



Abdominal multi-organ segmentation with organ-attention networks and statistical fusion

Yan Wang^{a,1}, Yuyin Zhou^{a,1}, Wei Shen^{b,a}, Seyoun Park^{c,*}, Elliot K. Fishman^c, Alan L. Yuille^{d,a}

^a Department of Computer Science, Johns Hopkins University, Baltimore, MD, USA

^b Key Laboratory of Specialty Fiber Optics and Optical Access Networks, Shanghai University, Shanghai, China

^c Department of Radiology and Radiological Science, Johns Hopkins University, Baltimore, MD, USA

^d Department of Cognitive Science, Johns Hopkins University, Baltimore, MD, USA

ARTICLE INFO

Article history:

Received 3 May 2018

Revised 5 April 2019

Accepted 17 April 2019

Available online 18 April 2019

Keywords:

Multi-organ segmentation

Organ-attention network

Reverse connection

Statistical label fusion

Abdominal CT

ABSTRACT

Accurate and robust segmentation of abdominal organs on CT is essential for many clinical applications such as computer-aided diagnosis and computer-aided surgery. But this task is challenging due to the weak boundaries of organs, the complexity of the background, and the variable sizes of different organs. To address these challenges, we introduce a novel framework for multi-organ segmentation of abdominal regions by using organ-attention networks with reverse connections (OAN-RCs) which are applied to 2D views, of the 3D CT volume, and output estimates which are combined by statistical fusion exploiting structural similarity. More specifically, OAN is a two-stage deep convolutional network, where deep network features from the first stage are combined with the original image, in a second stage, to reduce the complex background and enhance the discriminative information for the target organs. Intuitively, OAN reduces the effect of the complex background by focusing attention so that each organ only needs to be discriminated from its local background. RCs are added to the first stage to give the lower layers more semantic information thereby enabling them to adapt to the sizes of different organs. Our networks are trained on 2D views (slices) enabling us to use holistic information and allowing efficient computation (compared to using 3D patches). To compensate for the limited cross-sectional information of the original 3D volumetric CT, e.g., the connectivity between neighbor slices, multi-sectional images are re-constructed from the three different 2D view directions. Then we combine the segmentation results from the different views using statistical fusion, with a novel term relating the structural similarity of the 2D views to the original 3D structure. To train the network and evaluate results, 13 structures were manually annotated by four human raters and confirmed by a senior expert on 236 normal cases. We tested our algorithm by 4-fold cross-validation and computed Dice–Sørensen similarity coefficients (DSC) and surface distances for evaluating our estimates of the 13 structures. Our experiments show that the proposed approach gives strong results and outperforms 2D- and 3D-patch based state-of-the-art methods in terms of DSC and mean surface distances.

© 2019 Published by Elsevier B.V.

1. Introduction

Segmentation of the internal structures, like body organs, in medical images is an essential task for many clinical applications such as computer-aided diagnosis (CAD), computer-aided surgery (CAS) and radiation therapy (RT). However, despite intensive studies of automatic or semi-automatic segmentation methods, there remain challenges which need to be overcome before these meth-

ods can be applied to clinical environments. In particular, detailed abdominal organ segmentation on CT is a challenging task both for manual human annotation and for automatic segmentation algorithms for various reasons including the morphological complexity of the structures, the large variations between inter- and intra-subjects, and image characteristics such as low contrast of soft tissues.

Early studies of abdominal organ segmentation focused on specific single organs, for example relatively large isolated structures such as the liver (Heimann et al., 2009; Mharib et al., 2012; Li et al., 2015) or critical structures such as blood vessels (Kirbas and Quek, 2004; Lesage et al., 2009). However, most of the algo-

* Corresponding author.

E-mail address: spark139@jhmi.edu (S. Park).

¹ These authors equally contributed to the work.

rithms were based on specific features of the target organ, and so extensibility to the simultaneous segmentation of multiple organs was limited. For multi-organ segmentation, atlas-based approaches were adopted for many applications (Iglesias and Sabuncu, 2015; Asman and Landman, 2013; Chu et al., 2013; Wolz et al., 2013; Kada et al., 2015; Zhuang and Shen, 2016; Karasawa et al., 2017). The general framework of atlas-based segmentations is to deformably register selected atlas images with segmented structures to the target image. Critical issues for this approach, which affect performance accuracy, include proper atlas selection, accurate deformable image registration, and label fusion. In particular, for the abdominal region, inter-subject variations are relatively large compared with other parts of the body (e.g., the brain) so the segmentation results are dependent on deformable registration between inter-subjects from the limited set of atlases, which is a challenging problem that critically affects the final accuracies. In addition, computational time is strongly dependent on the number of atlases. Therefore, selection of the proper number and types of atlases is a critical factor for both of the accuracy and efficiency.

Recently, learning-based approaches exploiting large datasets have been applied to the segmentation of medical images (Dou et al., 2016; Çiçek et al., 2016; Milletari et al., 2016; Nascimento and Carneiro, 2016; Setio et al., 2016; Chen et al., 2017a; Rajchl et al., 2017; Havaei et al., 2017; Jamnitsas et al., 2017; Zu et al., 2017). In particular, deep convolutional neural networks (CNN) have been very successful (Dou et al., 2016; Çiçek et al., 2016; Roth et al., 2015; 2016; Milletari et al., 2016; Setio et al., 2016; Chen et al., 2017a; Rajchl et al., 2017; Havaei et al., 2017; Jamnitsas et al., 2017). Targets include regions in the brain (Chen et al., 2017a; Havaei et al., 2017; Jamnitsas et al., 2017), chest (Setio et al., 2016), and abdomen (Dou et al., 2016; Roth et al., 2015; 2016). The performance results of CNNs for organs (and even tumors) reach, or outperform, alternative state-of-the-art methods. Unlike multi-atlas-based approaches, deep networks do not require selecting a specific atlas or require deformable registration from training sets to a target image. In this study, we apply deep network approaches to abdominal organ segmentation.

Most studies based on deep networks, however, focused on a single structure segmentation, particularly for abdominal regions, and there are few studies of multi-organ segmentation partly due to technical challenges discussed later. We note that fully convolutional networks (FCNs) (Long et al., 2015) have been generally accepted for organ segmentations on CT scans (Çiçek et al., 2016; Zhou et al., 2017; Roth et al., 2018) partly because they give state-of-the-art performance for semantic segmentation of natural images (Long et al., 2015; Chen et al., 2017b). But there are three major characteristics of abdominal CT which we must address in order to obtain strong performance on multi-organ segmentation.

First, many abdominal organs have weak boundaries between spatially adjacent structures on CT, e.g. between the head of the pancreas and the duodenum. In addition, the entire CT volume includes a large variety of different complex structures. Morphological and topological complexity includes anatomically connected structures such as the gastrointestinal (GI) track (stomach, duodenum, small bowel and colon) and vascular structures. The correct anatomical borders between connected structures may not be always visible in CT, especially in sectional images (i.e., 2D slices), and may be indicated only by subtle texture and shape change, which causes uncertainty even for human experts. This makes it hard for deep networks to distinguish the target organs from the complex background.

Second, there are large variations in the relative sizes of different target organs, e.g. the liver compared to the gallbladder. This causes problems when applying deep networks to multi-organ segmentation because lower layers typically lack semantic information when segmenting small structures. The same problem has

been observed in semantic segmentation of natural images where the segmentation performance on small regions is typically much worse than on large regions, motivating the need to introduce mechanisms which attend to the scale (Chen et al., 2016).

Thirdly, although CT scans are high-resolution three-dimensional volumes, most current deep network methods were designed for 2D images. To overcome the limitations of using 2D CNNs for 3D images, Setio et al. (2016) used multiple 2D patches reconstructed from 9 different directions around the target region for the task of pulmonary nodule detection. Zhuang and Shen (2016) used 2D axial, coronal, and sagittal slices for pancreas detection at the coarse level and also for segmentation at the finer level. More recently, there are studies which use 3D deep networks (Çiçek et al., 2016; Milletari et al., 2016; Roth et al., 2018; Jamnitsas et al., 2017; Roth et al., 2018). These, however, are not networks that act on the entire 3D CT volume but instead are local patch-based approaches (due to complex challenges of 3D deep networks discussed later in this paragraph). To address the problems caused by restricting to image patches, Roth et al. (2018) and Jamnitsas et al. (2017) used a hierarchical approach with multi-resolutions, which reduces the dimension of the whole volume for initial detection and focuses on smaller regions at the finer resolution. But this strategy is best suited to a single target structure. Roth et al. (2018) applied a bigger patch size to deal with the whole dense pancreatic volume, but this was also for single pancreas segmentation and hard to extend to the whole abdominal region. In general, 3D deep networks face far greater complex challenges than 2D deep networks. Both approaches rely heavily on graphics processing units (GPUs) but these GPUs have limited memory size which makes it difficult when dealing with full 3D CT volumes compared to 2D CT slices (which require much less memory). In addition, 3D deep networks typically require many more parameters than 2D deep networks and hence require much more training data, unless they are restricted to patches. But there is limited training data for abdominal CT images, because annotating them is challenging and requires expert human radiologists, which makes it particularly difficult to apply 3D deep networks to abdominal multi-organ segmentation. We have, however, implemented a 3D patch based approach for comparison.

To deal with the technical difficulties for abdominal multi-organ segmentation on CT, we introduce a novel framework of an organ-attention 2D deep networks with reverse connections (OAN-RC) followed by statistical fusion to combine the information from the three different views exploiting structural similarity using local isotropic 3D patches. OAN is a two-stage deep network, which computes an organ-attention map (OAM) from typical probability map of labels for input images in the first stage and combines OAM to the original input image for the second stage. This two-stage strategy effectively reduces the complexity of the background while enhancing the discriminative information of target structures (by concentrating attention close to the target structures). By training OAM with additional deep network, uncertainties and errors from the first stage are adjusted and the fidelity of the final probability map is improved. In this procedure, we apply reverse connections (Kong et al., 2017) to the first stage so that we can localize organ information at different scales by assisting the lower layers with semantic information.

More specifically, we apply OAN-RC to each sectional slice, which is an extreme form of anisotropic local patches but include the whole semantic (i.e. volume) information from one viewing direction. This yields segmentation information from separate sets of multi-sectional images (axial, coronal, and sagittal planes in this study similarly to most of medical image platforms for 2D visualization). We statistically fuse the three sources of information using local isotropic 3D patches based

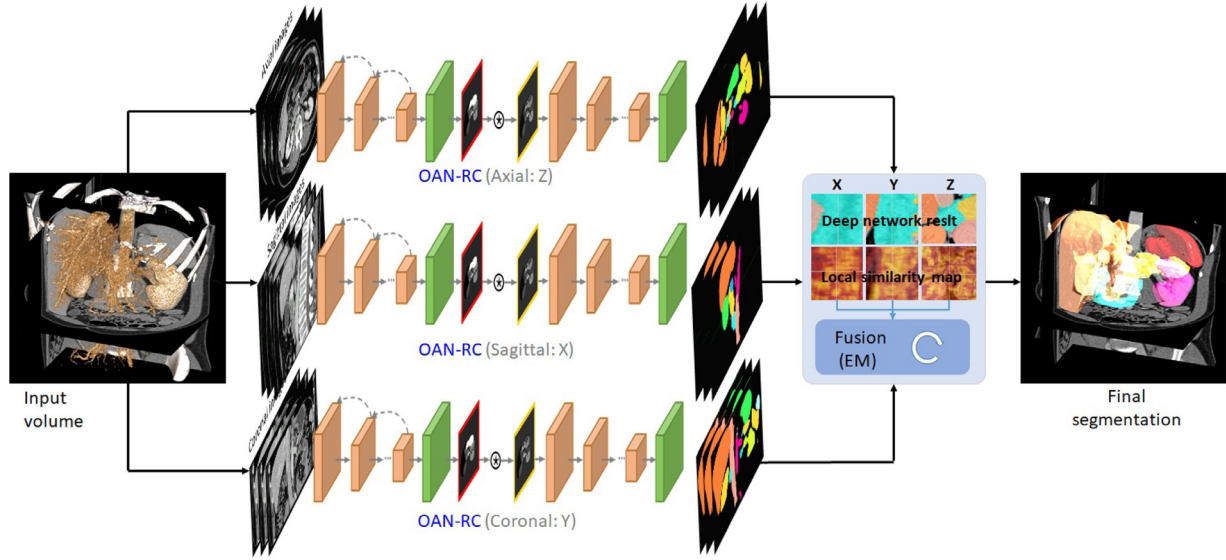


Fig. 1. The overall framework.

on direction-dependent local structural similarity. The basic fusion framework uses expectation-maximization (EM) similar to Warfield et al. (2004) and Asman and Landman (2013). But, unlike typical statistical fusion methods used for atlas-based segmentation, the input volumes and the target volumes for segmentation in our problem are the same. But different structures and texture patterns, from different viewing directions, will often generate nonidentical segmentations in 3D. Our strategy is to exploit structural similarity by computing a direction-dependent local property at each voxel. This models the structural similarity from the 2D images to the original 3D structure (in the 3D volume) by local weights. This structural statistical fusion improves our overall performance by combining the information from the three different views in a principled manner and also imposing local structure.

Fig. 1 describes the graphical concept of our framework. Our proposed algorithm was tested on 236 abdominal CT scans of normal cases collected as a part of FELIX project for pancreatic cancer research (Lugo-Fagundo et al., 2018). By experiments, our method showed robust and high fidelities to the ground-truth for all target structures with smooth boundaries. It outperformed 3D patch-based algorithms as well as 2D-based in terms of DICE-similarity coefficient and average surface distance with memory and computational efficiency.

2. Organ-attention networks with reverse connections

Given a 3D volume of interest (VOI) of a scanned CT image $V \subset \mathbb{R}^3$, our goal is to find the label of each voxel $v \in V$. The target structures (i.e., the labeled structures) are restricted to be organs which do not overlap with each other, so every voxel v should be assigned to a label in a finite set $\mathcal{L}$. In this section we introduce our proposed organ-attention networks with reverse connections (annotated as OAN-RC) which is run separately on three different views, and then in the next section we describe our novel structural similarity statistical fusion method which combines the segmentation results obtained from the OAN-RCs on the three different views.

2.1. Two-stage organ attention network

We first introduce the OAN, which is composed of two jointly optimized stages. The first stage (stage-I) transforms the organ

segmentation probability map to provide spatial attention to the second stage (stage-II), so that the segmentation network trained in stage-II is more discriminative for segmenting organs (because it only has to deal with local context). To assist the lower layers in stage-I with more semantic information, we employ reverse connections (Section 2.2), which pass semantic information down from high layers to low layers. The OAN is trained in an end-to-end fashion to enhance the learning ability of all stages.

There are previous studies using attention maps. Xu et al. (2015) used soft-attention map for text captioning purpose from natural images which can be different from pixelwise semantic segmentation of organs. Especially in medical images, the locations and sizes of organs are more structured than natural image whose context can be random. Therefore, it is possible and useful to get and use more accurate attention map in our pixelwise segmentation. Otkay et al. (2018) applied attention gates into the U-net structure for pancreas detection. In our approach, we use explicit separate 2-stages with supervisions for accurate multi-organ segmentation.

The input images to our OAN are reconstructed 2D slices from axial, sagittal and coronal directions. Based on the normal vector directions of the sagittal (X), coronal (Y) and axial (Z) planes, we denote the 2D images by $\mathbf{I}_i^X$, $\mathbf{I}_j^Y$ and $\mathbf{I}_k^Z$ respectively, where $i = 1, \dots, n_x$, $j = 1, \dots, n_y$, $k = 1, \dots, n_z$ and n_x, n_y, n_z are the numbers of slices for the three directions, respectively, and $\bigcup_i \mathbf{I}_i^X = \bigcup_j \mathbf{I}_j^Y = \bigcup_k \mathbf{I}_k^Z = V$. Following the work of Zhou et al. (2017), we train an individual OAN for each direction.

Fig. 2 illustrates the conceptual architecture of our organ-attention-network. The whole operations (Fig. A1) and the full architecture (Fig. A2) can be found in Appendix A. The network consists of two stages, where each stage is a segmentation network. For notational simplicity, we denote an input 2D slice by $\mathbf{I} \in \mathbb{R}^{H \times W}$ and its corresponding label map by $\mathbf{T} = \{t_i\}_{i=1, \dots, H \times W}$. Stage-I outputs a probability map $\mathbf{P}^{(1)} = f(\mathbf{I}; \Theta^{(1)}) \in \mathbb{R}^{H \times W \times |\mathcal{L}|}$ for each label at every pixel, where the probability density function $f(\cdot; \Theta^{(1)})$ is a segmentation network parameterized by $\Theta^{(1)}$. We use FCN (Long et al., 2015) with reverse connections, which is explained in Section 2.2, as $\Theta^{(1)}$. FCN is the backbone network throughout the paper. Each element $p_{i,l}^{(1)} \in \mathbf{P}^{(1)}$ is the probability that the i th pixel in the input slice belongs to label l , where $l = 0$ is the background, and $l = 1, \dots, |\mathcal{L}|$ are target organs. We define $p_{i,l}^{(1)} = \sigma(a_{i,l}^{(1)}) =$

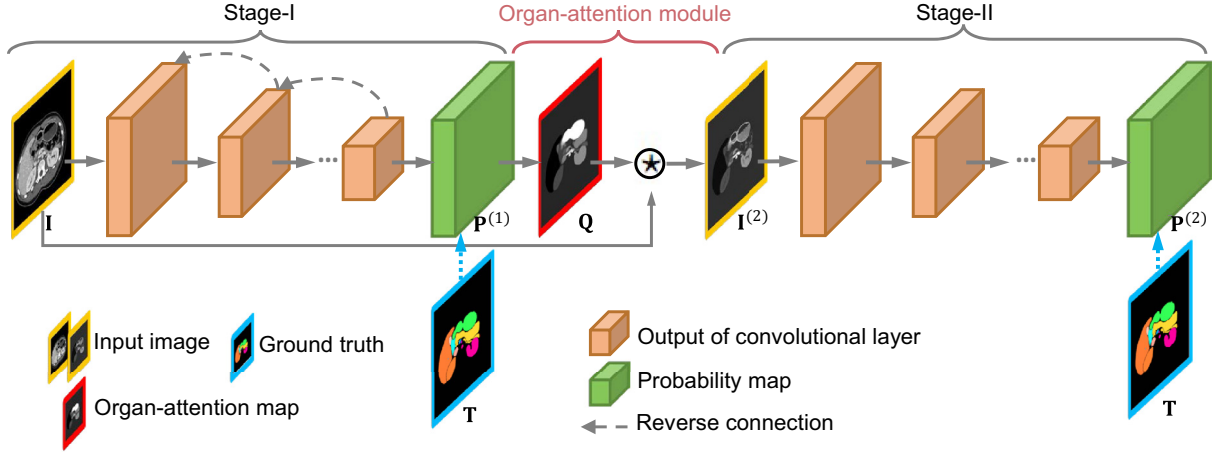


Fig. 2. The architecture of our two-stage organ-attention network with reverse connections. The organ-attention network (OAN) is composed of two jointly optimized stages, where the first stage (stage-I) transforms the organ segmentation probability map by spatial attention to the second stage (stage-II). Hence the organ segmentation map generated in the organ-attention module guides the latter computation. The reverse connections, described in Section 2.2, modify the first stage of OAN as shown by dashed lines.

$\frac{\exp(a_{i,l}^{(1)})}{\sum_{t=0}^{|L|} \exp(a_{i,t}^{(1)})}$, where $a_{i,l}^{(1)}$ is the activation value of the i th pixel on the l th channel dimension. Let $\mathbf{A}^{(1)} = \{a_{i,l}^{(1)}\}_{i=1, \dots, H \times W, l=0, \dots, |L|}$ be the activation map. The objective function to minimize for $\Theta^{(1)}$ is given by

$$\mathcal{J}^{(1)}(\Theta^{(1)}) = -\frac{1}{H \times W} \left[\sum_{i=1}^{H \times W} \sum_{l=0}^{|L|} \mathbf{1}(t_i = l) \log p_{i,l}^{(1)} \right], \quad (1)$$

where $\mathbf{1}(\cdot)$ is an indicator function.

Using a preliminary organ segmentation map to guide the computation of a better organ segmentation can be thought as employing an attentional mechanism. Toward this end, we propose an organ-attention module by

$$\mathbf{Q} = \mathbf{W} * \mathbf{P}^{(1)} + \mathbf{b}, \quad (2)$$

where $*$ denotes the convolution operator, $\mathbf{W}$ indicates the convolutional filters whose dimension is $(5 \times 5 \times |L|)$ in this study, and $\mathbf{b}$ is the bias. Eq. (2) embeds cross-organ information into a single organ-attention map, $\mathbf{Q} \in \mathbb{R}^{H \times W}$, which learns discriminative spatial attention for different organs automatically. By combining $\mathbf{Q}$ with the original input $\mathbf{I}$, we get an image which emphasizes each organ by

$$\mathbf{I}^{(2)} = \mathbf{I} * \mathbf{Q}, \quad (3)$$

where $*$ is the element-wise product operator. We apply $\mathbf{I}^{(2)} \in \mathbb{R}^{H \times W}$ to the input of stage-II, and the probability of stage-II then becomes $\mathbf{P}^{(2)} = f(\mathbf{I}^{(2)}; \Theta^{(2)})$.

In order to drive stage-II to focus on organ regions without needing to deal with complicated non-local background, we define a selection function, $\mathbf{1}(\mathbf{P}_0^{(1)} \leq \rho)$ where $\mathbf{P}_0^{(1)} = \{p_{i,0}^{(1)}\}_{i=1, \dots, H \times W}$ is the probability map provided by stage-I. In stage-II, we only accept the region if $p_{i,0}^{(1)} > \rho$ and do not back-propagate it to stage-I. We selected $\rho = 0.98$ in a conservative way in our experiments to keep the possible foreground region as much as possible from the stage-I so that the remaining boundary uncertainty can be dealt at the stage-II and possible noise and inconsistency can be treated by the following statistical fusion step. The value between $0.95 < \rho < 0.98$ did not sensitively affect to the final results. The loss function for stage-II is formulated as

$$\mathcal{J}^{(2)}(\Theta^{(2)}, \mathbf{W}, \mathbf{b}) = -\frac{1}{H \times W} \left[\sum_{i=1}^{H \times W} \sum_{l=0}^{|L|} \mathbf{1}(p_{i,0}^{(1)} \leq \rho) \cdot \mathbf{1}(t_i = l) \log p_{i,l}^{(2)} \right]. \quad (4)$$

To jointly optimize stage-I and stage-II, we define a loss function aiming at estimating parameters $\Theta^{(1)}$, $\Theta^{(2)}$, $\mathbf{W}$, and $\mathbf{b}$ by optimizing

$$\mathcal{J} = h^{(1)} \mathcal{J}^{(1)}(\Theta^{(1)}) + h^{(2)} \mathcal{J}^{(2)}(\Theta^{(2)}, \mathbf{W}, \mathbf{b}), \quad (5)$$

where $h^{(1)}$ and $h^{(2)}$ are the fusion weights. By experiments, a stronger weight of the finer network, stage-II, than the coarse network, stage-I, i.e. $h^{(1)} < h^{(2)}$, showed better performance compared to $h^{(1)} \geq h^{(2)}$, and $h^{(1)} = 0.5$ and $h^{(2)} = 1.5$ were initially set empirically and fixed during training.

2.2. Reverse connections

FCNs (Long et al., 2015) have shown good segmentation results in recent studies, especially for single organ segmentation. However, for multi-organ segmentation, lower layers typically lack semantic information, which may lead to inaccurate segmentation particularly for smaller structures. Therefore, we propose reverse connections which feed coarse-scale (high) layer information backward to fine-scale (low) layer for semantic segmentation of multi-scale structures, inspired by Kong et al. (2017). This enables us to connect abstract high-level semantic information to the more detailed lower layers so that all the target organs have similar levels of details and abstract information at the same layer. The reverse connections framework for stage-I is shown in Fig. 3. Fig. 4 illustrates a reverse connection block. Let $\mathbf{R}_n$ denote the reverse connection map of the n th convolutional layer in the backbone network, i.e. FCN in this study, where $\mathbf{C}_n$ is the output of the n th convolutional layer. A convolutional layer (with 512 channels by 3×3 kernels) is added after $\mathbf{C}_n$, and a deconvolutional layer (with 512 channels by 4×4 kernels) is applied after $\mathbf{R}_{n+1}$. $\mathbf{R}_n$ is then obtained via an element-wise summation of these two maps. $\mathbf{R}_7$ is the output of a convolutional layer (with 512 channels by 2×2 kernels) grafted onto $\mathbf{C}_7$. Let $\mathbf{w}^n$ denote the corresponding weights for obtaining $\mathbf{R}_n$. Following Kong et al. (2017), we add reverse connections from $\mathbf{C}_4$ to $\mathbf{C}_7$.

With these learnable reverse connections, the semantic information of the lower layers can be enriched. In order to drive learned reverse connection maps to produce segmentation results approaching the ground-truth, we make each reverse connection map associate with a classifier. As the side-output layers proposed in Kong et al. (2017) are designed for detection purposes, they are not suitable for our task. Instead we follow the side-outputs used

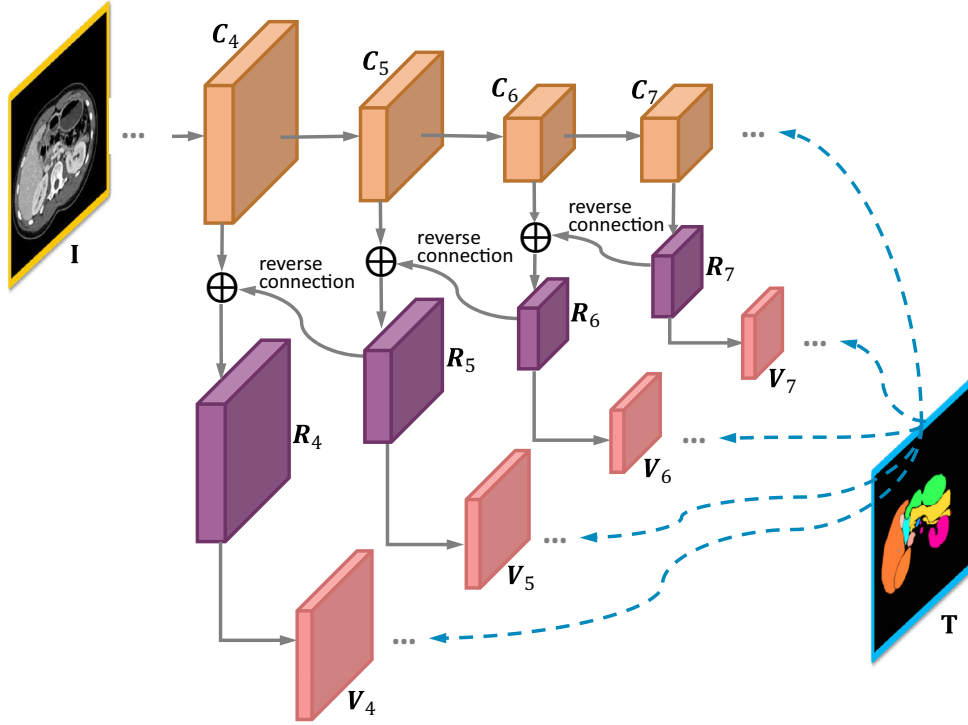


Fig. 3. The reverse connections architecture of OAN stage-I. The network has reverse connections to the output of convolutional layers. In the training step, both backbone network and reverse connection side-outputs are supervised by the ground-truth. Finally, all reverse connection side-outputs and the output of backbone network are fused and made to approach ground-truth.

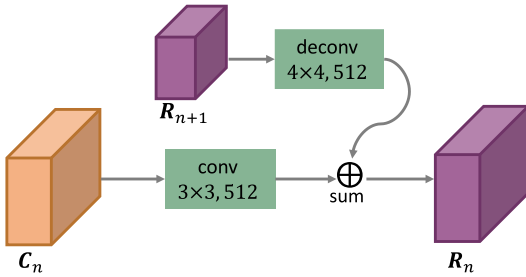


Fig. 4. A reverse connection block.

in Xie and Tu (2015). More specifically, a convolutional layer (with $|\mathcal{L}|$ channels by 1×1 kernels) is added on top of R_n , whose output is denoted as V_n , and followed by a deconvolutional layer (with $|\mathcal{L}|$ channels), and each side-output layer is supervised by the ground-truth as shown in Fig. 3. This side-output layer also brings differences from the original deep supervision approach (Wang et al., 2015) by allowing additional convolution between different level of side-output layers which makes multi-scale fusion of the layers. We denote the weights of the n th side-output layer by θ^n . Therefore, the loss function for side-output layers $\mathcal{J}^{(s,1)}$ is defined as

$$\mathcal{J}^{(s,1)}(\Theta^{(1)}, \mathbf{w}, \theta) = \sum_{n=4}^7 h_n^{(s,1)} \ell_n^{(s,1)}(\Theta^{(1)}, \mathbf{w}^n, \theta^n), \quad (6)$$

where $\ell_n^{(s,1)} = -\frac{1}{H \times W} \left[\sum_{i=1}^{H \times W} \sum_{l=0}^{|\mathcal{L}|} \mathbf{1}(t_i = l) \log p_{i,l}^{(s,1)} \right]$ and $p_{i,l}^{(s,1)}$ is the probability output of the n th side-output layer. Each $h_n^{(s,1)}$ was set as 0.5 and fixed during training.

In order to combine the learned reverse connection maps of fine layers and coarse layers, we add up the predictions (i.e., V_n) of the reverse connection maps from high layer to low layer gradually. First, V_6 is fused with a $2 \times$ upsampling of V_7 by an element-wise addition. Then we follow the same strategy and gradually merge V_5 and V_4 , as shown in Fig. 5. To obtain a fused activation map $\mathbf{A}^{(f,1)} = \{a_{i,l}^{(f,1)}\}_{i=1, \dots, H \times W, l=0, \dots, |\mathcal{L}|}$ from the activation map of both side-outputs (i.e., $\mathbf{A}^{(s,1)}$) and convolutional layers in the backbone network (i.e., $\mathbf{A}^{(b,1)}$), a scale function is adopted followed by an element-wise addition by

$$\mathbf{A}_l^{(f,1)} = h_l^{(s,1)} \mathbf{A}_l^{(s,1)} + h_l^{(b,1)} \mathbf{A}_l^{(b,1)}, \quad l = 0, \dots, |\mathcal{L}| \quad (7)$$

where $\mathbf{A}_l$ indicates the l th channel of the activation map. $h_l^{(s,1)}$ and $h_l^{(b,1)}$ are fusion weights, and we set $h_l^{(s,1)} = 0.75$ and $h_l^{(b,1)} = 0.15$ for each l th channel of the activation map and updated them by gradient descent during training. Then the fused probability map, $\mathbf{P}^{(f,1)} = \{p_{i,l}^{(f,1)}\}_{i=1, \dots, H \times W, l=0, \dots, |\mathcal{L}|}$, can be obtained by $p_{i,l}^{(f,1)} = \sigma(a_{i,l}^{(f,1)})$. The final objective function for stage-I is defined by

$$\mathcal{J}^{(1)}(\Theta^{(1)}, \mathbf{w}, \theta) = h^{(b,1)} \mathcal{J}^{(b,1)}(\Theta^{(1)}) + h^{(s,1)} \mathcal{J}^{(s,1)}(\Theta^{(1)}, \mathbf{w}, \theta) + h^{(f,1)} \mathcal{J}^{(f,1)}(\Theta^{(1)}, \mathbf{w}, \theta), \quad (8)$$

where $h^{(b,1)}$, $h^{(s,1)}$ and $h^{(f,1)}$ are fusion weights, and

$$\mathcal{J}^{(f,1)}(\Theta^{(1)}, \mathbf{w}, \theta) = -\frac{1}{H \times W} \left[\sum_{i=1}^{H \times W} \sum_{j=0}^{|\mathcal{L}|} \mathbf{1}(t_i = j) \log p_{i,j}^{(f,1)} \right]. \quad (9)$$

Note that in our full system with the two-stage organ-attention network and reverse connections, all the parameters are optimized simultaneously by standard back-propagation

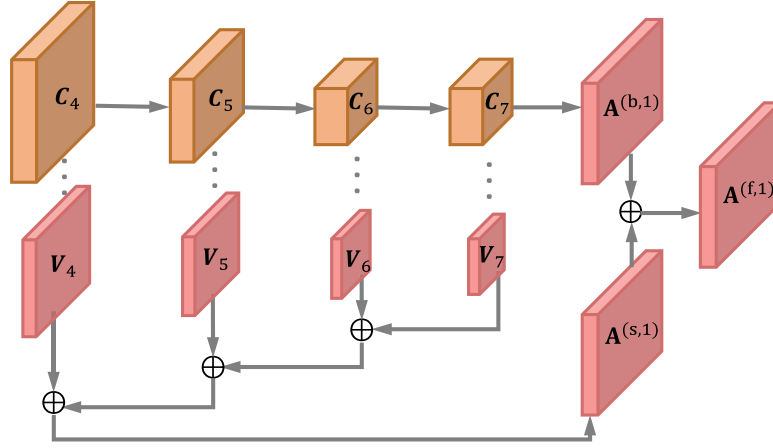


Fig. 5. Feature fusion strategy. A deep-to-shallow refinement is adopted for multi-scale side-output features. The final activation map ($A^{(f,1)}$) for stage-I is an element-wise addition of the side-output activation map ($A^{(s,1)}$) and the backbone network activation map ($A^{(b,1)}$).

$$(\hat{\Theta}^{(1)}, \hat{\mathbf{w}}, \hat{\theta}, \hat{\Theta}^{(2)}, \hat{\mathbf{W}}, \hat{\mathbf{b}}) = \arg \min \{h^{(1)} \mathcal{J}^{(1)}(\Theta^{(1)}, \mathbf{w}, \theta) + h^{(2)} \mathcal{J}^{(2)}(\Theta^{(2)}, \mathbf{W}, \mathbf{b})\}. \quad (10)$$

Algorithm 1 shows a step-by-step forward pass iteration of our OAN-RC.

Algorithm 1 One forward pass iteration of the organ attention network with reverse connections

Input: initialize network parameters, $\Theta^{(1)}, \Theta^{(2)}, \mathbf{w}, \mathbf{W}, \theta$ randomly

input training slice $\mathbf{I}$ and its label map $\mathbf{T}$.

Output: updated network parameters, $\Theta^{(1)}, \Theta^{(2)}, \mathbf{w}, \mathbf{W}, \theta, \mathbf{b}$.

- 1: Obtain $\mathbf{P}^{(1)}$ by the forward pass $\mathbf{P}^{(1)} = f(\mathbf{I}; \Theta^{(1)}, \mathbf{w}, \theta)$
- 2: Compute the organ-attention map $\mathbf{Q}$ by Eq. 2
- 3: Calculate the input to stage-II $\mathbf{I}^{(2)}$ by Eq. 3
- 4: Obtain $\mathbf{P}^{(2)}$ by the forward pass $\mathbf{P}^{(2)} = f(\mathbf{I}^{(2)}; \Theta^{(2)})$
- 5: Update parameters $\Theta^{(1)}, \Theta^{(2)}, \mathbf{w}, \mathbf{W}, \theta, \mathbf{b}$ by standard back-propagation, Eq. 10

2.3. Testing phase

In the testing stage, given a slice $\mathbf{I}$, we obtain the stage-I and stage-II probability map by

$$\begin{aligned} \mathbf{P}^{(1)} &= f(\mathbf{I}; \hat{\Theta}^{(1)}, \hat{\mathbf{w}}, \hat{\theta}) \\ \mathbf{P}^{(2)} &= f(\mathbf{I}; \hat{\Theta}^{(2)}, \hat{\mathbf{W}}, \hat{\mathbf{b}}), \end{aligned} \quad (11)$$

where $f(\cdot, \cdot)$ is the network functions defined in Section 2.1. A fused probability map of $\mathbf{P}^{(1)}$ and $\mathbf{P}^{(2)}$ is then given by

$$\mathbf{P} = \mathbf{P}^{(1)} \circ \mathbf{1}(\mathbf{P}_0^{(1)} > \rho) + \mathbf{P}^{(2)} \circ \mathbf{1}(\mathbf{P}_0^{(1)} \leq \rho). \quad (12)$$

The final label map $\mathbf{S} = \{s_i\}_{i=1, \dots, H \times W}$ is determined by $s_i = \arg \min_{l \in \mathcal{L}} p_{i,l}$.

3. Statistical label fusion based on local structural similarity

As described in Section 1, our OAN-RC is based on 2D images which is an extreme case of 3D anisotropic patches. In this section, we propose to fuse anisotropic information obtained from different viewing directions using isotropic 3D local patches to estimate the final segmentation. Let us denote the segmentation results by $\mathbf{S}^j$, ($j = 1, \dots, M = 3$), which are obtained as described in

Section 2.3 from the axial (Z), sagittal (X), and coronal (Y) OAN-RCs. Depending on the viewing directions, sectional images contain different structures and may have different texture patterns in the same organs. These differences can cause nonidentical segmentations by the deep network as shown in Fig. 6 in 3D. In addition, there is no guarantee of connectivity between neighbor slices by independent use of slices for training and testing. Possible naïve approaches for determining the final segmentation in 3D from the OAN-RC results can be boolean operations such as union or intersection. Majority voting (MV) is another candidate for efficient fusion, however, these approaches assume the same global weights of OAN-RC results. From the observations that the performance level of segmentation, e.g. sensitivity, can be different from viewing directions for each organ, we set the performance level to be an unknown variable when computing the probability of labeling. This concept is similar to the label fusion algorithms using expectation-maximization (EM) framework such as STAPLE (simultaneous truth and performance level estimation) and its extensions (Warfield et al., 2004; Asman and Landman, 2012; 2013).

Let us denote the true label of the V by $\mathbf{T}$, which is unknown, and the unknown performance level parameter of segmentation by θ . The segmentations from the deep networks $\mathbf{S} = \{\mathbf{S}^j | j = 1, \dots, M\}$ are observed values. Under this condition, the basic EM framework is performed by following two steps in an iterative manner: (1) to compute $Q^0(\theta | \theta^{(k)}) = E_{\mathbf{T}}[\ln L(\theta | \mathbf{S}, \mathbf{T}) | \mathbf{S}, \theta^{(k)}]$ which is the expected value of the log likelihood, $\ln L(\theta | \mathbf{S}, \mathbf{T}) = \ln P(\mathbf{S}, \mathbf{T} | \theta)$, under the current estimate of the parameters $\theta^{(k)}$ at k th iteration, and (2) to find the parameter $\theta^{(k+1)}$ which maximizes $Q^0(\theta | \theta^{(k)})$.

The maximization step can be written as

$$\begin{aligned} \theta^{(k+1)} &= \arg \max_{\theta} E_{\mathbf{T}}[\ln P(\mathbf{S}, \mathbf{T} | \theta) | \mathbf{S}, \theta^{(k)}] \\ &= \arg \max_{\theta} E_{\mathbf{T}}[\ln P(\mathbf{S} | \mathbf{T}, \theta) P(\mathbf{T}) | \mathbf{S}, \theta^{(k)}] \\ &= \arg \max_{\theta} \sum_{\mathbf{T}} \ln \{P(\mathbf{S} | \mathbf{T}, \theta) P(\mathbf{T})\} P(\mathbf{T} | \mathbf{S}, \theta^{(k)}) \\ &= \arg \max_{\theta} \sum_{\mathbf{T}} \{\ln P(\mathbf{S} | \mathbf{T}, \theta) + \ln P(\mathbf{T})\} P(\mathbf{T} | \mathbf{S}, \theta^{(k)}). \end{aligned} \quad (13)$$

By assuming independence between $\mathbf{T}$ and θ in our problem, the second term $\sum_{\mathbf{T}} \ln P(\mathbf{T}) P(\mathbf{T} | \mathbf{S}, \theta^{(k)})$ in Eq. (13) becomes free of θ and the maximization step can be written as

$$\begin{aligned} \theta^{(k+1)} &= \arg \max_{\theta} \sum_{\mathbf{T}} \ln P(\mathbf{S} | \mathbf{T}, \theta) P(\mathbf{T} | \mathbf{S}, \theta^{(k)}) \\ &= \arg \max_{\theta} E_{\mathbf{T}}[\ln P(\mathbf{S} | \mathbf{T}, \theta) | \mathbf{S}, \theta^{(k)}]. \end{aligned} \quad (14)$$

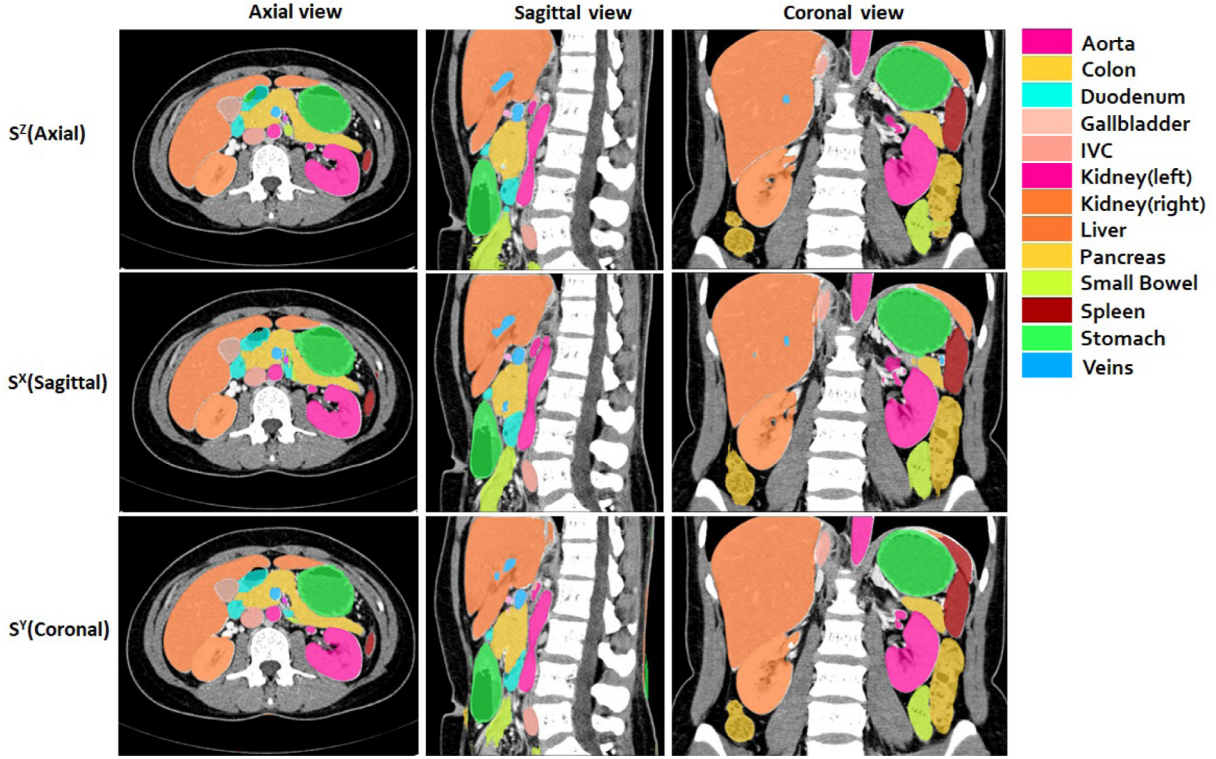


Fig. 6. An example of multi-planar reconstruction view of OAN-RC estimations.

Therefore, we redefine $Q^0(\theta|\theta^{(k)})$ as $Q(\theta|\theta^{(k)}) = E_T[\ln P(\mathbf{S}|\mathbf{T}, \theta)|\mathbf{S}, \theta^{(k)}]$.

The performance level parameter in this framework is a global property representing the overall confidence of deep network segmentation for the whole volume. However, it can also vary according to the voxel spatial locations via the local and neighbor structures as we use 2D slices for the initial segmentation. Therefore, we propose to combine local structural similarity shown from a specific viewing direction to the original 3D volume and the global performance level, conceptually similar to local weighted voting (Sabuncu et al., 2010). We compute the probability of correspondence between 2D images and the 3D volume by structural similarity (SSIM) (Wang et al., 2004) by

$$\alpha_i^j = P(\ell_2(I_i^j)|\ell_3(V_i)) \equiv \text{SSIM}(\ell_2(I_i^j), \ell_3(V_i))$$

$$= \frac{(2\mu_{\ell_2(I_i^j)}\mu_{\ell_3(V_i)} + c_1)(2\sigma_{\ell_2(I_i^j)\ell_3(V_i)} + c_2)}{(\mu_{\ell_2(I_i^j)}^2 + \mu_{\ell_3(V_i)}^2 + c_1)(\sigma_{\ell_2(I_i^j)}^2 + \sigma_{\ell_3(V_i)}^2 + c_2)}, \quad (15)$$

where α_i^j is the SSIM from the j th viewing direction at the i th voxel. c_1 and c_2 are user-defined constants, and $\ell_2(I_i)$ and $\ell_3(V_i)$ represent local 2D and 3D patches centered at the i th voxel, respectively. μ_ℓ and σ_ℓ are the average and standard deviation of the patch ℓ , respectively, and $\sigma_{\ell_2(I_i)\ell_3(V_i)}$ is the covariance of $\ell_2(I_i)$ and $\ell_3(V_i)$. Fig. 7 shows an example of the structural similarity computed on different viewing directions as a color map.

Considering the local image properties, the expectation of log likelihood function in our problem becomes

$$Q(\theta|\theta^{(k)}) = E[\ln P(\mathbf{S}, I|\mathbf{T}, V, \theta)|\mathbf{S}, I, V, \theta^{(k)}]$$

$$= \sum_{\mathbf{T}} \ln P(\mathbf{S}, I|\mathbf{T}, V, \theta) P(\mathbf{T}|\mathbf{S}, I, V, \theta^{(k)}). \quad (16)$$

The global underlying performance level parameters of the deep network segmentations is defined as

$$\theta_{js's} \equiv P(\mathbf{S}_i^j = s'|\mathbf{T}_i = s, \theta_{js's}^{(k)}), \quad (17)$$

where $\theta_{js's}$ is the probability of the voxel labeled as s' from the j th deep network with the current estimated performance value $\theta_{js's}^{(k)}$, when the true label is s .

To make the problem simple, we assume conditional independence between labeling and the original volume intensities. The labeling probability with the target image intensity then becomes

$$P(\mathbf{S}_i^j = s', \ell_2(I_i^j)|\mathbf{T}_i = s, \ell_3(V_i), \theta_{js's}^{(k)})$$

$$= P(\mathbf{S}_i^j = s'|\mathbf{T}_i = s, \theta_{js's}^{(k)}) P(\ell_2(I_i^j)|\ell_3(V_i))$$

$$= \theta_{js's} \alpha_i^j. \quad (18)$$

3.1. E-step

In the expectation step (E-step), we estimate the probability of voxelwise labels. Let us denote the probability that the true label of i th voxel is $s \in \mathcal{L}$ at the k th iteration by $\omega_{si}^{(k)}$. When the deep network segmentations $\mathbf{S}$ and performance level parameters at the k th iteration $\theta^{(k)}$ are given, $\omega_{si}^{(k)}$ can be then described as

$$P(\mathbf{T}_i = s|\mathbf{S}, I, V, \theta^{(k)}) \equiv \omega_{si}^{(k)}, \quad (19)$$

where $\theta \in \mathbb{R}^{N \times |\mathcal{L}| \times |\mathcal{L}|}$ is the vector of all $(\theta_{js's})^T$. From the independence between $\mathbf{S}^X, \mathbf{S}^Y$, and $\mathbf{S}^Z$, we apply Bayesian theorem to Eq. (19).

$$\omega_{si}^{(k)} = \frac{P(\mathbf{T}_i = s) \prod_j P(\mathbf{S}_i^j = s', \ell_2(I_i^j)|\mathbf{T}_i = s, \ell_3(V_i), \theta_j^{(k)})}{\sum_n P(\mathbf{T}_i = n) \prod_j P(\mathbf{S}_i^j = s', \ell_2(I_i^j)|\mathbf{T}_i = n, \ell_3(V_i), \theta_j^{(k)})}, \quad (20)$$

where $P(\mathbf{T}_i = s)$ is a priori of the i th voxel. By applying Eq. (18) to Eq. (20), we then obtain the probability of voxelwise labeling as

$$\omega_{si}^{(k)} = \frac{P(\mathbf{T}_i = s) \prod_j \theta_{js's}^{(k)} \alpha_i^j}{\sum_n P(\mathbf{T}_i = n) \prod_j \theta_{js'n}^{(k)} \alpha_i^j}. \quad (21)$$

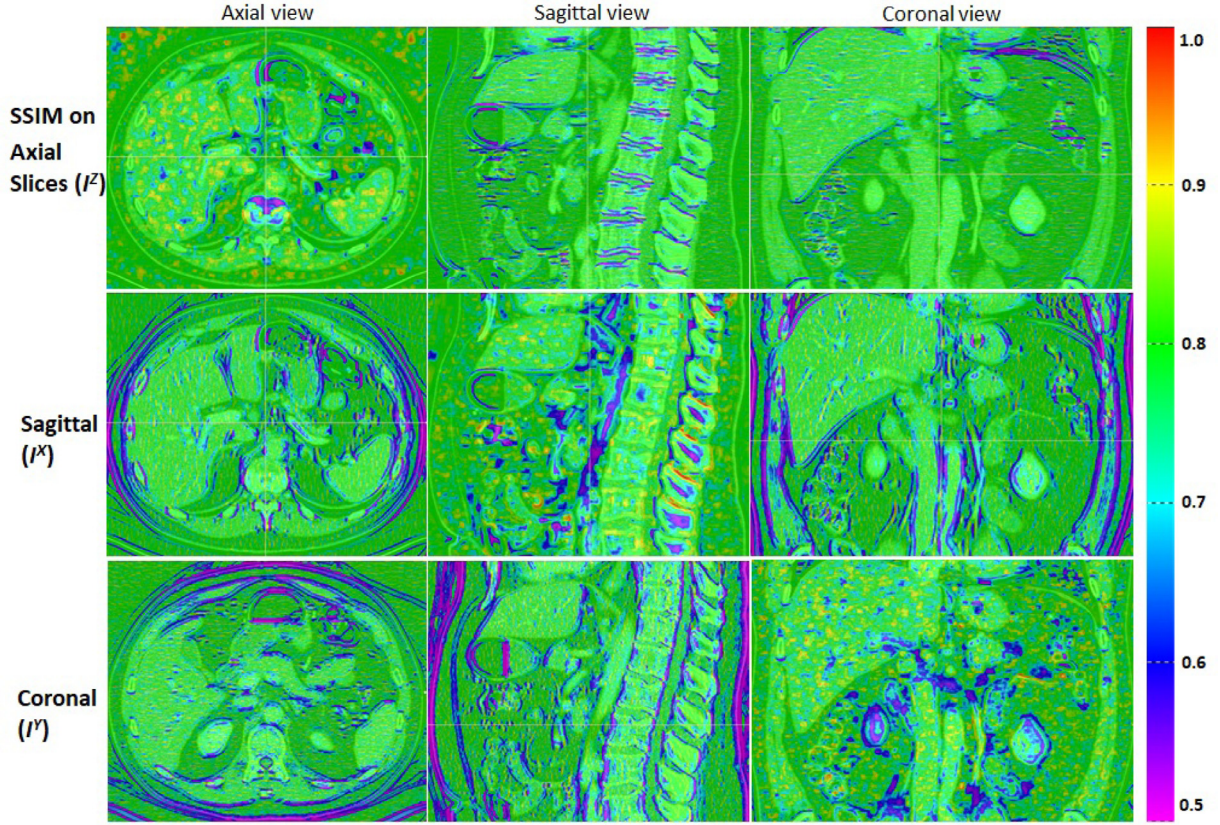


Fig. 7. The local structural similarity map between 2D slices and the 3D volume. Each row is captured from the same similarity map computed on one viewing direction. Each column shows the captures images at the same location computed from different viewing directions. (For interpretation of the references to colour in this figure legend, the reader is referred to the web version of this article.)

3.2. M-step

In the maximization step (*M-step*), the goal is to find the performance parameters, θ , which maximize Eq. (16) with the current given parameters. Considering each $\mathbf{S}^j$ and θ_j independently, the expectation of log likelihood function in Eq. (16) can be expressed with the estimated voxelwise probability in *E-step*. Then the performance parameter of each segmentation can be formulated to find the solution which maximizes the summation of voxelwise probability as

$$\theta_j^{(k+1)} = \arg \max_{\theta_j} Q(\theta_j | \theta_j^{(k)}) = \arg \max_{\theta_j} \sum_i Q_i(\theta_j | \theta_j^{(k)}), \quad (22)$$

where $Q_i = E[\ln P(\mathbf{S}_i, \ell_2(I_i) | \mathbf{T}_i, \ell_3(V_i), \theta^{(k)}) | \mathbf{S}, I, V, \theta^{(k)}]$ at *i*th voxel. By applying Eq. (19) and Eq. (18), Eq. (22) becomes

$$\begin{aligned} \theta_j^{(k+1)} &= \arg \max_{\theta_j} \sum_i \sum_s P(\mathbf{T}_i = s | \mathbf{S}, I, V, \theta^{(k)}) \\ &\quad \times \ln P(\mathbf{S}_i^j, \ell_2(I_i^j) | \mathbf{T}_i = s, \ell_3(V_i), \theta_j^{(k)}) \\ &= \arg \max_{\theta_j} \sum_i \sum_s \omega_{si}^{(k)} \ln P(\mathbf{S}_i^j, \ell_2(I_i^j) | \mathbf{T}_i = s, \ell_3(V_i), \theta_j^{(k)}) \\ &= \arg \max_{\theta_j} \sum_{s'} \sum_{i: \mathbf{S}_i^j = s'} \sum_s \omega_{si}^{(k)} \\ &\quad \times \ln P(\mathbf{S}_i^j = s', \ell_2(I_i^j) | \mathbf{T}_i = s, \ell_3(V_i), \theta_j^{(k)}) \\ &= \arg \max_{\theta_j} \sum_{s'} \sum_{i: \mathbf{S}_i^j = s'} \sum_s \omega_{si}^{(k)} \ln \theta_{js's} \alpha_i^j. \end{aligned} \quad (23)$$

From the definition of θ in Eq. (17), the summation of probability mass function, $\sum_{s'} \theta_{js's}^{(k)}$ must be 1, and Eq. (22) becomes a con-

strained optimization problem which can be solved by introducing a Lagrange multiplier, λ . We then obtain the optimal solution by making the first gradient zero as

$$0 = \frac{\partial}{\partial \theta_{js's}} \left[Q(\theta_j | \theta_j^{(k)}) + \lambda \sum_{s'} \theta_{js's} \right]. \quad (24)$$

By applying the derivation of Q in Eqs. (16), (22) and (23), Eq. (24) becomes

$$\begin{aligned} 0 &= \frac{\sum_{i: \mathbf{S}_i^j = s'} \omega_{si}^{(k)} \alpha_i^j}{\theta_{js's}} + \lambda \\ \theta_{js's}^{(k+1)} &= \frac{\sum_{i: \mathbf{S}_i^j = s'} \omega_{si}^{(k)} \alpha_i^j}{-\lambda}. \end{aligned} \quad (25)$$

By substituting the constraint of $\sum_{s'} \theta_{js's}^{(k)} = 1$, we can obtain the final optimal solution as

$$\theta_{js's}^{(k+1)} = \frac{\sum_{i: \mathbf{S}_i^j = s'} \alpha_i^j \omega_{si}^{(k)}}{\sum_i \omega_{si}^{(k)}}. \quad (26)$$

The two steps, Eqs. (21) and (26), are then computed alternatively in the EM iterations until they converge. From the final values of Eq. (21), the final segmentation can be computed by graph-based approaches such as Boykov et al. (2001).

3.3. Parallel computing using GPUs

The fusion step can be efficiently computed in a parallel way on a GPU. The local structural similarity α_i^j of *i*th voxel in *j*th deep network and *priori* $P(\mathbf{T}_i)$ can be computed for each voxel and saved

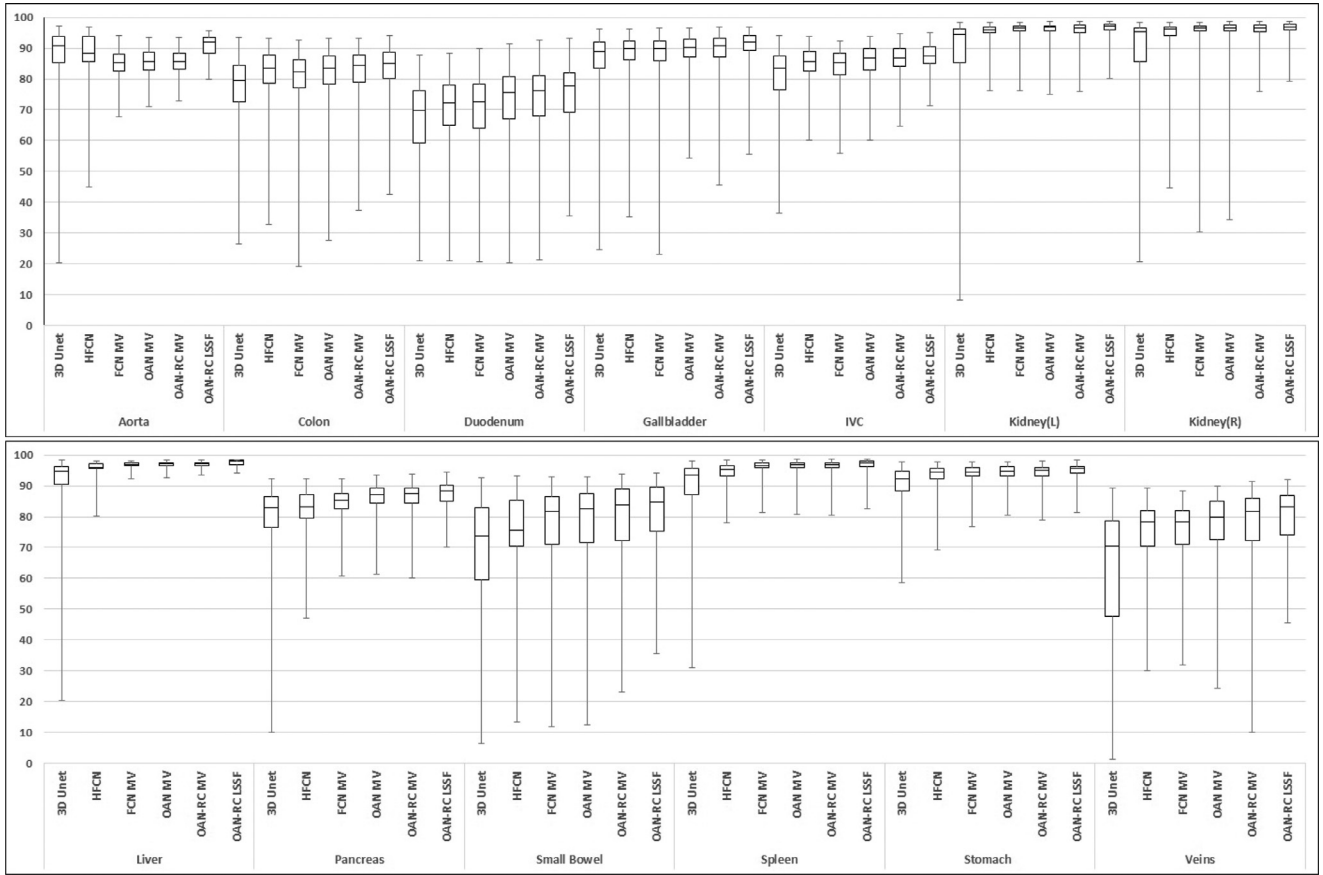


Fig. 8. Box plots of the Dice-Sørensen similarity coefficients of 13 structures to compare performance. As in typical box plots, the box represents the first quartile, median, and the third quartile from the lower border, middle and the upper boarder, respectively, and the lower and the upper whiskers show the minimum and the maximum values. (LSF: local similarity-based statistical fusion.)

as a pre-processing step. In the EM iterations, as shown in Eq. (21), the probability can be computed and updated for each structure at each voxel. In our implementation, a GPU thread is logically allocated for each voxel. However, to reduce the used memory and computation cost, the target volume of interest (VOI) for each structure s is computed in an extended region as $\delta = 4$ voxels for each direction from $V(\bigcup_j S^j = s)$ in our implementation. For parallel computing, one CPU thread is allocated to a structure and launches a kernel of one GPU to compute EM iteration for each structure.

4. Experimental results

We evaluated our methods on 236 abdominal CT images of normal cases under an IRB (institutional review board) approved protocol in Johns Hopkins Hospital as a part of the FELIX project for pancreatic cancer research (Lugo-Fagundo et al., 2018). CT images were obtained by Siemens Healthineers (Erlangen, Germany) SOMATOM Sensation and Definition CT scanners. CT scans are composed of (319–1051) slices of (512×512) images, and have voxel spatial resolution of $([0.523 - 0.977] \times [0.523 - 0.977] \times 0.5) \text{ mm}^3$. All CT scans are contrast enhanced images and obtained in the portal venous phase.

A total of 13 structures for each case were segmented by four human annotators/raters, one case by one person, and confirmed by an independent senior expert. The structures include the aorta, colon, duodenum, gallbladder, interior vena cava (IVC), kidney (left, right), liver, pancreas, small bowel, spleen, stomach, and large

veins. Vascular structures were segmented only outside of the organs in order to make the structures exclusive to each other (i.e. no overlaps).

As explained in Section 2, we used OAN-RCs for multi-organ segmentation whose backbone FCNs had been pre-trained by *PascalVOC* dataset (Everingham et al., 2015). From the possible variants of FCNs (e.g., FCN-32s, FCN-16s, and FCN-8s), which depend on how they combine the fine detailed predictions (Shelhamer et al., 2017). FCN-32s upsample stride 32 predictions back to pixels in a single step, FCN-16s combine predictions from both the final layer and the 4th pooling layer, at stride 16, and FCN-8s combine additional predictions from the 3rd pooling layer, at stride 8. We selected FCN-8s in this study because it captures very fine details in the 3rd and 4th pooling layer, and keeps high-level semantic contextual information from the final layer. Our algorithm was implemented and tested on a workstation with Intel i7-6850K CPU, NVidia TITAN X (PASCAL) GPU. With 236 cases, the initial segmentations using OAN-RCs were tested by four-fold cross-validation. All the input images of OAN-RCs are 1.5 times enlarged by upsampling, which lead to improved performance in our experiments.

In the fusion step, the average probability of S^X , S^Y , S^Z are taken as a *priors* in Eq. (21) and the initial performance levels $\theta_{js's}^{(0)}$ were computed by randomly selecting 5 cases and by comparing them to the ground-truth. To compute the local patch-based structural similarity in Eq. (15), patches of $(4.5 \times 4.5 \times 4.5) \text{ mm}^3$ size cubes were used for 3D volume. Since CT voxels are not always isotropic and spatial resolutions can be different between scan volumes, we re-sampled the 3D patch with 0.5 mm length cubic voxels so that

Table 1Dice–Sørensen similarity coefficient (DSC, %) of thirteen segmented organs (mean $\pm$ standard deviation of 236 cases).

Structure	3D U-net	HFCN	FCN MV	OAN MV	OAN-RC MV	OAN-RC LSSF
Aorta	87.0 $\pm$ 12.3	88.3 $\pm$ 8.8	85.0 $\pm$ 4.2	85.5 $\pm$ 4.2	85.3 $\pm$ 4.1	91.8 $\pm$ 3.5
Colon	77.0 $\pm$ 11.0	79.3 $\pm$ 9.2	80.3 $\pm$ 9.1	81.5 $\pm$ 9.4	82.0 $\pm$ 8.8	83.0 $\pm$ 7.4
Duodenum	66.8 $\pm$ 12.8	70.3 $\pm$ 10.4	70.2 $\pm$ 11.3	72.6 $\pm$ 11.4	73.4 $\pm$ 11.1	75.4 $\pm$ 9.1
Gallbladder	85.4 $\pm$ 10.3	87.9 $\pm$ 7.5	87.8 $\pm$ 8.3	88.9 $\pm$ 6.2	89.4 $\pm$ 6.1	90.5 $\pm$ 5.3
IVC	80.8 $\pm$ 10.2	84.7 $\pm$ 5.9	84.0 $\pm$ 6.0	85.6 $\pm$ 5.8	86.0 $\pm$ 5.5	87.0 $\pm$ 4.2
Kidney(L)	83.9 $\pm$ 22.4	95.2 $\pm$ 2.6	96.1 $\pm$ 2.0	96.2 $\pm$ 2.2	95.9 $\pm$ 2.3	96.8 $\pm$ 1.9
Kidney(R)	88.0 $\pm$ 14.4	95.6 $\pm$ 4.5	95.8 $\pm$ 4.9	95.9 $\pm$ 4.9	96.0 $\pm$ 2.5	98.4 $\pm$ 2.1
Liver	91.4 $\pm$ 9.9	95.7 $\pm$ 1.8	96.8 $\pm$ 0.8	97.0 $\pm$ 0.9	97.0 $\pm$ 0.8	98.0 $\pm$ 0.7
Pancreas	79.3 $\pm$ 11.7	81.4 $\pm$ 10.8	84.3 $\pm$ 4.9	86.2 $\pm$ 4.5	86.6 $\pm$ 4.3	87.8 $\pm$ 3.1
Small bowel	69.9 $\pm$ 17.3	71.1 $\pm$ 15.0	76.9 $\pm$ 14.0	78.0 $\pm$ 13.8	79.0 $\pm$ 13.4	80.1 $\pm$ 10.2
Spleen	89.6 $\pm$ 9.5	93.1 $\pm$ 2.1	96.3 $\pm$ 1.9	96.4 $\pm$ 1.9	96.4 $\pm$ 1.7	97.1 $\pm$ 1.5
Stomach	90.1 $\pm$ 7.2	93.2 $\pm$ 5.4	93.9 $\pm$ 3.2	94.2 $\pm$ 2.9	94.2 $\pm$ 3.0	95.2 $\pm$ 2.6
Veins	60.7 $\pm$ 23.7	74.5 $\pm$ 10.5	74.8 $\pm$ 10.7	76.8 $\pm$ 11.2	77.4 $\pm$ 12.1	80.7 $\pm$ 9.3

The highest DSC for each structure is highlighted as bold.

Table 2Average surface distances of thirteen segmented organs for all 236 cases (mean $\pm$ standard deviation of average surface distances in mm).

Structure	3D U-net	HFCN	FCN MV	OAN MV	OAN-RC MV	OAN-RC LSSF
Aorta	0.44 $\pm$ 1.01	0.42 $\pm$ 0.58	0.56 $\pm$ 0.47	0.47 $\pm$ 0.42	0.44 $\pm$ 0.28	0.39 $\pm$ 0.21
Colon	6.75 $\pm$ 9.01	6.35 $\pm$ 8.12	6.27 $\pm$ 7.44	5.65 $\pm$ 7.25	4.07 $\pm$ 5.72	3.59 $\pm$ 4.17
Duodenum	2.01 $\pm$ 2.46	1.70 $\pm$ 2.18	1.71 $\pm$ 2.25	1.49 $\pm$ 1.87	1.54 $\pm$ 1.43	1.36 $\pm$ 1.31
Gallbladder	1.31 $\pm$ 0.76	1.21 $\pm$ 0.50	1.22 $\pm$ 0.52	1.12 $\pm$ 0.50	1.05 $\pm$ 0.41	0.95 $\pm$ 0.37
IVC	1.57 $\pm$ 1.53	1.15 $\pm$ 1.05	1.26 $\pm$ 1.08	1.16 $\pm$ 1.38	1.12 $\pm$ 1.24	1.08 $\pm$ 1.03
Kidney(L)	0.77 $\pm$ 1.04	0.41 $\pm$ 0.42	0.36 $\pm$ 0.47	0.34 $\pm$ 0.47	0.30 $\pm$ 0.33	0.30 $\pm$ 0.30
Kidney(R)	1.39 $\pm$ 2.01	1.03 $\pm$ 1.68	1.05 $\pm$ 1.74	0.74 $\pm$ 1.32	0.54 $\pm$ 1.09	0.45 $\pm$ 0.89
Liver	1.89 $\pm$ 3.21	1.60 $\pm$ 0	1.61 $\pm$ 2.98	1.39 $\pm$ 2.64	1.32 $\pm$ 1.74	1.23 $\pm$ 1.52
Pancreas	1.78 $\pm$ 1.05	1.51 $\pm$ 0.80	1.41 $\pm$ 0.88	1.19 $\pm$ 0.82	1.17 $\pm$ 0.72	1.05 $\pm$ 0.65
Small bowel	4.21 $\pm$ 5.78	4.01 $\pm$ 6.01	3.91 $\pm$ 6.05	3.20 $\pm$ 4.05	3.37 $\pm$ 5.48	3.01 $\pm$ 3.35
Spleen	0.98 $\pm$ 0.56	0.59 $\pm$ 0.37	0.60 $\pm$ 0.36	0.56 $\pm$ 0.40	0.47 $\pm$ 0.27	0.42 $\pm$ 0.25
Stomach	2.78 $\pm$ 5.89	2.50 $\pm$ 5.02	2.51 $\pm$ 5.13	2.36 $\pm$ 5.65	1.88 $\pm$ 1.64	1.68 $\pm$ 1.55
Veins	2.31 $\pm$ 4.51	1.75 $\pm$ 3.51	1.69 $\pm$ 3.61	1.92 $\pm$ 6.48	1.40 $\pm$ 3.61	1.21 $\pm$ 3.05

The highest DSC for each structure is highlighted as bold.

the same size of ($9 \times 9 \times 9$) 3D patches and (9×9) 2D patches from all directions can be used for all cases in our experiments.

The final segmentation results using OAN-RC with local structural similarity-based statistical fusion (LSSF) were compared with the 3D-patch based state-of-the-art approaches, 3D U-net (Çiçek et al., 2016) and hierarchical 3D FCN (HFCN) (Roth et al., 2018) as well as 2D-based FCN, OAN and OAN-RC with majority voting (MV). For a quantitative comparison, we computed the well-known Dice–Sørensen similarity coefficient (DSC) and the surface distances based on the manual annotations as ground-truth. For a structure s , DSC is computed as $\frac{2V(S \cap T)}{V(S) + V(T)}$ where S is the estimated segmentation and T is the ground-truth, i.e. manual annotations in this study. The surface distance was computed from each vertex of the ground-truth and to the estimates of our algorithms. Fig. 8 shows comparison results by box plots, while Tables 1 and 2 represent the mean and standard deviations for all the 236 cases.

As shown in Fig. 8, the basic OAN-RC outperforms other state-of-the-art approaches and our local structural similarity-based fusion improves the results even more. We note that although DSC shows the relative overall volume similarity, it does not quantify the boundary smoothness or the boundary noise of the results. But evaluating the surface distances, see below, shows that our method works effectively for both the whole volumes and the boundaries of the organs.

Tables 1 and 2 represent the mean and standard deviations of performance measures for 13 critical organs. Similar to the box plots, they show that our OAN-RCs with statistical fusion improves the overall mean performance and also reduces the standard deviations significantly. Fig. 9 shows an example generated by our proposed OAN-RC with LSSF, which is visually in-

distinguishable from manual segmentation for almost all target structures.

The OAN-RC training and testing can be computed in parallel for each view direction. In our experiments, the training took 40 hours for 120,000 iterations for 177 training cases and the average testing time for each volume was 76.73 s. The fusion time depended on the volume of the target structure, and the average computation time for 13 organs was 6.87s.

5. Discussion

Multi-organ segmentation using OAN-RCs alone, without the statistical fusion, gave similar or better performance compared with the state-of-the-art approaches as shown in Tables 1 and 2. Comparisons to non-deep network based algorithms can be found in Karasawa et al. (2017). Table 3 shows a statistical test results in terms of p -values for difference between different networks structures, denoted as Φ , and our OAN-RC, represented as Φ^* . A ranking test introduced by a difference zone $\delta \in \mathbb{R}$ is proper to test $DSC_{\Phi^*} > DSC_{\Phi}$, which equals $d_{\Phi} = DSC_{\Phi^*} - DSC_{\Phi} > \delta (\delta > 0)$. We tested the null hypothesis as $d = \delta$ versus the alternative hypothesis $d > \delta$. Based on $\delta = 0.1\%$, the test result shows $p < 0.001$ in most of the cases which explains that OAN-RC is superior to state-of-the-art methods including our OAN without reverse connection in terms of DSC. Among target organs, our performance on structures such as gallbladder and pancreas, whose sizes are relatively small and have particularly weak boundaries improves significantly from using basic FCNs or using OANs without reverse connections.

Moreover, as shown in Section 4, our statistical fusion based on local structural similarity improves the overall segmentation accuracies in terms of both DSC and average surface distances. In

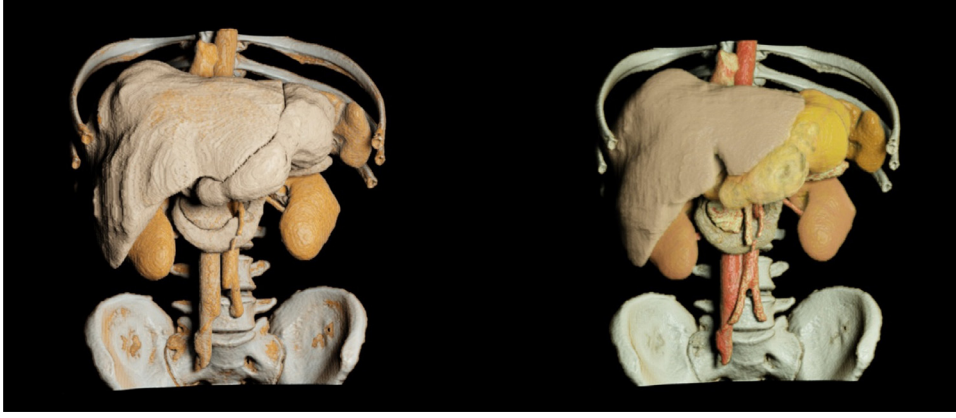


Fig. 9. 3D photo-realistic volume rendering (Kores et al., 2012) of the ground-truth (left) and the results from OAN-RC with statistical fusion (right). The aorta, duodenum, IVC, liver, kidneys, pancreas, duodenum, spleen, and stomach are rendered. The difference between our results and the ground-truth are almost visually indistinguishable. The rib and spine bones are also rendered as a reference frame to show the spatial location and relationship of the internal organs for the qualitative comparison. To differentiate adjacent organs and from manual segmentation, different color setting were applied to the our methods results. (For interpretation of the references to colour in this figure legend, the reader is referred to the web version of this article.)

Table 3

Statistical difference test between OAN-RC with majority voting (MV) and the following approaches. p -values are reported. The bold values represent the test result of $-d_{\Phi} = \text{DSC}_{\Phi} - \text{DSC}_{\Phi^*} = 0$ for the cases with $\text{DSC}_{\Phi} > \text{DSC}_{\Phi^*}$.

Structure	3D U-net	HFCN	FCN MV	OAN-RC MV(Φ^*)
Aorta	0.1825	0.0139	0.1792	0.1061
Colon	< 0.0001	< 0.0001	< 0.0001	0.0528
Duodenum	< 0.0001	< 0.0001	< 0.0001	< 0.0001
Gallbladder	< 0.0001	0.0002	0.0002	0.0027
IVC	< 0.0001	0.0001	< 0.0001	0.0002
Kidney (L)	< 0.0001	0.2045	0.4967	0.0025
Kidney (R)	< 0.0001	0.2143	0.3134	0.4255
Liver	< 0.0001	< 0.0001	0.0002	0.174
Pancreas	< 0.0001	< 0.0001	< 0.0001	0.0381
Small bowel	< 0.0001	< 0.0001	0.0024	0.0001
Spleen	< 0.0001	< 0.0001	0.4107	0.3974
Stomach	< 0.0001	0.0002	0.0031	0.3027
Veins	< 0.0001	< 0.0001	< 0.0001	< 0.0001

particular, there are significant performance improvements for the minimum values as shown in Fig. 8, which helps explain the robustness of the algorithm.

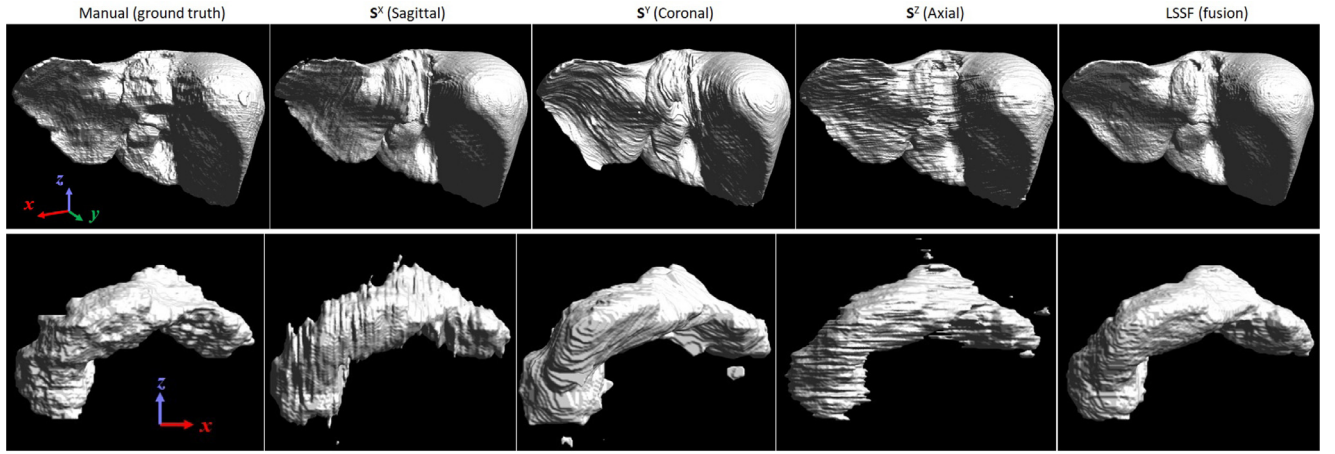
The differences can be depicted more clearly by visualizing the 3D surfaces as shown in Figs. 10 and 11. The noise of the deep network segmentations is distributed over large regions, without much connectivity, and occasionally they show significantly different patterns. But our fusion step exploits structural similarity which outputs clean and smooth boundaries by effectively combining different information based on the local structure of the original 3D volume.

Table 4 summarizes the number of network parameters used in the experiment and average computation time in seconds per a test case. Note that the time of FCN, OAN, and OAN-RC is for all three directions, which can be parallelized into 3 threads. Table 5 shows three types of additional experiments between the FCN and a wider FCN, FCN LSSF, and comparison of three fusion approaches including MV, STAPLE, and LSSF. The number of parameters of the wider FCN was 269M, which is almost double of the number of parameters of the FCN and similar to the OAN's. However, as shown in Table 5, the performance was similar to or slightly lower than the FCN. This shows our architecture is effectively designed for multi-organ segmentation not only by increasing number of parameters.

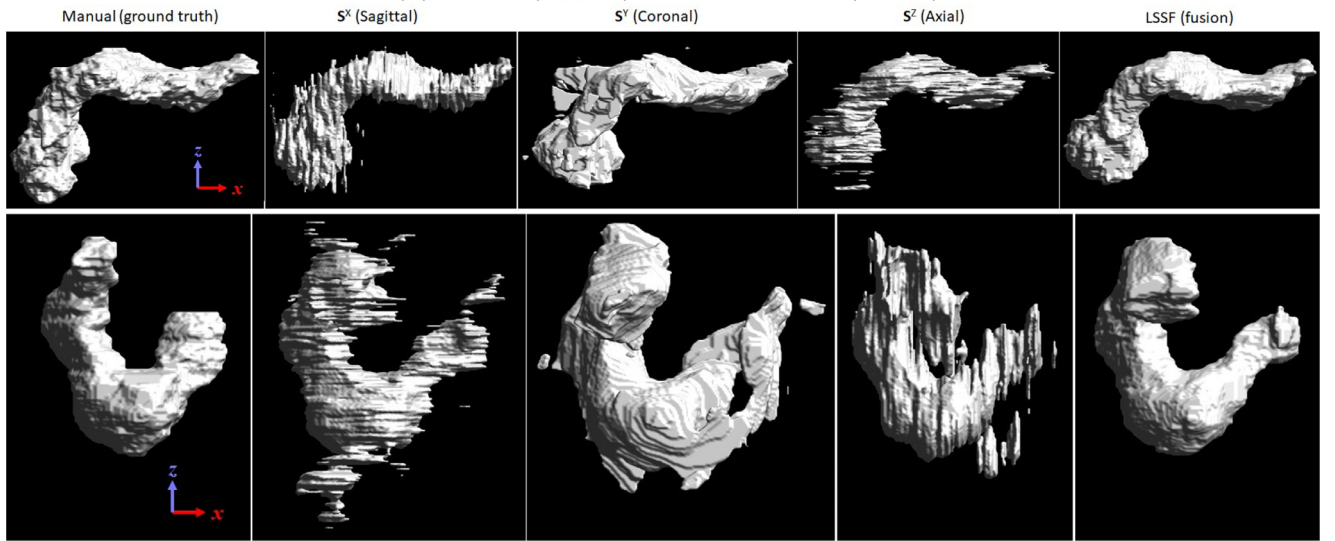
When applying our proposed method and interpreting the evaluation results, we must address several considerations. As shown in our experiments, our proposed algorithm also outperforms 3D patch based approaches. But 3D (isotropic) patch-based approaches have several issues which make it hard to apply to this problem. To make bigger patch size, they require more parameters and hence require more training data or, if this is not available, significant data augmentation (e.g., by scaling, rotation, and elastic deformation). In addition, there can be practical memory limitation on GPUs which restricts the expandable patch size. The limited patch size means that the deep networks receptive field sizes contains only limited local information which is problematic for multi-organ segmentation and the discontinuities between the patches also raises problems. It is possible that solutions to these three problems may make 3D patch based methods work better in the future. Unlike 3D approaches, the local structure-similarity used in our fusion method effectively combine the information from anisotropic patches to 3D at each voxel.

The ground-truth used in this study for training and evaluation was specified using manual annotations by human observers. It is well known that there can be significant inter-/intra-observer variations in manual segmentation. But, as explained before, the ground-truth was created by four human observers and checked by experts in a visual way, and we randomly divided testing groups in our 4-fold cross-validation to avoid biased comparison. However, it is still possible that inaccuracies due to human variability may affect the evaluation as well as the training. This can be further intensively explored as separate experiments.

Another possible consideration when applying the proposed approach is the image quality which can affect both of manual annotations and deep network segmentation results. Various factors such as spatial resolution, level of artifacts and reconstruction kernels should be considered. The dataset used in this study has been collected in the same institute with control over the scanning parameters. CT images used in this study is scanned from 2005 to 2009 from adult kidney donors to guarantee that the cases are pathologically proven as normal. Two different types of machines were used in this study and 4 different parameters of reconstruction kernels were used. As explained in Section 4, the CT protocol is the portal venous phase and the spatial resolution is almost isotropic. But different scanning parameters and artifacts may affect our algorithms performance when applied to other datasets. Proper data collection with various scanning parameters and with



(a) Liver (upper) and Pancreas (lower)



(b) Pancreas (Upper) and Duodenum (lower)

Fig. 10. Effects of local structural similarity-based statistical fusion (LSSF) for estimating 3D surfaces. From left to right, the manual segmentation (ground-truth), initial segmentations from OAN-RCs with X, Y, Z slices, and the results of our proposed algorithm with statistical fusion. (a) When S^X , S^Y , and S^Z show similar result, statistical fusion produces smoother and less-noisy boundaries. (b) Surface estimation examples when initial OAN-RCs give differing results. But our approach effectively fuses the information, exploiting the local structural similarity.

Table 4
Number of parameters between different networks.

Structure	3D U-net	HFCN	FCN	OAN	OAN-RC
No. parameters	19M	38M	134M	269M	313M
Average computation time(s)/case	1012	1105	146	213	230

enough distribution of samples including multi-cohorts is another issue to apply this approach to clinical environment.

The same issues about manual segmentations and image qualities can be raised in general segmentation and evaluations. Specifically for our proposed approach, especially in the fusion step, the way of computing *priori*, $P(T)$, used in Eq. (21) can in practice affect the final segmentation. But considering that the deep network segmentation results from different viewing-directions are independently obtained, the mean can be accepted in general. However, if the deep network segmentations show clear tendencies toward over-estimation or under-estimation, then different types of mod-

els for *priors* may need to be used in order to improve the final result for practical applications.

One of the main advantages of our algorithm is the efficient computation time. Table 4 shows the comparison of number of parameters and computation time between different deep network architectures. Considering the computation for three directions can be parallelized, the segmentation of 13 organs of the whole volume takes similar to or less than 1.5 minute with better performance reported than the state-of-the-art methods (Karasawa et al., 2017). Hence our approach can be practically useful in clinical environments.

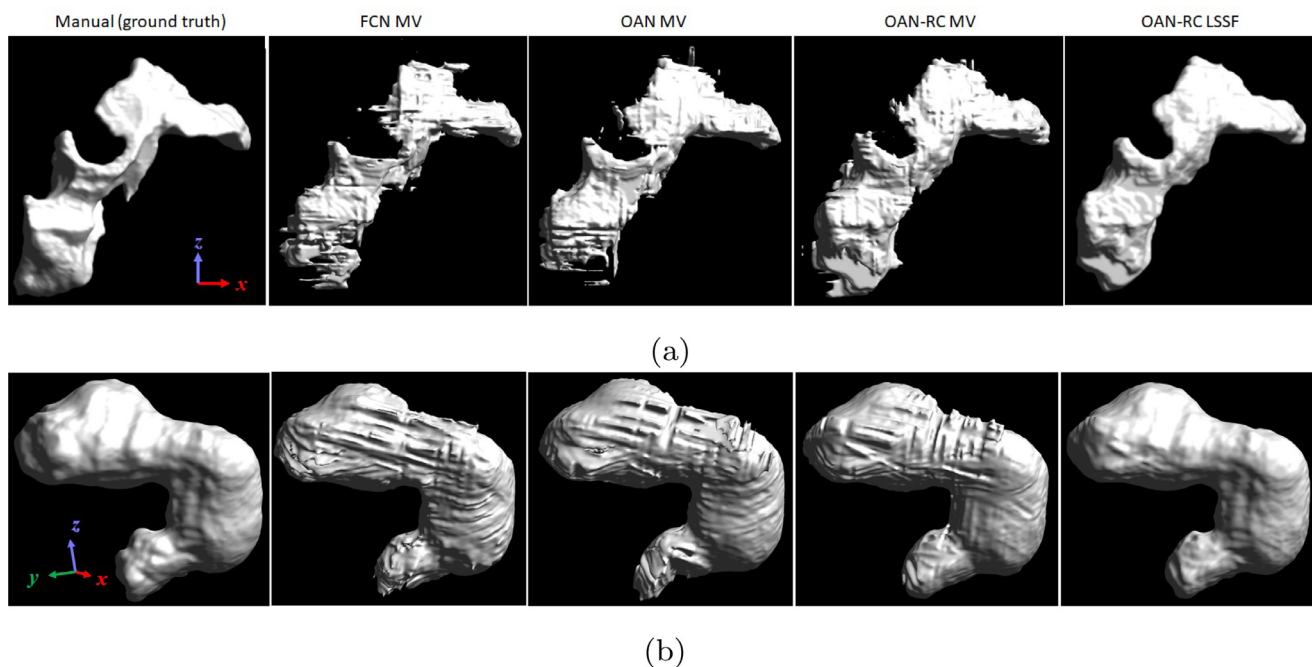


Fig. 11. Examples of FCN, OAN, OAN-RC, and OAN-RC LSSF. The manual segmentation (ground-truth), FCN MV, OAN MV, OAN-RC MV, OAN-RC LSSF (from left to right). (a) Pancreas: DSC(%) and surface distances (mean $\pm$ standard deviation in mm) to the ground-truth are 72.5 and 2.13 ± 1.74 (FCN MV), 77.2 and 1.90 ± 1.77 (OAN MV), 82.4 and 1.33 ± 1.31 (OAN-RC MV), and 85.5 and 0.71 ± 0.81 (OAN-RC LSSF), respectively. (b) Stomach: DSC(%) and surface distances (mean $\pm$ standard deviation in mm) to the ground-truth are 92.5 and 2.44 ± 1.27 (FCN MV), 93.6 and 1.63 ± 1.14 (OAN MV), 94.9 and 2.25 ± 1.30 (OAN-RC MV), and 97.1 and 1.26 ± 0.88 (OAN-RC LSSF), respectively.

Table 5

Comparison of different FCN variants and fusion algorithms. DSC for thirteen segmented organs are presented.

Structure	FCNs and fusion			OAN-RC fusion		
	FCN MV	Wider FCN MV	FCN LSSF	OAN-RC MV	OAN-RC STAPLE	OAN-RC LSSF
Aorta	85.0 $\pm$ 4.2	85.4 $\pm$ 4.4	90.9 $\pm$ 3.5	85.3 $\pm$ 4.1	88.2 $\pm$ 3.8	91.8 $\pm$ 3.5
Colon	80.3 $\pm$ 9.1	80.0 $\pm$ 9.5	81.5 $\pm$ 8.5	82.0 $\pm$ 8.8	82.2 $\pm$ 8.2	83.0 $\pm$ 7.4
Duodenum	70.2 $\pm$ 11.3	69.3 $\pm$ 11.7	72.6 $\pm$ 9.9	73.4 $\pm$ 11.1	74.0 $\pm$ 10.5	75.4 $\pm$ 9.1
Gallbladder	87.8 $\pm$ 8.3	87.2 $\pm$ 9.6	88.4 $\pm$ 7.3	89.4 $\pm$ 6.1	89.6 $\pm$ 6.0	90.5 $\pm$ 5.3
IVC	84.0 $\pm$ 6.0	83.3 $\pm$ 6.3	85.1 $\pm$ 5.4	86.0 $\pm$ 5.5	86.2 $\pm$ 5.0	87.0 $\pm$ 4.2
Kidney(L)	96.1 $\pm$ 2.0	95.9 $\pm$ 1.9	96.9 $\pm$ 1.9	95.9 $\pm$ 2.3	96.2 $\pm$ 2.2	96.8 $\pm$ 1.9
Kidney(R)	95.8 $\pm$ 4.9	95.9 $\pm$ 2.4	97.9 $\pm$ 0.8	96.0 $\pm$ 2.5	97.9 $\pm$ 2.3	98.4 $\pm$ 2.1
Liver	96.8 $\pm$ 0.8	96.8 $\pm$ 0.9	97.8 $\pm$ 0.8	97.0 $\pm$ 0.8	97.8 $\pm$ 0.7	98.0 $\pm$ 0.7
Pancreas	84.3 $\pm$ 4.9	83.7 $\pm$ 5.1	85.1 $\pm$ 3.5	86.6 $\pm$ 4.3	87.0 $\pm$ 4.0	87.8 $\pm$ 3.1
Small bowel	76.9 $\pm$ 14.0	75.9 $\pm$ 15.6	77.8 $\pm$ 12.5	79.0 $\pm$ 13.4	79.2 $\pm$ 12.4	80.1 $\pm$ 10.2
Spleen	96.3 $\pm$ 1.9	96.1 $\pm$ 2.8	97.0 $\pm$ 1.6	96.4 $\pm$ 1.7	96.9 $\pm$ 1.5	97.1 $\pm$ 1.5
Stomach	93.9 $\pm$ 3.2	93.6 $\pm$ 3.6	94.5 $\pm$ 2.6	94.2 $\pm$ 3.0	94.7 $\pm$ 2.9	95.2 $\pm$ 2.6
Veins	74.8 $\pm$ 10.7	72.0 $\pm$ 17.8	77.1 $\pm$ 10.1	77.4 $\pm$ 12.1	78.9 $\pm$ 10.8	80.7 $\pm$ 9.3

The highest DSC for each structure is highlighted as bold.

6. Conclusion

In this paper, we proposed a novel framework for multi-organ segmentation using OAN-RCs with statistical fusion exploiting structural similarity. Our two-stage organ-attention network reduces uncertainties at weak boundaries, focuses attention on organ regions with simple context, and adjusts FCN error by training the combination of original images and OAMs. Reverse connections deliver abstract level semantic information to lower layers so that hidden layers can be assisted to contain more semantic information and give good results even for small organs. The results are improved by the statistical fusion, based on local structural similarity, which smooths our noise and removes biases leading to better overall segmentation performance in terms of DSC and surface distances. We showed that our performance is better than previous state of the art algorithms. Our framework is not specific to any

particular body region, but gives high quality and robust results for abdominal CTs, which are typically challenging regions due to their low contrast, large intra-/inter-variations, and different scales. In addition, the efficient computational time of our algorithm makes our approach practical for clinical environments such as CAD, CAS or RT.

Conflict of interest

The authors have no conflicts of interest.

Supplementary material

Supplementary material associated with this article can be found, in the online version, at doi:[10.1016/j.media.2019.04.005](https://doi.org/10.1016/j.media.2019.04.005).

Appendix A. The full architecture of the 2-stage network

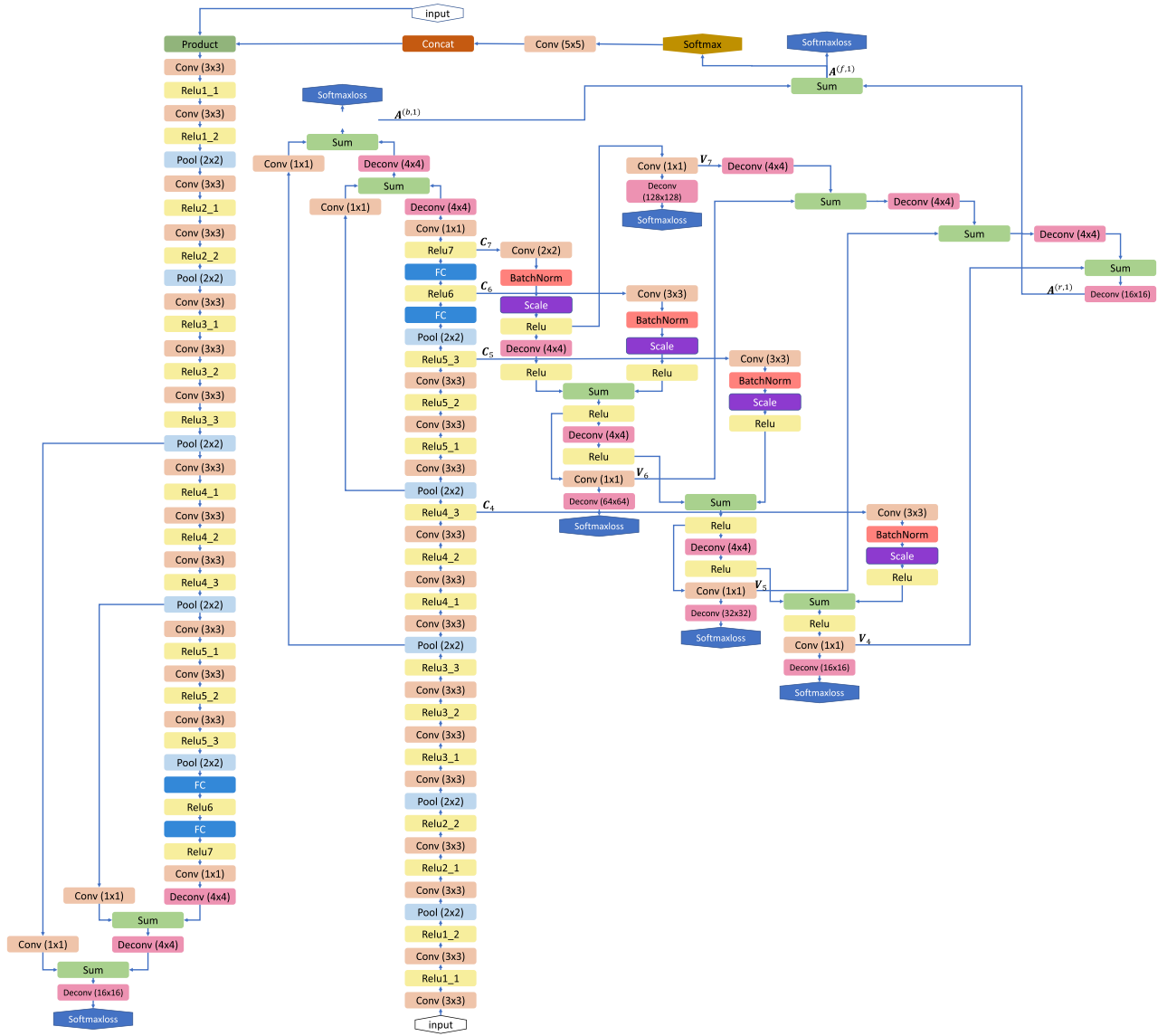


Fig. A1. The full operation chart of the 2-stage OAN-RC.

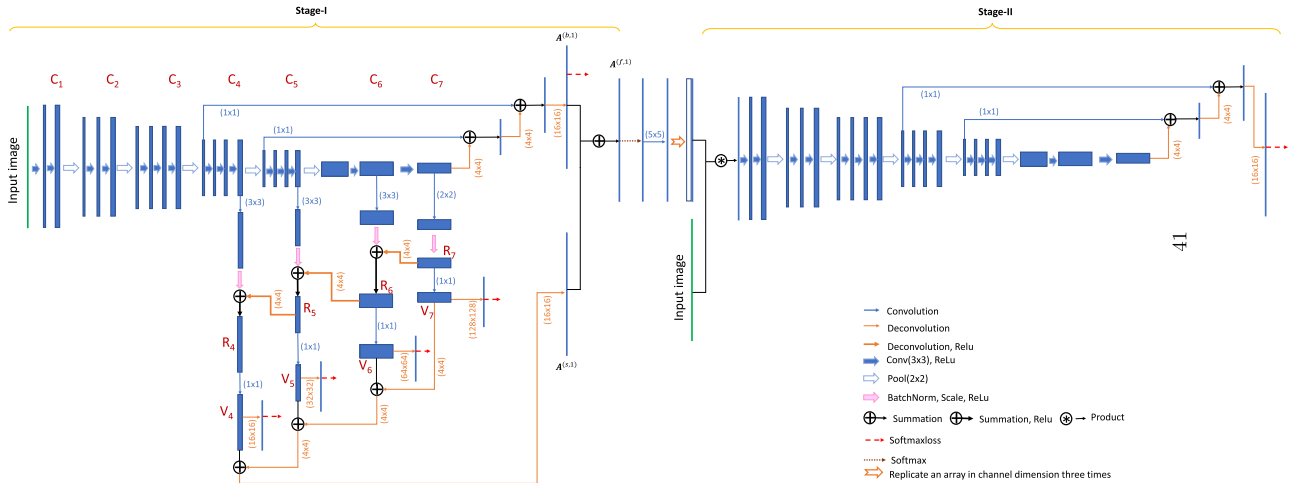


Fig. A2. The full architecture of the 2-stage OAN-RC.

References

- Asman, A.J., Landman, B.A., 2012. Formulating spatially varying performance in the statistical fusion framework. *IEEE Trans. Med. Imaging* 31 (6), 1326–1336. doi:[10.1109/TMI.2012.2190992](#).
- Asman, A.J., Landman, B.A., 2013. Non-local statistical label fusion for multi-atlas segmentation. *Med. Image Anal.* 17 (2), 194–208. doi:[10.1016/j.media.2012.10.002](#).
- Boykov, Y., Veksler, O., Zabih, R., 2001. Fast approximate energy minimization via graph cuts. *IEEE Trans. Pattern Anal. Mach. Intell.* 23 (11), 1222–1239. doi:[10.1109/34.969114](#).
- Chen, H., Dou, Q., Yu, L., Qin, J., Heng, P.-A., 2017a. Voxresnet: deep voxelwise residual networks for brain segmentation from 3D MR images. *Neuroimage* doi:[10.1016/j.neuroimage.2017.04.041](#).
- Chen, L., Papandreou, G., Kokkinos, I., Murphy, K., Yuille, A.L., 2017b. Deeplab: semantic image segmentation with deep convolutional nets, atrous convolution, and fully connected crfs. *IEEE Trans. Pattern Anal. Mach. Intell.* 40, 834–848.
- Chen, L., Yang, Y., Wang, J., Xu, W., Yuille, A.L., 2016. Attention to scale: Scale-aware semantic image segmentation. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 3640–3649.
- Chu, C., Oda, M., Kitasaka, T., Misawa, K., Fujiwara, M., Hayashi, Y., Nimura, Y., Rueckert, D., Mori, K., 2013. Multi-organ segmentation based on spatially-divided probabilistic atlas from 3D abdominal CT images. *Lect. Notes Comput. Sci.* 8150 LNCS (part 2), 165–172. doi:[10.1007/978-3-642-40763-5_21](#).
- Çiçek, Ö., Abdulkadir, A., Lienkamp, S.S., Brox, T., Ronneberger, O., 2016. 3D U-Net: learning dense volumetric segmentation from sparse annotation. *Lect. Notes Comput. Sci.* 9901, 424–432. doi:[10.1007/978-3-319-46723-8_49](#).
- Dou, Q., Chen, H., Jin, Y., Yu, L., Qin, J., Heng, P.A., 2016. 3D deeply supervised network for automatic liver segmentation from CT volumes. *Lect. Notes Comput. Sci.* 9901 LNCS, 149–157. doi:[10.1007/978-3-319-46723-8_18](#).
- Everingham, M., Ali Eslami, S., Van Gool, L., Williams, C.K.I., Winn, J., Zisserman, A., 2015. The PASCAL visual object classes challenge: A retrospective. *Int. J. Comput. Vis.* 111, 98–136.
- Havaei, M., Davy, A., Warde-Farley, D., Biard, A., Courville, A., Bengio, Y., Pal, C., Jodoin, P.-M., Larochelle, H., 2017. Brain tumor segmentation with deep neural networks. *Med. Image Anal.* 35, 18–31. doi:[10.1016/j.media.2016.05.004](#).
- Heimann, T., van Ginneken, B., Styner, M., Arzhaeva, Y., Aurich, V., Bauer, C., Beck, A., Becker, C., Beichel, R., Bekes, G., Bello, F., Binnig, G., Bischof, H., Bornik, A., Cashman, P., Chi, Y., Cordova, A., Dawant, B., Fidrich, M., Furst, J., Furukawa, D., Grenacher, L., Hornegger, J., Kainmiller, D., Kitney, R., Kobatake, H., Lamecker, H., Lange, T., Lee, J., Lennon, B., Li, R., Li, S., Meinzer, H., Nemeth, G., Raicu, D., Rau, A., van Rikxoort, E., Rousson, M., Rusko, L., Saddi, K., Schmidt, G., Seghers, D., Shimizu, A., Slagmolen, P., Sorantin, E., Soza, G., Susomboon, R., Waite, J., Wimmer, A., Wolf, I., 2009. Comparison and evaluation of methods for liver segmentation from CT datasets. *IEEE Trans. Med. Imaging* 28, 1251–1265. doi:[10.1109/TMI.2009.2013851](#).
- Iglesias, J.E., Sabuncu, M.R., 2015. Multi-atlas segmentation of biomedical images: a survey. *Med. Image Anal.* 24 (1), 205–219. doi:[10.1016/j.media.2015.06.012](#).
- Jamnitsas, K., Ledig, C., Newcombe, V.F., Simpson, J.P., Kane, A.D., Menon, D.K., Rueckert, D., Glocker, B., 2017. Efficient multi-scale 3D CNN with fully connected CRF for accurate brain lesion segmentation. *Med. Image Anal.* 36, 61–78. doi:[10.1016/j.media.2016.10.004](#).
- Kada, T., Linguraru, M.G., Hori, M., Summers, R.M., Tomiyama, N., Sato, Y., 2015. Abdominal multi-organ segmentation from CT images using conditional shape-location and unsupervised intensity priors. *Med. Image Anal.* 26 (1), 1–18. doi:[10.1016/j.media.2015.06.009](#).
- Karasawa, K., Oda, M., Kitasaka, T., Misawa, K., Fujiwara, M., Chu, C., Zheng, G., Rueckert, D., Mori, K., 2017. Multi-atlas pancreas segmentation: atlas selection based on vessel structure. *Med. Image Anal.* 39, 18–28. doi:[10.1016/j.media.2017.03.006](#).
- Kirbas, C., Quek, F., 2004. A review of vessel extraction techniques and algorithms. *ACM Comput. Surv.* 36, 81–121.
- Kong, T., Sun, F., Yao, A., Liu, H., Lu, M., Chen, Y., 2017. RON: reverse connection with objectness prior networks for object detection. In: *IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*.
- Kores, T., Post, F.H., Botha, C.P., 2012. Exposure render: an interactive photo-realistic volume rendering framework. *PLoS One* 7, e38586:1–10.
- Lesage, D., Angelini, E.D., Bloch, I., Funka-Lea, G., 2009. A review of 3d vessel lumen segmentation techniques: models, features and extraction schemes. *Med. Image Anal.* 13, 819–845. doi:[10.1016/j.media.2009.07.011](#).
- Li, G., Chen, X., Shi, F., Zhu, W., Tian, J., Xiang, D., 2015. Automatic liver segmentation based on shape constraints and deformable graph cut in CT images. *IEEE Trans. Image Process.* 24, 5315–5329.
- Long, J., Shelhamer, E., Darrell, T., 2015. Fully convolutional networks for semantic segmentation. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 3431–3440.
- Lugo-Fagundo, C., Vogelstein, B., Yuille, A., Fishman, E.K., 2018. Deep learning in radiology: now the real work begins. *J. Am. Coll. Radiol.* 15, 364–367.
- Mharib, A.M., Ramli, A.R., Mashohor, S., Mahmood, R.B., 2012. Survey on liver CT image segmentation methods. *Artif. Intell. Rev.* 37. doi:[10.1007/s10462-011-9220-3](#).
- Milletari, F., Navab, N., Ahmadi, S.A., 2016. V-Net: fully convolutional neural networks for volumetric medical image segmentation. In: *Proceedings - 2016 4th International Conference on 3D Vision, 3DV 2016*, pp. 565–571. doi:[10.1109/3DV.2016.79](#).
- Nascimento, J., Carneiro, G., 2016. Multi-atlas segmentation using manifold learning with deep belief networks. In: *Proceedings - International Symposium on Biomedical Imaging*, June doi:[10.1109/ISBI.2016.7493403](#).
- Otkay, O., Schlemper, J., Folgoc, L. L., Lee, M., Heinrich, M., Misawa, K., Mori, K., McDonagh, S., Hammerla, N. Y., Kainz, B., Glocker, B., Rueckert, D., 2018. Attention U-net: learning where to look for the pancreas. *arXiv:1804.03999*.
- Rajchl, M., Lee, M.C., Oktay, O., Kamnitsas, K., Passerat-Palmbach, J., Bai, W., Damodaram, M., Rutherford, M.A., Hajnal, J.V., Kainz, B., Rueckert, D., 2017. Deepcut: object segmentation from bounding box annotations using convolutional neural networks. *IEEE Trans. Med. Imaging* 36 (2), 674–683. doi:[10.1109/TMI.2016.2621185](#).
- Roth, H., Oda, M., Shimizu, N., Oda, H., Hayashi, Y., Kitasaka, T., Fujiwara, M., Misawa, K., Mori, K., 2018. Towards dense volumetric pancreas segmentation in ct using 3D fully convolutional networks. In: *Proceedings of SPIE Medical Imaging 2018*, p. 10574.
- Roth, H.R., Farag, A., Lu, L., Turkbey, E.B., Summers, R.M., 2015. Deep convolutional networks for pancreas segmentation in CT imaging. In: *Proceedings of SPIE Medical Imaging 2015*, p. 9413. doi:[10.1117/12.2081420](#).
- Roth, H.R., Lu, L., Farag, A., Sohn, A., Summers, R.M., 2016. Spatial aggregation of holistically-nested networks for automated pancreas segmentation. In: *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 9901, pp. 451–459. doi:[10.1007/978-3-319-46723-8_52](#).
- Roth, H.R., Oda, H., Zhou, X., Shimizu, N., Yang, Y., Hayashi, Y., Oda, M., Fujiwara, M., Misawa, K., Mori, K., 2018. An application of cascaded 3D fully convolutional networks for medical image segmentation. *Comput. Med. Imaging Graph.* 66, 90–99.
- Sabuncu, M., Yeo, B., van Leemput, K., Fische, B., Golland, P., 2010. A generative model for image segmentation based on label fusion. *IEEE Trans. Med. Imaging* 29, 1714–1729.
- Setio, A.A.A., Ciompi, F., Litjens, G., Gerke, P., Jacobs, C., Van Riel, S.J., Wille, M.M.W., Naqibullah, M., Sanchez, C.I., Van Ginneken, B., 2016. Pulmonary nodule detection in CT images: false positive reduction using multi-View convolutional networks. *IEEE Trans. Med. Imaging* 35 (5), 1160–1169. doi:[10.1109/TMI.2016.2536809](#).
- Shelhamer, E., Long, J., Darrell, T., 2017. Fully convolutional networks for semantic segmentation. *IEEE Trans. Pattern Anal. Mach. Intell.* 39 (4), 640–651.
- Wang, L., Lee, C.-Y., Tu, Z., Lazebnik, S., 2015. Training deeper convolutional networks with deep supervision. *arXiv:1505.02496*.
- Wang, Z., Bovik, A.C., Sheikh, H.R., Simoncelli, E.P., 2004. Image quality assessment: from error visibility to structural similarity. *IEEE Trans. Image Process.* 13 (4), 600–612. doi:[10.1109/TIP.2003.819861](#).
- Warfield, S.K., Zou, K.H., Wells, W.M., 2004. Simultaneous truth and performance level estimation (STAPLE): an algorithm for the validation of image segmentation. *IEEE Trans. Med. Imaging* 23 (7), 903–921. doi:[10.1109/TMI.2004.828354](#).
- Wolz, R., Chu, C., Misawa, K., Fujiwara, M., Mori, K., Rueckert, D., 2013. Automated abdominal multi-organ segmentation with subject-specific atlas generation. *IEEE Trans. Med. Imaging* 32 (9), 1723–1730. doi:[10.1109/TMI.2013.2265805](#).
- Xie, S., Tu, Z., 2015. Holistically-nested edge detection. In: *IEEE International Conference on Computer Vision (ICCV)*.
- Xu, K., Ba, J.L., Kiros, R., Cho, K., Courville, A., Salakhutdinov, R., Zemel, R.S., Bengio, Y., 2015. Show, attend and tell: Neural image caption generation with visual attention. In: *Proceedings of the 32nd International Conference on Machine Learning*, pp. 2048–2057.
- Zhou, Y., Xie, L., Shen, W., Wang, Y., Fishman, E.K., Yuille, A.L., 2017. A fixed-point model for pancreas segmentation in abdominal CT scans. *Medical Image Computing and Computer Assisted Intervention (MICCAI)*.
- Zhuang, X., Shen, J., 2016. Multi-scale patch and multi-modality atlases for whole heart segmentation of MRI. *Med. Image Anal.* 31, 77–87. doi:[10.1016/j.media.2016.02.006](#).
- Zu, C., Wang, Z., Zhang, D., Liang, P., Shi, Y., Shen, D., Wu, G., 2017. Robust multi-atlas label propagation by deep sparse representation. *Pattern Recognit.* 63, 511–517. doi:[10.1016/j.patcog.2016.09.028](#).