithm.

1 Introduction

Segmentation of 4D echocardiography plays an important role in the quantitative analysis that provides important cardiac functional parameters such as ejection fraction and strain. Due to gross intensity inhomogeneities, characteristic artifacts, and poor contrast, automatic segmentation of the left ventricle is particularly challenging in echocardiography. The inherent spatiotemporal coherence of echocardiographic data provides useful constraints. The key observation is that the inherent spatio-temporal consistencies regarding image appearance (e.g., speckle pattern) and shape over the sequence can be exploited to guide cardiac border estimation. Statistical models have received considerable attention. Following the seminal work of Cootes et al. on statistical shape modeling [1],

* This work was supported by NIH RO1HL082640.

a number of statistical models [2–5] have been proposed for learning spatiotemporal priors offline from a database. The main limitation of these methods is that the high level spatiotemporal patterns in routine clinical images, especially for disease cases, may deviate from the priors learned from a database.

Exploiting individual data coherence through online learning overcomes this limitation. It is particularly attractive when a database is inapplicable or unavailable. Sparse representation and dictionary learning have recently been successfully applied to modelling local image appearance and segmenting left ventricular borders in 4D echocardiography [6, 7]. Dictionary learning on the fly exploits the spatiotemporal coherence inherent to individual data and achieves promising segmentation results [6]. However, these methods use classical second-order batch procedures for dictionary learning. The batch algorithm assumes a fixed-size dataset and accesses the whole training set at each iteration. It is memory-consuming and computationally expensive. It can be impractical when the training set is large. Every time new data is added to the training set, the dictionary needs to be retrained on the new complete training set in order to incorporate the new information, which makes the batch algorithm inefficient for dynamically changing data and online learning. In [6, 7], the appearance dictionaries are trained only on the last segmented frame rather than all the previous frames. This accelerates error accumulation and compromises the segmentation accuracy and reliability, especially for endocardial borders.

To overcome these limitations, we present a stochastic online dictionary learning approach for segmenting left ventricular borders from 4D echocardiography. It utilizes a stochastic optimization technique and processes a mini-batch of samples at a time, which results in lower memory consumption and lower computational cost than classical second-order batch algorithms. In contrast to the previous methods, our stochastic online learning procedure optimizes the dictionaries and the corresponding weights by aggregating the information of all the past frames while adapting the dictionaries to the latest segmented frame. The past information is carried forward by sufficient statistics. We weight the past information to control the rate at which the past information is updated by the new information. This updating rate varies with appearance scale to maintain a balance between old and new information.

2 Methods

2.1 Segmentation Framework

We employ a frame-by-frame sequential segmentation procedure interlaced with dictionary learning on the fly introduced in [6, 7]. Multiscale appearance dictionaries are dynamically updated each time a new frame is segmented. In a maximum a posteriori (MAP) framework, we estimate the shape S_t in frame I_t given the knowledge of $\hat{S}_{1:t-1}$ and $I_{1:t}$:

$$\hat{S}_t = \arg \max_{S_t} p(S_t | \hat{S}_{1:t-1}, I_{1:t}). \quad (1)$$

It is approximated by a decomposition of information into intensity I_t , local appearance discriminant R_t , and shape prediction S_t^* :

$$\hat{S}_t \approx \arg \max_{S_t^*} p(S_t^*|S_t)p(R_t|S_t)p(I_t|S_t)p(S_t). \quad (2)$$

The discriminant R_t summarizing multiscale local appearance dominates the estimation. It is predicted by multiscale appearance dictionaries D_t that are derived from $\hat{S}_{1:t-1}$ and $I_{1:t-1}$ through sparse representation and dictionary learning. In [6, 7], the dictionaries D_t are trained only on $\hat{S}_{t-1}$, I_{t-1} . The knowledge of the previous information is not fully utilized. This paper focuses on computing D_t more efficiently and reliably and achieving more accurate and reliable discriminant R_t . Further details of solving (2) can be found in [6].

2.2 Multiscale Sparse Representation

Let Ω denote the 3D image domain. We describe a pixel $\mathbf{u} \in \Omega$ in frame I_t with a series of appearance vectors $\mathbf{y}_t^k(\mathbf{u}) \in \mathbb{R}^n$ at different appearance scales $k = 1, \dots, J$. $\mathbf{y}_t^k(\mathbf{u})$ is constructed by concatenating orderly the pixels in a local block centered at $\mathbf{u}$ and normalized to unit length. Complementary multiscale appearance information is extracted at different levels of Gaussian pyramid. A shape S_t in I_t is represented by a level set function $\Phi_t(\mathbf{u})$. The regions of interest are two band regions $\Omega_t^1 = \{\mathbf{u} \in \Omega : 0 \leq \Phi_t(\mathbf{u}) < \psi_2\}$ and $\Omega_t^2 = \{\mathbf{u} \in \Omega : 0 > \Phi_t(\mathbf{u}) > -\psi_1\}$ which form two appearance classes. Let $\{\mathbf{D}_t^1, \mathbf{D}_t^2\}_k$ denote two dictionaries adapted to appearance classes Ω_t^1 and Ω_t^2 respectively at scale k . Under a sparse linear model, an appearance vector $\mathbf{y} \in \mathbb{R}^n$ can be decomposed as a sparse linear combination of the atoms from a dictionary $\mathbf{D} \in \mathbb{R}^{n \times K}$ which encodes the typical patterns of a corresponding appearance class. That is, $\mathbf{y} \approx \mathbf{D}\mathbf{x}$, and $\|\mathbf{x}\|_0$ is small. How well $\mathbf{y}_t^k(\mathbf{u})$ is sparsely represented by the appearance dictionary $\{\mathbf{D}_t^c\}_k$ is measured by the reconstruction residue:

$$\{R_t^c(\mathbf{u})\}_k = \|\mathbf{y}_t^k(\mathbf{u}) - \{\mathbf{D}_t^c \hat{\mathbf{x}}_t^c(\mathbf{u})\}_k\|_2 \quad (3)$$

$\forall k \in \{1, \dots, J\}$ and $c \in \{1, 2\}$, where

$$\{\hat{\mathbf{x}}_t^c(\mathbf{u})\}_k = \arg \min_{\mathbf{x}} \|\mathbf{y}_t^k(\mathbf{u}) - \{\mathbf{D}_t^c\}_k \mathbf{x}\|_2^2 \text{ s.t. } \|\mathbf{x}\|_0 \leq T, \quad (4)$$

where T is a sparsity factor. The residue indicates the likelihood $\mathbf{u}$ is in class c . Combining the multiscale information, we define the discriminant as

$$R_t(\mathbf{u}) = \sum_{k=1}^J \left[\left(\log \frac{1}{\beta_t^k} \right) \text{sgn}(\{R_t^2(\mathbf{u})\}_k - \{R_t^1(\mathbf{u})\}_k) / \sum_{j=1}^J \left(\log \frac{1}{\beta_t^j} \right) \right], \quad (5)$$

$\forall \mathbf{u} \in \Omega$, where β_t^k 's are the weighting parameters of the J appearance scales.

2.3 Stochastic Online Dictionary Learning

Learning a dictionary $\mathbf{D} \in \mathbb{R}^{n \times K}$ from a finite training set $\mathbf{Y} = [\mathbf{y}_1, \dots, \mathbf{y}_M] \in \mathbb{R}^{n \times M}$ is to solve a joint optimization problem with respect to the dictionary $\mathbf{D}$ and the sparse representation coefficients $\mathbf{X} = [\mathbf{x}_1, \dots, \mathbf{x}_M] \in \mathbb{R}^{K \times M}$:

$$\min_{\mathbf{D}, \mathbf{X}} \frac{1}{2} \|\mathbf{Y} - \mathbf{DX}\|_2^2 + \lambda \sum_{i=1}^M \|\mathbf{x}_i\|_q, \quad (6)$$

where $\|\mathbf{x}\|_q$ is a sparsity-inducing regularization that can be ℓ_0 pseudo norm or ℓ_1 norm. Classic algorithms for dictionary learning are second-order iterative batch algorithms such as the K-SVD [8] algorithm that is used in [6, 7]. The batch algorithm accesses the whole training set at each iteration and is memory consuming and computationally expensive. It may become impractical in the case of large training sets. This problem is aggravated when the data is dynamically changing over time like echocardiography, since the dictionary needs to be retrained on the new complete dataset each time new data is available. In [6, 7], the appearance dictionaries are updated each time a new frame is segmented, but they are only optimized on the newly segmented frame rather than all the previous frames. This accelerates accumulation of errors, especially at endocardial borders where there are often large deformations.

Stochastic online learning technique proposed in [9] can be used to overcome these limitations. It has recently been applied to shape modeling [10]. It processes one element of the training set at a time, which particularly suits applications with large training sets or image sequence analysis. It alternates classic sparse coding steps with dictionary update steps where the new dictionary $\mathbf{D}_m$ at m th iteration minimizes a surrogate for the empirical cost (6):

$$\mathbf{D}_m = \arg \min_{\mathbf{D}} \frac{1}{m} \sum_{i=1}^m \left(\frac{1}{2} \|\mathbf{y}_i - \mathbf{D}\mathbf{x}_i\|_2^2 + \lambda \|\mathbf{x}_i\|_1 \right) \quad (7)$$

where sufficient statistics $\mathbf{x}_i$ computed during the previous steps aggregate the past information. The past information is carried forward in matrices:

$$\mathbf{A}_m = \mathbf{A}_{m-1} + \mathbf{x}_m \mathbf{x}_m^T \text{ and } \mathbf{B}_m = \mathbf{B}_{m-1} + \mathbf{y}_m \mathbf{x}_m^T, \quad (8)$$

which enables optimizing dictionaries on the past information without accessing the past data again. Then the dictionary update step (7) is reduced to solving (9) with initialization $\mathbf{D}_{m-1}$. This procedure leads to faster performance and better dictionaries than classical batch algorithms [9]. It converges almost surely to a stationary point of the cost function and scales up gracefully to large datasets [9]. For dynamic data, the dictionary is dynamically updated by the new data while optimized on the whole dataset. Here we use a variant of [9] as summarized in Algorithm 1. We use a mini-batch extension that accesses a mini-batch of η samples per iteration to accelerate convergence. We assign weights ϱ to the past training data to control the rate of updating out-of-date information.

We introduce a stochastic online learning process supervised in a boosting framework [11] as detailed in Algorithm 2. Algorithm 1 is invoked to enforce the

Algorithm 1. Stochastic Online Dictionary Learning

Require: training set $\mathbf{y} \in \mathbb{R}^n \sim p(\mathbf{y})$, sparsity weight λ , initial dictionary $\mathbf{D}_{m_I} \in \mathbb{R}^{n \times K}$, initial iteration number m_I and terminal iteration number m_T , mini-batch size η , weight ϱ , and initial matrices $\mathbf{A}_{m_I}$ and $\mathbf{B}_{m_I}$.

for $m = m_I$ to m_T **do**

 Draw η samples $\mathbf{Y}_m = \{\mathbf{y}_{m,i}\}_{i=1}^\eta$ from $p(\mathbf{y})$

Sparse coding: $\mathbf{x}_{m,i} = \arg \min_{\mathbf{x} \in \mathbb{R}^K} \|\mathbf{y}_{m,i} - \mathbf{D}_{m-1}\mathbf{x}\|_2^2 + \lambda \|\mathbf{x}\|_1, \forall i \in \{1, \dots, \eta\}$

$\mathbf{A}_m = \varrho \mathbf{A}_{m-1} + \frac{1}{\eta} \sum_{i=1}^\eta \mathbf{x}_{m,i} \mathbf{x}_{m,i}^T$, $\mathbf{B}_m = \varrho \mathbf{B}_{m-1} + \frac{1}{\eta} \sum_{i=1}^\eta \mathbf{y}_{m,i} \mathbf{x}_{m,i}^T$.

Update dictionary: compute $\mathbf{D}_m$ with $\mathbf{D}_{m-1}$ as initialization

$$\mathbf{D}_m = \arg \min_{\mathbf{D}} \frac{1}{m} \left(\frac{1}{2} \text{Tr}(\mathbf{D}^T \mathbf{D} \mathbf{A}_m) - \text{Tr}(\mathbf{D}^T \mathbf{B}_m) \right). \quad (9)$$

end for

return dictionary $\mathbf{D}_{m_T}$, and matrices $\mathbf{A}_{m_T}$ and $\mathbf{B}_{m_T}$.

reconstructive property of the dictionaries. The boosting supervision strengthens the discriminative property and optimizes the weighting of multiscale information. At each time point t , the series of multiscale appearance dictionary pairs $\{\mathbf{D}_t^1, \mathbf{D}_t^2\}_k$, matrices $\mathbf{A}_t^k$ and $\mathbf{B}_t^k$, and the corresponding weighting parameters β_t^k , $k = 1, \dots, J$, are updated by the latest segmented frame $t-1$: training samples of appearance vectors belonging to two classes $\{\mathbf{Y}_{t-1}^1\}_k = \{\mathbf{y}_{t-1}^k(\mathbf{u}) : \mathbf{u} \in \Omega_{t-1}^1\}$ and $\{\mathbf{Y}_{t-1}^2\}_k = \{\mathbf{y}_{t-1}^k(\mathbf{u}) : \mathbf{u} \in \Omega_{t-1}^2\}$. In contrast to [6, 7] where $\{\mathbf{D}_t^1, \mathbf{D}_t^2\}_k$ and β_t^k depend only on frame $t-1$, we optimize $\{\mathbf{D}_t^1, \mathbf{D}_t^2\}_k$ and β_t^k by aggregating the information of all the preceding frames (stored in $\mathbf{A}_{t-1}^k$, $\mathbf{B}_{t-1}^k$, and β_{t-1}^k). If an error occurs in one frame, it can be compensated by the information of the previous frames. The propagation of errors is alleviated. The rate of updating the past information varies with appearance scale. Let l_k be the axial width in millimeter of the local image at scale k , we set $\varrho_k = a l_k^{-2}$ where $a \in \mathbb{R}^+$. Higher ϱ 's are assigned to finer appearance scales to incorporate more past information. Lower ϱ 's are assigned to coarser appearance scales to put more emphasis on the latest information, since the coarse appearance scale is more sensitive to cardiac deformation. The stochastic online learning procedure can be initialized either by offline learning from a suitable database or by a manual tracing.

3 Results

We validated our method on 26 4D canine open-chest echocardiographic images acquired from both healthy and post-infarct animals using Phillips iE33 and an X7-2 array probe. Each image sequence spanned a cardiac cycle and contained about 25 – 30 volumes. The sequential segmentation was initialized with a manual tracing of the end-diastole volume. 100 volumes were randomly selected for expert manual segmentation and quality assessment. We evaluated automatic results against manual tracings using the following segmentation quality metrics: Hausdorff Distance (HD), Mean Absolute Distance (MAD), and Dice coefficient (DICE). We compared the proposed method to [6] that uses the batch dictionary learning technique K-SVD. The two algorithms shared the same set of relevant

Algorithm 2. Boosted Multiscale Online Dictionary Learning

Require: training sets $\{\mathbf{Y}_{t-1}^1\}_k = \{\mathbf{y}_{1,i}^k\}_{i=1}^{M_1}$ and $\{\mathbf{Y}_{t-1}^2\}_k = \{\mathbf{y}_{2,j}^k\}_{j=1}^{M_2}$, initial dictionaries $\{\mathbf{D}_{t-1}^1, \mathbf{D}_{t-1}^2\}_k$, matrices $\mathbf{A}_{t-1}^k$ and $\mathbf{B}_{t-1}^k$, accumulated # of previous iterations N_{t-1} , weighting parameters $\beta_{t-1}^k, \varrho_k, k = 1, \dots, J$, mini-batch size η , # of iterations N , and sparsity factor T .

$\mathbf{w}_1^1 = \{w_{1,i}^1\}_{i=1}^{M_1} = \mathbf{1}, \mathbf{w}_2^1 = \{w_{2,j}^1\}_{j=1}^{M_2} = \mathbf{1}.$

for $k = 1$ to J **do**

Dictionary Learning: Apply Algorithm 1 for N iterations to adapt $\{\mathbf{D}_t^1, \mathbf{D}_t^2\}_k$

to $\{\mathbf{Y}_{t-1}^1\}_k \sim \mathbf{p}_1^k = \{p_{1,i}^k\}_{i=1}^{M_1} = \frac{\mathbf{w}_1^k}{\sum_{i=1}^{M_1} w_{1,i}^k}$ and $\{\mathbf{Y}_{t-1}^2\}_k \sim \mathbf{p}_2^k = \{p_{2,j}^k\}_{j=1}^{M_2} = \frac{\mathbf{w}_2^k}{\sum_{j=1}^{M_2} w_{2,j}^k}$. Use $\varrho = \varrho_k$ for the first iteration and $\varrho = 1$ for the rest.

Sparse Coding: $\forall \mathbf{y} \in \{\mathbf{Y}_{t-1}^1, \mathbf{Y}_{t-1}^2\}_k$, solve (4) for sparse representations w.r.t. $\{\mathbf{D}_t^1\}_k$ and $\{\mathbf{D}_t^2\}_k$ and get residues $R(\mathbf{y}, \mathbf{D}_t^1)_k$ and $R(\mathbf{y}, \mathbf{D}_t^2)_k$.

Hypothesis $h_k: \mathbf{y} \in \{\mathbf{Y}_{t-1}^1, \mathbf{Y}_{t-1}^2\}_k \rightarrow \{0, 1\}$: $h_k(\mathbf{y}) = \text{Heaviside}(R(\mathbf{y}, \mathbf{D}_t^2)_k - R(\mathbf{y}, \mathbf{D}_t^1)_k)$. Calculate the error of h_k : $\epsilon_k = \frac{1}{1+\varrho_k} [\sum_{i=1}^{M_1} p_{1,i}^k |h_k(\mathbf{y}_{1,i}^k) - 1| + \sum_{j=1}^{M_2} p_{2,j}^k |h_k(\mathbf{y}_{2,j}^k)|] + \frac{\varrho_k}{1+\varrho_k} (\frac{\beta_{t-1}^k}{1+\beta_{t-1}^k})$. Set $\beta_t^k = \epsilon_k / (1 - \epsilon_k)$.

Weight Update: $w_{1,i}^{k+1} = w_{1,i}^k \beta_t^{k-1-|h_k(\mathbf{y}_{1,i}^k)-1|}, w_{2,j}^{k+1} = w_{2,j}^k \beta_t^{k-1-h_k(\mathbf{y}_{2,j}^k)}$.

end for

return $\{\mathbf{D}_t^1, \mathbf{D}_t^2\}_k, \mathbf{A}_t^k, \mathbf{B}_t^k, \beta_t^k, k = 1, \dots, J$, and $N_t = N_{t-1} + N$.

parameters. We used the following parameter setting: $J = 10, T = 2, K = 1.5n, N = 10(t > 2)$ or $20(t = 2), \eta = 2048, \lambda = 0.8$, and $a = 100$.

Figure 1 shows representative segmentation results for frames at end-systole when it is easiest to access error accumulation. In the top row, the batch method [6] resulted in more errors in end-systolic segmentations, since it learned appearance dictionaries only on the latest segmented frame and did not fully leverage the information carried in all the previous frames. The segmentation error of a frame is likely to propagate to the following frames. Images in the bottom row show the improved segmentation results by employing our new stochastic learning procedure. Since we optimize the dictionaries on all the previous frames, the error in a given frame is compensated by the information of the other frames. Figure 2 presents the quality measure curves from end-diastole to end-systole for the endocardial segmentations of a healthy sequence and a post-infarct sequence. DICE decays and HD and MAD rise from end-diastole to end-systole due to accumulation of errors. Compared to the batch method, our method resulted in flattened curves, which suggests our method effectively alleviates error accumulation and improves segmentation performance for both healthy and post-infarct images. For epicardial segmentation, the improvement was not significant, since the baseline accuracy of [6] was already very high (97% in DICE). Table 1 summarizes the statistics of segmentation quality measures and computational efficiency achieved by the two algorithms in segmenting endocardial borders. The proposed method achieved smaller mean MAD, smaller mean HD, larger mean DICE, and smaller standard deviations of all the measures. The overall segmentation accuracy and robustness were

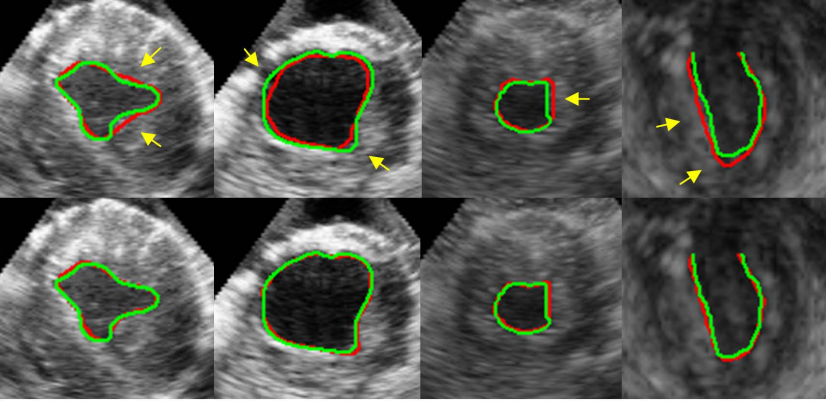


Fig. 1. Comparisons of segmentation results by the batch method (top row) and our method (bottom row). Green: Manual segmentation. Red: Automatic segmentation.

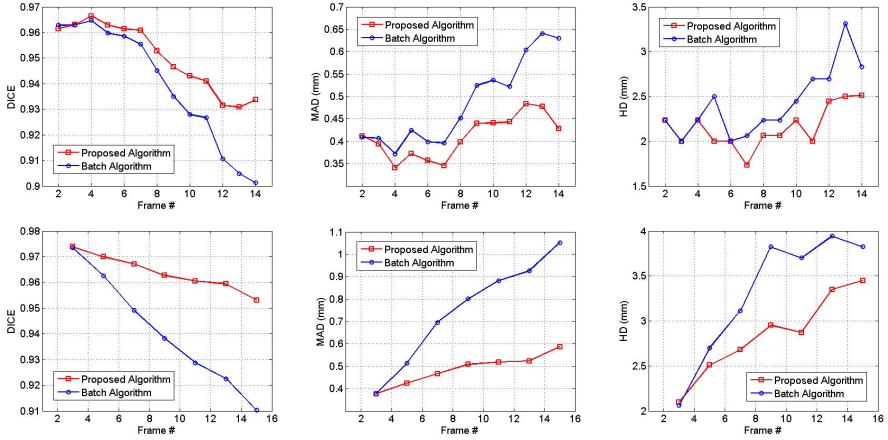


Fig. 2. Segmentation quality measures at different frames of two example sequences (healthy (top row) and post-infarct (bottom row)) from end-diastole to end-systole. Blue: the batch method. Red: the proposed method.

Table 1. Sample means $\pm$ standard deviations of the quality measures and dictionary learning time per frame for the segmentation of endocardial borders

	DICE (%)	MAD (mm)	HD (mm)	Time (s)
Batch Algorithm [6]	93.6 ± 2.49	0.57 ± 0.14	2.95 ± 0.62	~ 45
Proposed Algorithm	94.6 ± 2.17	0.48 ± 0.11	2.83 ± 0.53	~ 25

effectively improved using our stochastic online learning procedure. We tested the two algorithms on a laptop with Intel quad-core 2.2 GHz CPU and 8 GB memory. Both algorithms were implemented with a mixture of MATLAB and C++. The batch algorithm took about 45 seconds per frame for dictionary learning. The proposed algorithm took only about 25 seconds per frame.

4 Conclusion

We have presented an approach for segmenting left ventricular borders from 4D echocardiography using stochastic online dictionary learning. It is based on a stochastic optimization technique resulting in lower memory consumption and computational cost than classical batch algorithms. We optimize the dictionaries and their weights on all the preceding frames while adapting them to the latest segmented frame. The rate of updating the past information is controlled and varies with appearance scale. Our method effectively improved the accuracy and robustness of endocardial segmentation and computational efficiency compared to the previous batch methods. Future work will include automating the dictionary initialization through offline learning. The stochastic learning procedure is suitable for both offline and online learning. A database that is too large for batch methods can be gracefully handled by our method which avoids accessing the database during online learning. Our method can ultimately be used to build an integrated offline and online learning framework.

References

1. Cootes, T.F., Edwards, G.J., Taylor, C.J.: Active appearance models. *IEEE TPAMI* 23(6), 681–685 (2001)
2. Bosch, J.G., Mitchell, S.C., Lelieveldt, B.P.F., Nijland, F., Kamp, O., Sonka, M., Reiber, J.H.C.: Automatic segmentation of echocardiographic sequences by active appearance motion models. *IEEE TMI* 21(11), 1374–1383 (2002)
3. Jacob, G., Noble, J.A., Behrenbruch, C.P., Kelion, A.D., Banning, A.P.: A shape-space based approach to tracking myocardial borders and quantifying regional left ventricular function applied in echocardiography. *IEEE TMI* 21(3), 226–238 (2002)
4. Sun, W., Çetin, M., Chan, R., Reddy, V., Holmvang, G., Chandar, V., Willsky, A.S.: Segmenting and tracking the left ventricle by learning the dynamics in cardiac images. In: Christensen, G.E., Sonka, M. (eds.) *IPMI 2005*. LNCS, vol. 3565, pp. 553–565. Springer, Heidelberg (2005)
5. Zhu, Y., Papademetris, X., Sinusas, A.J., Duncan, J.S.: A dynamical shape prior for LV segmentation from RT3D echocardiography. In: Yang, G.-Z., Hawkes, D., Rueckert, D., Noble, A., Taylor, C. (eds.) *MICCAI 2009, Part I*. LNCS, vol. 5761, pp. 206–213. Springer, Heidelberg (2009)
6. Huang, X., Dione, D.P., Compas, C.B., Papademetris, X., Lin, B.A., Sinusas, A.J., Duncan, J.S.: A Dynamical Appearance Model Based on Multiscale Sparse Representation: Segmentation of the Left Ventricle from 4D Echocardiography. In: Ayache, N., Delingette, H., Golland, P., Mori, K. (eds.) *MICCAI 2012, Part III*. LNCS, vol. 7512, pp. 58–65. Springer, Heidelberg (2012)
7. Huang, X., Lin, B.A., Compas, C.B., Sinusas, A.J., Staib, L.H., Duncan, J.S.: Segmentation of left ventricles from echocardiographic sequences via sparse appearance representation. In: *MMBIA*, pp. 305–312 (2012)
8. Aharon, M., Elad, M., Bruckstein, A.: K-SVD: An algorithm for designing overcomplete dictionaries for sparse representation. *IEEE TSP* 54(11), 4311–4322 (2006)

9. Mairal, J., Bach, F., Ponce, J., Sapiro, G.: Online dictionary learning for sparse coding. In: ICML, p. 87 (2009)
10. Zhang, S., Zhan, Y., Zhou, Y., Uzunbas, M., Metaxas, D.N.: Shape prior modeling using sparse representation and online dictionary learning. In: Ayache, N., Delingette, H., Golland, P., Mori, K. (eds.) MICCAI 2012, Part III. LNCS, vol. 7512, pp. 435–442. Springer, Heidelberg (2012)
11. Freund, Y., Schapire, R.: A decision-theoretic generalization of on-line learning and an application to boosting. In: Vitényi, P.M.B. (ed.) EuroCOLT 1995. LNCS, vol. 904, pp. 23–37. Springer, Heidelberg (1995)