

Robust PCA Unrolling Network for Super-Resolution Vessel Extraction in X-Ray Coronary Angiography

Binjie Qin¹, Member, IEEE, Haohao Mao¹, Yiming Liu¹, Jun Zhao¹, Member, IEEE, Yisong Lv¹, Yueqi Zhu, Song Ding, and Xu Chen²

Abstract—Although robust PCA has been increasingly adopted to extract vessels from X-ray coronary angiography (XCA) images, challenging problems such as inefficient vessel-sparsity modelling, noisy and dynamic background artefacts, and high computational cost still remain unsolved. Therefore, we propose a novel robust PCA unrolling network with sparse feature selection for super-resolution XCA vessel imaging. Being embedded within a patch-wise spatiotemporal super-resolution framework that is built upon a pooling layer and a convolutional long short-term memory network, the proposed network can not only gradually prune complex vessel-like artefacts and noisy backgrounds in XCA during network training but also iteratively learn and select the high-level spatiotemporal semantic information of moving contrast agents flowing in the XCA-imaged vessels. The experimental results show that the proposed method significantly outperforms state-of-the-art methods, especially in the imaging of the vessel network and its distal vessels, by restoring the intensity and geometry profiles of heterogeneous vessels against complex and dynamic backgrounds. The source code is available at <https://github.com/Binjie-Qin/RPCA-UNet>

Index Terms—Algorithm unrolling, RPCA unrolling network, X-ray coronary angiography, vessel extraction, sparse feature selection, super-resolution.

Manuscript received 14 March 2022; revised 9 May 2022; accepted 21 May 2022. Date of publication 23 May 2022; date of current version 27 October 2022. This work was supported in part by the Science and Technology Commission of Shanghai Municipality under Grant 19dz1200500 and Grant 19411951507; in part by the National Natural Science Foundation of China under Grant 61271320 and Grant 82070477; in part by the Shanghai Shenkang Hospital Development Center under Grant SHDC12019X12; and in part by the Interdisciplinary Program of Shanghai Jiao Tong University under Grant ZH2018ZDA19, Grant YG2021QN122, and Grant YG2021QN99. (Corresponding author: Binjie Qin.)

Binjie Qin, Haohao Mao, Yiming Liu, and Jun Zhao are with the School of Biomedical Engineering, Shanghai Jiao Tong University, Shanghai 200240, China (e-mail: bj Qin@sjtu.edu.cn; oo_mhh@sjtu.edu.cn; m201812235@sjtu.edu.cn; junzhao@sjtu.edu.cn).

Yisong Lv is with the School of Continuing Education, Shanghai Jiao Tong University, Shanghai 200240, China (e-mail: yslv@sjtu.edu.cn).

Yueqi Zhu is with the Department of Radiology, Shanghai Jiao Tong University Affiliated Sixth People's Hospital, Shanghai Jiao Tong University, Shanghai 200233, China (e-mail: zhuyueqi@hotmail.com).

Song Ding is with the Department of Cardiology, Ren Ji Hospital, School of Medicine, Shanghai Jiao Tong University, Shanghai 200127, China (e-mail: dingsong@renji.com).

Xu Chen is with the Center for Advanced Neuroimaging, University of California at Riverside, Riverside, CA 92521 USA (e-mail: xu.chen@ucr.edu).

Digital Object Identifier 10.1109/TMI.2022.3177626

I. INTRODUCTION

CARDIOVASCULAR diseases (CVDs) threaten human health worldwide [1]. Percutaneous coronary intervention (PCI) is very important for the diagnosis and treatment of CVDs, during which X-ray coronary angiography (XCA) is a primary technique for imaging morphological and functional information about blood vessels. Due to X-ray beams being attenuated by varying amounts when they pass through tissues with different densities along the projection path of XCA imaging, the XCA sequence displays heterogeneous blood vessels that overlap with various anatomical structures (such as bones, lungs and diaphragms), mixed Poisson-Gaussian noise [2], [3], and respiratory and cardiac motions. It is very difficult for surgeons to clearly identify blood vessels, let alone extract vessels for the quantitative analysis of a vessel's structure and function. Vessel extraction algorithms [4], [5] are usually built upon tube-like feature representation, which is very sensitive to noisy and dynamic background artefacts. Recently, XCA vessel extraction [6]–[12] was regarded as the separation of foreground vessels and background structures within sparse and low-rank modelling via robust principal component analysis (RPCA) [13] to achieve state-of-the-art performance. However, the extracted results in these studies still have some noisy artefacts.

Moreover, RPCA-based vessel extraction consumes a large amount of storage and time. Therefore, a deep neural network called a convolutional robust PCA (CORONA) [14] unfolds the RPCA algorithm for ultrasonic vascular imaging. Algorithm unrolling [15] or unfolding was first introduced in [16], the result of which being that the time efficiency of the unfolded deep network is greatly improved compared with that of the original iterative algorithm. However, the problems caused by the dynamic background and complex noise patterns still remain in the CORONA solution when it is used to extract vessels from XCA images. Furthermore, the mixed Gaussian-Poisson noise in XCA is complex and heterogeneous in different XCA imaging machines and locally affects the extraction of distal vessels with low contrast and low SNR. RPCA-based methods and CORONA globally implement foreground/background decomposition and cannot effectively cope with the local interference caused by mixed noise and heterogeneous artefacts.

To solve these problems, we propose a patch-wise spatiotemporal super-resolution (SR) module to refine the vessel features outputted by the deep unfolded RPCA layer. Specifically, a feature pooling layer inputs the original data to the unfolded RPCA layer and applies patch-wise sparse feature selection in the SR module to eliminate redundant vessel-like artefacts while retaining the useful vessel features [17]. Different from current deep-learning-based SR neural networks [18] that automatically extract features for a non-linear low-resolution to high-resolution mapping and cannot select local/non-local sparse features from a single image or video, the proposed RPCA unrolling network, called RPCA-UNet, implements a patch-wise spatiotemporal SR module with sparse feature selection that is based on a residual module and a convolutional long short-term memory (CLSTM) network [19]. RPCA-UNet can effectively enhance patch-wise vessel features by extracting not only the heterogeneous grey level information but also the geometrical structures of XCA vessels in a spatiotemporally consistent way. Specifically, the residual module is first applied to extract deep features through multiple convolutional layers and transfer these features along with the original features via the residual operation to the subsequent CLSTM network. By saving complementary features of pervious frame in an XCA sequence, the CLSTM network integrates the features of the current frame into the complementary features of previous frames. This feature aggregation establishes a spatiotemporal evolution for accurately extracting both the image grey values and geometrical features of XCA vessels. The main contribution of this work is threefold:

- 1) A novel RPCA unrolling (or unfolded RPCA) network with a patch-wise SR module is proposed to iteratively extract XCA vessels with a certain time and space efficiency. The unrolling network in each iteration/layer has a pooling layer as the preprocessing layer and a patch-wise SR module as the postprocessing layer that consists of a residual module and a CLSTM network. The proposed RPCA-UNet can not only achieve uninformative feature pruning and Gaussian-Poisson denoising but also selectively learn sparse vessel features from complex and dynamic backgrounds. To the best of our knowledge, this is the first use a RPCA unrolling network with sparse feature selection to extract vessels from XCA images. Experiments show that the proposed method significantly outperforms state-of-the-art methods in both vessel extraction and vessel segmentation.
- 2) We apply CLSTM network to the proposed SR module that can not only learn sparse features selectively from the current frame of XCA sequence but also preserve the high-level spatiotemporal semantic detail of moving contrast agents in the whole XCA sequence. CLSTM network in a patch-wise SR network is proven to boost the performance of vessel extraction by significantly improving the distal vessel detection accuracy and spatiotemporal consistency in the XCA sequence.
- 3) The proposed RPCA-UNet is implemented as a weakly supervised learning method such that grey value vessel labelling is automatically generated by our vessel

extraction method, called VRBC (vessel region background completion) [9], and our training data and testing data comprise heterogeneous XCA images that are collected from different machines. This weakly supervised learning in a heterogeneous environment overcomes the need of expensive and time-consuming manual annotation and improves a generalization ability of the proposed network.

II. RELATED WORKS

A. XCA Vessel Extraction

Compared with other imaging modalities reviewed in recent survey studies [4], [5], [20], such as computed tomography angiography, magnetic resonance angiography and retinal fundus images, few studies on extracting vessels from XCA images have been conducted. XCA vessel extraction methods [21] can be divided into the following four categories that corporately transform XCA images into segmentation results: vessel enhancement, deformable model, vessel tracking, and machine learning. Vessel enhancement approaches [22]–[25] aggregate compact image patches in local/non-local filtering to enhance the tube-like vessel features and increase the contrast between the foreground vessels and background structures. For example, Hessian-based multiscale local or non-local [26] filtering yielding geometrical features [27] regarding both vesselness and direction information was incorporated into an iterative region growing [24], a statistical region merging [25], and multiscale superpixels [23] to group enhanced pixels into correct clusters of arteries and background. However, Hessian-matrix-based segmentation is largely dependent on the optimal scale selection of the major vessel radius and highly sensitive to mixed Gaussian-Poisson noise in the spatial domain. In the frequency domain, single-scale Gabor filters with optimized parameters [22] and multiscale Gabor filters with optimized response thresholding [28] are developed for XCA vessel segmentation, but the thresholding techniques in the Gabor and wavelet domains cannot easily distinguish the vascular structures from many vessel-like artefacts. Inspired by phase congruency, which has stability in the presence of noise and is invariant to changes in contrast, some detail-preserving image enhancement methods have exploited phase-congruency-based feature indicator called phase symmetry or phase asymmetry [29], [30] to detect edge- and ridge-like features such as 2D/3D vessels [29], [31]. Nevertheless, vessel enhancement methods can simultaneously enhance the vessel-like background structures in XCA images.

Another common method is using deformable models to segment vessels, which have parametric and geometric deformable models. Parametric models such as active contour model [32] directly represent the target curves or surfaces during deformation, resulting in an efficient and lower computational cost segmentation, but are not suitable for XCA vessels with complex topologies. In contrast, geometric deformable models are implemented as an energy minimization within a level set framework [33], [34], which can be adapted to changing vessel topologies and effectively extract thin vessels and complex vessel branches. Currently, integrating not only the

edge and centerline information [35] but also the region [36] and shape prior constraints [37] into the optimization model can lead to more precise XCA vessel segmentation. However, deformable models have strong dependence on initialization, high sensitivity to irregular vessel shapes with inhomogeneous intensity and low contrast, and high computational cost.

Additionally, vessel-tracking methods also attract much attention. Vessel-tracking methods usually place initial seed points and drive the growth process with specific constraints to segment the vessel area. They are generally divided into two categories by different tracking constraints: model-based [38] and minimal path [39]–[41] methods. Model-based techniques track vessels by searching for and matching a predefined vessel model with different shapes and thicknesses, but their matching performance decreases sharply on images with high noise and inhomogeneous intensity as well as many non-vascular structures. Minimal path methods [39], [40] can efficiently extract the XCA vessel centreline by finding the path with minimal accumulated cost between two given end-points via centreline evolution over a filter-response-derived vectorial multiscale feature image [39] or via the backtracking operation [40]. The work in [41] has extracted the complete vessel lumen within the framework of backtracked minimal path propagation. Due to a lack of global context feature selection, the centreline extraction method has difficulty in avoiding under-segmentation in some clear gaps between vessel structures with low-contrast-intensity inhomogeneities or in some weak vessel structures such as distal vessels.

Machine learning methods such as RPCA- [6]–[12] and graph-based [38], [42] methods treat segmentation as an optimized classification to distinguish between foreground and background pixels [5]. However, the globally optimal solutions have several unsolved problems such as discriminative feature representation, spatiotemporal regularization, and mixed Gaussian-Poisson noise removal. With the ability to perform featurization and classification of big data, deep-learning-based methods, especially the convolutional neural network (CNN) combined with image enhancement [43] for preprocessing and graph-based vessel connection [44] for post-processing as well as pyramid pooling and the convolving of multiscale features with small sample transfer learning [45], have proven effective in XCA segmentation but still have several unsolved problems related to spatiotemporal and semantic context modelling. Using an encoder-decoder architecture equipped with skipping connections, U-Net [46] and fully convolutional networks (FCNs) combine high-level semantic information with low-level appearance details to efficiently achieve end-to-end semantic segmentation of entire vessel trees [45], [47]–[50]. For example, SVS-net [48] embedding channel attention mechanism for suppressing noisy backgrounds and the spatiotemporal FCN [49] integrating interframe information with influence layers are proposed to extract multiscale features for segmenting entire vessels from XCA sequence. However, deep networks of a certain depth for accumulating multiscale feature have limitations in extracting more features and handling details such as distal vessels.

Unfortunately, most deep learning techniques have limitations in not only efficiently extracting more spatiotemporal

features in a sequential way but also discriminatively selecting sparse vessel features from vessel-like and signal-dependent noisy backgrounds. To the best of our knowledge, none of the current vessel extraction methods can fully restore the intensity and geometry profiles of entire heterogeneous XCA vessels, except VRBC method [9].

B. Unrolling Neural Network

The unrolling neural network was first proposed by Gregor and LeCun [16] to approximate the iterative soft-threshold algorithm (ISTA) for sparse coding. The unfolded network called the learned ISTA (LISTA) achieves great performance, being nearly 20 times faster than the traditional accelerated ISTA. The success of the LISTA shows the significant computational benefits of the deep unfolding algorithm. Moreover, by considering each iteration of an iterative algorithm as a layer of an unrolling network and then concatenating a few of these layers, one needs only a few iterations of training to achieve a dramatic improvement in convergence.

Recently, algorithm unrolling has attracted significant attention in signal and image processing [15], where the collection of sufficient data is expensive and difficult to achieve and the performance of conventional networks is limited. Solomon *et al.* [14] proposed CORONA to separate the blood vessels and background tissues from an ultrasound signal. CORONA was trained on simulated data, and then the resulting network was trained on in vivo data. This hybrid policy can not only improve the network performance but also achieve a fully automated network, in which all the regularization parameters are also learned. Moreover, by exploiting spatial invariance and facilitating the training process, one can reduce the number of learnable parameters dramatically through the use of convolutional layers.

Algorithm unrolling shows great potential in solving inverse problems in biomedical imaging. Xiang *et al.* [51] unfolded the fast ISTA (FISTA) framework into FISTA-Net, which achieves great performance in different imaging tasks. Algorithm unrolling has been expanded to the graph domain and designed as an interpretable architecture from a signal processing perspective [52]. The graph unrolling networks [52] are trained through unsupervised learning, where the input noisy measurements are used to supervise the neural network training. The network output does not overfit the noisy input in most cases, indicating that the unrolling network can carry out implicit graph regularization and thus avoid overfitting.

C. Feature Selection

Feature selection [53] can improve learning accuracy while defying the curse of dimensionality of high-dimensional data in an efficient way. For example, PCI needs an accurate real-time navigation system to visualize and navigate inside the vessel network, which presents structure-functional information about the cardiac perfusion; thus, an XCA image could be categorized into vessel and non-vessel features in deep-learning-based image analysis. Therefore, feature selection [53] can be used to find the most appropriate lightweight feature subset that preserves relevant vessel information while discarding the non-vessel and artefact features. Generally,

feature selection techniques can be classified into four main categories: filter, wrapper, embedded, and hybrid methods.

Filter methods evaluate feature relevance in discriminating different classes according to predefined criteria without using any learning algorithm. The criteria include information theoretic criteria such as mutual information [54] and multivariate joint entropy [55]. Filter methods are fast, but their selected subset is usually not an optimal feature subset from the perspective of classification performance. Recently, feature interactions among multiple variables [55] and views [56] in multisource heterogeneous data environments were studied to increase the classification accuracy.

Wrapper methods select features to achieve the best performance of a specific learning algorithm [57], [58]. Although various algorithms are used to accelerate the combinatorial optimization for maximizing the relevancy to the target class and minimizing the redundancy of selected features, the computational cost of wrapper methods is still too high, especially when the number of selected features greatly increases. Therefore, hybrid approaches [59] that use filter methods to remove the irrelevant features and then select important features from a candidate subset by wrapper models are developed to achieve the best possible performance by a particular learning algorithm with time complexity similar to that of the filter methods.

Because feature selection is embedded in the training of a learning machine, embedded methods are better than other methods in jointly achieving high classification performance and computational efficiency. Typically, the popular sparse learning models implement embedded feature selection [60] by minimizing an empirical error penalized by a regularization term such as the $l_{r,p}$ -norm regularizer. For example, the traditional RPCA-based and CORONA methods usually apply the $l_{1,2}$ -norm to select sparse features and eliminate redundant features. However, these approaches select solely individual sparse features and ignore the possible interaction between different features. Therefore, group sparse feature selection was recently developed [61] to model the strengths of interactions between different features with graph-based edge weights and to partition the graph into different groups in terms of their strengths. However, how to combine sparse feature selection with neural networks [62], [63] is still a poorly understood and unexplored research area. While deep neural networks can automatically extract features appropriate for the target task and use an attention mechanism [48], [64] to weigh the different features to increase the classification performance, they usually cannot check important input signals and select sparse features based on some predefined criteria, which leads to a lack of interpretability. To solve this problem, an attempt at feature selection is conducted in the proposed RPCA-UNet to gain prediction accuracy and computational efficiency with existing XCA data. Feature selection for pruning neural network and reinforcement learning is beyond the scope of our paper, we refer the interested reader to the related works [65], [66].

III. METHOD

The overall architecture of each iteration/layer from RPCA-UNet is shown in Fig. 1 for decomposing a given XCA

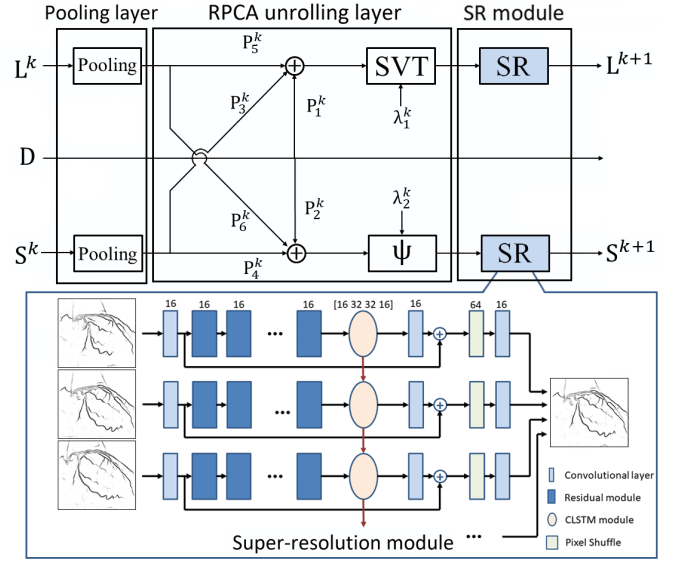


Fig. 1. The architecture of a single iteration/layer of RPCA-UNet for decomposing XCA data D into vessel (S) and background (L) components, which consists of a pooling layer, an RPCA unrolling layer, and an SR module. The SR module is mainly built upon the convolutional layer, residual module and CLSTM network.

data D into the sum of a vessel (S) and a background (L) component. It is difficult for data-driven neural networks to build models for removing the underlying mixed noises and artefacts. We then focus on sparse feature selection, which plays an important role in RPCA-UNet. Specifically, RPCA-UNet in each layer has a feature selection module that combines a pooling layer as the subsampling-based pre-processing layer with a patch-wise SR module as the post-processing layer, which consists of a convolutional layer, a residual module and a CLSTM network.

A. RPCA Modelling

The contrast agents move quickly in vessels and can be regarded as sparse foreground anomalies. Then, background structures in slow motions can be treated as the low-rank component. Therefore, RPCA is defined to decompose a given XCA data D into a sum of a vessel and a background component:

$$\min \|L\|_* + \lambda \|S\|_1 \text{ s.t. } D = L + S \quad (1)$$

where L is the slowly changing background layer, which can be described as a low-rank matrix, and S is the foreground layer, which can be described as a sparse matrix. $\|\cdot\|_*$ is the nuclear norm (which is the sum of its singular values), $\|\cdot\|_1$ is the l_1 -norm regularization, and λ is a regularizing parameter to control the detection sensitivity to a number of outliers. The RPCA problem is further expanded into a more general form with the input image D being defined as [14]:

$$D = H_1 L + H_2 S + N \quad (2)$$

where H_1 and H_2 are the measurement matrices of L and S , respectively, representing the physical acquisition mechanisms which are determined by the characteristics of measurement

devices (in XCA images, $H_1 = H_2 = I$), and N is the additive noise. The RPCA is then formulated in a Lagrangian form as:

$$\min \frac{1}{2} \|M - H_1 L - H_2 S\|_F^2 + \lambda_1 \|L\|_* + \lambda_2 \|S\|_{1,2} \quad (3)$$

where $\|\cdot\|_{1,2}$ is the mixed $l_{1,2}$ -norm and λ_1 and λ_2 are the regularizing parameters of L and S , respectively. The mixed $l_{1,2}$ -norm is applied in the cardiovascular extraction task because the structure and position of the vessels change slowly between adjacent frames. We can define this as:

$$X = \begin{bmatrix} L \\ S \end{bmatrix}, P_1 = \begin{bmatrix} I \\ 0 \end{bmatrix}, P_2 = \begin{bmatrix} 0 \\ I \end{bmatrix}, A = \begin{bmatrix} H_1 \\ H_2 \end{bmatrix} \quad (4)$$

Then, Equation (3) can be rewritten as

$$\min_{L,S} \frac{1}{2} \|D - AX\|_F^2 + h(X) \quad (5)$$

where $h(X) = \lambda_1 \|P_1 X\|_* + \lambda_2 \|P_2 X\|_{1,2}$. Thus, the minimization problem (5) can be regarded as a regularized least-squares problem, which can be solved by the iterative shrinkage/thresholding algorithm, where L and S are iteratively updated until the formula reaches its minimum. L^{k+1} and S^{k+1} at iteration $k+1$ can be updated [14] via

$$L^{k+1} = SVT_{\lambda_1/L_f}(I - \frac{1}{L_f} H_1^H H_1) L^k - H_1^H H_2 S^k + H_1^H D \quad (6)$$

$$S^{k+1} = \psi_{\lambda_2/L_f}(I - \frac{1}{L_f} H_2^H H_2) L^k - H_2^H H_1 S^k + H_2^H D \quad (7)$$

where SVT_{λ_1/L_f} is the singular-value thresholding operator, ψ_{λ_2/L_f} is the soft-thresholding operator, and L_f is the Lipschitz constant.

B. RPCA Unrolling Network

Traditional iterative algorithm can be unfolded into a deep neural network, where each layer of the network is represented as one iteration of the algorithm [15]. Thus, passing through the layers of the deep unfolded network can be viewed as calculation via the iterative algorithm a finite number of times. Following the principle of deep unfolding, the matrices dependent on H_1 and H_2 in equations (6) and (7) can be replaced with convolutional kernels. We form a deep network by employing convolutional layers $P_1, P_2, \dots, P_6$ to replace the matrices dependent on H_1 and H_2 . Convolutional layers are applied rather than fully connected layers, aimed at reducing the number of learned parameters to improve the time and space efficiency of the network. Then, the equations for the k th layer in the unfolded network are computed as follows:

$$L^{k+1} = SVT_{\lambda_1}^k P_5^k * L^k + P_3^k * S^k + P_1^k * D \quad (8)$$

$$S^{k+1} = \psi_{\lambda_2}^k P_6^k * S^k + P_4^k * L^k + P_2^k * D \quad (9)$$

where $*$ denotes a convolutional operator. The diagram of a single layer of the unfolded network is shown in Fig. 1. Here, the convolutional layers $P_1^k, \dots, P_6^k$, regularization parameters, and λ_1^k and λ_2^k are learned during the training process.

C. Patch-Wise Super-Resolution Module

Unrolling RPCA directly to extract XCA vessels faces limitations since it ignores the influence of additive noise N , in which a dynamic background with complex variations and mixed Gaussian-Poisson noise in XCA images largely affects the foreground/background decomposition. Although the regularization parameters λ_1^k and λ_2^k can be changed to adjust the number of foreground components in the final result to reduce noise to some extent, it is very difficult for the unfolded and traditional RPCA methods to eliminate noisy background disturbances while simultaneously preserving the entire vessel structure, especially the distal vessels.

Considering that SR network can effectively extract structural features of target object and selectively enhance these features without introducing much noise for image segmentation [67], we assume that this SR network's ability to select features can be explored for our work. In addition, inspired by the fact that the mixed Gaussian-Poisson noise locally corrupts the detailed information of vessel branches and can be successfully removed in a patch-wise Gaussian denoising [2], [3], we propose a patch-wise SR module with sparse feature selection in RPCA-UNet to extract vessels and eliminate artefacts simultaneously.

The patch-wise SR module is embedded into each iteration/layer of RPCA-UNet to gradually refine vessel extraction and simultaneously remove background artefacts. Inspired by [17], we introduce a feature pooling layer at the beginning of each iteration to downsample the input signal first, which can reduce the influence of redundant information. In RPCA-UNet, the motion artefacts and complex Gaussian-Poisson noise, which are described as N in equation (2), can then be locally discarded to a large extent such that the input of the unfolded RPCA algorithm can be approximately regarded as consisting of only sparse components and low-rank components.

At the end of each iteration, the residual module and CLSTM network are introduced to iteratively learn the high-level spatiotemporal semantic information of sparsely distributed vessels and refine the output of the deep unfolded RPCA layer. Specifically, the residual module is first applied to extract deep features through multiple convolutional layers. After this feature extraction, the extracted features along with the original features are transferred to the next step via the residual operation. Then, the CLSTM network is applied to combine high-level spatiotemporal semantic details in the whole XCA sequence and selectively adjust the weights of these features.

Different from traditional LSTM network, which uses full connections, CLSTM network replaces the multiplication operations with convolutional layers. This modification enables the CLSTM network to propagate spatiotemporal features in the training process of deep network [19]. The key CLSTM network formulation is:

$$\begin{aligned} i_t &= \sigma(W_{xi} * X_t + W_{hi} * h_{t-1} + W_{ci} \circ c_{t-1} + b_i) \\ f_t &= \sigma(W_{xf} * X_t + W_{hf} * h_{t-1} + W_{cf} \circ c_{t-1} + b_f) \end{aligned}$$

$$\begin{aligned}
c_t &= f_t \circ c_{t-1} + i_t \circ \tanh(W_{xc} * X_t + W_{hc} * h_{t-1} + b_c) \\
o_t &= \sigma(W_{xo} * X_t + W_{ho} * h_{t-1} + W_{co} \circ c_t + b_o) \\
h_t &= o_t \circ \tanh(c_t)
\end{aligned} \tag{10}$$

where $*$ denotes the convolutional operator and $\circ$ denotes the Hadamard product. The memory cell c_t can be used to store the spatiotemporal information of previous frames. The information in the memory cell can be written, cleared and propagated by controlling gates i_t , f_t , and o_t , respectively. The gate o_t that is regarded as a selector can select features from the complementary spatiotemporal information of previous frames to enhance deep features. Then, h_t is the final output, which is determined by the current input and the spatiotemporal information in the memory cell to aid sparse feature selection for better prediction.

Usually, the CLSTM network can be inserted at different stages of the SR module, such as at the beginning of the module, at the end of the module or during the feature extraction of the residual module. We choose to embed the CLSTM network in the feature extraction such that the weights of extracted features can be selectively adjusted through the spatiotemporal information in memory cells. This embedded feature selection via CLSTM network is assumed to be the core mechanism that enables the sparse feature selection for patch-wise SR vessel extraction and non-vessel artefact removal in RPCA-UNet.

Finally, the output of the CLSTM network is transferred to a sub-pixel convolution layer, which is often used in the image SR task to upscale the output for enhancing detailed information, such as distal and branch vessels in an XCA sequence.

D. Automatic Vessel Labelling

RPCA-UNet aims to extract both the geometrical features and image grey values of XCA vessels, which are unrealistic to be labelled manually for the deep learning of RPCA-UNet. Therefore, RPCA-UNet implemented weakly supervised learning through an automatic vessel labelling with tensor-completion-based VRBC method [9], which is the only method available to accurately and automatically recover vessel intensity information with rarely introducing background component. Specifically, VRBC first extracts vessel structures from complex and noisy backgrounds by RPCA-based vessel extraction. An accurate binary mask of vessel is then finely generated via Radon-like feature filtering with spatially adaptive thresholding. Subsequently, vessel-masked background regions are recovered to complete background layers by implementing tensor completion with the spatiotemporal consistency of whole background regions. Finally, the layers containing vessels' greyscale values can be accurately extracted by subtracting the completed background layers from the overall XCA images.

IV. EXPERIMENTAL RESULTS

A. Experimental Materials

Our experiments collected 43 sequences of real clinical XCA images from Renji Hospital of Shanghai Jiao Tong University. The length of each sequence ranges from 30 to 140 frames. Images from these XCA sequences were manually

annotated by three experts to obtain the vessel mask ground truth for evaluation. To eliminate differences in size, these frames were resized to 512×512 resolution with 8 bits per pixel. It is worth noting that these sequences are heterogeneous since they are collected from different machines, including a medical angiography X-ray system from Philips and the 800 mAH digital silhouette angiography X-ray machine from Siemens. Therefore, the noise distribution and the pixel grey level range of each sequence are very different.

B. Experiment Settings and RPCA-UNet Training

RPCA-UNet consists of 4 layers. The first two layers use convolutional kernels of size = 5 with stride = 1, padding = 2 and a bias, and the other two layers use convolutional kernels of size = 3 with stride = 1, padding = 1 and a bias. We choose the ADAM optimizer with a learning rate of 0.0001. In the feature selection module, the average pooling layer with pooling window = 2 and stride = 2 is selected. The upscaling rate of the SR module is set to 2.

RPCA-UNet is trained using back-propagation in a weakly-supervised manner. Training pairs of vessel/background labelling are generated by the VRBC method [9] and the training images are divided into $64 \times 64 \times 20$ patches with a certain overlap (50% between two neighboring patches). A total of 15 sequences containing 900 samples are used in the experiment and the total amount of patches used in dataset is 20000. Then, the dataset is randomly divided into training, validation, and test datasets at a ratio of approximately 0.6:0.2:0.2, respectively. The patches in the output are spliced with their grey values being the average of overlapping patches. The loss function is chosen as the sum of the mean square errors between the predicted vessel/background values of the network and the corresponding vessel/background labels.

C. Comparison Methods

We used the VRBC [9] and several state-of-the-art RPCA-based methods for comparison, which include ALF-RPCA [68], MoG-RPCA [69], our previous MCR-RPCA [6] and CORONA [14]. Moreover, to evaluate the performance of our network on the vessel segmentation task, we compared the vessel segmentation results with those of several other vessel segmentation algorithms, including Coye [70], Frangi *et al.* [27] results and those of the deep-learning-based SVS-net [48] and CS²-Net [71]. The parameters of these segmentation algorithms were tuned to achieve the best performance.

D. Visual Evaluation on Experimental Results

Vessel extraction results are shown in Fig. 2, in which three RPCA-based methods can extract major vessels relatively well but obvious vessel-like residuals with considerable noises still remain (see Fig. 2(b)-(d)). Moreover, distal vessels are hardly extracted because they are completely submerged in the noisy background. Although CORONA performs better in distal vessel extraction (see Fig. 2(e)), the noisy background still has

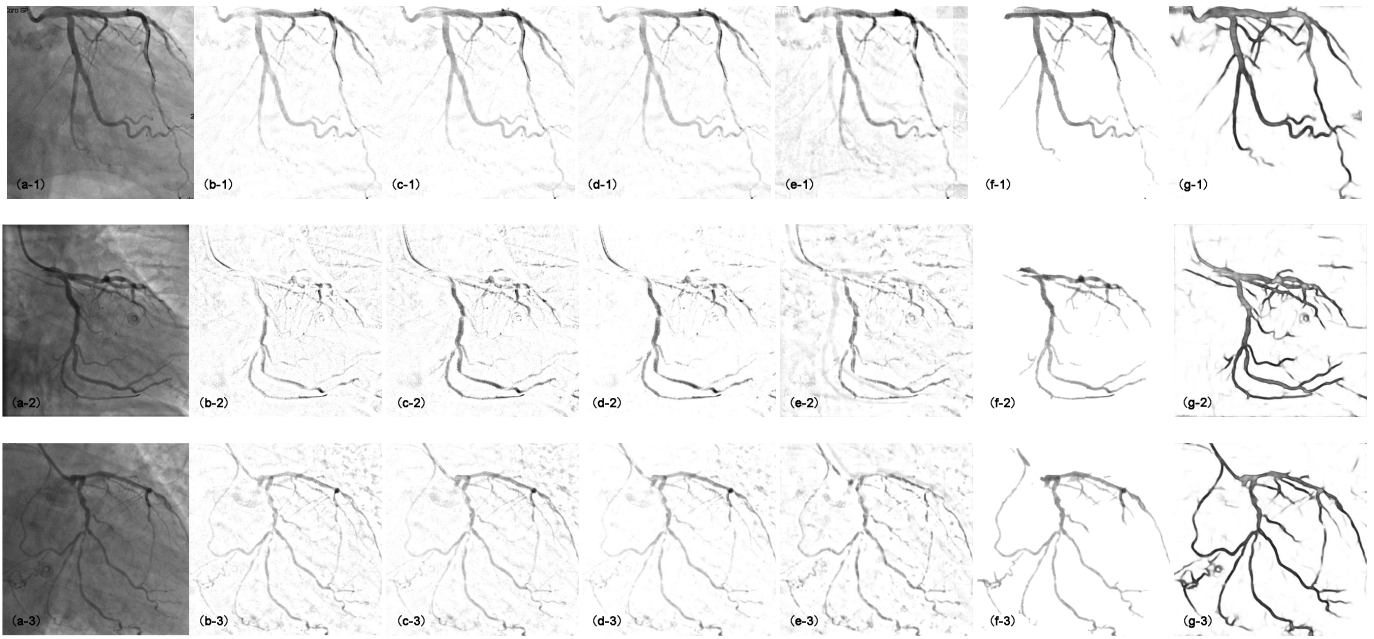


Fig. 2. XCA vessel extraction results. (a) Original XCA image; (b) ALF-RPCA; (c) MoG-RPCA; (d) MCR-RPCA [6]; (e) CORONA [14]; (f) VRBC [9]; (g) RPCA-UNet.

a large influence on the extraction results, as for the traditional RPCA-based methods. The VRBC framework extracts much better grey value vessels than the above methods, with most of the background artefacts being removed (see Fig. 2(f)). However, the performance of VRBC in the extraction of vessel branches and distal vessels is still not satisfactory. Compared to these methods, RPCA-UNet greatly improves the vessel extraction performances since the extracted vessel tree structure is clearer and more complete, especially for the vessel branches and distal vessels (see Fig. 2(g)).

It is worth noting that the visual contrast of vessel extraction (see Fig. 2) of RPCA-UNet is obviously enhanced compared with that of the VRBC method [9]. Specifically, in the recovered profiles of vessel intensity, the grey level of distal vessels is lower and that of large vessels is increasingly higher. These results are entirely consistent with the X-ray attenuation coefficients of various structures imaged in the XCA images. During low-dose XCA imaging, the amount of contrast agent in large vessels is greater than that of distal vessels, which makes the grey level of large vessels higher than that of distal vessels. Therefore, the recovered grey levels of vessels achieved by RPCA-UNet is in high fidelity relative to the real distribution of contrast agent in XCA vessels, which is helpful for quantitatively analysing the structure-functional characterization of cardiac perfusion.

In vessel segmentation evaluation, we use different colours to label the pixels of segmentation results by comparing with ground truth vessel mask (see Fig. 3(b)), in which green pixels represent the true positive pixels that are correctly classified as vessels, blue pixels represent false negative pixels that are vessel pixels but wrongly classified as backgrounds, red pixels are false positive pixels that are wrongly classified as vessels but practically belonging to the backgrounds. The segmentation results in Fig. 3 show that Coye's and Frangi's

methods detect either too few vessels or too much noises (see Fig. 3(c)-(d)). These traditional methods have poor performance in the foreground and background areas with similar grey values. SVS-net can detect most of the vessels and effectively suppress the background noise (see Fig. 3(e)). However, in some distal vessel areas, the detection result appears to have missing or discontinuous segments. CS²-Net has a great segmentation performance on distal vessels while in some cases it may introduce vessel-like noisy artefacts (see Fig. 3(f)). Moreover, its comprehensive performance seems to be unstable since discontinuities appear in some major blood vessels. RPCA-UNet performs best in the segmentation experiments since almost all the major vessels and most of the distal vessels can be detected (see Fig. 3(g)). Although some background residue still exists in the detection results, it does not have a significant impact on the visual effect.

E. Quantitative Evaluation of Vessel Extraction

The vessel visibility can be quantitatively evaluated by using the contrast-to-noise ratio (CNR) [14] of the vessel layer images. The CNR measures the contrast between the vessels and backgrounds, where a larger CNR value means a better vessel visibility. The CNR can be calculated by:

$$CNR = \frac{|\mu_V - \mu_B|}{\sqrt{\sigma_B^2 + \sigma_V^2}} \quad (11)$$

where μ_V and μ_B are the pixel intensity means in the vessel and background regions, respectively, and σ_V and σ_B are the standard deviation of the pixel intensity values in the vessel regions and background regions, respectively.

To further evaluate vessel visibility, we define global and local background regions to cover all the image regions except

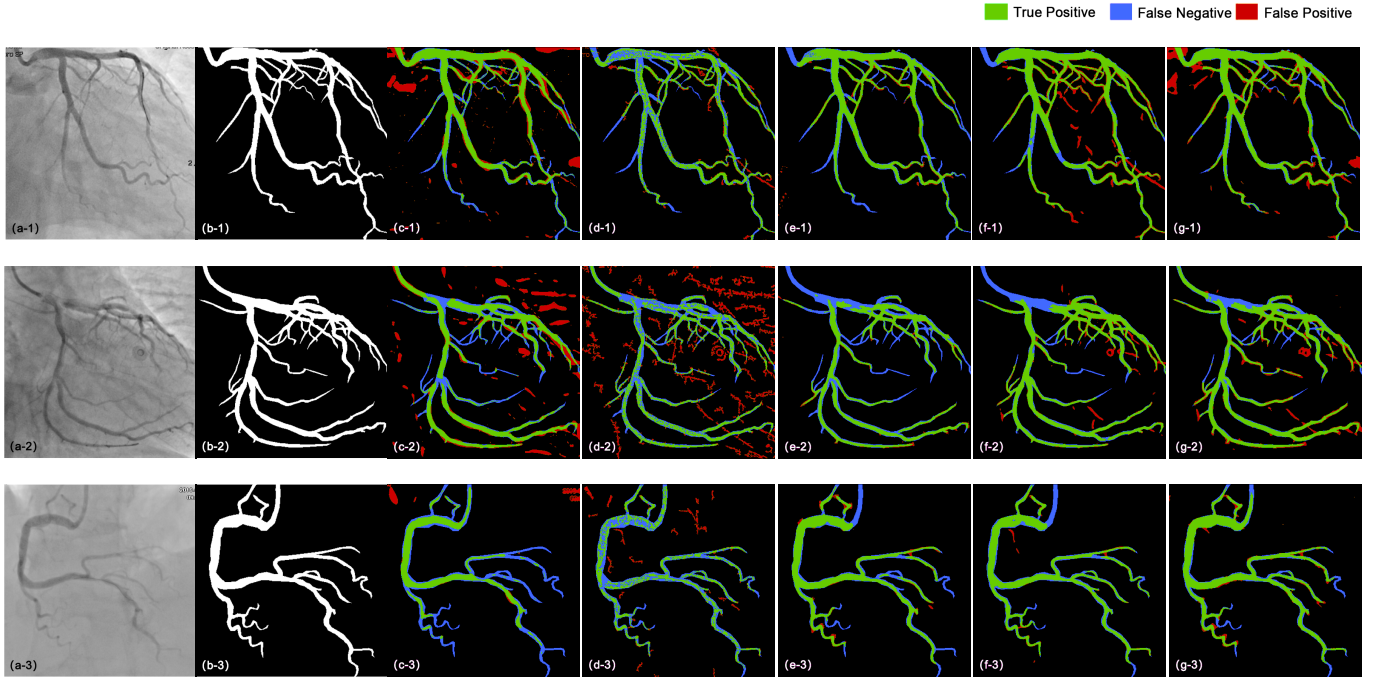


Fig. 3. XCA vessel segmentation results. Pixels labelled with green, blue, and red colours represent true positive pixels, false negative pixels, and false positive pixels, respectively. (a) Original XCA image; (b) Ground-truth vessel mask; (c) Frangi's; (d) Coye's; (e) SVS-net; (f) CS²-Net; (g) RPCA-UNet.

TABLE I

PERFORMANCE OF DIFFERENT VESSEL EXTRACTION METHODS IN TERMS OF CNR VALUES (MEAN $\pm$ STANDARD DEVIATION)

Method	Global CNR	Local CNR	Time(s)
MCR-RPCA	1.01 $\pm$ 0.19	1.00 $\pm$ 0.20	23.19
MoG-RPCA	1.06 $\pm$ 0.22	1.06 $\pm$ 0.21	1.03
ALF-RPCA	0.93 $\pm$ 0.25	0.95 $\pm$ 0.24	0.20
CORONA	0.96 $\pm$ 0.18	1.01 $\pm$ 0.183	0.61
VRBC	1.04 $\pm$ 0.14	1.02 $\pm$ 0.14	24.68
RPCA-UNet	1.78 $\pm$ 0.25	1.65 $\pm$ 0.20	0.92

the vessel regions and the 7-pixel-wide neighbourhood regions surrounding the vessel regions, respectively. The CNR calculation results are shown in TABLE I and Fig. 4. The results show that RPCA-UNet achieves the highest global and local CNRs, which indicates that RPCA-UNet greatly improves the vessel visibility both globally and locally because it achieves excellent vessel extraction and noise suppression.

To evaluate the time efficiency of RPCA-UNet, we calculate the average running time per image of the above vessel extraction methods and the results are shown in TABLE I. RPCA-UNet has relatively fast speed even though it divides each image into patches with 50% overlapping to eliminate the influence of mixed noise, which will increase the amount of calculation by about 4 times. Such sacrifice in speed for better extraction results is assumed to be worthwhile in meeting the clinical demand. Moreover, the parameter sizes of RPCA-UNet and SVS-net is 0.76M and 10.06M, respectively. This indicates that the parameter size of RPCA-UNet is quite small and its storage efficiency is remarkable.

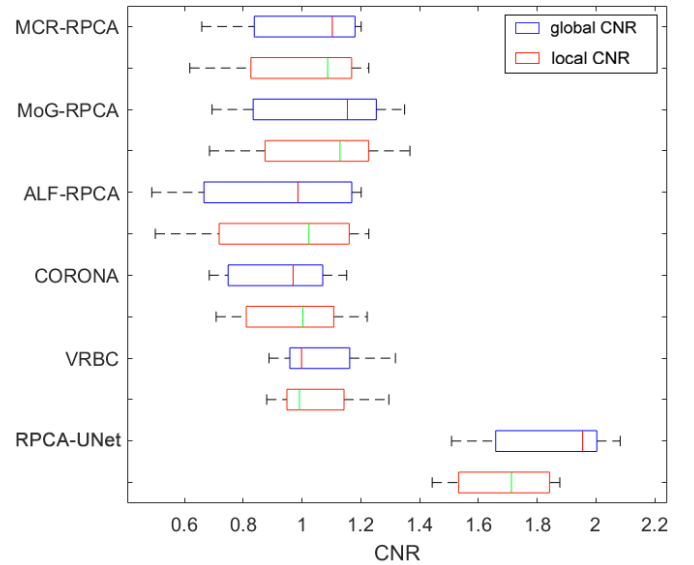


Fig. 4. Performance of vessel extraction methods in terms of global and local CNRs for twelve real XCA images.

F. Quantitative Evaluation of Vessel Segmentation

To evaluate the performances of the segmentation results of the proposed method, the detection rate (DR), precision (P), and F-measure (F) are employed. We also compare RPCA-UNet with other state-of-the-art methods. The above-mentioned metrics can be calculated as follows:

$$DR = \frac{TP}{TP + FN}, P = \frac{TP}{TP + FP}, F = \frac{2 \times DR \times P}{DR + P} \quad (12)$$

TABLE II
MEANS AND STANDARD DEVIATIONS OF THE DR, P,
AND F METRICS ON TWELVE XCA IMAGES

Method	Detection Rate	Precision	F-measure
Coye's	0.592 ± 0.087	0.810 ± 0.134	0.675 ± 0.065
Frangi's	0.577 ± 0.120	0.686 ± 0.181	0.617 ± 0.126
SVS-net	0.635 ± 0.090	0.948 ± 0.030	0.757 ± 0.063
CS ² -Net	0.747 ± 0.119	0.842 ± 0.089	0.778 ± 0.063
RPCA-UNet	0.810 ± 0.057	0.774 ± 0.161	0.783 ± 0.086

where TP is the total number of true positive pixels, FP indicates the total number of false positive pixels, and TN and FN indicate the total numbers of true negative and false negative pixels that are correctly classified as background pixels and wrongly predicted as background pixels in the segmentation output, respectively. The DR represents the proportion between the correctly classified vessel pixels and the total vessel pixels in the ground truth, P represents the ratio of the true positive pixels among all positive pixels, and F comprehensively considers both the P and DR metrics and indicates the overall segmentation performance. All these metrics range from 0 to 1, where higher values mean a better segmentation performance.

In the quantitative evaluation of the segmentation experiments, 12 images selected from different sequences are manually annotated as the ground truth. The DR, P, and F-measure of these 12 images are measured and displayed in TABLE II. RPCA-UNet generally obtains the highest DR and F scores. RPCA-UNet achieves relatively lower P value than some other methods do, which represents the proportion of positive examples that are actually positive. This is because that RPCA-UNet tends to extract complete vessel information as much as possible and may inevitably cause an increase in false positives, which leads to a relative low P. However, the improvement of vessel detection makes RPCA-UNet achieve better performance in terms of both DR and F-measure. RPCA-UNet performs best in term of F-measure that represents the comprehensive performance of the P and the DR. Moreover, the XCA images typically selected in our experiments contain many distal vessel branches with a very low contrast to the background components such that it is quite difficult to detect all distal vessels correctly. Therefore, the DR and F results of state-of-the-art methods and RPCA-UNet are generally low. Under these experimental settings, RPCA-UNet still obtained relatively high results due to the larger improvement in its capability to detect vessel branches and distal vessels. Therefore, we believe RPCA-UNet performs better than other methods even if its P is not the highest.

G. Ablation Study

Several ablation experiments are tested to validate the effectiveness of RPCA-UNet architecture. We compare the vessel extraction results of RPCA-UNet using different iterations/layers with different results from CORONA [14], SR network, RPCA-UNet with SR module but without CLSTM network. The layers of RPCA-UNet can gradually extract the moving contrast agents and prune complex vessel-like

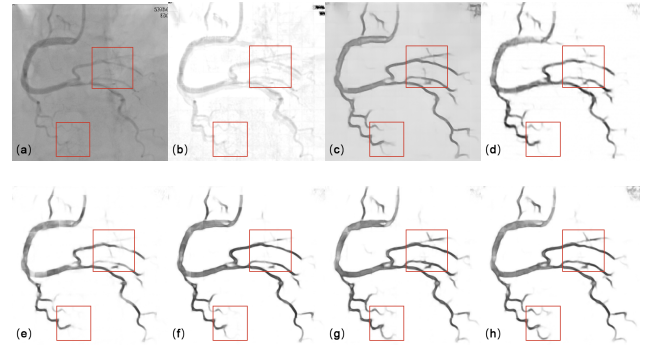


Fig. 5. The results of ablation study. (a) original XCA image; (b) CORONA [14]; (c) SR network; (d) RPCA-UNet with SR but without CLSTM network; (e)-(h) RPCA-UNet with 1-4 layers.

TABLE III
MEAN SQUARED ERRORS OF RPCA-UNET WITH DIFFERENT LAYERS

Layer Number	1	2	3	4
MSE	8.29e-03	6.68e-03	6.56e-03	6.26e-03

artefacts. As in [14], for each layer number, we construct RPCA-UNet with that number of layers. These networks are trained for 50 epochs on the same training dataset.

Fig. 5 shows the comparison results of ablation study. CORONA [14] can extract vessel grey value information while the result (see Fig. 5(b)) is severely disturbed by noisy artefacts. The output of SR network shows clear vessel contour (see Fig. 5(c)) while the grey value information is almost lost since SR network is mainly designed for extracting structure information. The RPCA unrolling network embedded with SR module but without CLSTM network (see Fig. 5(d)) improves the vessel extraction result significantly. However, such combined network has an imperfect performance in extracting distal vessels with low contrast and strong background noises.

The second row of Fig. 5 shows the comparison between different versions of RPCA-UNet with different layers (see Fig. 5(e)-(h)). The results show that integrating our CLSTM-network-based SR module into RPCA-UNet can extract vessel information iteratively and gradually prune complex vessel-like artefacts and noisy backgrounds. We also compare the mean squared error (MSE) of RPCA-UNet and the results in TABLE III show that the MSE decreases as the number of layers increases.

H. Coarse Versus Fine Labels for Weakly Supervised Learning

Due to our weakly supervised method replacing manual annotation with automatic vessel labelling by VRBC method [9], it would be interesting to see whether this replacement with automatic vessel labelling had influenced the results. Although obtaining a large number of manually annotated vessel masks with grey values is almost impossible, we have manually annotated binary labels for our proposed SVS-net [48], whose output is then used as a binary-mask in VRBC method for automatic vessel labelling (refer to the results in Fig. 2(f)). Therefore, we can provide coarse and fine

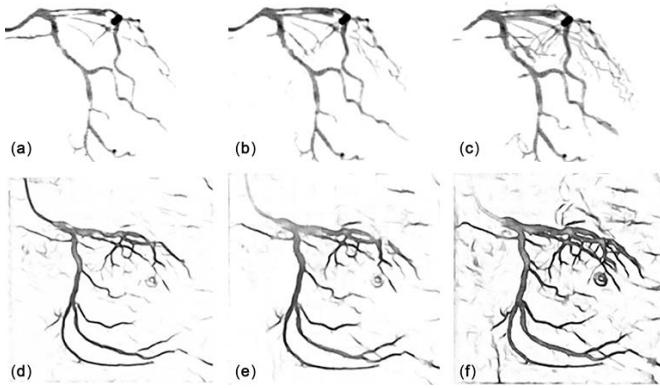


Fig. 6. The effect of coarse versus fine vessel label on the result of weakly supervised learning. The first row are the coarse and fine grey value labels automatically generated by the VRBC combined with different binary vessel mask segmentations, i.e. from left to right being original segmentation method in the VRBC, SVS-net with training data generated by the original segmentation method, SVS-net with training data by manual annotation; the second row of results are test cases of the corresponding networks trained with different grey value labels.

vessel labels to evaluate the impact of labelling quality on the vessel extraction results.

To assess the impact of coarsely labelled versus finely labelled data on weakly supervised vessel extraction, we have automatically generated three types of grey value labels using VRBC method with their corresponding binary masks being first segmented in different ways: original segmentation method [8] adopted in the VRBC method, SVS-net with training data being generated by the original segmentation method, SVS-net with training data being manually annotated. With these different binary segmentation masks, the grey value labels generated by the VRBC method are displayed in the first row of Fig. 6. We assume that the quality of fine grey value labels generated by the VRBC plus SVS-net with manual annotations is close to that of manual annotations. The vessel extraction results via RPCA-UNet are shown in the second row of Fig. 6. The two networks trained by the first two types of grey value labels similarly achieve great performances while the network trained by the third type of grey value labels, where the labelling quality is the highest and is close to manual annotations, introduces more background impurities in some test cases. We believe the reason for this phenomenon is that too fine labels will cause overfitting in the trained neural network that may have poor generalization ability, thereby erroneously identifying some background components as vessels in some noisy XCA images.

V. CONCLUSION AND DISCUSSION

To efficiently remove background artefacts and mixed Gaussian-Poisson noises for XCA vessel extraction, we propose a detail-preserving RPCA-UNet with a patch-wise spatiotemporal SR via sparse feature selection, which can not only achieve uninformative feature pruning and Gaussian-Poisson denoising but also selectively enhance vessel features from the backgrounds. The experimental results show superior perfor-

mance in both vessel extraction and vessel segmentation in an accurate and efficient way.

To the best of our knowledge, RPCA-UNet is the first neural network to implement an automatic weakly supervised vessel recovery from dynamic and complex backgrounds in XCA. Specifically, the heterogeneous grey value vessel layers automatically produced by the VRBC method [9] are used as the training data. Such grey value labels contain the main branches of vessels, enabling RPCA-UNet to learn the greyscale and motion information of the whole vessel network. After that, RPCA-UNet can combine the information provided by the grey value labels and the characteristics of RPCA-UNet to achieve a great effect of vessel extraction. Moreover, we compared different training strategies with fine grey value labels where almost all the distal branches are annotated, and coarse grey value labels where only the major vessels and relatively thick vessel branches are annotated. The comparison results show that RPCA-UNet trained by coarse labels perform better than that trained by fine labels. Specifically, the RPCA-UNet trained by fine labels introduces significantly more noises, which is assumed to result from the overfitting in the trained neural network. Therefore, the proposed weakly supervised learning can not only largely reduce the labour and time spent on labelling data, but also improve the generalization ability of RPCA-UNet.

To achieve a better detail-preserving vessel extraction, future research can explore more effective pooling layers [17] and interpretable [72] CLSTM network in the patch-wise spatiotemporal SR module for selecting sparse feature to improve the restoration of heterogeneous vessel profiles. For distal vessel detection, applying a self-attention mechanism to improve the inter-class discrimination and intra-class aggregation abilities [71] can help unrolling network in accurately classifying the vessel pixels in the easily confused regions between the distal branches and the background. Furthermore, there is still room for improvement to develop better evaluation metrics for vessel extraction performance. Most current studies on vessel extraction employ CNR metric for visual effect evaluation. CNR solely quantifies the visibility of extracted vessels via the contrast between foreground vessels and backgrounds while ignoring the heterogeneity of extracted blood vessels. In addition to CNR, we can apply patch-based principal component analysis for intensity-dependent variation analysis [3] in vessel regions to evaluate the restoration of heterogeneous vessel profiles. Due to contrast agent particle moving in a blood stream causing transient changes of dose rate near its flow path, an evaluation metric for dynamic heterogeneity [73] can also be explored to consider the variation of blood flow capacity and calculate the blood flow geometry in terms of distribution in blood vessel length and connected structure of the vessel paths in sequential XCA imaging.

ACKNOWLEDGMENT

The authors thank all the cited authors for providing the source codes used in this work and an anonymous reviewers for their valuable comments on the manuscript.

REFERENCES

- [1] U. Sechtem, D. Brown, S. Godo, G. A. Lanza, H. Shimokawa, and N. Sidik, "Coronary microvascular dysfunction in stable ischaemic heart disease (non-obstructive coronary artery disease and obstructive coronary artery disease)," *Cardiovascular Res.*, vol. 116, no. 4, pp. 771–786, Mar. 2020.
- [2] P. Irrera, I. Bloch, and M. Delplanque, "A flexible patch based approach for combined denoising and contrast enhancement of digital X-ray images," *Med. Image Anal.*, vol. 28, pp. 33–45, Feb. 2016.
- [3] W. Zhao, Q. Liu, Y. Lv, and B. Qin, "Texture variation adaptive image denoising with nonlocal PCA," *IEEE Trans. Image Process.*, vol. 28, no. 11, pp. 5537–5551, Nov. 2019.
- [4] S. Moccia, E. De Momi, S. E. Hadji, and L. S. Mattos, "Blood vessel segmentation algorithms—Review of methods, datasets and evaluation metrics," *Comput. Methods Programs Biomed.*, vol. 158, pp. 71–91, May 2018.
- [5] D. Jia and X. Zhuang, "Learning-based algorithms for vessel tracking: A review," *Computerized Med. Imag. Graph.*, vol. 89, Apr. 2021, Art. no. 101840.
- [6] M. Jin, R. Li, J. Jiang, and B. Qin, "Extracting contrast-filled vessels in X-ray angiography by graduated RPCA with motion coherency constraint," *Pattern Recognit.*, vol. 63, pp. 653–666, Mar. 2017.
- [7] H. Ma, A. Hoogendoorn, E. Regar, W. J. Niessen, and T. van Walsum, "Automatic online layer separation for vessel enhancement in X-ray angiograms for percutaneous coronary interventions," *Med. Image Anal.*, vol. 39, pp. 145–161, Jul. 2017.
- [8] M. Jin, D. Hao, S. Ding, and B. Qin, "Low-rank and sparse decomposition with spatially adaptive filtering for sequential segmentation of 2D+t vessels," *Phys. Med. Biol.*, vol. 63, no. 17, 2018, Art. no. 17LT01.
- [9] B. Qin *et al.*, "Accurate vessel extraction via tensor completion of background layer in X-ray coronary angiograms," *Pattern Recognit.*, vol. 87, pp. 38–54, Mar. 2019.
- [10] H. Fang *et al.*, "Topology optimization using multiple-possibility fusion for vasculature extraction," *IEEE Trans. Circuits Syst. Video Technol.*, vol. 30, no. 2, pp. 442–456, Feb. 2020.
- [11] S. Song *et al.*, "Spatio-temporal constrained online layer separation for vascular enhancement in X-ray angiographic image sequence," *IEEE Trans. Circuits Syst. Video Technol.*, vol. 30, no. 10, pp. 3558–3570, Oct. 2020.
- [12] S. Xia *et al.*, "Vessel segmentation of X-ray coronary angiographic image sequence," *IEEE Trans. Biomed. Eng.*, vol. 67, no. 5, pp. 1338–1348, May 2020.
- [13] T. Bouwmans, S. Javed, H. Zhang, Z. Lin, and R. Otazo, "On the applications of robust PCA in image and video processing," *Proc. IEEE*, vol. 106, no. 8, pp. 1427–1457, Aug. 2018.
- [14] O. Solomon *et al.*, "Deep unfolded robust PCA with application to clutter suppression in ultrasound," *IEEE Trans. Med. Imag.*, vol. 39, no. 4, pp. 1051–1063, Apr. 2020.
- [15] V. Monga, Y. Li, and Y. C. Eldar, "Algorithm unrolling: Interpretable, efficient deep learning for signal and image processing," *IEEE Signal Process. Mag.*, vol. 38, no. 2, pp. 18–44, Mar. 2021.
- [16] K. Gregor and Y. LeCun, "Learning fast approximations of sparse coding," in *Proc. 27th Int. Conf. Mach. Learn.*, 2010, pp. 399–406.
- [17] P. Singh, P. Raj, and V. P. Namboodiri, "EDS pooling layer," *Image Vis. Comput.*, vol. 98, Jun. 2020, Art. no. 103923.
- [18] Z. Wang, J. Chen, and S. C. H. Hoi, "Deep learning for image super-resolution: A survey," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 43, no. 10, pp. 3365–3387, Mar. 2021.
- [19] X. Shi, Z. Chen, H. Wang, D.-Y. Yeung, W.-K. Wong, and W.-C. Woo, "Convolutional LSTM network: A machine learning approach for precipitation nowcasting," 2015, *arXiv:1506.04214*.
- [20] M. R. K. Mookiah *et al.*, "A review of machine learning methods for retinal blood vessel segmentation and artery/vein classification," *Med. Image Anal.*, vol. 68, Feb. 2021, Art. no. 101905.
- [21] B. Qin, M. Jin, and S. Ding, "Extracting heterogeneous vessels in X-ray coronary angiography via machine learning," in *Cardiovascular and Coronary Artery Imaging*. Amsterdam, The Netherlands: Elsevier, 2022, pp. 89–127.
- [22] F. Cervantes-Sanchez, I. Cruz-Aceves, A. Hernandez-Aguirre, S. Solorio-Meza, T. Cordova-Fraga, and J. G. Aviña-Cervantes, "Coronary artery segmentation in X-ray angiograms using Gabor filters and differential evolution," *Appl. Radiat. Isot.*, vol. 138, pp. 18–24, Aug. 2018.
- [23] H. R. Fazlali *et al.*, "Vessel segmentation and catheter detection in X-ray angiograms using superpixels," *Med. Biol. Eng. Comput.*, vol. 56, no. 9, pp. 1515–1530, Sep. 2018.
- [24] A. Kerkeni, A. Benabdallah, A. Manzanera, and M. H. Bedoui, "A coronary artery segmentation method based on multiscale analysis and region growing," *Computerized Med. Imag. Graph.*, vol. 48, pp. 49–61, Mar. 2016.
- [25] T. Wan, X. Shang, W. Yang, J. Chen, D. Li, and Z. Qin, "Automated coronary artery tree segmentation in X-ray angiography using improved Hessian based enhancement and statistical region merging," *Comput. Methods Programs Biomed.*, vol. 157, pp. 179–190, Apr. 2018.
- [26] Y. Qian, Z. Wang, L. Chen, and Z. Huang, "Vascular enhancement with structure preservation from noisy X-ray angiogram images by employing non-local Hessian-based filter," *Optik*, vol. 232, Apr. 2021, Art. no. 166523.
- [27] A. F. Frangi, W. J. Niessen, K. L. Vincken, and M. A. Viergever, "Multiscale vessel enhancement filtering," in *Proc. Int. Conf. Med. Image Comput. Comput.-Assist. Intervent.* Berlin, Germany: Springer, 1998, pp. 130–137.
- [28] I. Cruz-Aceves, F. Oloumi, R. M. Rangayyan, J. G. Aviña-Cervantes, and A. Hernandez-Aguirre, "Automatic segmentation of coronary arteries using Gabor filters and thresholding based on multiobjective optimization," *Biomed. Signal Process. Control*, vol. 25, pp. 76–85, Mar. 2016.
- [29] Y. Zhao *et al.*, "Automatic 2-D/3-D vessel enhancement in multiple modality images using a weighted symmetry filter," *IEEE Trans. Med. Imag.*, vol. 37, no. 2, pp. 438–450, Feb. 2018.
- [30] K. Mei, B. Hu, B. Fei, and B. Qin, "Phase asymmetry ultrasound despeckling with fractional anisotropic diffusion and total variation," *IEEE Trans. Image Process.*, vol. 29, pp. 2845–2859, 2020.
- [31] R. Reisenhofer and E. J. King, "Edge, ridge, and blob detection with symmetric molecules," *SIAM J. Imag. Sci.*, vol. 12, no. 4, pp. 1585–1626, Jan. 2019.
- [32] M. Kass, A. Witkin, and D. Terzopoulos, "Snakes: Active contour models," *Int. J. Comput. Vis.*, vol. 1, no. 4, pp. 321–331, 1988.
- [33] S. Osher and J. A. Sethian, "Fronts propagating with curvature-dependent speed: Algorithms based on Hamilton-Jacobi formulations," *J. Comput. Phys.*, vol. 79, no. 1, pp. 12–49, 1988.
- [34] L. Zou *et al.*, "A survey on regional level set image segmentation models based on the energy functional similarity measure," *Neurocomputing*, vol. 452, pp. 606–622, Sep. 2021.
- [35] T. Lv *et al.*, "Vessel segmentation using centerline constrained level set method," *Multimedia Tools Appl.*, vol. 78, no. 12, pp. 17051–17075, Jun. 2019.
- [36] K. Sun, Z. Chen, and S. Jiang, "Local morphology fitting active contour for automatic vascular segmentation," *IEEE Trans. Biomed. Eng.*, vol. 59, no. 2, pp. 464–473, Feb. 2012.
- [37] S. Ge, Z. Shi, G. Peng, and Z. Zhu, "Two-steps coronary artery segmentation algorithm based on improved level set model in combination with weighted shape-prior constraints," *J. Med. Syst.*, vol. 43, no. 7, pp. 1–10, 2019.
- [38] H. Fang *et al.*, "Greedy soft matching for vascular tracking of coronary angiographic image sequences," *IEEE Trans. Circuits Syst. Video Technol.*, vol. 30, no. 5, pp. 1466–1480, May 2020.
- [39] O. Wink, W. J. Niessen, and M. A. Viergever, "Multiscale vessel tracking," *IEEE Trans. Med. Imag.*, vol. 23, no. 1, pp. 130–133, Jan. 2004.
- [40] Y. Chen *et al.*, "Curve-like structure extraction using minimal path propagation with backtracking," *IEEE Trans. Image Process.*, vol. 25, no. 2, pp. 988–1003, Feb. 2016.
- [41] G. Yang *et al.*, "Vessel structure extraction using constrained minimal path propagation," *Artif. Intell. Med.*, vol. 105, May 2020, Art. no. 101846.
- [42] X. Liu, F. Hou, H. Qin, and A. Hao, "Robust optimization-based coronary artery labeling from X-ray angiograms," *IEEE J. Biomed. Health Inform.*, vol. 20, no. 6, pp. 1608–1620, Nov. 2015.
- [43] E. Nasr-Esfahani *et al.*, "Segmentation of vessels in angiograms using convolutional neural networks," *Biomed. Signal Process. Control*, vol. 40, pp. 240–251, Feb. 2018.
- [44] S. Y. Shin, S. Lee, I. D. Yun, and K. M. Lee, "Deep vessel segmentation by learning graphical connectivity," *Med. Image Anal.*, vol. 58, Dec. 2019, Art. no. 101556.
- [45] X. Zhu, Z. Cheng, S. Wang, X. Chen, and G. Lu, "Coronary angiography image segmentation based on PSPNet," *Comput. Methods Programs Biomed.*, vol. 200, Mar. 2021, Art. no. 105897.
- [46] O. Ronneberger, P. Fischer, and T. Brox, "U-Net: Convolutional networks for biomedical image segmentation," in *Proc. Int. Conf. Med. Image Comput. Comput.-Assist. Intervent.* Cham, Switzerland: Springer, 2015, pp. 234–241.

- [47] J. Fan *et al.*, "Multichannel fully convolutional network for coronary artery segmentation in X-ray angiograms," *IEEE Access*, vol. 6, pp. 44635–44643, 2018.
- [48] D. Hao *et al.*, "Sequential vessel segmentation via deep channel attention network," *Neural Netw.*, vol. 128, pp. 172–187, Aug. 2020.
- [49] T. Wan, J. Chen, Z. Zhang, D. Li, and Z. Qin, "Automatic vessel segmentation in X-ray angiogram using spatio-temporal fully-convolutional neural network," *Biomed. Signal Process. Control*, vol. 68, Jul. 2021, Art. no. 102646.
- [50] P. M. Samuel and T. Veeramalai, "VSSC Net: Vessel specific skip chain convolutional network for blood vessel segmentation," *Comput. Methods Programs Biomed.*, vol. 198, Jan. 2021, Art. no. 105769.
- [51] J. Xiang, Y. Dong, and Y. Yang, "FISTA-Net: Learning a fast iterative shrinkage thresholding network for inverse problems in imaging," *IEEE Trans. Med. Imag.*, vol. 40, no. 5, pp. 1329–1339, May 2021.
- [52] S. Chen, Y. C. Eldar, and L. Zhao, "Graph unrolling networks: Interpretable neural networks for graph signal denoising," *IEEE Trans. Signal Process.*, vol. 69, pp. 3699–3713, 2021.
- [53] J. Cai, J. Luo, S. Wang, and S. Yang, "Feature selection in machine learning: A new perspective," *Neurocomputing*, vol. 300, pp. 70–79, Jul. 2018.
- [54] R. Battiti, "Using mutual information for selecting features in supervised neural net learning," *IEEE Trans. Neural Netw.*, vol. 5, no. 4, pp. 537–550, Jul. 1994.
- [55] S. Yu, L. G. S. Giraldo, R. Jenssen, and J. C. Principe, "Multivariate extension of matrix-based Rényi's α -order entropy functional," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 42, no. 11, pp. 2960–2966, Nov. 2020.
- [56] R. Zhang, F. Nie, X. Li, and X. Wei, "Feature selection with multi-view data: A survey," *Inf. Fusion*, vol. 50, pp. 158–167, Oct. 2019.
- [57] O. Tarkhaneh, T. T. Nguyen, and S. Mazaheri, "A novel wrapper-based feature subset selection method using modified binary differential evolution algorithm," *Inf. Sci.*, vol. 565, pp. 278–305, Jul. 2021.
- [58] L. Jiang, G. Kong, and C. Li, "Wrapper framework for test-cost-sensitive feature selection," *IEEE Trans. Syst., Man, Cybern., Syst.*, vol. 51, no. 3, pp. 1747–1756, Mar. 2021.
- [59] A. Got, A. Moussaoui, and D. Zouache, "Hybrid filter-wrapper feature selection using whale optimization algorithm: A multi-objective approach," *Expert Syst. Appl.*, vol. 183, Nov. 2021, Art. no. 115312.
- [60] J. Gui, Z. Sun, S. Ji, D. Tao, and T. Tan, "Feature selection based on structured sparsity: A comprehensive study," *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 28, no. 7, pp. 1490–1507, Jul. 2017.
- [61] X. Li, Y. Wang, and R. Ruiz, "A survey on sparse learning models for feature selection," *IEEE Trans. Cybern.*, vol. 52, no. 3, pp. 1642–1660, Mar. 2022.
- [62] L. Zhao, Q. Hu, and W. Wang, "Heterogeneous feature selection with multi-modal deep neural networks and sparse group LASSO," *IEEE Trans. Multimedia*, vol. 17, no. 11, pp. 1936–1948, Nov. 2015.
- [63] F. Farokhmanesh and M. T. Sadeghi, "Deep neural networks regularization using a combination of sparsity inducing feature selection methods," *Neural Process. Lett.*, vol. 53, no. 1, pp. 701–720, Feb. 2021.
- [64] H. Dong, L. Zhang, D. Lu, and B. Zou, "Attention-based polarimetric feature selection convolutional network for PolSAR image classification," *IEEE Geosci. Remote Sens. Lett.*, vol. 19, pp. 1–5, 2022.
- [65] T. Hoefer, D. Alistarh, T. Ben-Nun, N. Dryden, and A. Peste, "Sparsity in deep learning: Pruning and growth for efficient inference and training in neural networks," *J. Mach. Learn. Res.*, vol. 22, pp. 1–124, Sep. 2021.
- [66] K. Liu, Y. Fu, L. Wu, X. Li, C. Aggarwal, and H. Xiong, "Automated feature selection: A reinforcement learning perspective," *IEEE Trans. Knowl. Data Eng.*, early access, Sep. 24, 2021, doi: 10.1109/TKDE.2021.3115477.
- [67] L. Wang, D. Li, Y. Zhu, L. Tian, and Y. Shan, "Dual super-resolution learning for semantic segmentation," in *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit. (CVPR)*, Jun. 2020, pp. 3774–3783.
- [68] S. Wang, Y. Wang, Y. Chen, P. Pan, Z. Sun, and G. He, "Robust PCA using matrix factorization for background/foreground separation," *IEEE Access*, vol. 6, pp. 18945–18953, 2018.
- [69] Q. Zhao, D. Meng, Z. Xu, W. Zuo, and L. Zhang, "Robust principal component analysis with complex noise," in *Proc. Int. Conf. Mach. Learn.*, 2014, pp. 55–63.
- [70] T. Coyle. (2022). Novel Retinal Vessel Segmentation Algorithm: Fundus Images. MATLAB Central File Exchange. Accessed: May 24, 2022. [Online]. Available: <https://www.mathworks.com/matlabcentral/fileexchange/50839-novelretinal-vessel-segmentation-algorithm-fundus-images>
- [71] L. Mou *et al.*, "CS²-Net: Deep learning segmentation of curvilinear structures in medical imaging," *Med. Image Anal.*, vol. 67, Jan. 2021, Art. no. 101874.
- [72] Y. Huang, J. Li, X. Gao, Y. Hu, and W. Lu, "Interpretable detail-fidelity attention network for single image super-resolution," *IEEE Trans. Image Process.*, vol. 30, pp. 2325–2339, 2021.
- [73] G. M. Shook and K. M. Mitchell, "A robust measure of heterogeneity for ranking earth models: The F PHI curve and dynamic Lorenz coefficient," in *Proc. SPE Annu. Tech. Conf. Exhib.*, 2009, pp. 1–13, Paper SPE124625.