

JointVesselNet: Joint Volume-Projection Convolutional Embedding Networks for 3D Cerebrovascular Segmentation

Yifan Wang¹, Guoli Yan¹, Haikuan Zhu¹, Sagar Buch², Ying Wang², Ewart Mark Haacke², Jing Hua¹, and Zichun Zhong¹

¹ Department of Computer Science, Wayne State University, Detroit MI, USA
zichunzhong@wayne.edu

² Department of Radiology, Wayne State University, Detroit MI, USA

Abstract. In this paper, we present an end-to-end deep learning method, *JointVesselNet*, for robust extraction of 3D sparse vascular structure through embedding the image composition, generated by maximum intensity projection (MIP), into the 3D magnetic resonance angiography (MRA) volumetric image learning process to enhance the overall performance. The MIP embedding features can strengthen the local vessel signal and adapt to the geometric variability and scalability of vessels. Therefore, the proposed framework can better capture the small vessels and improve the vessel connectivity. To our knowledge, this is the first time that a deep learning framework is proposed to construct a joint convolutional embedding space, where the computed joint vessel probabilities from 2D projection and 3D volume can be integrated synergistically. Experimental results are evaluated and compared with the traditional 3D vessel segmentation methods and the state-of-the-art in deep learning, by using both public and real patient cerebrovascular image datasets.

Keywords: Deep neural network · 3D cerebrovascular segmentation · Maximum intensity projection (MIP) · Joint embedding

1 Introduction

A growing body of evidence in animal and human studies shows that micro cerebrovascular abnormalities are the source of many vascular diseases and neurologic disorders, e.g., hypertension, arteriosclerosis, cerebral amyloid angiopathy, diabetes [1–3]. There is a pressing need for better extracting and understanding of vascular abnormalities in vivo at the micro-level [4]. However, compared with traditional organ segmentations, segmenting the cerebrovascular structure from magnetic resonance angiography (MRA) is very challenging due to difficulties, such as complex geometry and topology variations as well as data sparseness (artifacts, noises, low signal-to-noise ratio).

In recent decades, there have been many automatic model-driven vessel segmentation approaches proposed, such as multiscale filtering [5], region growing techniques [6], active contours [7], statistical and shape models [8, 9], particle filtering [10], geometric flow [11], path tracing [12], level-set approach [13], etc. However, these approaches are

easily overwhelmed by tons of low-level handcrafted features and complicated manual parameter adjustments to overcome aforementioned difficulties and subject variations.

Recently, data-driven approaches have been proposed to robustly investigate the correlations between different objects / instances without relying on hard-coded metrics. In medical image processing, several deep learning-based methods have been proposed to extract vessels from 2D retinal images, such as DeepVessel [14], cross-modality learning approach [15], multi-level deep supervised networks [16], DNN-based method [17], unified convolutional neural network (CNN) and graph neural network (GNN) [18], etc. These methods can well perform 2D vessel segmentation tasks, but are far from success in 3D vessel scenarios. There are still very few dedicated deep learning architectures for 3D vessel segmentation. For instance, Uception [19] presents a network inspired by 3D U-Net [20] and Inception modules [21]. DeepVesselNet [22] and VesselNet [23] propose 2D orthogonal cross-hair filters in three planes on each voxel to obtain the 3D contextual information at a reduced computational burden and less memory. The major challenges of 3D cerebrovascular segmentation are complicated vessel geometry and topology variations, high sparseness of vessel data in a large-sized 3D volume, and the limited resource of 3D vasculature datasets. Existing methods are not specifically designed for solving these challenges.

To fill the gap in the high-fidelity 3D cerebrovascular segmentation, we present an end-to-end deep learning method, *JointVesselNet*. A multi-stream CNN framework is adopted to effectively learn the 3D volume and multislice composited 2D maximum intensity projection (MIP) [24] feature vectors respectively. It then explores the inter-dependencies between 3D and 2D embedded feature vectors in a joint volume-projection embedding space by backprojecting the 2D feature vectors, learned from MIP, into the 3D volume embedding space. MIP is a widely-used scientific method for visualizing and analyzing 3D vasculature structure in MRA diagnosis by domain scientists. It can enhance the local vessel signal by canceling out the random noises and adapt to geometric variability and scalability. Through the novel integrative learning of both 3D volume and 2D MIPs, the proposed framework can better capture the small vessels and improve the subtle vessel connectivity. To our knowledge, this is the first time that a deep learning framework is proposed to construct a joint convolutional embedding space, where the computed joint vessel probabilities from 2D projection and 3D volume can be integrated synergistically. The key motivation of the proposed network is to integrate the trustworthy auxiliary from learned 2D MIP features into the 3D volume segmentation network, instead of using more complicated networks empirically. Experimental results are evaluated and compared with the traditional 3D vessel segmentation methods and the state-of-the-art in deep learning by using both public and real patient cerebrovascular image datasets. The application of this accurate segmentation of sparse and complicated 3D vascular structure facilitated by our method demonstrates the potential in improving MRA diagnosis of vascular diseases.

2 Method

In this section, we introduce the components of the JointVesselNet model: dataset preparation and generation, network architecture, and loss function.

2.1 Dataset Preparation and Generation

In this work, we use two different real patient datasets to evaluate our proposed JointVesselNet method. The first dataset is the public TubeTK Toolkit MRA dataset from University of North Carolina at Chapel Hill ³, acquired by a Siemens Allegra head-only 3T MR system. There are 42 patient cases in the whole dataset, which have the manual-labeled vessel segmentation masks. The voxel spacing of the MRA images is $0.5 \times 0.5 \times 0.8 \text{ mm}^3$ with a volume size of $448 \times 448 \times 128$ voxels.

The second dataset is provided by domain experts, who are neurologists and radiologists under our collaboration. 11 healthy volunteers were scanned in midbrain regions with a dual echo susceptibility weighted imaging (SWI) sequence at four time points: the first was acquired pre-contrast and the remaining three were acquired post-contrast during a gradual increase in dose delivered over the time frame of 20 mins (final concentration = 4 mg / kg); with the imaging parameters: echo time (TE)1 / echo time (TE)2 / repetition time (TR) = 7.5 / 15 / 27 ms, bandwidth = 180 Hz / pxl, flip angle = 15° (pre-contrast and final post-contrast data) and 20° (first and second post-contrast data). The voxel spacing is $0.22 \times 0.22 \times 1 \text{ mm}^3$ with a volume size of $832 \times 1024 \times 48$ voxels. This protocol enables short and long TE magnitude data to produce MRA. Then, the MRA is calculated using a non-linear subtraction (MRAnls) method [25], which is employed for selective MRA enhancement utilizing the flow rephased and dephased images. Finally, the ground truth vessel labels are obtained by integrating an enhanced angiography map from the computed MRAnls to set a more reasonable threshold for the initial mask, followed by domain experts' post-manual labeling refinement using a developed cerebrovascular labeling and visualization tool.

2.2 Network Architecture

The proposed JointVesselNet mainly consists of a dual-stream component (i.e., a 3D volume segmentation stream and a 2D composited MIP segmentation stream) and the bi-directional operations between these two streams (i.e., 3D-to-2D projection and 2D-to-3D backprojection). The overall architecture is demonstrated in Fig. 1.

In this work, the two-stream segmentation component learns vessel feature vectors in 3D volume and corresponding multiple 2D MIPs (enhanced and dense depiction of 3D relationships via a 3D-to-2D projection computation) contexts, respectively. We use a 3D U-Net [20] as the 3D volume segmentation branch and a half 2D U-Net [26] (in terms of feature channel numbers) as the 2D composited MIP segmentation branch, respectively. Since U-Net-like networks are the most commonly-used and robust medical imaging segmentation neural networks across different data modalities for varying organ / tissue geometries, it is suitable for our work to justify the benefits from the 2D-to-3D backprojection and joint embedding of 3D volume and 2D composited MIP. A U-Net-like network is essentially a convolutional encoder-decoder network. In Fig. 1, the layer output feature channel numbers are denoted in the corresponding blocks and layer input spatial dimensions are shown in the horizontal levels of every block.

³ <https://public.kitware.com/Wiki/TubeTK/Data>

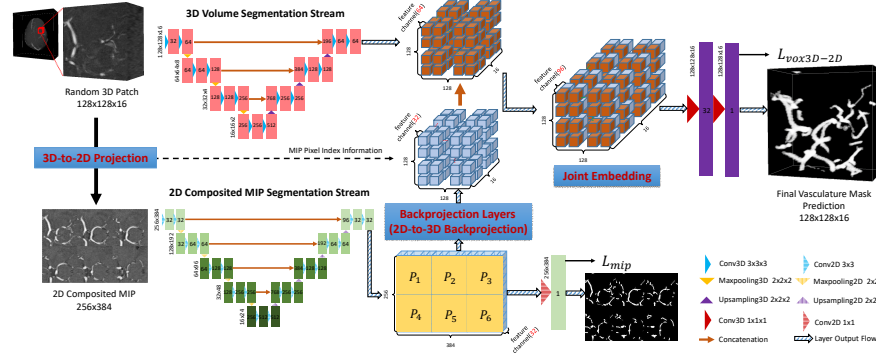


Fig. 1: The architecture of JointVesselNet. The major procedure includes composing MIPs via 3D-to-2D projection, dual-stream segmentation learning for 3D volume and 2D composited MIP feature vectors, mapping 2D composited MIP feature vectors back to 3D volume feature space via 2D-to-3D backprojection, building a joint embedding for learning the final vasculature mask.

Due to the limited data availability and the large volume size in cerebrovascular image datasets, we choose to train the network patch-wisely. Specifically, from the observation that most brain MRAs have much higher resolutions in axial plane than other planes, we adaptively train our network using none-cubic patches, which have larger dimension sizes across axial plane, instead of resizing the data into the uniform voxel spacing through an interpolation before the network training, to avoid potential segmentation inaccuracy. As shown in Fig. 1, the key step in our network is the effective integration of the features from two different streams / domains. Accordingly, two main challenges need to be overcome in this work. The first one is the effective format of the corresponding 2D composited MIPs from a randomly-extracted 3D volume patch that is suitable for simultaneous dual-stream learning design. The second one is the effective approach for backprojecting the feature vectors extracted from the composited MIP image plane pixels (in a dimension-reduced 2D space) back to the corresponding 3D volume spacing voxels. More details are introduced in the following.

3D-to-2D Projection in Dual-Stream Design. The major motivation for projecting the 3D volume space into the 2D MIP space is to enhance the local vessel probability (sparseness) as well as the signal-to-noise ratio. Given a randomly-extracted 3D volume patch V of the size $K_1 \times K_2 \times K_3$ (e.g., we use $128 \times 128 \times 16$ in our experiments) and K_3 along vertical axis, we compute s -sliced (e.g., $s = 5$ in our experiments as suggested by domain experts) MIPs of V along vertical axis with overlapping coverage every t slice interval. Consequently we can get a set of m consecutive MIPs, i.e., $\mathbf{P} = \{P_1, P_2, \dots, P_k, \dots, P_{m-1}, P_m\}$, in which P_k is the MIP across the $[(k-1)t+1]^{th}$ slice to the $[(k-1)t+s]^{th}$ slice in V . It is noted that in a 2D MIP, only one voxel with the maximum intensity among the s voxels along the vertical axis in V will be recorded, which is prone to an information loss, considering the segmentation

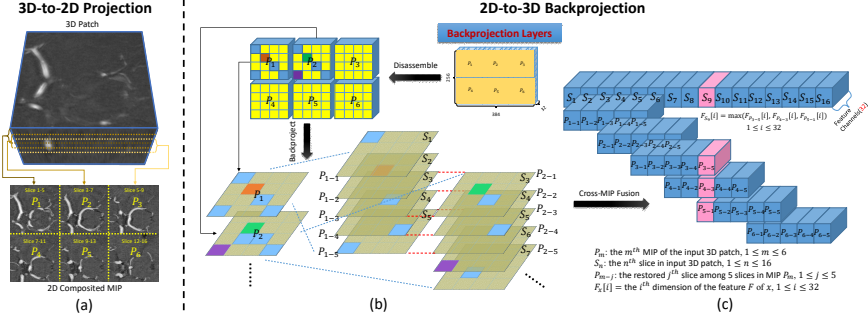


Fig. 2: (a) Illustration of the 3D-to-2D projection in the spatial domain for computing a 2D composited MIP from a 3D volume patch. (b) and (c) Illustration of the detailed 2D-to-3D restorations within backprojection layers in the embedded feature domain. Backprojected pixel-voxel feature pair examples can be traced via the pairing colors (e.g., green, orange, purple, and blue) in (b).

task needing information of every voxel. Consequently, we set $t = 2$ as a trade-off between computation cost and information completeness / denseness. We can get m MIPs of size $K_1 \times K_2$ for V , where the MIP number m is computed as:

$$m = \left\lfloor \frac{1}{t} (K_3 - s) \right\rfloor + 1. \quad (1)$$

A MIP conveys denser vessel information and is naturally suitable for 2D convolution. However, we now have m different MIPs and need to feed them to our network in the MIP stream, in company with the 3D volume stream V as an input pair to our entire network. The information from the m MIPs is equally important. In order to avoid stacking them to a $K_1 \times K_2 \times m$ volume such that the 2D CNN would essentially treat it as a 2D input of a spatial dimension $K_1 \times K_2$ with m different properties (feature channels), we convert the m MIPs to a tiled MIP with a larger 2D spatial size, such as $0.5mK_1 \times 2K_2$. In this case, the 2D convolution operates equally across the 2D composited MIP plane domain. The slice indices from where the MIP pixels are selected in the original V are also recorded so as to effectively restore the pixel-wise information extracted from MIP to the 3D volume space, which will be used in the 2D-to-3D backprojection transformation in the following process. The format of the 2D composited MIP (e.g., six consecutive MIPs) computed from a 3D patch is shown in Fig. 2 (a).

2D-to-3D Backprojection for Joint Embedding. Once the 3D volume and 2D MIP streams learn their segmentation features respectively, we intend to integrate them in a unified joint hidden feature embedding space to yield the final 3D segmentation prediction. In order to achieve this, we conduct several operations within our network to backproject the pixel features extracted from the composited MIP back to their corresponding 3D voxel feature space.

The final-stage hidden feature from 2D composited MIP segmentation branch has the size $0.5mK_1 \times 2K_2$ with C_1 channels ($C_1 = 32$ as shown in Fig. 1), which is

the input of the backprojection layers. We first disassemble it to restore m C_1 -channel features for the corresponding MIPs (e.g., $P_1, P_2, \dots, P_{m-1}, P_m$, where $m = 6$ as illustrated in Fig. 2 b). Then we use the recorded index information to map the MIP pixel features back to where they are selected from V during the 2D composited MIP generation. Fig. 2 (b) shows how the feature vectors of two consecutive MIPs (P_1 and P_2) are disassembled from the composited MIP. They project their pixel feature space (P_{m-j} , i.e., the j -th slice among 5-sliced MIP P_m , $1 \leq j \leq 5$) back to the voxel feature space (S_n , i.e., the n -th slice in the input 3D patch, $1 \leq n \leq 16$). It is noted that the feature dimension is reduced from 3D to 2D for a convenient illustration in Fig. 2 (b) (i.e., without considering the 32 feature channels).

For the features of overlapping slices (from the consecutive MIPs) derived from the enhanced vessel probability / features, which are covered by multiple MIPs, we take the element-wise maximum value across the overlapping restoration through the feature channels, i.e., $F_{S_n}[i] = \max(F_{P_{1-1}}[i], \dots, F_{P_{6-5}}[i])$, $1 \leq i \leq 32$, where $F_{S_n}[i]$ represents the i -th dimensional feature (channel) F at the n -th slice in the 3D patch. For example, the feature F_{S_9} is computed across the overlapping slices of $P_{3-5}, P_{4-3}, P_{5-1}$ as highlighted in the pink color in Fig. 2 (c). The whole process of the cross-MIP fusion in the feature channels of the 3D volume feature space is shown in Fig. 2 (c) in detail. Now, in this 3D volume feature space, the backprojected 2D MIP features and 3D volume features from two streams are integrated together through concatenation, constructing a unified high-dimensional joint convolutional embedding for predicting the final vessel segmentation.

2.3 Loss Function

The major learning objective of our JointVesselNet network is to extract the sparse 3D vasculature structure from the 3D MRA volume image using a 3D segmentation network supplemented by information from the denser and more connected multiple 2D MIPs. Consequently the network loss function consists of two terms:

$$L = L_{vox_{3D-2D}} + \lambda L_{mip}, \quad (2)$$

where $L_{vox_{3D-2D}}$ is the joint 3D-2D segmentation Dice loss defined as:

$$L_{vox_{3D-2D}} = -\frac{2\sum_{x \in V} p(x)g(x) + \delta}{\sum_{x \in V} p(x) + \sum_{x \in V} g(x) + \delta}, \quad (3)$$

where $p(x)$ and $g(x)$ are the predicted voxel-wise vessel probability maps and ground truth binary labels within the query volume patch V , respectively. δ is a smooth constant. L_{mip} acts as a regularization term during training, which is also a Dice loss function defined (similarly to $L_{vox_{3D-2D}}$) within the 2D composited MIP and supervised by the ground truth MIP vessel binary labels. λ is the constant coefficient of L_{mip} , which is set to be 0.2 for our best experiment performance.

3 Experiments and Results

For both datasets, we first apply the MR-based skull-stripping method [27] to extract the pure brain from each MRA image. 3D training patches with the imbalanced dimen-

sions are randomly-extracted with overlapping focusing on the brain area in a whole 3D MRA, e.g., 80 patches for each TubeTK case and 440 patches for each collaborative clinical case. The random training / validation / testing case split is 33 / 3 / 6 and 6 / 2 / 3 for TubeTK dataset and the clinical dataset, respectively. The testing accuracy is computed in the full brain patched with no overlap.

Our JointVesselNet network adopts the Adam optimizer with an initial learning rate of 0.0001 with 0.5 as the learning decay factor and 10 epochs as the learning patience. The network is implemented with the TensorFlow framework and the total training time is about 10 hours on two NVIDIA GeForce GTX 1080 GPUs with 8 GB GDDR5X memory. The source code of our method and datasets will be made available later.

We first compare our JointVesselNet performance on TubeTK dataset with four state-of-the-art deep learning based methods (i.e., 3D U-Net [20], 2D U-Net [26], DeepVesselNet [22], and Uception [19]) and one classical parametric intensity-based method (i.e., vesselness algorithm [5, 11]) in 3D vessel segmentation. All deep learning methods in comparison are trained until convergence by using the same dataset split or using the results reported from their original publication (such as Uception). For 2D U-Net, we train it with 128×128 2D patches, whose amount is over 10 times of the amount of 3D patches extracted for the 3D CNN based methods with on-the-fly data augmentation for a fair data acquisition. For DeepVesselNet, we have tried different combinations of their data pre-processing processes and chosen the image intensity clipping for obtaining an optimal performance on TubeTK dataset. For the evaluation on the clinical dataset from our medical collaboration, we use the state-of-the-art model-driven MRAnIs method [25] mentioned in Sec. 2.1 to extract the vessels for comparison.

Four quantitative metrics, i.e., Dice Similarity (Dice), Sensitivity, Precision, and False Positive Rate (FPR), are used for numerical evaluation. The performance comparison of different methods on TubeTK dataset is shown in Tab. 1. ‘—’ means ‘not applicable’ due to the lack of their implementations or results. The best results in the table are shown in bold font. From Tab. 1, we can see that our JointVesselNet has the best overall performance among all the methods on TubeTK dataset. With the 2D composited MIP feature integration, our network performs better than a pure 3D U-Net [20] over all four different metrics. The qualitative comparison of MIP-wise (e.g., 5-sliced) segmentation results and 3D global vessel segmentation results between our JointVesselNet and 3D U-Net (one of the most robust state-of-the-art deep learning based methods for biomedical image segmentation) is shown in Fig. 3. With the 2D composited MIP complementary information, the final vessel segmentation shows better connectivity and better small vessel capturing compared to 3D U-Net, as marked in red circles (3D global vessel segmentation visualization) and green circles (2D MIP vessel segmentation visualization). Another observation is that 3D U-Net greatly outperforms 2D U-Net [26] since the former method is able to capture the cross-slice continuity. That is why 3D CNN should be involved in sparse 3D object segmentation with complex topology. DeepVesselNet [22] fails to yield a good performance as reported in their dataset, which could result from the lack of the pre-training procedure and the instability of their loss function, as well as over-simplified network (e.g., five convolutional layers). Our method also performs better than the best Uception result in [19] on TubeTK dataset with even less data pre-processing procedure. We also employ the vesselness

algorithm [5, 11], a widely-used approach to segment cylindrical vessel structures in the medical field, as a traditional benchmark method for comparison. Last but not least, all deep learning methods greatly outperform vesselness method on TubeTK dataset.

Table 1: Quantitative performance evaluation of different methods.

Metrics / Methods	Ours	3D U-Net	Uception	DeepVesselNet	2D U-Net	Vesselness
Dice (%) $\uparrow$	71.81	71.01	67.01	64.12	65.10	37.71
Sensitivity (%) $\uparrow$	79.33	75.99	66.02	63.20	73.93	65.34
Precision (%) $\uparrow$	76.66	74.00	—	63.75	70.05	47.69
FPR (%) $\downarrow$	0.0821	0.0958	—	0.1465	0.1041	0.1393

Note: Accuracy metric is not included here, since it is always very high (e.g., $\geq 99\%$) in highly sparse vessel data.

As mentioned in Sec. 2.1, our collaborative medical domain experts collect clinical MRA datasets and use the state-of-the-art non-linear subtraction (MRAnls) method [25] to extract the clean midbrain vessels, which tend to be a good physical indicator of several cerebral diseases. MRAnls requires different data modalities and tedious manual parameter-tuning; however, as shown in Fig. 3, its segmentation result still fails to be free from location-dependent interference, such as superior sagittal sinus (red dotted circles in 3D visualization and green dotted circles in MIP visualization) and some random noises (red solid circles). In order to improve the segmentation performance with limited data inventory (11 cases in total), we fine-tune our network pre-trained on TubeTK dataset with six cases in the clinical dataset. Our network only needs TE1 pre-contrast SWI (a single-modal MRA) data as input. The quantitative comparison results between ours and MRAnls method are **82.98%** / 80.60%, **83.77%** / 82.26%, **83.69%** / 82.07%, **0.0337%** / 0.0354% for Dice, Sensitivity, Precision, and FPR, respectively. Our numeric evaluation outperforms MRAnls method on all metrics. Fig. 3 further shows that our vessel segmentation results are more continuous and cleaner than MRAnls results.

4 Conclusion

In this work, we have proposed a deep neural network method, JointVesselNet, to segment high-fidelity 3D cerebrovascular structure from MRA images. By backprojecting the learned multislice composited 2D MIP feature vectors into the 3D volume embedding space, the proposed framework can strengthen the sparse 3D vascular representation by better capturing the small vessels as well as improving the vessel connectivity, which outperforms the state-of-the-art classical and deep learning based methods. In medical practice, this work can be used as the key functions for real-time segmentation and visualization of sparse and complicated 3D vasculature to improve MRA diagnosis.

Acknowledgements. We would like to thank the reviewers for their valuable comments. We are grateful to Yongsheng Chen from Neurology for the early discussion of this work, Pavan K. Jella from Radiology for preparing and collecting the clinical

datasets, and Michelle Hua from Cranbrook Schools for pre-processing the datasets and proofreading. This work was partially supported by the NSF under Grant Numbers IIS-1816511, CNS-1647200, OAC-1657364, OAC-1845962, OAC-1910469, the Wayne State University Subaward 4207299A of CNS-1821962, NIH 1R56AG060822-01A1, NIH 1R44HL145826-01A1, ZJNSF LZ16F020002, and NSFC 61972353.

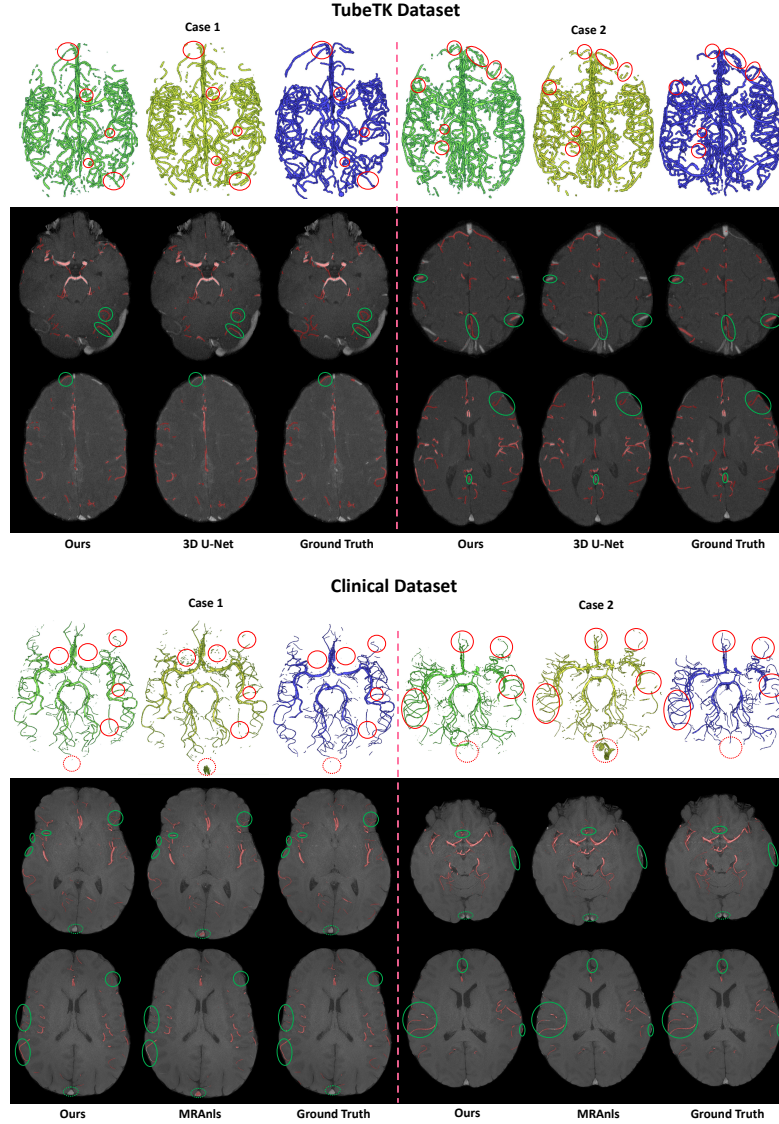


Fig. 3: Some qualitative comparison results from two datasets. 3D global vessel segmentations are shown from superior direction. Two MIP segmentations are visualized by 5-sliced MRA images.

References

1. Brown, W., Thore, C.: Cerebral microvascular pathology in ageing and neurodegeneration. *Neuropathology and Applied Neurobiology* **37** (1), 56–74 (2011)
2. Dorr, A., Sahota, B., et al.: Amyloid- β -dependent compromise of microvascular structure and function in a model of Alzheimer's disease. *Brain* **135** (10), 3039–3050 (2012)
3. Gouw, A., Seewann, A., Van Der Flier, W., Barkhof, F., et al.: Heterogeneity of small vessel disease: a systematic review of MRI and histopathology correlations. *Journal of Neurology, Neurosurgery & Psychiatry* **82** (2), 126–135 (2011)
4. Mott, M., Pahigiannis, K., Koroshetz, W.: Small blood vessels: big health problems: National Institute of Neurological Disorders and Stroke update. *Stroke* **45** (12), e257–e258 (2014)
5. Frangi, A., Niessen, W., et al.: Multiscale vessel enhancement filtering. *International Conference on Medical Image Computing and Computer-Assisted Intervention*, 130–137 (1998)
6. Martínez-Pérez, M., Hughes, A., Stanton, A., et al.: Retinal blood vessel segmentation by means of scale-space analysis and region growing. *International Conference on Medical Image Computing and Computer-Assisted Intervention*, 90–97 (1999)
7. Nain, D., Yezzi, A., Turk, G.: Vessel segmentation using a shape driven flow. *International Conference on Medical Image Computing and Computer-Assisted Intervention*, 51–59 (2004)
8. Chung, A., et al.: Statistical 3D vessel segmentation using a Rician distribution. *International Conference on Medical Image Computing and Computer-Assisted Intervention*, 82–89 (1999)
9. Liao, W., Rohr, K., Wörz, S.: Globally optimal curvature-regularized fast marching for vessel segmentation. *International Conference on Medical Image Computing and Computer-Assisted Intervention*, 550–557 (2013)
10. Florin, C., Paragios, N., Williams, J.: Globally optimal active contours, sequential Monte Carlo and on-line learning for vessel segmentation. *European Conference on Computer Vision*, 476–489 (2006)
11. Descoteaux, M., Collins, D., Siddiqi, K.: A geometric flow for segmenting vasculature in proton-density weighted MRI. *Medical Image Analysis*, **12** (4), 497–513 (2008)
12. Wang, S., Peplinski, B., Lu, L., et al.: Sequential Monte Carlo tracking for marginal artery segmentation on CT angiography by multiple cue fusion. *International Conference on Medical Image Computing and Computer-Assisted Intervention*, 518–525 (2013)
13. Forkert, N., Schmidt-Richberg, A., Fiehler, J., et al.: 3D cerebrovascular segmentation combining fuzzy vessel enhancement and level-sets with anisotropic energy weights. *Magnetic Resonance Imaging*, **31** (2), 262–271 (2013)
14. Fu, H., Xu, Y., Lin, S., Wong, D., Liu, J.: DeepVessel: retinal vessel segmentation via deep learning and conditional random field. *International Conference on Medical Image Computing and Computer-Assisted Intervention*, 132–139 (2016)
15. Li, Q., Feng, B., Xie, L., et al.: A cross-modality learning approach for vessel segmentation in retinal images. *IEEE Transactions on Medical Imaging*, **35** (1), 109–118 (2015)
16. Mo, J., Zhang, L.: Multi-level deep supervised networks for retinal vessel segmentation. *International Journal of Computer Assisted Radiology and Surgery*, **12** (12), 2181–2193 (2017)
17. Liskowski, P., Krawiec, K.: Segmenting retinal blood vessels with deep neural networks. *IEEE Transactions on Medical Imaging*, **35** (11), 2369–2380 (2016)
18. Shin, S., Lee, S., Yun, I., Lee, K.: Deep vessel segmentation by learning graphical connectivity. *Medical Image Analysis*, **58**, 101556 (2019)
19. Sanches, P., Meyer, C., et al.: Cerebrovascular network segmentation of MRA images with deep learning. *IEEE International Symposium on Biomedical Imaging*, 768–771 (2019)
20. Çiçek, Ö., Abdulkadir, A., Lienkamp, S., et al.: 3D U-Net: learning dense volumetric segmentation from sparse annotation. *International Conference on Medical Image Computing and Computer-Assisted Intervention*, 424–432 (2016)

21. Szegedy, C., Ioffe, S., et al.: Inception-v4, inception-resnet and the impact of residual connections on learning. The AAAI Conference on Artificial Intelligence, (2017)
22. Tetteh, G., Efremov, V., Forkert, N., Schneider, M., et al.: DeepVesselNet: vessel segmentation, centerline prediction, and bifurcation detection in 3D angiographic volumes. arXiv preprint arXiv:1803.09340 (2018)
23. Kitrungsakul, T., Han, X., Iwamoto, Y., Lin, L., Foruzan, A., et al.: VesselNet: a deep convolutional neural network with multi pathways for robust hepatic vessel segmentation. *Computerized Medical Imaging and Graphics*, **75**, 74–83 (2019)
24. Napel, S., Marks, M., Rubin, G., et al.: CT angiography with spiral CT and maximum intensity projection. *Radiology*, **185** (2), 607–610 (1992)
25. Ye, Y., Hu, J., Wu, D., Haacke, E.: Noncontrast-enhanced magnetic resonance angiography and venography imaging with enhanced angiography. *Journal of Magnetic Resonance Imaging*, **38** (6), 1539–1548 (2013)
26. Ronneberger, O., Fischer, P., Brox, T.: U-Net: Convolutional networks for biomedical image segmentation. *International Conference on Medical Image Computing and Computer-Assisted Intervention*, 234–241 (2015)
27. Jenkinson, M., Pechaud, M., et al.: BET2: MR-based estimation of brain, skull and scalp surfaces. *Annual Meeting of the Organization for Human Brain Mapping*, (2005)