



# Rethinking the detail-preserved completion of complex tubular structures based on point cloud: A dataset and a benchmark<sup>☆</sup>

Yaolei Qi<sup>a,1</sup>, Yikai Yang<sup>a,1</sup>, Wenbo Peng<sup>a</sup>, Shumei Miao<sup>c,a</sup>, Yutao Hu<sup>a,\*</sup>, Guanyu Yang<sup>a,b</sup><sup>ID</sup>,\*\*

<sup>a</sup> Key Laboratory of New Generation Artificial Intelligence Technology and Its Interdisciplinary Applications (Southeast University), Ministry of Education, No.2, Sipai Lou, Xuanwu District, Nanjing, 210096, China

<sup>b</sup> Jiangsu Province Joint International Research Laboratory of Medical Information Processing, Nanjing, 210096, China

<sup>c</sup> The First Affiliated Hospital of Nanjing Medical University, Nanjing, 210029, China

## ARTICLE INFO

### Keywords:

Tubular structure reconnection  
Topology preservation  
Point cloud completion  
Benchmark  
Dataset

## ABSTRACT

Complex tubular structures are essential in medical imaging and computer-assisted diagnosis, where their integrity enhances anatomical visualization and lesion detection. However, existing segmentation algorithms struggle with structural discontinuities, particularly in severe clinical cases such as coronary artery stenosis and vessel occlusions, which leads to undesired discontinuity and compromising downstream diagnostic accuracy. Therefore, it is imperative to reconnect discontinuous structures to ensure their completeness. In this study, we explore the tubular structure completion based on point cloud for the first time and establish a Point Cloud-based Coronary Artery Completion (PC-CAC) dataset, which is derived from real clinical data. This dataset provides a novel benchmark for tubular structure completion. Additionally, we propose TSRNet, a Tubular Structure Reconnection Network that integrates a detail-preserved feature extractor, a multiple dense refinement strategy, and a global-to-local loss function to ensure accurate reconnection while maintaining structural integrity. Comprehensive experiments on our PC-CAC and two additional public datasets (PC-ImageCAS and PC-PTR) demonstrate that our method consistently outperforms state-of-the-art approaches across multiple evaluation metrics, setting a new benchmark for point cloud-based tubular structure reconstruction. Our benchmark is available at <https://github.com/YaoleiQi/PCCAC>.

## 1. Introduction

Complex tubular structures, such as blood vessels, constitute essential components of human physiology and play a critical role in medical imaging and computer-assisted diagnosis (Li et al., 2022). Preserving their structural integrity is essential, as it enables the accurate visualization for radiologists and facilitates the detection of lesions (Ko et al., 2017), supporting more informed and precise treatment decisions (Gharleghi et al., 2022). Existing segmentation algorithms (Qi et al., 2023; Kong et al., 2020; Zhang et al., 2023, 2022) have made progress in analyzing tubular structures. However, in more clinically significant and severe cases (Fig. 1(a)), such as coronary artery stenosis (Madhavan et al., 2014; Pijls and Sels, 2012), blockage (Garje et al.,

2015), and myocardial bridging (Fig. 2(a)) (Lee and Chen, 2015), the continuity of segmentation is often compromised, resulting in fractures that negatively affect downstream diagnosis. Therefore, it is imperative to design an effective and robust reconstruction framework to preserve the anatomical integrity of complex tubular structures.

Recently, some studies (Mou et al., 2019; Qiu et al., 2023; Han et al., 2022; Weng et al., 2023) have investigated the reconnection of fractured tubular structures using voxel-based approaches. However, these methods encounter notable limitations in handling the following complex scenarios. As shown in Figs. 1(a) and 2(a), lesions such as coronary artery stenosis cause blurry and unclear image representations, reducing the visibility of critical vascular structures. In cases of complete occlusion (Thanvi and Robinson, 2007), this issue may further

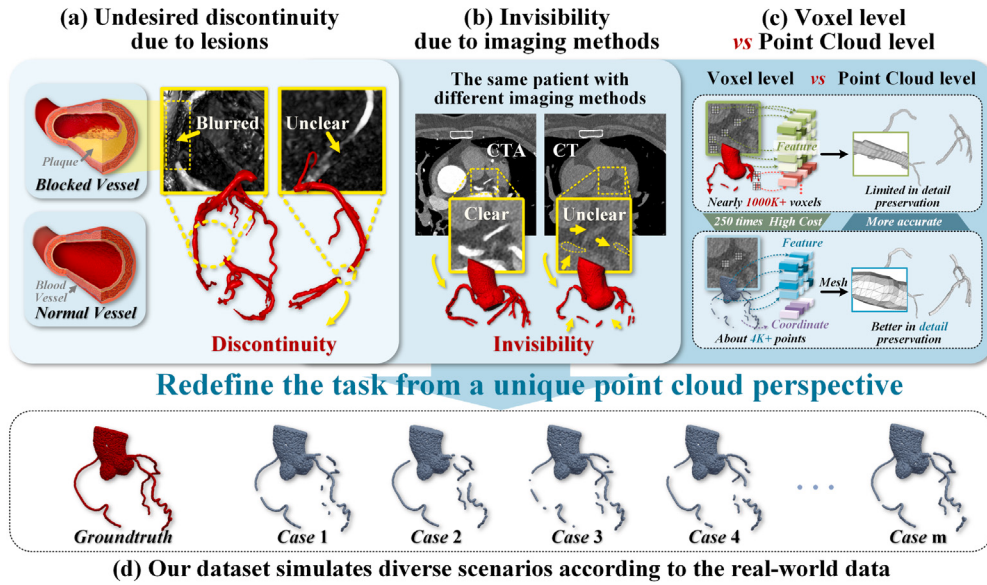
<sup>☆</sup> This research was supported by the National Key Research and Development Program of China (2024YFC2417803), National Natural Science Foundation of China (82441021), Jiangsu Province Cutting-edge Technology R&D Program (BF2025628), in part by the Funded by Basic Research Program of Jiangsu under Grant (BK20251279), the Postgraduate Research & Practice Innovation Program of Jiangsu Province, the Fundamental Research Funds for the Central Universities, China (KYCX22\_0239), and the Big Data Computing Center of Southeast University.

\* Corresponding author.

\*\* Corresponding author at: Key Laboratory of New Generation Artificial Intelligence Technology and Its Interdisciplinary Applications (Southeast University), Ministry of Education, No.2, Sipai Lou, Xuanwu District, Nanjing, 210096, China.

E-mail address: [yang.list@seu.edu.cn](mailto:yang.list@seu.edu.cn) (G. Yang).

<sup>1</sup> Yaolei Qi and Yikai Yang contributed equally to this work.



**Fig. 1.** Motivation. (a) shows the undesired discontinuity caused by unclear image representation in structures affected by lesions. (b) shows vessels visualized across different modalities, where thin tubular structures are prone to being obscured, making them difficult to observe without the use of contrast agents. (c) shows the differences in structural reconnection from the perspectives of both voxel-based and point cloud representation. (d) our work redefines the reconnection task from a point cloud perspective and builds a dataset that is simulated from real clinical data.

lead to missing information in local regions, resulting in undesired discontinuity. Such disruptions compromise the structural integrity of tubular anatomy and pose significant challenges for accurate reconnection. Similarly, as depicted in Fig. 1(b), variations across imaging modalities also present additional challenges. For patients allergic to contrast agents, non-contrast imaging presents substantial obstacles in analyzing tubular structures due to reduced visibility and lack of enhancement. These fine vascular features are often difficult to distinguish, which hampers visual assessment and complicates structural reconnection when relying solely on voxel-based methods.

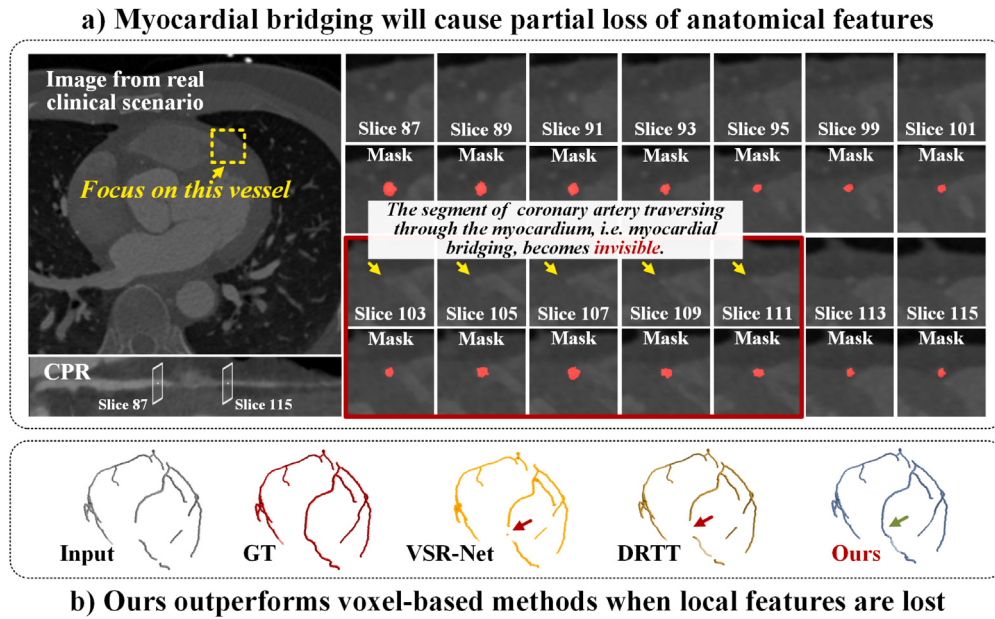
Fortunately, point clouds, as a structured data representation that emphasizes coherence and consistency (Fei et al., 2022; Chen et al., 2021; Bernard et al., 2017; Beetz et al., 2023), provide advantages that are absent at the voxel level. As shown in Fig. 1(c), point cloud-based methods offer notable advantages in computational efficiency and accuracy. Voxel-based methods process over 1000K voxels, resulting in high computational cost, whereas point cloud-based methods operate on around 4K points. Additionally, the reconnected point cloud can then be directly utilized to construct a high-resolution sub-voxel mesh model, enabling accurate anatomical reconstruction and simulation without the need to revert to the original voxel representation. For instance, in coronary arteries, this enables the creation of a more accurate physical model for downstream tasks such as hemodynamic simulation (Ko et al., 2017). Thus, we redefine the reconnection of complex tubular structures from a point cloud perspective (Fig. 1(d)) for the first time.

However, the direct application of point cloud completion to fractured tubular structures presents several key challenges: (a) **Imbalanced point distribution.** The number and density of points vary greatly between large anatomical regions (e.g., the aorta) and slender elongated segments (e.g., coronary vessels). This extreme disparity biases the model towards larger structures and hinders its ability to capture fine structural details, ultimately degrading performance on thin tubular structures. (b) **Complex global topology.** Tubular structures exhibit intricate and delicate morphologies across individuals. Compared to objects with regular or simple structures, predicting missing segments in thin tubular structures is particularly difficult, especially when capturing intricate curvature details. As a result, the model tends to produce incomplete or topologically incorrect structures during point generation.

To tackle the above obstacles, as a first attempt, we propose a novel Point Cloud-based Coronary Artery Completion (PC-CAC) dataset and propose a robust, detail-preserved baseline for the precise reconnection of fractured tubular structures. (1) To enhance the integrity of complex tubular structures, we redefine the reconnection task from a point cloud perspective. To our best knowledge, we are the first to construct a dedicated dataset PC-CAC. Our dataset is derived from clinical data, incorporating segmentation results while preserving fractures, under-segmentation, and other imperfections. Furthermore, to simulate more challenging real-world scenarios, 10%–30% of the points from damage-prone regions within the main structures are removed from the input to better reflect the complexity of structural discontinuities. This design facilitates a more rigorous evaluation of completion algorithms and promotes the advancement of effective reconnection strategies. (2) To address the challenges of imbalanced point distribution and complex global topology, we propose a Tubular Structure Reconnection Network (TSRNet) designed for the accurate reconnection of fractured complex tubular structures, which comprises a detail-preserved feature extractor, a multiple dense refinement strategy, and a global-to-local loss function. These components cooperate to enhance detail preservation and effectively handle hard-to-represent regions. (3) To objectively evaluate our approach, experiments are conducted on our PC-CAC dataset and two public datasets. We select a coronary artery segmentation (ImageCAS) dataset (Zeng et al., 2023) and a pulmonary tree repairing (PTR) dataset (Weng et al., 2023). Then, we convert these two open-source datasets into point cloud format, namely PC-ImageCAS and PC-PTR, and conduct validation accordingly.

In summary, our contributions are as follows:

- **The first point cloud-based tubular structure reconnection dataset:** To our best knowledge, we construct the first point cloud-based coronary artery completion (PC-CAC) dataset from clinical data. Our benchmark provides a new perspective for tubular structure reconnection and fostering advancements in this field.
- **A novel exploration and high-performing baseline:** Our work represents the first attempt to explore tubular structure reconnection based on point cloud. We propose a baseline designed for accurately reconnecting fractured tubular structures, comprising



**Fig. 2.** (a) Illustration of myocardial bridging in a real clinical case, where the lesion leads to local structural loss that challenges voxel-based models. (b) Our proposed method shows superior reconnection performance, and better handles such fracture compared to existing voxel-based approaches.

a detail-preserved feature extractor, a multiple dense refinement strategy, and a global-to-local loss function. These methods cooperate to enhance detail preservation and effectively handle hard-to-represent regions.

- **A comprehensive evaluation across multiple datasets.** To objectively evaluate our approach, experiments are conducted on our PC-CAC dataset and two public datasets. Experimental results show that our method achieves state-of-the-art performance across multiple datasets.

## 2. Related works

### 2.1. Reconnection methods from voxel perspective

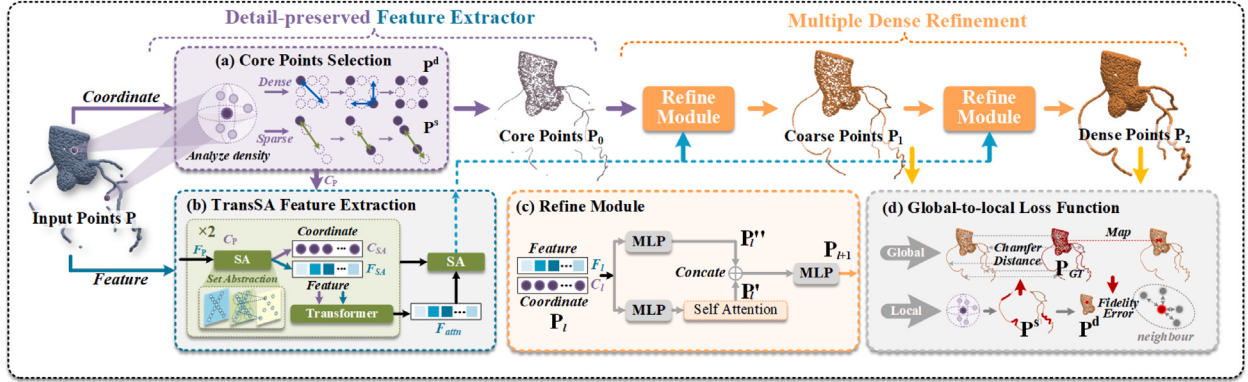
Many methods have been proposed to reconnect fractured vessels from the image perspective. (1) Traditional walk-based methods (Mou et al., 2019; Qiu et al., 2023; Gupta et al., 2023). Random walk classifies unlabeled pixels based on pre-annotated vessel (foreground) and background seed points, initially used directly for vessel segmentation (Grady, 2006; Li et al., 2015). Building on this, Mou et al. (2019) proposed the Probability Regularized Walk (PRW) algorithm for fractured vessel connection: labeling all connected regions via morphological operations, matching minimal-distance point pairs, and integrating local vessel direction modeling with segmentation network probability outputs to form “walking” connections. Qiu et al. (2023) extended the PRW algorithm to 3D and applied it to coronary arteries. Gupta et al. (2023) introduced a perturb-and-map mechanism to diversify path search and avoid local optima or deviations. These methods rely on local information while neglecting global morphology and are limited by weak voxel-level features, heavily depending on endpoint detection accuracy. As a result, they struggle to maintain structural continuity, often leading to misaligned connections, incomplete reconstructions, or even false reconnection in complex structures. These limitations hinder their applicability in real-world clinical scenarios, where missing or noisy data further exacerbate the challenges. However, our approach integrates global and local information through point cloud level completion without requiring explicit fracture endpoint identification.

### 2.2. Generic point cloud completion methods

PointNet (Qi et al., 2017a) and PointNet++ (Qi et al., 2017b) advanced point cloud completion. Numerous deep learning-based methods (Xie et al., 2020; Yu et al., 2021; Zhou et al., 2022; Xiang et al., 2022; Wen et al., 2022; Chen et al., 2023; Rong et al., 2024; Wang et al., 2024) have emerged, typically using encoder-decoder architectures to map partial inputs to complete shapes. Xie et al. (2020) regularized unordered point clouds via 3D grids as intermediate representations, leveraging structural and contextual neighbor information. Yu et al. (2021) adopted a Transformer encoder-decoder architecture, converting point clouds into point proxies and designing geometry-aware modules to exploit 3D geometric inductive biases. Zhou et al. (2022) introduced Patch Seeds to capture global structures and local patterns. Wen et al. (2022) predicted unique point displacement paths under distance constraints for strict correspondence learning. Wang et al. (2024) proposed geometric detail perception and self-feature augmentation units to establish structural relationships via attention mechanisms. These methods prioritize large surface optimization over slender tubular structures, thus failing to capture the intricate geometric details. The absence of dedicated mechanisms for preserving thin tubular features tends to result in disrupted topology, or over-reconstructions, ultimately compromising diagnostic reliability and downstream clinical applications. However, our approach specifically targets elongated vascular features through clinical data-driven training and continuity constraints for vascular topology.

## 3. Methodology

As shown in Fig. 3, our TSRNet comprises two main stages: a Detail-preserved Feature Extractor (Section 3.1) and a Multiple Dense Refinement strategy (Section 3.2). These are jointly optimized using a Global-to-local Loss Function (Section 3.3) to ensure the structural consistency. To leverage geometric sparsity, the entire framework operates directly on point cloud inputs. The procedure for converting voxel-based segmentation results with anatomical discontinuities into point clouds is detailed in the next section (Section 4) and included in our open-source implementation.



**Fig. 3.** Framework. Our framework comprises two main stages: a Detail-preserved Feature Extractor and a Multiple Dense Refinement strategy. These are jointly optimized using a Global-to-local Loss Function to ensure the structural consistency. (a) Core Points Selection analyzes the density to extract critical points that are easily overlooked. (b) TransSA Feature Extraction utilizes a Transformer-based set abstraction to capture both local and global features. (c) Refine Module progressively refines results, generating coarse-to-fine reconstructions. (d) Global-to-local Loss constrains the reconstruction by enforcing global consistency via Chamfer Distance and enhancing local structure with Fidelity Error.

### 3.1. Detail-preserved feature extractor

In this section, we discuss how to perform the Detail-preserved Feature Extractor. Given the input point cloud as  $\mathbf{P} = \{\mathbf{p}_i | i = 1, 2, \dots, N\} \in \mathbb{R}^{N \times 3}$ , where  $N$  represents the total number of points, and each  $\mathbf{p}_i$  contains 3D coordinates  $(x, y, z)$ . Within our network architecture, the extractor captures both global shape and local geometric context of the input point cloud  $\mathbf{P}$  and downsamples it to support the subsequent coarse-to-fine refinement process.

#### 3.1.1. Core points selection module

Progressive generation aims to achieve global densification of sparse point clouds. Thus, we first derive a sparse core point cloud  $\mathbf{P}_0$  from the original  $\mathbf{P}$  to integrate the reconnection of fracture locations into the global densification process. To address the challenge of imbalanced point distribution, which arises from the small spatial proportion occupied by complex tubular structures in the overall point cloud, we propose the Core Points Selection (CPS) module. CPS dynamically adjusts the attention of the model to different point types during training (Fig. 3(a)). By analyzing density variations and employing adaptive selection strategies, CPS effectively identifies critical points, ensuring a well-balanced point distribution and mitigating the impact of sparsity.

Specifically, CPS first analyzes the density distribution of the input point cloud  $\mathbf{P}$  to separate dense  $\mathbf{P}^d$  and sparse regions  $\mathbf{P}^s$ . Then, the Farthest Point Selection (Qi et al., 2017a) is employed as the fundamental approach for core point selection, while we introduce adaptive constraints tailored to different density conditions to optimize the selection process in both dense and sparse regions. To ensure balanced and effective critical points selection, CPS incorporates three main constraints: (1) **Simultaneous selection in separate dense and sparse regions.** CPS operates separately on both dense and sparse regions, ensuring that the selection process adapts to local point distribution features without biasing towards high-density areas. (2) **Region-restricted point selection.** Points are selected strictly within their respective regions, meaning that points in high-density regions cannot be chosen from points in low-density regions. This constraint prevents structural distortion caused by disproportionate sampling across density variations. (3) **Priority on end and isolated points.** CPS prioritizes the retention of endpoints and isolated points, which are often crucial for maintaining continuity. These points play a key role in guiding the refinement process and preventing structural loss in later stages.

Under the condition of adhering to the principles, CPS starts from a single point and iteratively selects the most suitable point in different

situations from the current point until the designated number  $N_0$  of points is reached by  $\mathbf{P}_0 = \{\mathbf{P}^d, \mathbf{P}^s\} = CPS(\mathbf{P}, N_0)$ , where  $\mathbf{P}^d$  refers to the dense regions, and  $\mathbf{P}^s$  refers to the sparse regions.  $N_0$  is set to one quarter of the complete point cloud to facilitate high-quality refinement and to retain as much information from  $\mathbf{P}$  as possible.  $\mathbf{P}_0$  selected by CPS represent core points containing crucial information, exhibiting a uniform overall distribution, and preserving attention on the complex vascular structures.

#### 3.1.2. TransSA feature extraction module

Inspired by Zhou et al. (2022), which applies patch seeds to obtain rich information, we propose the TransSA module, which integrates Set Abstraction (SA) and Transformer (Zhao et al., 2021) to effectively extract features from irregular structures.

As shown in Fig. 3(b) and Eq. (1), the SA block concatenates the input point coordinates  $C_P$  with their corresponding features  $F_P$ , selects the centroid  $C_{SA}$  within the local region, and identifies neighboring points of  $C_{SA}$  via K-Nearest Neighbors (KNN) (Zhang et al., 2017) to construct local representation. Subsequently, PointNet (Qi et al., 2017a) is applied to extract refined feature points  $F_{SA}$ . The Transformer module further processes the extracted features  $F_{SA}$  and coordinates  $C_{SA}$ , mitigating the unordered nature of point cloud data while leveraging self-attention mechanisms for enhanced feature extraction, denoted as  $F_{attn}$ . Additionally, the SA module utilizes  $F_{attn}$  to capture higher-dimensional representations, enriching the feature extraction process. By iteratively integrating Transformer with SA, the TransSA module effectively enhances feature discrimination, ensuring robust representation learning for complex structures.

$$\begin{aligned} F_{SA}, C_{SA} &= SA(F_P, C_P), \\ F_{attn} &= Transformer(F_{SA}, C_{SA}). \end{aligned} \quad (1)$$

### 3.2. Multiple dense refinement strategy

As shown in Fig. 3, the progressive network is designed to faithfully restore details from the sparse point cloud, gradually guiding the model to focus on details of completion. Multiple refine modules (Fig. 3(c)) are designed to sequentially reconstruct the core point cloud  $\mathbf{P}_0$  into the coarse point cloud  $\mathbf{P}_1$ , and then further refine  $\mathbf{P}_1$  into the dense point cloud  $\mathbf{P}_2$ , which constitutes the final output. Features  $F_\ell$  extracted by the TransSA module are incorporated into each refine module to better preserve the fidelity of the data. Within each refine module, the coordinates  $C_\ell$  is fused with the feature  $F_\ell$ , which undergoes

processing through an MLP followed by three rounds of Self-Attention, generating an intermediate representation  $\mathbf{P}'_\ell$ . Simultaneously,  $\mathbf{P}_\ell$  is processed through another MLP to produce  $\mathbf{P}''_\ell$ . The final refined output  $\mathbf{P}_{\ell+1}$  is obtained by applying an MLP to the concatenation of  $\mathbf{P}'_\ell$  and  $\mathbf{P}''_\ell$ . The first refine module focuses on initial completion by addressing major discontinuities, while the second refine module further enhances the reconstructed details, emphasizing minute vascular structures to ensure finer structural integrity.

### 3.3. Global-to-local loss function

Considering the progress of the coarse-to-fine training process, it becomes imperative to impose stage-specific constraints, gradually guiding the model optimization towards the desired direction. Therefore, the overall loss function  $\mathcal{L}$  encompasses two components, global  $\mathcal{L}_g$  and local  $\mathcal{L}_l$ , which are tailored to the differences in the outputs across different stages. In detail,  $\ell_1$  Chamfer Distances ( $CD^{\ell_1}$ ) (Fan et al., 2017) is employed as the primary metric, which utilizes  $\ell_1$ -norm to compute the distance between two sets of points, ensuring the precise global structural alignment. Furthermore, to maintain a high degree of similarity between the reconstructed output and the original input, and enhance the focus of the model on local regions, we integrate Fidelity Error ( $FE$ ) (Yuan et al., 2018) into the loss formulation. Since  $FE$  is defined as the average distance from each point in the input to its nearest neighbor in the output, effectively measuring how well the original structure is preserved during completion.

In our study, loss calculation is divided into global and local perspectives. Specifically, the global perspective (Eq. (2)) focuses on the overall structural restoration, ensuring the coherence of the reconstructed shape. Meanwhile, the local perspective (Eq. (3)) leverages the CPS module mentioned above to analyze point set density and processes different density regions separately. This approach dynamically balances the representation of sparse and dense point sets, enhancing the model's attention to sparse structures. The loss from global  $\mathcal{L}_g$  becomes:

$$\mathcal{L}_g = \underbrace{CD^{\ell_1}(\mathbf{P}_1, \mathbf{P}_{GT}) + FE(\mathbf{P}_1, \mathbf{P}_{GT})}_{\mathcal{L}_{coarse}} + \underbrace{CD^{\ell_1}(\mathbf{P}_2, \mathbf{P}_{GT})}_{\mathcal{L}_{fine}}, \quad (2)$$

where  $CD^{\ell_1}$  represents the  $\ell_1$  Chamfer Distance, which measures the discrepancy between two point sets.  $\mathcal{L}_g$  denotes the loss function from global perspective, consisting of  $\mathcal{L}_{coarse}$  and  $\mathcal{L}_{fine}$ , corresponding to different stages.  $\mathbf{P}_1$  and  $\mathbf{P}_2$  represent the reconstructed point clouds at the coarse and fine stages, respectively, while  $\mathbf{P}_{GT}$  denotes the ground truth point cloud.  $FE$  refers to Fidelity Error, which evaluates how well the reconstructed structure preserves the original input, focusing on fine-grained details.

From the local perspective, the CPS module mentioned above is utilized to analyze the density distribution of the input point set  $\mathbf{P}$ , and separate it into two regions, including  $\mathbf{P}^d$  for high-density areas and  $\mathbf{P}^s$  for low-density areas, express as:  $\{\mathbf{P}^d, \mathbf{P}^s\} = CPS(\mathbf{P}, N_0)$ . Then  $CPS(\cdot)$  is performed simultaneously on the output point sets  $\mathbf{P}_1$ ,  $\mathbf{P}_2$ , and the ground truth  $\mathbf{P}_{GT}$ . Then the following loss calculation exclusively on  $\mathbf{P}^s$  region:

$$\mathcal{L}_l = \underbrace{CD^{\ell_1}(\mathbf{P}_1^s, \mathbf{P}_{GT}^s) + FE(\mathbf{P}_1^s, \mathbf{P}_{GT}^s)}_{\mathcal{L}_{coarse}} + \underbrace{CD^{\ell_1}(\mathbf{P}_2^s, \mathbf{P}_{GT}^s)}_{\mathcal{L}_{fine}}, \quad (3)$$

where  $CD^{\ell_1}$  and  $FE$  are consistent with Eq. (2). Specifically,  $\mathbf{P}_1^s$ ,  $\mathbf{P}_2^s$  and  $\mathbf{P}_{GT}^s$  represent the sparse regions within the point sets  $\mathbf{P}_1$ ,  $\mathbf{P}_2$  and  $\mathbf{P}_{GT}$ , respectively.

Finally, the global-to-local training process is regulated by a threshold  $\varepsilon$ , ensuring a stage-aware transition between loss components. The complete formulation of the loss function  $\mathcal{L}$  is defined as follows:

$$\mathcal{L} = \begin{cases} \mathcal{L}_g, & \mathcal{L}_g > \varepsilon, \\ (1 - \gamma) \mathcal{L}_g + \gamma \mathcal{L}_l, & \text{otherwise,} \end{cases} \quad (4)$$

where  $\mathcal{L}_g$  represents the global loss, ensuring overall structural consistency, while  $\mathcal{L}_l$  accounts for finer details by focusing on local reconstruction accuracy. The threshold  $\varepsilon$  acts as a phase transition parameter, determining whether the loss function should prioritize global structural alignment or incorporate local refinements. The trade-off parameter  $\gamma$  dynamically adjusts the weighting between global and local losses when the global loss  $\mathcal{L}_g$  is below the threshold, allowing a gradual shift in emphasis as training progresses. This adaptive loss strategy ensures that early training stages prioritize overall geometric consistency, while later stages refine fine details, leading to progressive, high-fidelity reconstructions of complex tubular structures.

## 4. Experimental setting

### 4.1. Datasets description

As shown in Table 1, to objectively evaluate our approach, experiments are conducted on our PC-CAC dataset and two public datasets. We select a coronary artery segmentation (ImageCAS) dataset (Zeng et al., 2023) and a pulmonary tree repairing (PTR) dataset (Weng et al., 2023). Then, we convert these two open-source datasets into point cloud format, namely PC-ImageCAS and PC-PTR, and conduct validation accordingly.

**PC-CAC Dataset.** We construct the Point Cloud-based Coronary Artery Completion (PC-CAC) dataset, derived from real clinical coronary artery data, comprising a total of 427 patients (3416 cases). These images were scanned on a SOMATOM Definition Flash and the contrast media was injected during the scanning process. The x/y-resolution of these CT images is between 0.25~0.57 mm/voxel and the slice thickness is between 0.75 to 3 mm/voxel. The x/y-size of the images is 512 voxels and the z-size is between 128~994 voxels. For pre-processing, we first resample the resolution of these images to 1 mm × 1 mm × 1 mm for a unified resolution, then threshold their grayscale value to [0, 2048] and normalize them to [0, 1] for unified intensity.

The dataset is partitioned into 300 patients (2400 cases) for training, 40 patients (320 cases) for validation, and 87 patients (696 cases) for testing. Each point cloud consists of 4096 points, representing the aorta and the left and right coronary arteries, where the aorta contains 3072 points, and the coronary arteries contain 1024 points. The input point clouds in our dataset originate from real segmentation results (Isensee et al., 2021; Qi et al., 2023) using Alg. 1. To further simulate real-world challenges, we introduce synthetic fractures at high-risk locations, including bifurcations and tortuous structures, where discontinuities are most likely to occur. Additionally, we incorporate random fractures and noise to enhance robustness and generalization. To ensure a diverse and challenging reconstruction task, we generate 8 distinct input cases per patient, capturing a wide range of vascular conditions. To simulate fractures, 10%–30% of the points in the coronary arteries are removed, while the aortic structure remains unchanged, ensuring a realistic and challenging reconstruction scenario.

**PC-ImageCAS Dataset.** To further enhance the evaluation of our proposed approach, we introduce a public coronary artery segmentation dataset ImageCAS (Zeng et al., 2023), and convert it into a point cloud format called PC-ImageCAS dataset. The dataset consists of 8000 cases, derived from 1000 patients, where each complete point cloud sample captures detailed anatomical structures. The dataset is partitioned into 700 patients (5600 cases) for training, 100 patients (800 cases) for validation, and 200 patients (1600 cases) for testing, maintaining consistency in dataset division. This structured partitioning ensures a comprehensive assessment across varying clinical conditions, reinforcing the reliability of our framework in real-world applications.

**PC-PTR Dataset.** To further assess the generalization capability of our proposed baseline, we utilize the pulmonary arteries in Pulmonary Tree Repairing (PTR) dataset (Weng et al., 2023) and convert them into point clouds format, referred to as PC-PTR. The dataset comprises

**Table 1**

The details of our dataset from the real clinical scenario and two public available datasets in our verification tasks.

| (a) The details of two public available datasets in our task    |                                           |                                           |                                                                                                                                                                                     |                                                                                                                                                                                                                                                                                                                                      |
|-----------------------------------------------------------------|-------------------------------------------|-------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Name                                                            | Target dataset                            | Train/Val/Test                            | Detail information                                                                                                                                                                  | Pre-processing                                                                                                                                                                                                                                                                                                                       |
| PC-ImageCAS                                                     | ImageCAS (Zeng et al., 2023) <sup>b</sup> | 700/100/200<br>5600/800/1600 <sup>a</sup> | 1. Scanner: Siemens 128-slice dual-source<br>2. Planar resolution: 0.29~0.43 mm <sup>2</sup><br>3. Slice thickness: 0.25~0.45 mm<br>4. x/y-size: 512 voxels, z-size: 206~275 voxels | 1. Resample the resolution to 1 mm <sup>3</sup><br>2. Normalize via $\frac{\max(\min(0,x),2048)}{2048}$<br>3. Obtain segmentation results<br>4. Extract <b>aorta</b> and main <b>coronary</b> branches<br>5. Generate point cloud based on surface of <b>aorta</b><br>6. Generate point cloud based on centerline of <b>coronary</b> |
| PC-PTR                                                          | PTR (Weng et al., 2023) <sup>c</sup>      | 599/80/160<br>4472/640/1280 <sup>a</sup>  | 1. Scan from multiple medical centers<br>2. Resolution: already be processed to 1 mm <sup>3</sup><br>3. x/y-size: 512 voxels, z-size: 177~798 voxels                                | Generate point cloud based on centerline of vessels                                                                                                                                                                                                                                                                                  |
| (b) The details of our proposed dataset from real clinical data |                                           |                                           |                                                                                                                                                                                     |                                                                                                                                                                                                                                                                                                                                      |
| Name                                                            | Target dataset                            | Train/Val/Test                            | Detail information                                                                                                                                                                  | Pre-processing                                                                                                                                                                                                                                                                                                                       |
| PC-CAC                                                          | PC-CAC                                    | 300/40/87<br>2400/320/696 <sup>a</sup>    | 1. Scanner: SOMATOM Definition Flash<br>2. x/y-resolution: 0.25~0.57 mm/voxel<br>3. Slice thickness: 0.75~3 mm/voxel<br>4. x/y-size: 512 voxels, z-size: 128~994 voxels             | 1. Resample the resolution to 1 mm <sup>3</sup><br>2. Normalize via $\frac{\max(\min(0,x),2048)}{2048}$<br>3. Obtain segmentation results<br>4. Extract <b>aorta</b> and main <b>coronary</b> branches<br>5. Generate point cloud based on surface of <b>aorta</b><br>6. Generate point cloud based on centerline of <b>coronary</b> |

<sup>a</sup> To ensure a diverse and challenging reconstruction task, each patient generates 8 distinct input cases, capturing a wide range of conditions.

<sup>b</sup> ImageCAS: <https://github.com/XiaoweiXu/ImageCAS-A-Large-Scale-Dataset-and-Benchmark-for-Coronary-Artery-Segmentation-based-on-CT>.

<sup>c</sup> PTR: <https://github.com/M3DV/pulmonary-tree-repairing>.

799 patients (6392 cases), with each sample consisting of 8192 points, where 5%–15% of the points are deliberately removed to simulate structural degradation. The dataset is partitioned into 559 patients (4472 cases) for training, 80 patients (640 cases) for validation, and 160 patients (1280 cases) for testing, following the same as the original PTR dataset.

#### 4.2. Details of voxel-to-point clouds process

This section discusses how to perform the voxel-to-point cloud conversion process. Using our proposed coronary artery dataset as an example, it consists of segmentation results (using segmentation models such as nnU-Net Isensee et al., 2021; Qi et al., 2023) that preserve fractures and under-segmentation, along with data that simulate challenging real-world scenarios. Regarding the segmentation results, the skeleton is extracted from the voxels, followed by a point-by-point mapping to form a point cloud. For the simulation of fractures, the complete procedure is illustrated in Alg. 1. The operations on the segmentation results are similar to those of this algorithm, except that simulated fractures are not required. The input point cloud simulation construction process involves four steps: obtaining segmentation or annotation, converting into point clouds, normalizing, and simulating fractures. To better simulate real-world scenarios, operations such as removing cavities from voxels and extracting the main trunk from point clouds are also performed.

#### 4.3. Experimental configurations

**Evaluation Settings.** Our evaluation is divided into two main parts: (1) comparisons with point cloud-based methods (Section 5.1), and (2) comparisons with voxel-based methods (Section 5.2). For both parts, we conduct comprehensive performance evaluations as well as ablation studies to validate the effectiveness of our proposed TSRNet. For **point cloud-based methods**, we select classic point cloud completion algorithms GRNet (Xie et al., 2020) and PointTR (Yu et al., 2021), as well as the methods such as SeedFormer (Zhou et al., 2022) and SnowflakeNet (Xiang et al., 2022) proposed in 2022, AnchorFormer (Chen et al., 2023) proposed in 2023, PointAttN (Wang et al., 2024) and CRA-PCN (Rong et al., 2024) proposed in 2024. For **voxel-based methods**, we adopt DRTT (Li et al., 2021) and VSRNet (Ye et al., 2025) as representative methods, both of which have been used in

volumetric vascular reconstruction tasks. All models are trained on the same datasets under identical implementation settings to ensure a fair comparison. All models were trained independently on each dataset, and no transfer learning was used across datasets.

**Evaluation Metrics.** All metrics are calculated for each case and averaged. (1)  $\ell_1$  Chamfer Distances ( $CD^{\ell_1}$ ) (Fan et al., 2017). It uses  $\ell_1$ -norm to calculate the sum of the average closest distance between two point clouds. (2) F1-Score (F1) (Tatarchenko et al., 2019). It evaluates the distance between the surfaces of objects. (3) Fidelity Error (Yuan et al., 2018). It is used to measure how well the input is preserved, which calculates the average distance between the point and the nearest neighbor. To better assess whether the vessels are functionally reconnected, we additionally include topology-aware metrics. (4) Overlap (OV) (Schaap et al., 2009). It quantifies the structural overlap between the predicted result and the ground truth, reflecting vessel coverage after reconstruction. (5) Bottleneck Distance (BD) (Cohen-Steiner et al., 2005; Kerber et al., 2017). It measures the discrepancy between the persistence diagrams of the prediction and the ground truth, and thus characterizes topological consistency. (6) Betti Error (BE) (Hu et al., 2019; Wyburd et al., 2024). It evaluates topological differences in terms of Betti numbers, capturing discrepancies in connected components and loop structures.

**Implementation Details.** Experiments are implemented using PyTorch on GeForce RTX 2080 Ti. Adam optimizer is used to train the models with an initial learning rate of  $10^{-4}$ . The threshold parameter  $\epsilon$  for the loss function  $\mathcal{L}$  is set to  $6 \times 10^{-3}$  and the trade-off parameter  $\gamma$  is set to 0.5. To ensure a more objective evaluation and to ignore the effect of inflated metrics caused by the high number of points in the aorta, the **coronary arteries are assessed independently**.

## 5. Results and analysis

### 5.1. Compared with point cloud-based methods

In this section, we present a comprehensive comparison between our TSRNet and the state-of-the-art point cloud-based completion methods. The evaluation is conducted on three benchmark datasets: PC-CAC, PC-ImageCAS, and PC-PTR. We report both quantitative and qualitative results, followed by detailed ablation studies to assess the contribution of each component in our framework. The analysis focuses on three key performance indicators.

---

**Algorithm 1:** Conversion of Coronary Artery Voxels into Corresponding Point Cloud from Clinical Data.

---

**input :** The annotation of voxel-based **aorta**  $y^a$  and **coronary**  $y^c$ ;  
 Numbers of patients  $K$ , numbers of points  $N$ , number of cases for each simulation  $M$ .  
**output:** The ground truth of point cloud-based vessels (including **aorta** and **coronary**)  $P_{GT}$ ;  
 The input of point cloud-based vessels (including **aorta** and **coronary**)  $P$ .

- 1 Eliminate cavities in  $y^c$ ;
- 2 Define the spherical structuring element  $e$  to control the extent of elimination;
- 3 Apply morphological operations (Dilation, Erosion) with structuring element  $e$ ;
- 4 Obtain centerline point cloud  $y^{lc}$  of **coronary**;
- 5 Obtain surface point cloud  $y^{sa}$  of **aorta**;
- 6 **for**  $k = 1 \dots K$  **do**
- 7    $y^a = MC(y^{sa})$ ;  $\triangleright$  MC represents the Marching Cubes.  $y^a$  refers to the point cloud surface of **aorta**.
- 8    $N_a = N - \text{PointNums}(y^{lc})$
- 9    $y^{sa} = \text{FPS}(y^a, N_a)$   $\triangleright$  FPS represents the Farthest Point Sampling.
- 10 Obtain main trunk point cloud  $y^t$  of **coronary**;  $\triangleright$  Remove irrelevant branches.
- 11 **for**  $k = 1 \dots K$  **do**
- 12   Calculate all pairs of endpoints  $(t_1, t_2)$  for  $y^{lc}$ ;
- 13   **for** each endpoint pair  $(t_1, t_2)$  **do**
- 14     Perform Depth-First Search from  $t_1$  to  $t_2$ , updating the shortest path;
- 15      $i = \text{GetIndex}(t_1, t_2)$ ;  $\triangleright$  Assign index  $i$  for each pair.
- 16     Record the shortest path  $y_i^t$ .
- 17    $y^t = \text{Union}\{y_i^t\}$ .
- 18 Normalize the point clouds  $y^t$  and  $y^{sa}$ ;
- 19 **for**  $k = 1 \dots K$  **do**
- 20   Calculate the centroid of  $\text{Union}\{y^t, y^{sa}\}$ ;
- 21   **for** each point  $p \in \text{Union}\{y^t, y^{sa}\}$  **do**
- 22     Subtract the centroid from point  $p$ ;
- 23     Scale the points to fit within the range  $[-1, 1]$ .
- 24 Obtain ground truth point cloud  $P_{GT}$  and simulated fractured input point clouds  $P$ ;
- 25 **for**  $k = 1 \dots K$  **do**
- 26    $P_{GT} = \text{Union}\{y^{sa}, y^t\}$ ,  $P = \emptyset$ ;
- 27   **for**  $m = 1 \dots M$  **do**
- 28      $R = \text{Random}(\text{min\_ratio}, \text{max\_ratio})$ ;  $\triangleright$  Proportion of points to be removed.
- 29      $B = \text{Random}(\text{min\_break}, \text{max\_break})$ ;  $\triangleright$  Number of regions to break.
- 30     Randomly partition  $y^t \times R$  points into  $B$  groups and store them in  $\text{remove}[B]$ ;
- 31      $y_{\text{copy}}^t = \text{Copy}(y^t)$ ;
- 32     **for**  $b = 1 \dots B$  **do**
- 33        $\text{connect} = \text{GetConnectedComponents}(y_{\text{copy}}^t)$ ;
- 34       Randomly delete a connected component of size  $\text{remove}[b]$  from  $y_{\text{copy}}^t$ ;
- 35        $\text{connect}' = \text{GetConnectedComponents}(y_{\text{copy}}^t)$ ;
- 36       **if**  $\text{connect}' == \text{connect}$  **then**
- 37          Retry the current deletion;  $\triangleright$  Did not cause a fracture, retry.
- 38      $P = P \cup \{y_{\text{copy}}^t\}$ ;
- 39 End, ground truth  $P_{GT}$  and  $M$  sets of input point clouds  $P$  are obtained.

---

### 5.1.1. Quantitative evaluation

**Comparative Results.** The advantages of our TSRNet on each metric are demonstrated in Table 2, and the results show that our framework consistently outperforms existing approaches across all three datasets: PC-CAC, PC-ImageCAS, and PC-PTR, achieving state-of-the-art

performance in  $CD^{\ell_1}$ , F1, and Fidelity Error. Specifically, our method attains the lowest  $CD^{\ell_1}$ , demonstrating its effectiveness in accurately reconstructing fine-grained structural details. Additionally, our model achieves the highest F1 on PC-CAC and PC-PTR, while maintaining the second-best F1 on PC-ImageCAS. Furthermore, our approach consistently outperforms other methods in Fidelity Error, with the lowest values observed across all datasets. Notably, the most striking improvement is in Fidelity Error in PC-PTR which consists of dense and numerous elongated tubular structures, where our model achieves 0.102, marking a **97.2% reduction** compared to the second-best model SeedFormer, demonstrating unparalleled structural preservation. This confirms that our model preserves the original structure more faithfully while reducing unwanted distortions. Compared to other recent methods such as AnchorFormer, PointAttN and CRA-PCN, our framework achieves more stable and better performance, particularly in complex vascular structures from PC-PTR, where detail preservation is critical. In addition to these geometric metrics, the topology-aware results in Table 3 further verify that TSRNet preserves structural connectivity more effectively. Specifically, our method achieves the highest OV on all three datasets, the lowest BE on all three datasets, and the best BD on two out of three datasets, indicating more complete vessel coverage and better topological consistency than competing methods. These results validate the effectiveness of our detail-preserved feature extractor, multiple dense refinement strategy, and global-to-local loss function, collectively contributing to the superior performance of our method.

**Ablation Study.** The results in Table 4 demonstrate that each module contributes to the overall performance, with the best results achieved when all three components are combined (Row D). By comparing Row A (baseline) and Row B (with the Multiple Dense Refinement Strategy enabled), we observe a substantial improvement in  $CD^{\ell_1}$  from 9.137 to 4.763, and an increase in F1 from 73.90% to 93.27%. This confirms that the Multiple Dense Refinement Strategy effectively enhances the completeness and precision of the reconstructed structures by progressively refining missing regions. Enabling the Global-to-Local Loss Function further improves performance, as seen in Row C, where  $CD^{\ell_1}$  decreases to 3.877, F1 rises to 94.39%, and Fidelity Error is also notably reduced, indicating that this loss function helps balance global structural consistency while preserving fine details. The integration of the Detail-Preserved Feature Extractor (Row D) leads to the best overall performance, achieving the lowest  $CD^{\ell_1}$  of 3.783, lowest Fidelity Error of 3.318, and highest F1 of 94.58%, which demonstrates that this module plays a crucial role in accurately reconstructing fine-grained tubular structures.

**Hyperparameter Ablation.** To evaluate the impact of key hyperparameters in TSRNet, we perform a comprehensive ablation study on the PC-CAC dataset, as summarized in Table 5. Specifically, we vary the global loss threshold  $\epsilon$ , the trade-off parameter  $\gamma$ , and the number of stacked SA+Transformer (S+T) blocks. Notably, Row 7 achieves the best overall performance with  $\epsilon = 2.7$  and  $\gamma = 0.5$ . We also explored the number of SA+Transformer blocks, the choice of 2 achieves the best result, and is thus adopted in our model.

**Loss Ablation.** To assess whether the use of Fidelity Error in our loss formulation introduces metric bias, we clarify that all competing baselines in the main comparisons were trained using the loss functions and optimization settings reported in their original papers, without the Fidelity Error term used in TSRNet. We further retrained the three top-performing baselines in our main comparison, namely SeedFormer, SnowflakeNet, and CRA-PCN, on PC-ImageCAS with the same Fidelity-related loss. As shown in Table 6, TSRNet still achieves the best overall performance, indicating that its advantage is not solely due to the use of Fidelity Error in training.

**Sensitivity Analysis of Input Point Number and Core Point Ratio.** To further examine the robustness of the proposed Core Points Selection module, we conducted a sensitivity analysis on both the number of input points and the selection ratio  $N_0$ . As shown in Table

**Table 2**

Quantitative results of different point cloud-based methods on datasets PC-CAC, PC-ImageCAS and PC-PTR. **Bold** indicates the best results, and Underlined indicates the second-best results.

| Method                            | PC-CAC                              |                                    |                                     | PC-ImageCAS                         |                                    |                                     | PC-PTR                              |                                    |                                     |
|-----------------------------------|-------------------------------------|------------------------------------|-------------------------------------|-------------------------------------|------------------------------------|-------------------------------------|-------------------------------------|------------------------------------|-------------------------------------|
|                                   | $CD^{\downarrow}$ ( $10^{-3}$ )     | F1 (%) $\uparrow$                  | Fidelity ( $10^{-3}$ ) $\downarrow$ | $CD^{\downarrow}$ ( $10^{-3}$ )     | F1 (%) $\uparrow$                  | Fidelity ( $10^{-3}$ ) $\downarrow$ | $CD^{\downarrow}$ ( $10^{-3}$ )     | F1 (%) $\uparrow$                  | Fidelity ( $10^{-3}$ ) $\downarrow$ |
| GRNet (Xie et al., 2020)          | 32.683 $\pm$ 9.488                  | 32.60 $\pm$ 4.95                   | 52.549 $\pm$ 17.918                 | 28.953 $\pm$ 5.411                  | 46.12 $\pm$ 4.15                   | 46.292 $\pm$ 10.401                 | 11.641 $\pm$ 0.836                  | 58.81 $\pm$ 4.57                   | 13.900 $\pm$ 0.856                  |
| PoinTr (Yu et al., 2021)          | 4.452 $\pm$ 0.896                   | 87.98 $\pm$ 10.76                  | 4.335 $\pm$ 1.634                   | 4.861 $\pm$ 0.920                   | 85.83 $\pm$ 12.61                  | 4.796 $\pm$ 1.422                   | 6.129 $\pm$ 1.510                   | 80.80 $\pm$ 15.72                  | 5.203 $\pm$ 2.572                   |
| SeedFormer (Zhou et al., 2022)    | <u>4.139 <math>\pm</math> 1.286</u> | <u>93.41 <math>\pm</math> 2.84</u> | <u>3.598 <math>\pm</math> 1.368</u> | <u>2.799 <math>\pm</math> 0.712</u> | <b>97.47 <math>\pm</math> 1.88</b> | <u>2.563 <math>\pm</math> 0.825</u> | <u>4.803 <math>\pm</math> 1.356</u> | <u>88.82 <math>\pm</math> 7.02</u> | <u>3.666 <math>\pm</math> 0.914</u> |
| SnowflakeNet (Xiang et al., 2022) | 5.661 $\pm$ 1.661                   | 88.11 $\pm$ 3.37                   | 5.816 $\pm$ 2.365                   | 2.979 $\pm$ 0.893                   | 95.89 $\pm$ 2.54                   | 2.836 $\pm$ 1.288                   | 6.033 $\pm$ 1.327                   | 83.06 $\pm$ 6.97                   | 5.666 $\pm$ 1.009                   |
| AnchorFormer (Chen et al., 2023)  | 6.747 $\pm$ 0.871                   | 83.11 $\pm$ 7.32                   | 6.178 $\pm$ 0.845                   | 4.699 $\pm$ 0.435                   | 87.68 $\pm$ 11.62                  | 4.392 $\pm$ 0.693                   | 11.491 $\pm$ 1.150                  | 48.48 $\pm$ 7.51                   | 9.237 $\pm$ 1.107                   |
| PointAttN (Wang et al., 2024)     | 8.288 $\pm$ 1.348                   | 67.22 $\pm$ 4.47                   | 6.660 $\pm$ 2.488                   | 5.757 $\pm$ 1.422                   | 83.84 $\pm$ 4.01                   | 4.595 $\pm$ 2.587                   | 11.387 $\pm$ 1.777                  | 54.67 $\pm$ 7.60                   | 9.324 $\pm$ 1.657                   |
| CRA-PCN (Rong et al., 2024)       | 5.869 $\pm$ 1.100                   | 88.75 $\pm$ 3.87                   | 6.245 $\pm$ 2.171                   | 3.526 $\pm$ 0.749                   | 93.54 $\pm$ 4.44                   | 4.414 $\pm$ 1.650                   | 9.088 $\pm$ 0.982                   | 66.01 $\pm$ 7.20                   | 8.950 $\pm$ 2.114                   |
| TSRNet (Ours)                     | <b>3.783 <math>\pm</math> 1.456</b> | <b>94.58 <math>\pm</math> 2.28</b> | <b>3.318 <math>\pm</math> 1.768</b> | <b>2.395 <math>\pm</math> 0.716</b> | <u>96.83 <math>\pm</math> 1.94</u> | <b>2.062 <math>\pm</math> 0.910</b> | <b>2.382 <math>\pm</math> 0.533</b> | <b>96.31 <math>\pm</math> 1.83</b> | <b>0.102 <math>\pm</math> 0.031</b> |

**Table 3**

Quantitative results of different point cloud-based methods on datasets PC-CAC, PC-ImageCAS and PC-PTR, evaluated using topology-aware metrics. **Bold** indicates the best results, and Underlined indicates the second-best results.

| Method        | PC-CAC                             |                                     |                                     | PC-ImageCAS                        |                                     |                                     | PC-PTR                             |                                     |                                     |
|---------------|------------------------------------|-------------------------------------|-------------------------------------|------------------------------------|-------------------------------------|-------------------------------------|------------------------------------|-------------------------------------|-------------------------------------|
|               | OV (%) $\uparrow$                  | BD ( $10^{-2}$ ) $\downarrow$       | BE $\downarrow$                     | OV (%) $\uparrow$                  | BD ( $10^{-2}$ ) $\downarrow$       | BE $\downarrow$                     | OV (%) $\uparrow$                  | BD ( $10^{-2}$ ) $\downarrow$       | BE $\downarrow$                     |
| GRNet         | 68.83 $\pm$ 8.44                   | 4.028 $\pm$ 1.344                   | 5.475 $\pm$ 3.376                   | 70.21 $\pm$ 8.41                   | 4.095 $\pm$ 0.573                   | 9.420 $\pm$ 14.469                  | 89.58 $\pm$ 1.79                   | 2.539 $\pm$ 0.824                   | 87.275 $\pm$ 138.998                |
| PoinTr        | 86.74 $\pm$ 1.98                   | 5.060 $\pm$ 0.719                   | 5.400 $\pm$ 2.083                   | 92.48 $\pm$ 2.65                   | 4.110 $\pm$ 0.685                   | 2.910 $\pm$ 1.425                   | 91.50 $\pm$ 1.33                   | 2.444 $\pm$ 0.728                   | 34.875 $\pm$ 25.639                 |
| SeedFormer    | <u>97.93 <math>\pm</math> 1.87</u> | <b>2.064 <math>\pm</math> 1.219</b> | <u>0.188 <math>\pm</math> 0.470</u> | <u>98.56 <math>\pm</math> 1.53</u> | <u>1.737 <math>\pm</math> 1.115</u> | <u>0.140 <math>\pm</math> 0.375</u> | <u>95.40 <math>\pm</math> 1.55</u> | <u>1.944 <math>\pm</math> 0.566</u> | <u>1.031 <math>\pm</math> 2.008</u> |
| SnowflakeNet  | 97.13 $\pm$ 1.94                   | 2.817 $\pm$ 1.778                   | 0.425 $\pm$ 0.667                   | 98.58 $\pm$ 2.04                   | 1.602 $\pm$ 1.170                   | 0.205 $\pm$ 0.493                   | 94.91 $\pm$ 1.10                   | 1.956 $\pm$ 0.792                   | 3.000 $\pm$ 5.089                   |
| AnchorFormer  | 85.32 $\pm$ 2.60                   | 3.340 $\pm$ 1.474                   | 0.916 $\pm$ 0.940                   | 89.65 $\pm$ 2.53                   | 3.498 $\pm$ 1.255                   | 0.425 $\pm$ 0.587                   | 78.54 $\pm$ 2.53                   | 2.124 $\pm$ 0.538                   | 61.425 $\pm$ 84.174                 |
| PointAttN     | 86.49 $\pm$ 2.74                   | 5.256 $\pm$ 1.585                   | 8.538 $\pm$ 3.634                   | 87.14 $\pm$ 3.59                   | 4.325 $\pm$ 1.137                   | 4.313 $\pm$ 0.982                   | 75.16 $\pm$ 10.73                  | 2.576 $\pm$ 0.716                   | 156.613 $\pm$ 188.114               |
| CRA-PCN       | 96.74 $\pm$ 2.51                   | 2.728 $\pm$ 1.260                   | 0.500 $\pm$ 0.870                   | 96.27 $\pm$ 3.23                   | 2.728 $\pm$ 1.627                   | 0.490 $\pm$ 0.538                   | 93.17 $\pm$ 2.34                   | 2.058 $\pm$ 0.510                   | 1.109 $\pm$ 1.359                   |
| TSRNet (ours) | <b>98.13 <math>\pm</math> 2.03</b> | <u>2.236 <math>\pm</math> 1.253</u> | <b>0.172 <math>\pm</math> 0.439</b> | <b>99.12 <math>\pm</math> 1.76</b> | <b>1.448 <math>\pm</math> 0.962</b> | <b>0.080 <math>\pm</math> 0.271</b> | <b>96.78 <math>\pm</math> 1.05</b> | <b>1.821 <math>\pm</math> 0.579</b> | <b>0.719 <math>\pm</math> 1.068</b> |

**Table 4**

Quantitative results of ablation analysis of different modules on our PC-CAC. ① represents the Multiple Dense Refinement Strategy. ② represents the Global-to-Local Loss Function. ③ represents the Detail-Preserved Feature Extractor.

|   | ① | ② | ③ | $CD^{\downarrow}$ ( $10^{-3}$ ) $\downarrow$ | F1 (%) $\uparrow$                  | Fidelity ( $10^{-3}$ ) $\downarrow$ |
|---|---|---|---|----------------------------------------------|------------------------------------|-------------------------------------|
| A |   |   |   | 9.137 $\pm$ 7.993                            | 73.90 $\pm$ 31.36                  | 7.950 $\pm$ 7.577                   |
| B | ✓ |   |   | 4.763 $\pm$ 0.593                            | 93.27 $\pm$ 3.40                   | 4.809 $\pm$ 0.654                   |
| C | ✓ | ✓ |   | 3.877 $\pm$ 1.684                            | 94.39 $\pm$ 2.45                   | 3.351 $\pm$ 1.809                   |
| D | ✓ | ✓ | ✓ | <b>3.783 <math>\pm</math> 1.456</b>          | <b>94.58 <math>\pm</math> 2.28</b> | <b>3.318 <math>\pm</math> 1.768</b> |

7, the default setting of 4096 points with  $N_0 = 1/4N$  achieves the best overall performance across all evaluation metrics. When the input size is reduced to 2048 points, the performance drops markedly, indicating that insufficient point sampling weakens the representation of complex tubular structures. Although increasing the input size to 8192 points improves the results compared with 2048 points, it does not surpass the default configuration, suggesting that simply using more points does not necessarily lead to better reconstruction.

In addition, we evaluated two alternative selection ratios under the 4096-point setting. Using  $N_0 = 1/2N$  yields competitive performance but remains slightly inferior to the default setting, while  $N_0 = 1/8N$  leads to a clear performance degradation across all metrics. These results suggest that selecting too few core points may discard important structural information, whereas the default ratio provides a better balance between preserving representative geometry and maintaining effective coarse-to-fine refinement. Overall, the analysis confirms that the proposed model is robust to moderate variations in point number and selection ratio, and supports the use of 4096 points with  $N_0 = 1/4N$  as the default configuration.

### 5.1.2. Qualitative evaluation

As shown in Fig. 4, six representative cases from the PC-CAC, PC-ImageCAS, and PC-PTR datasets illustrate the superior completion performance of our TSRNet compared to existing approaches. Our framework demonstrates better structural integrity, finer detail preservation, and improved robustness in reconstructing missing regions.

In our PC-CAC dataset (A-1, A-2), our method effectively restores complex and thin vascular structures, reducing discontinuities and achieving a more accurate topology compared to other methods. Referring to A-1 and A-2, it can be observed that PointAttN and CRA-PCN struggle to reconstruct and complete curved structures, leading to unintended fractures in originally intact regions. While Pointr performs relatively well in structural restoration, it fails to focus on elongated tubular structures, making completion infeasible. SeedFormer, SnowflakeNet, and AnchorFormer show comparatively better results. However, they still suffer from structural distortions, incomplete connections, and even incorrect reconstructions, particularly in complex or thin structures. In contrast, our method demonstrates superior robustness and performance, effectively preserving structural integrity while accurately reconstructing both complex and elongated tubular structures. Similarly, in the PC-ImageCAS dataset (B-1, B-2), our approach better captures the global shape while maintaining fine anatomical details, whereas other methods suffer from missing branches and distortions. For the PC-PTR dataset (C-1, C-2), which involves highly intricate pulmonary vessels, our model achieves remarkable consistency with the ground truth, accurately reconstructing fine-grained structures while avoiding excessive artifacts. In contrast, alternative methods exhibit structural collapse, missing connections, or distorted reconstructions, as highlighted by the red-marked error regions. These qualitative results further confirm the effectiveness of our detail-preserved feature extractor and multi-stage refinement strategy, enabling our model to accurately reconstruct complex tubular structures while minimizing artifacts and preserving structural coherence.

To further investigate the performance under structurally extreme conditions, we highlight two particularly challenging cases with **highly tortuous and complex anatomies** from our PC-CAC dataset (A-3, A-4). As shown in Fig. 5, these examples are characterized by extensive curvature and disconnections, posing significant challenges to existing methods. It is evident that PointAttN and CRA-PCN struggle with reconstructing curved structures and may even introduce errors in previously intact regions. While Pointr shows a reasonable ability to restore structure, it lacks sensitivity to fine tubular details, leading to incomplete reconstructions. SeedFormer, SnowflakeNet, and AnchorFormer yield relatively better results but still exhibit structural distortions, missing connections, and sometimes incorrect linkages, particularly in complex

**Table 5**

Ablation study of key hyperparameters. We vary the global loss threshold  $\epsilon$ , the trade-off parameter  $\gamma$ , and the number of SA+Transformer (S+T) blocks. The best performance is highlighted in bold.

|   | $\epsilon$ | $\gamma$   | (S+T)    | $CD^{\ell_1}$ ( $10^{-3}$ ) ↓ | F1 (%) ↑            | Fidelity ( $10^{-3}$ ) ↓ |
|---|------------|------------|----------|-------------------------------|---------------------|--------------------------|
| 1 | 3.1        | 0.5        | 2        | 2.485 ± 0.718                 | 96.06 ± 2.11        | 2.254 ± 0.970            |
| 2 | 2.9        | 0.5        | 2        | 2.405 ± 0.716                 | 96.60 ± 1.97        | 2.084 ± 0.912            |
| 3 | 2.7        | 0.3        | 2        | 2.502 ± 0.719                 | 96.56 ± 2.00        | 2.081 ± 1.200            |
| 4 | 2.7        | 0.7        | 2        | 2.457 ± 0.718                 | 96.79 ± 1.94        | 2.120 ± 0.930            |
| 5 | 2.7        | 0.5        | 4        | 2.513 ± 0.712                 | 96.34 ± 1.96        | 2.277 ± 0.945            |
| 6 | 2.7        | 0.5        | 8        | 2.575 ± 0.717                 | 95.24 ± 2.72        | 2.409 ± 1.017            |
| 7 | <b>2.7</b> | <b>0.5</b> | <b>2</b> | <b>2.395 ± 0.716</b>          | <b>96.83 ± 1.94</b> | <b>2.062 ± 0.910</b>     |

**Table 6**

Analysis on PC-ImageCAS: quantitative comparison of representative competing methods retrained with the Fidelity-related loss. “w/o” denotes the original training loss reported in the corresponding paper, while “w/” denotes retraining with the added Fidelity-related loss.

| Method        | Fidelity loss | Metrics                       |                     |                          |                     |                      |                      |
|---------------|---------------|-------------------------------|---------------------|--------------------------|---------------------|----------------------|----------------------|
|               |               | $CD^{\ell_1}$ ( $10^{-3}$ ) ↓ | F1 (%) ↑            | Fidelity ( $10^{-3}$ ) ↓ | OV (%) ↑            | BD ( $10^{-2}$ ) ↓   | BE ↓                 |
| SeedFormer    | w/o           | 2.799 ± 0.712                 | 97.47 ± 1.88        | 2.563 ± 0.825            | 98.56 ± 1.53        | 1.737 ± 1.115        | 1.400 ± 0.375        |
| SeedFormer    | w/            | 2.537 ± 0.598                 | 97.79 ± 1.85        | 2.280 ± 0.735            | 99.06 ± 1.28        | 1.384 ± 1.217        | 0.950 ± 0.263        |
| SnowflakeNet  | w/o           | 2.979 ± 0.893                 | 95.89 ± 2.54        | 2.836 ± 1.288            | 98.58 ± 2.04        | 1.602 ± 1.170        | 2.050 ± 0.493        |
| SnowflakeNet  | w/            | 3.297 ± 0.968                 | 95.25 ± 2.48        | 3.406 ± 1.546            | 98.32 ± 1.63        | 1.893 ± 1.087        | 2.400 ± 0.427        |
| CRA-PCN       | w/o           | 3.526 ± 0.749                 | 93.54 ± 4.44        | 4.414 ± 1.650            | 96.27 ± 3.23        | 2.728 ± 1.627        | 4.900 ± 0.538        |
| CRA-PCN       | w/            | 3.374 ± 0.728                 | 92.96 ± 4.64        | 4.467 ± 1.542            | 98.47 ± 1.94        | 1.972 ± 1.388        | 1.850 ± 0.459        |
| TSRNet (ours) | w/            | <b>2.395 ± 0.716</b>          | <b>96.83 ± 1.94</b> | <b>2.062 ± 0.910</b>     | <b>99.12 ± 1.76</b> | <b>1.448 ± 0.962</b> | <b>0.800 ± 0.271</b> |

**Table 7**

Sensitivity analysis of the Core Points Selection module under different input point numbers and selection ratios.

| Num of points | Ratio         | $CD^{\ell_1}$ ( $10^{-3}$ ) ↓ | F1 (%) ↑            | Fidelity ( $10^{-3}$ ) ↓ | OV (%) ↑            | BD ( $10^{-2}$ ) ↓   | BE ↓                 |
|---------------|---------------|-------------------------------|---------------------|--------------------------|---------------------|----------------------|----------------------|
| 2048          | $N_0 = 1/4 N$ | 3.927 ± 0.451                 | 82.22 ± 4.62        | 4.167 ± 1.893            | 92.10 ± 3.46        | 3.876 ± 1.141        | 8.750 ± 0.980        |
| 8192          | $N_0 = 1/4 N$ | 2.561 ± 0.260                 | 94.98 ± 2.27        | 2.424 ± 1.192            | 97.59 ± 2.64        | 2.014 ± 0.887        | 2.300 ± 0.337        |
| 4096          | $N_0 = 1/2 N$ | 2.403 ± 0.507                 | 95.23 ± 1.96        | 2.354 ± 1.177            | 98.95 ± 2.78        | 1.828 ± 1.049        | 2.150 ± 0.510        |
| 4096          | $N_0 = 1/8 N$ | 3.980 ± 1.265                 | 91.30 ± 1.93        | 3.190 ± 1.215            | 92.10 ± 3.46        | 3.876 ± 1.141        | 8.750 ± 0.980        |
| 4096          | $N_0 = 1/4 N$ | <b>2.395 ± 0.716</b>          | <b>96.83 ± 1.94</b> | <b>2.062 ± 0.910</b>     | <b>99.12 ± 1.76</b> | <b>1.448 ± 0.962</b> | <b>0.800 ± 0.271</b> |

and elongated vascular regions. In contrast, our method consistently achieves more accurate, well-connected, and robust reconstructions, more closely aligning with the ground truth.

## 5.2. Compared with voxel-based methods

To further evaluate the effectiveness of our TSRNet, we compare it against voxel-based tubular structure completion methods. Specifically, we select two recent and representative baselines, DRTT and VSR-Net, which have demonstrated good performance in voxel-domain vascular reconstructions. Both quantitative and qualitative comparisons are conducted on the public PC-ImageCAS dataset for fairness.

### 5.2.1. Quantitative evaluation

**Comparative Results.** The quantitative results are presented in Table 8. Our TSRNet achieves the best performance across all three evaluation metrics. Specifically, our method obtains a  $CD^{\ell_1}$  of 2.395, which is significantly lower than that of VSR-Net (2.529) and DRTT (2.869), indicating superior accuracy in structural alignment. In addition, TSRNet achieves the highest F1-score of 96.83% and the lowest Fidelity Error of 2.062, reflecting its ability to preserve both geometric detail and global continuity. These results demonstrate that operating directly in the point cloud domain provides advantages over voxel-based methods, particularly in terms of structural fidelity and precise reconnection. In addition to reconstruction accuracy, we also compare the computational efficiency of TSRNet with the voxel-based baselines. TSRNet achieves substantially lower inference time and GPU memory usage than both DRTT and VSR-Net. Specifically, TSRNet requires only

0.860 s per case, compared with 13.003 s for DRTT and 15.97 s for VSR-Net. In terms of GPU memory usage, TSRNet consumes 242.4 M, whereas DRTT and VSR-Net require 334 M and 10,753 M, respectively. These results provide empirical support for the computational efficiency of the proposed point-cloud-based representation in fractured tubular structure reconstruction.

### 5.2.2. Qualitative evaluation

Visual comparisons of different methods are presented in Fig. 6, where five representative cases from the public PC-ImageCAS dataset are selected to evaluate completion performance under various structural challenges. We focus particularly on vascular regions that are difficult to reconnect. For instance, certain cases from PC-ImageCAS exhibit long-gap structural discontinuities that pose significant reconstruction difficulties for existing methods. Our TSRNet consistently outperforms other approaches in restoring continuity and preserving anatomical fidelity across all datasets. The visual results demonstrate that our framework achieves better structural integrity, finer detail preservation, and improved robustness in reconstructing missing regions.

In the PC-ImageCAS dataset (B-1 to B-5), each row shows a different anatomical structure with varying types and severities of vascular disconnection. DRTT relies on a centerline template as input guidance. When this centerline contains fractures, the model fails to generate any meaningful reconnection in the corresponding regions, as evident in cases such as B-1 to B-5. As a result, large disconnected segments remain unreconstructed. In contrast, VSR-Net is capable of handling some short-range disconnections and successfully reconnects smaller

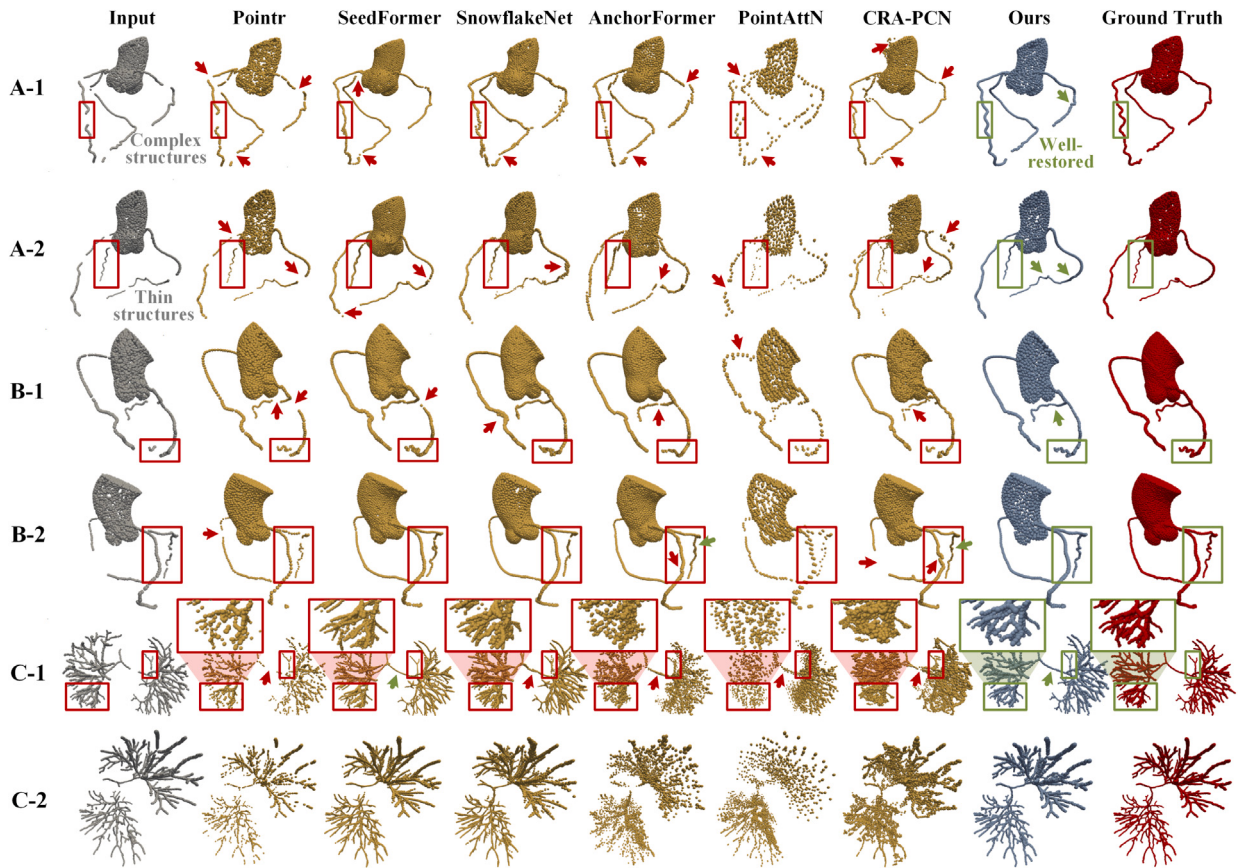


Fig. 4. Visual comparison of different point cloud-based models. A-1, A-2 are two cases from our PC-CAC dataset. B-1, B-2 are two cases from PC-ImageCAS dataset. C-1, C-2 are two cases from PC-PTR dataset.

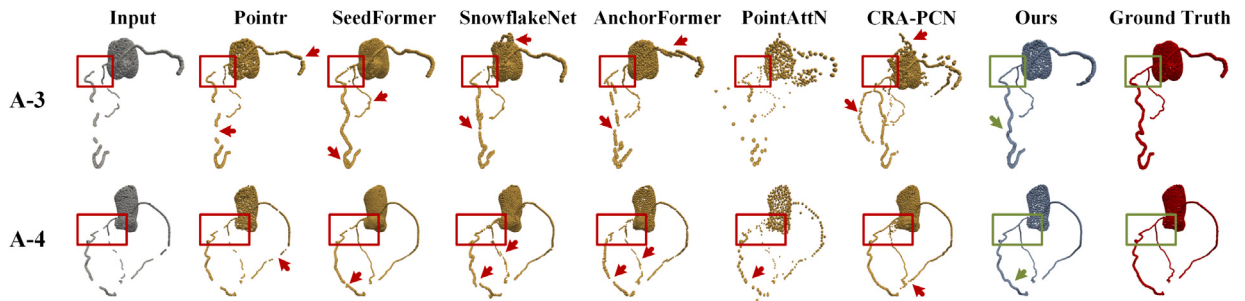


Fig. 5. Challenging cases with highly tortuous and complex anatomies. A-3,4 are two cases from our PC-CAC dataset.

fractures. However, it struggles in cases involving long-gap or high-curvature discontinuities (B-2 and B-4), where the restored structures remain incomplete or topologically inaccurate. Our method, by operating directly on the point cloud without requiring prior centerline input, demonstrates superior adaptability across all cases. In particular, in B-4, which includes both long-gap and short-gap fractures, our TSRNet successfully restores both segments and achieves a structurally consistent result that closely matches the ground truth. The qualitative results highlight the robustness of our model in restoring complex vascular geometry under varying levels of structural loss.

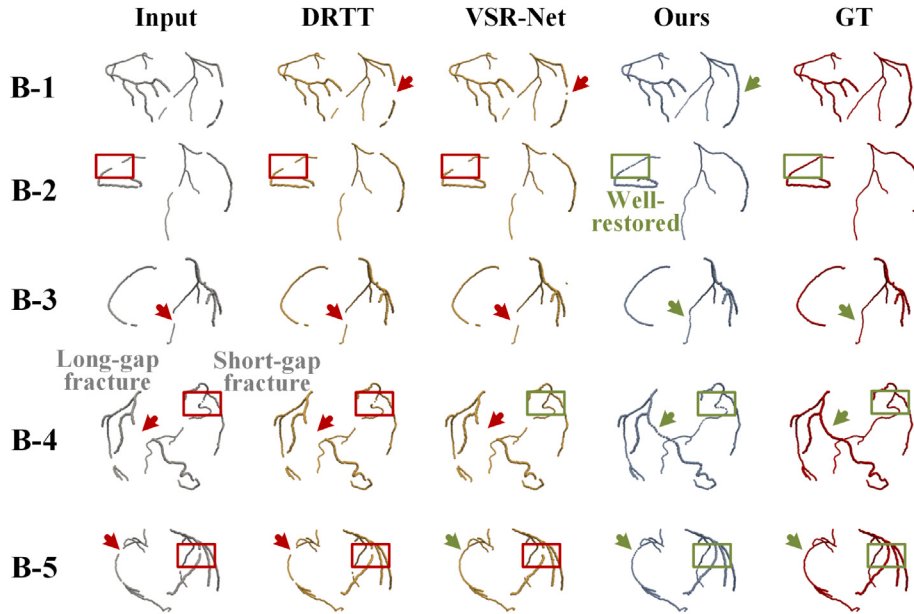
### 5.3. Model analysis – Completion capability under different fracture severity

To further explore the completion limits of each model and objectively evaluate the reconnection performance of different methods

under varying levels of structural disruption, we simulate fractures at different severities (10%, 20%, and 30%) and compare the completion results across multiple baseline methods. The bottom-right section of Fig. 7 illustrates a specific example simulating three different levels of fractures. At 10% fracture severity, a single fracture occurs in a complex and elongated region. At 20% fracture severity, multiple fractures emerge at critical locations such as thin structures and bifurcation points. At 30% fracture severity, a primary trunk is almost entirely fragmented, leaving only scattered points. Both the line and bar chart in Fig. 7 illustrate our model consistently outperforms others, achieving the lowest  $CD^{\dagger}$  and Fidelity Error, while maintaining the highest F1 across all fracture ratios. Notably, as the severity of fractures increases, competing methods exhibit a significant decline in performance, struggling with structural continuity and detail preservation. In contrast, our model remains robust, demonstrating superior completion accuracy and preserving fine structures even under severe fractures, highlighting

**Table 8**Quantitative results of different voxel-based methods on public dataset PC-ImageCAS. **Bold** indicates the best results.

|                           | $CD^{\ell_1}$ ( $10^{-3}$ ) ↓       | F1 (%) ↑                           | Fidelity ( $10^{-3}$ ) ↓            | OV (%) ↑                           | BD ( $10^{-2}$ ) ↓                  | BE ↓                                | Inference time ↓ | GPU usage ↓    |
|---------------------------|-------------------------------------|------------------------------------|-------------------------------------|------------------------------------|-------------------------------------|-------------------------------------|------------------|----------------|
| DRTT (Li et al., 2021)    | $2.869 \pm 0.630$                   | $96.41 \pm 1.03$                   | $3.765 \pm 0.627$                   | $95.75 \pm 1.60$                   | $3.391 \pm 0.375$                   | $2.800 \pm 0.872$                   | 13.003 s         | 334 M          |
| VSR-Net (Ye et al., 2025) | $2.529 \pm 0.719$                   | $96.54 \pm 1.15$                   | $3.205 \pm 0.910$                   | $96.34 \pm 2.18$                   | $3.530 \pm 0.463$                   | $2.300 \pm 1.100$                   | 15.97 s          | 10 753 M       |
| Ours                      | <b><math>2.395 \pm 0.716</math></b> | <b><math>96.83 \pm 1.94</math></b> | <b><math>2.062 \pm 0.910</math></b> | <b><math>99.12 \pm 1.76</math></b> | <b><math>1.448 \pm 0.962</math></b> | <b><math>0.080 \pm 0.271</math></b> | <b>0.860 s</b>   | <b>242.4 M</b> |

**Fig. 6.** Visual comparison of different voxel-based models. B-1 to B-5 are 5 cases from the public PC-ImageCAS dataset. We focus on **vascular regions** that are particularly difficult to reconnect. In particular, cases B-2 and B-4 involve long-gap structural discontinuities that pose significant challenges for reconstruction.

its strong generalization capability in reconstructing complex tubular structures.

#### 5.4. Model analysis – Impact of fracture type and severity

To further investigate the robustness of different models across diverse structural categories, we evaluate the reconnection performance on three representative types of anatomical regions: **opening**, **bifurcation**, and **terminal** segments. As shown in Fig. 8, we simulate fractures with increasing severity (Grade 1–3), corresponding to slight, moderate, and severe levels of disconnection. Each row represents a distinct structural region, while each column pair compares the performance of SeedFormer (the second best model) and our TSRNet under different fracture conditions. In **Grade 1 (Slight Fracture)**, all models generally perform well in reconnecting openings and bifurcation points, where disconnections are minor and structural cues remain clear. However, in terminal regions with highly tortuous geometries, SeedFormer often fails to maintain geometric fidelity, resulting in disrupted or overly smoothed reconstructions. In contrast, our TSRNet preserves the integrity of these curved structures more effectively. In **Grade 2 (Moderate Fracture)**, the fractures predominantly occur around complex bifurcation regions, where multiple branches intersect at varying angles. SeedFormer typically reconnects only a single dominant branch, failing to capture the full bifurcation topology. In contrast, our method accurately reconstructs all intersecting branches, maintaining topological completeness and structural coherence across the bifurcation. In **Grade 3 (Severe Fracture)**, as the degree of disconnection increases, the reconstruction task becomes significantly more challenging due to the loss of large structural segments and spatial cues. Under such extreme conditions, SeedFormer suffers from unstable predictions, including fragmented outputs. In contrast, our method maintains stable and coherent reconstructions, consistently producing anatomically plausible results that align more closely with the ground truth.

#### 5.5. Model analysis – External real-clinical cases

Although the benchmark is constructed using controlled simulated fractures to provide a standardized and reproducible setting, we further validated TSRNet on previously unseen real clinical CTA cases from 67 patients to assess its practical relevance. For the external CTA cases, we selected representative examples involving blurred vessel appearance, myocardial bridging, and imaging artifacts. The corresponding segmentation inputs were generated by nnU-Net (Isensee et al., 2021). These conditions often lead to local lumen narrowing, ambiguous vessel boundaries, and apparent discontinuities in clinical imaging. As shown in the added examples, TSRNet is able to reconnect incomplete vessel structures more consistently than the competing method, yielding reconstructions that are visually closer to the reference anatomy. In particular, the model preserves vessel continuity in regions where the segmentation output is fragmented, suggesting that the proposed framework can better handle clinically realistic incompleteness.

To further quantify the realism of the simulated benchmark, we directly tested the proposed method on these unseen clinical CTA cases without retraining (See Fig. 9). As shown in Table 9, the raw segmentation results exhibit limited structural completeness, and SeedFormer provides only modest improvement. In contrast, TSRNet achieves the best overall performance on most metrics, with the lowest  $CD^{\ell_1}$ , Fidelity Error, and BE, as well as the highest F1 and OV. These results suggest that the proposed framework can better reconnect real clinical discontinuities and preserve vessel continuity under practical imaging conditions, providing complementary evidence that the simulated benchmark remains relevant to clinically realistic scenarios.

From a clinical perspective, the proposed framework may be valuable in several aspects. **First**, as a point cloud-based post-processing module, TSRNet can be applied to vessel representations obtained from different acquisition settings or imaging types, potentially without

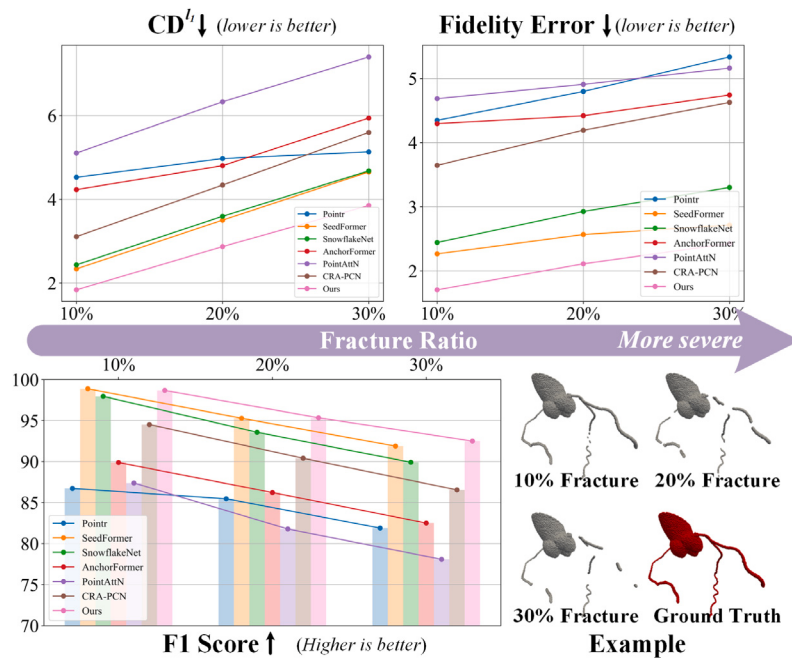


Fig. 7. Performance comparison under varying fracture ratios (10%, 20%, and 30%) in our PC-CAC dataset.

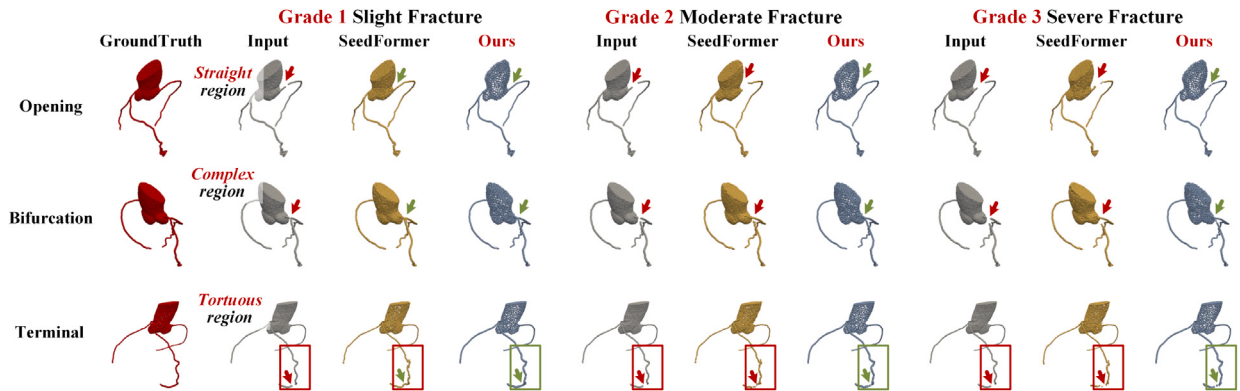


Fig. 8. Visual comparison of reconnection results across different fracture severities and structural types. We compare the performance of SeedFormer (the second best) and our method under three grades of fracture severity: slight (Grade 1), moderate (Grade 2), and severe (Grade 3).

task-specific retraining, making it suitable for downstream applications that require structurally complete vessels, such as centerline extraction, vessel tracking, branch-level analysis, and quantitative measurement. **Second**, in cases affected by imaging artifacts, low contrast, or generally low image quality, segmentation models are more likely to produce fragmented vessel outputs. In such scenarios, TSRNet can improve structural continuity and completeness, thereby providing a more reliable anatomical representation for subsequent analysis and interpretation. **Third**, the same framework is also potentially relevant to non-contrast CT, where the absence of vascular enhancement makes vessel delineation substantially more difficult. Since non-contrast CT is widely used in routine screening, baseline examination, and calcification-related assessment, the ability to recover incomplete vessel structures in this setting may further expand the practical utility of the proposed method.

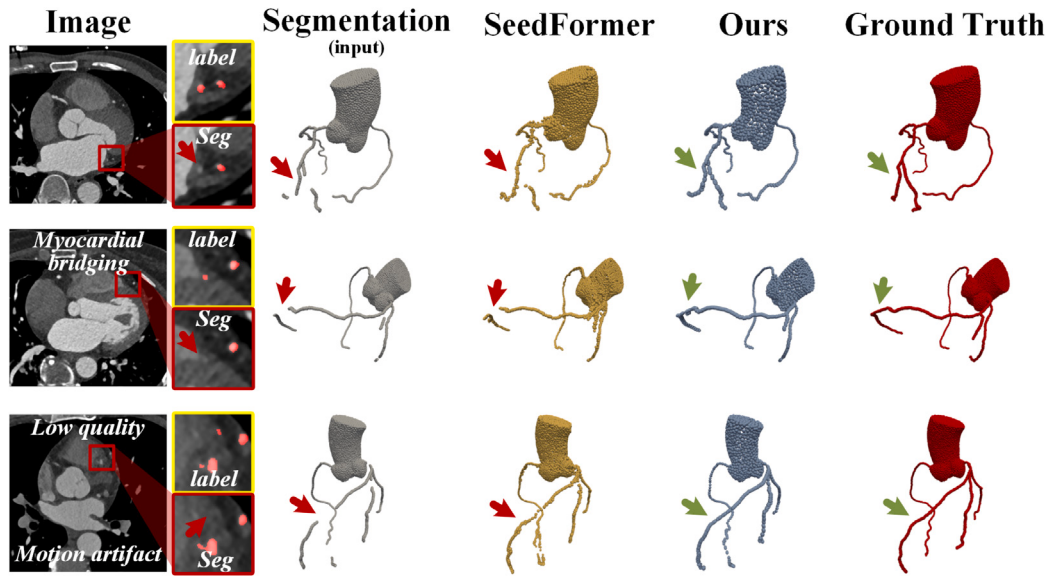
### 5.6. Future works

Looking ahead, point cloud-based tubular structure completion offers a strong foundation for preserving anatomical integrity. Building upon this, integrating point cloud representations with corresponding

voxel or image-based modalities holds great potential for further refining centerline estimation and improving structural accuracy. Such multi-modal fusion is expected to enhance model stability and provide more reliable reconstructions for downstream clinical applications such as diagnosis, simulation, or intervention planning. We will continue to release updates to the PC-CAC dataset, including the addition of paired image data, to promote ongoing research and foster broader advancements in this domain.

## 6. Discussion and conclusion

Structural discontinuity remains a fundamental challenge in medical image analysis of tubular structures, severely affecting the accuracy of downstream tasks such as flow simulation and structural quantification. Existing methods often fail to preserve anatomical coherence when dealing with complex morphologies or long-gap fractures. To address this, we propose a dedicated tubular structure completion framework based on point cloud modeling, which significantly improves the capacity to restore fine-grained topology across different types of structural degradation. In addition, we release a novel dataset PC-CAC and standardized benchmark to facilitate reproducible research



**Fig. 9.** Validation on external real cases. The figure shows three representative clinical scenarios, including blurred vessel appearance, myocardial bridging, and imaging artifacts. From left to right, each row presents the original image, the input segmentation, the reconstruction results of SeedFormer and our method, and the reference annotation.

**Table 9**

Quantitative comparison on unseen clinical CTA cases for external validation, evaluated without retraining on the external dataset. **Bold** indicates the best results and underlined indicates the second best results.

|               | $CD^d$ ( $10^{-3}$ ) ↓              | F1 (%) ↑                           | Fidelity ( $10^{-3}$ ) ↓            | OV (%) ↑                           | BD ( $10^{-2}$ ) ↓                  | BE ↓                                |
|---------------|-------------------------------------|------------------------------------|-------------------------------------|------------------------------------|-------------------------------------|-------------------------------------|
| Segmentation  | $5.293 \pm 3.609$                   | $92.42 \pm 4.35$                   | $3.283 \pm 5.386$                   | $98.36 \pm 1.25$                   | $2.314 \pm 1.359$                   | $0.657 \pm 1.100$                   |
| SeedFormer    | $5.075 \pm 2.164$                   | $91.19 \pm 3.99$                   | $3.941 \pm 1.480$                   | $98.48 \pm 1.18$                   | <b><math>1.827 \pm 1.490</math></b> | $0.299 \pm 0.547$                   |
| TSRNet (ours) | <b><math>4.872 \pm 2.206</math></b> | <b><math>92.83 \pm 3.98</math></b> | <b><math>3.177 \pm 1.673</math></b> | <b><math>98.51 \pm 1.13</math></b> | $1.972 \pm 1.254$                   | <b><math>0.283 \pm 0.541</math></b> |

and fair comparison in tubular structure reconnection tasks. Experimental results across three datasets (PC-CAC, PC-ImageCAS, and PC-PTR) confirm the robustness and generalizability of our approach in handling diverse anatomical patterns. Our TSRNet integrates a detail-preserved feature extractor, a progressive dense refinement strategy, and a global-to-local loss formulation, enabling accurate restoration of severely disconnected tubular structures while minimizing topological distortion. These findings underscore the importance of topology-aware modeling for medical imaging and suggest a promising direction for integrating structure-completion modules into broader diagnostic pipelines and patient-specific analysis systems.

In conclusion, this study introduces a novel approach for tubular structure reconnection from a unique point cloud perspective, addressing the limitations of traditional voxel-based methods in handling discontinuities within complex tubular structures. To our best knowledge, we construct the Point Cloud-based Coronary Artery Completion (PC-CAC) dataset, derived from real clinical data for the first time, providing a new benchmark for tubular structure reconnection tasks. Specifically, we propose the TSRNet, a dedicated Tubular Structure Reconnection Network, which integrates a detail-preserved feature extractor, a multiple dense refinement strategy, and a global-to-local loss function. These components collectively enhance structural integrity, detail preservation, and the reconstruction of missing regions. Experiments on PC-CAC and two additional public datasets (PC-ImageCAS and PC-PTR) validate the effectiveness of our method, demonstrating state-of-the-art performance. Our method consistently achieves the lowest Chamfer Distance and Fidelity Error while maintaining the highest F1-score, confirming its superiority in restoring fragmented tubular structures with high precision. Additionally, our model remains robust even under high-severity fractures, proving its strong capability. Overall, this work lays a solid step towards structure-aware modeling of fractured anatomy in medical imaging, paving the way for more effective and equitable applications in clinical diagnosis.

#### CRediT authorship contribution statement

**Yaolei Qi:** Writing – original draft, Software, Resources, Methodology, Investigation, Formal analysis, Data curation, Conceptualization. **Yikai Yang:** Visualization, Validation, Software, Methodology, Formal analysis, Data curation, Conceptualization. **Wenbo Peng:** Visualization, Validation, Conceptualization. **Shumei Miao:** Resources, Data curation. **Yutao Hu:** Writing – review & editing, Supervision, Resources, Project administration, Investigation, Formal analysis. **Guanyu Yang:** Writing – review & editing, Supervision, Project administration, Funding acquisition, Conceptualization.

#### Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

#### Acknowledgments

This research was supported by the National Key Research and Development Program of China (2024YFC2417803), National Natural Science Foundation of China (82441021), Jiangsu Province Cutting-edge Technology R&D Program (BF2025628), in part by the Funded by Basic Research Program of Jiangsu under Grant (BK20251279), the Postgraduate Research & Practice Innovation Program of Jiangsu Province, the Fundamental Research Funds for the Central Universities, China (KYCX22\_0239), and the Big Data Computing Center of Southeast University.

## References

- Beetz, M., Banerjee, A., Ossenbeng-Engels, J., Grau, V., 2023. Multi-class point cloud completion networks for 3D cardiac anatomy reconstruction from cine magnetic resonance images. *Med. Image Anal.* 90, 102975.
- Bernard, F., Salamanca, L., Thunberg, J., Tack, A., Jentsch, D., Lamecker, H., Zachow, S., Hertel, F., Goncalves, J., Gemmar, P., 2017. Shape-aware surface reconstruction from sparse 3D point-clouds. *Med. Image Anal.* 38, 77–89.
- Chen, Z., Long, F., Qiu, Z., Yao, T., Zhou, W., Luo, J., Mei, T., 2023. Anchorformer: Point cloud completion from discriminative nodes. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 13581–13590.
- Chen, X., Ravikumar, N., Xia, Y., Attar, R., Diaz-Pinto, A., Piechnik, S.K., Neubauer, S., Petersen, S.E., Frangi, A.F., 2021. Shape registration with learned deformations for 3D shape reconstruction from sparse and incomplete point clouds. *Med. Image Anal.* 74, 102228.
- Cohen-Steiner, D., Edelsbrunner, H., Harer, J., 2005. Stability of persistence diagrams. In: *Proceedings of the Twenty-First Annual Symposium on Computational Geometry*. pp. 263–271.
- Fan, H., Su, H., Guibas, L.J., 2017. A point set generation network for 3d object reconstruction from a single image. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*. pp. 605–613.
- Fei, B., Yang, W., Chen, W.-M., Li, Z., Li, Y., Ma, T., Hu, X., Ma, L., 2022. Comprehensive review of deep learning-based 3d point cloud completion processing and analysis. *IEEE Trans. Intell. Transp. Syst.* 23 (12), 22862–22883.
- Garje, A., Adhav, Y., Bodas, D., 2015. Design and simulation of blocked blood vessel for early detection of heart diseases. In: *2015 2nd International Symposium on Physics and Technology of Sensors. ISPTS, IEEE*. pp. 204–208.
- Gharleghi, R., Chen, N., Sowmya, A., Beier, S., 2022. Towards automated coronary artery segmentation: A systematic review. *Comput. Methods Programs Biomed.* 225, 107015.
- Grady, L., 2006. Random walks for image segmentation. *IEEE Trans. Pattern Anal. Mach. Intell.* 28 (11), 1768–1783.
- Gupta, S., Zhang, Y., Hu, X., Prasanna, P., Chen, C., 2023. Topology-aware uncertainty for image segmentation. *Adv. Neural Inf. Process. Syst.* 36, 8186–8207.
- Han, K., Jeon, J., Jang, Y., Jung, S., Kim, S., Shim, H., Jeon, B., Chang, H.-J., 2022. Reconnection of fragmented parts of coronary arteries using local geometric features in X-ray angiography images. *Comput. Biol. Med.* 141, 105099.
- Hu, X., Li, F., Samaras, D., Chen, C., 2019. Topology-preserving deep image segmentation. *Adv. Neural Inf. Process. Syst.* 32.
- Isensee, F., Jaeger, P.F., Kohl, S.A., Petersen, J., Maier-Hein, K.H., 2021. nnU-Net: A self-configuring method for deep learning-based biomedical image segmentation. *Nature Methods* 18 (2), 203–211.
- Kerber, M., Morozov, D., Nigmatov, A., 2017. Geometry helps to compare persistence diagrams.
- Ko, B.S., Cameron, J.D., Munnur, R.K., Wong, D.T., Fujisawa, Y., Sakaguchi, T., Hirohata, K., Hislop-Jambrich, J., Fujimoto, S., Takamura, K., Crossett, M., Leung, M., Kuganesan, A., Malaiapan, Y., Nasis, A., Troupis, J., Meredith, I.T., Seneviratne, S.K., 2017. Noninvasive CT-derived FFR based on structural and fluid analysis: A comparison with invasive FFR for detection of functionally significant stenosis. *JACC: Cardiovasc. Imaging* 10 (6), 663–673.
- Kong, B., Wang, X., Bai, J., Lu, Y., Gao, F., Cao, K., Xia, J., Song, Q., Yin, Y., 2020. Learning tree-structured representation for 3D coronary artery segmentation. *Comput. Med. Imaging Graph.* 80, 101688.
- Lee, M.S., Chen, C.-H., 2015. Myocardial bridging: An up-to-date review. *J. Invasive Cardiol.* 27 (11), 521.
- Li, Y., Gong, H., Wu, W., Liu, G., Chen, G., 2015. An automated method using hessian matrix and random walks for retinal blood vessel segmentation. In: *2015 8th International Congress on Image and Signal Processing. CISP, IEEE*. pp. 423–427.
- Li, H., Tang, Z., Nan, Y., Yang, G., 2022. Human treelike tubular structure segmentation: A comprehensive review and future perspectives. *Comput. Biol. Med.* 151, 106241.
- Li, Z., Xia, Q., Hu, Z., Wang, W., Xu, L., Zhang, S., 2021. A deep reinforced tree-traversal agent for coronary artery centerline extraction. In: *International Conference on Medical Image Computing and Computer-Assisted Intervention*. Springer, pp. 418–428.
- Madhavan, M.V., Tarigopula, M., Mintz, G.S., Maehara, A., Stone, G.W., G  n  reux, P., 2014. Coronary artery calcification: pathogenesis and prognostic implications. *J. Am. Coll. Cardiol.* 63 (17), 1703–1714.
- Mou, L., Chen, L., Cheng, J., Gu, Z., Zhao, Y., Liu, J., 2019. Dense dilated network with probability regularized walk for vessel detection. *IEEE Trans. Med. Imaging* 39 (5), 1392–1403.
- Pijls, N.H., Sels, J.-W.E., 2012. Functional measurement of coronary stenosis. *J. Am. Coll. Cardiol.* 59 (12), 1045–1057.
- Qi, Y., He, Y., Qi, X., Zhang, Y., Yang, G., 2023. Dynamic snake convolution based on topological geometric constraints for tubular structure segmentation. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision*. pp. 6070–6079.
- Qi, C.R., Su, H., Mo, K., Guibas, L.J., 2017a. Pointnet: Deep learning on point sets for 3d classification and segmentation. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*. pp. 652–660.
- Qi, C.R., Yi, L., Su, H., Guibas, L.J., 2017b. Pointnet++: Deep hierarchical feature learning on point sets in a metric space. *Adv. Neural Inf. Process. Syst.* 30, 5105–5114.
- Qiu, Y., Li, Z., Wang, Y., Dong, P., Wu, D., Yang, X., Hong, Q., Shen, D., 2023. Corsegrec: A topology-preserving scheme for extracting fully-connected coronary arteries from ct angiography. In: *International Conference on Medical Image Computing and Computer-Assisted Intervention*. Springer, pp. 670–680.
- Rong, Y., Zhou, H., Yuan, L., Mei, C., Wang, J., Lu, T., 2024. Cra-pcn: Point cloud completion with intra- and inter-level cross-resolution transformers. In: *Proceedings of the AAAI Conference on Artificial Intelligence*. Vol. 38, pp. 4676–4685.
- Schaap, M., Metz, C.T., van Walsum, T., van der Giessen, A.G., Weustink, A.C., Mollet, N.R., Bauer, C., Bogunovi  , H., Castro, C., Deng, X., et al., 2009. Standardized evaluation methodology and reference database for evaluating coronary artery centerline extraction algorithms. *Med. Image Anal.* 13 (5), 701–714.
- Tatarchenko, M., Richter, S.R., Ranftl, R., Li, Z., Koltun, V., Brox, T., 2019. What do single-view 3d reconstruction networks learn? In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 3405–3414.
- Thanvi, B., Robinson, T., 2007. Complete occlusion of extracranial internal carotid artery: Clinical features, pathophysiology, diagnosis and management. *Postgrad. Med. J.* 83 (976), 95–99.
- Wang, J., Cui, Y., Guo, D., Li, J., Liu, Q., Shen, C., 2024. Pointattn: You only need attention for point cloud completion. In: *Proceedings of the AAAI Conference on Artificial Intelligence*. Vol. 38, pp. 5472–5480.
- Wen, X., Xiang, P., Han, Z., Cao, Y.-P., Wan, P., Zheng, W., Liu, Y.-S., 2022. PMP-Net++: Point cloud completion by transformer-enhanced multi-step point moving paths. *IEEE Trans. Pattern Anal. Mach. Intell.* 45 (1), 852–867.
- Weng, Z., Yang, J., Liu, D., Cai, W., 2023. Topology repairing of disconnected pulmonary airways and vessels: baselines and a dataset. In: *International Conference on Medical Image Computing and Computer-Assisted Intervention*. Springer, pp. 382–392.
- Wyburd, M.K., Dinsdale, N.K., Jenkinson, M., Namburete, A.I., 2024. Anatomically plausible segmentations: Explicitly preserving topology through prior deformations. *Med. Image Anal.* 97, 103222.
- Xiang, P., Wen, X., Liu, Y.-S., Cao, Y.-P., Wan, P., Zheng, W., Han, Z., 2022. Snowflake point deconvolution for point cloud completion and generation with skip-transformer. *IEEE Trans. Pattern Anal. Mach. Intell.* 45 (5), 6320–6338.
- Xie, H., Yao, H., Zhou, S., Mao, J., Zhang, S., Sun, W., 2020. Grnet: Gridding residual network for dense point cloud completion. In: *European Conference on Computer Vision*. Springer, pp. 365–381.
- Ye, H., Zhang, X., Hu, Y., Fu, H., Liu, J., 2025. Vsr-net: Vessel-like structure rehabilitation network with graph clustering. *IEEE Trans. Image Process.* 1090–1105.
- Yu, X., Rao, Y., Wang, Z., Liu, Z., Lu, J., Zhou, J., 2021. PointR: Diverse point cloud completion with geometry-aware transformers. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision*. pp. 12498–12507.
- Yuan, W., Khot, T., Held, D., Mertz, C., Hebert, M., 2018. Pcn: Point completion network. In: *2018 International Conference on 3D Vision. 3DV, IEEE*. pp. 728–737.
- Zeng, A., Wu, C., Lin, G., Xie, W., Hong, J., Huang, M., Zhuang, J., Bi, S., Pan, D., Ullah, N., Khan, K.N., Wang, T., Shi, Y., Li, X., Xu, X., 2023. ImageCAS: A large-scale dataset and benchmark for coronary artery segmentation based on computed tomography angiography images. *Comput. Med. Imaging Graph.* 109, 102287.
- Zhang, S., Li, X., Zong, M., Zhu, X., Cheng, D., 2017. Learning k for knn classification. *ACM Trans. Intell. Syst. Technol. (TIST)* 8 (3), 1–19.
- Zhang, X., Sun, K., Wu, D., Xiong, X., Liu, J., Yao, L., Li, S., Wang, Y., Feng, J., Shen, D., 2023. An anatomy- and topology-preserving framework for coronary artery segmentation. *IEEE Trans. Med. Imaging* 43 (2), 723–733.
- Zhang, X., Zhang, J., Ma, L., Xue, P., Hu, Y., Wu, D., Zhan, Y., Feng, J., Shen, D., 2022. Progressive deep segmentation of coronary artery via hierarchical topology learning. In: *International Conference on Medical Image Computing and Computer-Assisted Intervention*. Springer, pp. 391–400.
- Zhao, H., Jiang, L., Jia, J., Torr, P.H., Koltun, V., 2021. Point transformer. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision*. pp. 16259–16268.
- Zhou, H., Cao, Y., Chu, W., Zhu, J., Lu, T., Tai, Y., Wang, C., 2022. Seedformer: Patch seeds based point cloud completion with upsample transformer. In: *European Conference on Computer Vision*. Springer, pp. 416–432.