

Topological Shape Representations for Aneurysm–Bifurcation Discrimination

Akshay Gokhale

*Computer Science and Engineering
Sardar Patel Institute of Technology
Mumbai, India*

AKSHAY.GOKHALE22@SPIT.AC.IN

Mansi Dhamne

*Computer Engineering
Sardar Patel Institute of Technology
Mumbai, India*

MANSI.DHAMNE22@SPIT.AC.IN

Abstract

Automated detection of intracranial aneurysms (IAs) from CT angiography (CTA) is severely hindered by high false-positive rates. Traditional convolutional neural networks (CNNs) rely fundamentally on local pixel intensities, causing them to systematically confuse true saccular aneurysms with healthy vascular bifurcations. This geometric ambiguity is especially catastrophic for small, clinically critical lesions (<3 mm), where standard detection sensitivity often plummets below 60%.

To resolve this, we propose a plug-and-play, topology-aware false-positive reduction framework. We evaluate the Smooth Euler Characteristic Transform (SECT)—a directional mathematical representation that encodes global 3D vascular geometry independently of intensity—against standard persistence-based summaries (Persistence Images and Landscapes). The framework is rigorously stress-tested on a curated, heavily stratified subset of the multi-institutional RSNA 2025 dataset, explicitly designed to challenge models with anatomically plausible bifurcation mimics.

SECT demonstrates exceptional, classifier-agnostic discriminative power, achieving an AUC of 0.943, substantially outperforming direction-agnostic persistence methods (AUC ~ 0.68). Crucially, our topological filter exhibits a clinical performance inversion: it excels on the most challenging sub-3 mm cohort, maintaining a 0.943 AUC and 78.5% sensitivity even under a strict 95% specificity constraint. Furthermore, the representation proves highly resilient to hardware-specific artifacts, maintaining a 0.927 mean AUC under strict leave-one-scanner-out (LOGO) validation across four distinct manufacturers.

By explicitly capturing asymmetric geometric invariants rather than intensity profiles, directional topological representations reliably resolve the primary structural confounder in IA detection. This establishes SECT as a highly robust, scanner-agnostic downstream filter ready for integration into hybrid deep-learning diagnostic pipelines.

1. Introduction

The automated detection of Intracranial Aneurysms (IAs) remains a critical challenge in machine learning for healthcare. IAs are abnormal, localized dilations of cerebral arteries affecting approximately 3% to 7% of the adult population (Vlak et al., 2011; Joo, 2025). If left untreated, these weakened vessel walls can rupture, causing a fatal subarachnoid

hemorrhage. While Digital Subtraction Angiography (DSA) remains the gold standard, Computed Tomography Angiography (CTA) is the primary non-invasive screening modality used in clinical practice (Hsu et al., 2025). However, reliably identifying IAs on CTA is notoriously difficult even for skilled radiologists. This difficulty is magnified for small lesions (< 3 mm in diameter), where sensitivity frequently drops to between 64% and 74% (Bizjak and Špiclin, 2023). Developing an automated diagnostic aid is paramount for early intervention, yet current deep learning systems generate an overwhelming number of false alarms, severely hindering their clinical viability.

Addressing this problem is challenging due to the inherent geometric limitations of standard convolutional neural networks (CNNs). CTA natively suffers from lower spatial resolution and low vessel-to-background contrast, making small vascular structures visually indistinct because of surrounding tissue or noise. More fundamentally, modern CNNs rely heavily on local pixel intensities and texture statistics. Because a saccular aneurysm and a healthy vascular bifurcation exhibit nearly identical local intensity distributions in a CTA volume, CNNs consistently confuse the two. This geometric ambiguity—treating a complex branching intersection exactly like a bulging sac—translates into a prohibitive number of false positives per scan. Standard intensity-based CNNs do not explicitly model global structural invariants; they cannot mathematically differentiate the asymmetric, spherical protrusion of an aneurysm dome from the symmetric branching of a healthy vessel bifurcation. Consequently, previous deep learning architectures, such as 3D U-Nets and ResNets, often exhibit lesion-level sensitivities for small aneurysms of merely 56% and fail to generalize across different scanner acquisition protocols (Bizjak and Špiclin, 2023; Joo, 2025).

In this work, we introduce a topologically-aware false-positive reduction framework that explicitly models global structural invariants to differentiate aneurysms from complex vessel bifurcations. Rather than replacing existing deep learning architectures, our module is designed as a plug-and-play false-positive reduction filter intended to sit downstream of high-sensitivity CNN candidate generators. Drawing on Persistent Homology (PH) from Topological Data Analysis (TDA), we capture coordinate-independent geometric descriptors that track the topological evolution of a shape across multiple scales. However, we hypothesize that standard direction-agnostic topological summaries are insufficient for cerebrovascular structures. Instead, we evaluate the Smooth Euler Characteristic Transform (SECT), a directional topological representation that captures the asymmetric spatial configuration of aneurysms relative to their parent vessels. We validate our framework on a curated, heavily stratified subset of the multi-institutional RSNA 2025 dataset (Radiological Society of North America (RSNA), 2025), explicitly designed to stress-test geometric discrimination against anatomically plausible bifurcation mimics.¹

Generalizable Insights about Machine Learning in the Context of Healthcare

Our work highlights a critical lesson for healthcare machine learning: models relying exclusively on local visual patterns (e.g., intensity and texture) are inherently brittle in clinical scenarios where healthy and pathological structures share similar density profiles. We demonstrate that incorporating explicit mathematical representations of global shape and spatial organization can resolve this geometric ambiguity in a principled manner. Further-

1. To support reproducibility, the full codebase can be made available on request.

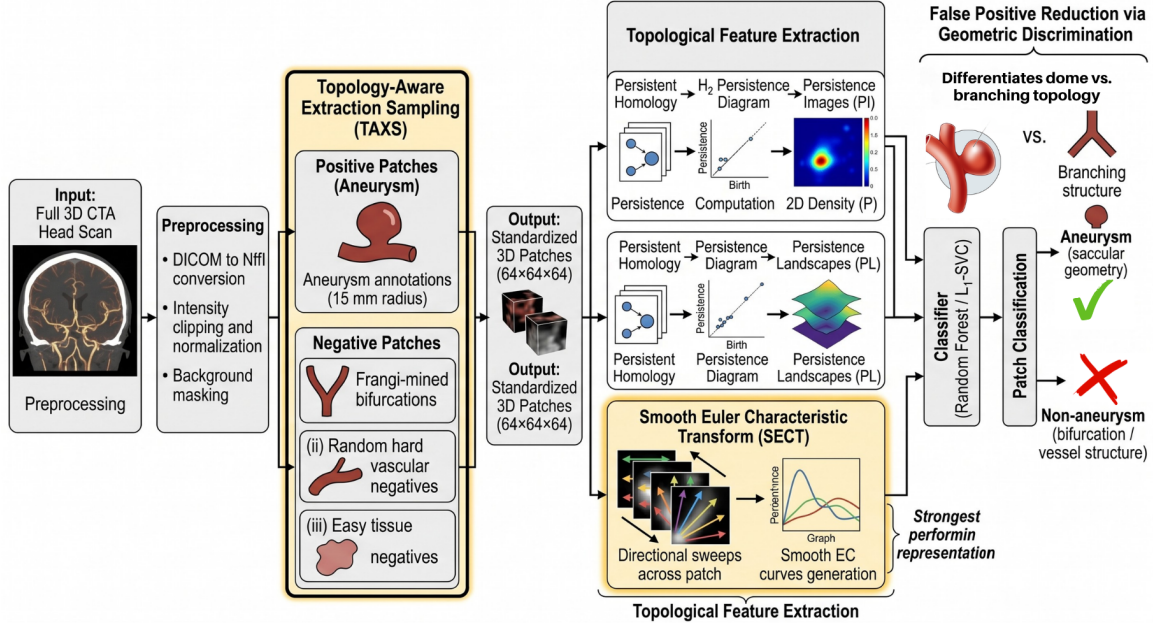


Figure 1: **Topology-Aware False Positive Reduction Pipeline.** Input CTA volumes undergo intensity standardization before localized 3D patches are extracted via the Topology-Aware Extraction Sampling (TAXS) strategy. Patches are heavily stratified into positive aneurysms and highly confounding negative classes (e.g., Frangi-mined bifurcations). Topological features are extracted via Persistent Homology (PI, PL) and directional filtration (SECT). Acting as a downstream filter, SECT provides the strongest representation, enabling the classifier to accurately discriminate saccular geometry from complex branching structures based purely on shape rather than pixel intensity.

more, by focusing on topological invariants rather than raw pixel values, we show that structural properties remain stable across varied imaging conditions. This provides a pathway to improving generalizability across disparate hospital systems and scanner manufacturers without requiring exhaustive retraining, emphasizing that integrating clinical intuition about anatomy with geometric machine learning leads to more robust, trustworthy, and deployable diagnostic systems.

2. Related Work

2.1. Deep Learning for Intracranial Aneurysm Detection

Automated intracranial aneurysm (IA) detection has predominantly been approached as a segmentation or object detection task using architectures such as 3D U-Nets, ResNets, and GLIA-Net (Bo et al., 2021; Bizjak and Špiclin, 2023; Hsu et al., 2025). While these systems achieve high sensitivity for large aneurysms, they fundamentally rely on localized

pixel intensities and texture statistics. Consequently, they falter on small lesions (< 3 mm), where sensitivity plummets to roughly 56%, and produce an overwhelming number of false positives (FPs) by confusing saccular aneurysms with healthy vascular bifurcations (Bizjak and Špiclin, 2023). Furthermore, as noted in recent scoping reviews, most studies lack external validation, multi-institutional data, and size-stratified metrics, yielding optimistic aggregate performance that proves brittle in clinical deployment (Joo, 2025).

2.2. Geometry-Aware and FP-Reduction Methods

To address the high FP rates of pure intensity-based CNNs, modern computer-aided detection (CAD) systems frequently adopt two-stage pipelines: a high-sensitivity candidate generator followed by a dedicated FP-reduction classifier. This paradigm has proven highly effective in both pulmonary nodule and aneurysm detection architectures (Setio et al., 2017). Within these secondary filters, explicitly modeling vascular geometry is critical. Techniques ranging from classical multi-scale vesselness filters (e.g., Frangi) to modern deep-learning skeletonization are frequently employed to extract centerlines and suppress tubular structures (Frangi et al., 1998). More recently, Graph Neural Networks (GNNs) have been applied directly to extracted vascular trees to enable complex bifurcation modeling and artery-vein classification (Vos et al., 2024; Hampe et al., 2024). However, these graph- and centerline-based approaches rely heavily on accurate preliminary segmentation, which is highly error-prone in the low-contrast, noisy regime of small aneurysms.

2.3. Topological Data Analysis in Medical Imaging

To overcome the limitations of localized texture classifiers and brittle skeletonization algorithms, Topological Data Analysis (TDA) offers a mathematically rigorous framework for extracting robust global geometries (Singh et al., 2023; Singh and Quaia, 2025). TDA investigates the underlying shape of a scalar field through algebraic topology, tracking the presence of connected components (H_0), loops (H_1), and voids (H_2) across multiple scales. The core mechanism, Persistent Homology (PH) (Otter et al., 2017; Clough et al., 2019), provides mathematical summaries that are inherently invariant to deformations and the localized scanner noise prevalent in CTA. While TDA has been successfully applied to characterize general tumor morphology and pulmonary tissue by separating meaningful anatomical structures from transient noise (Boehm et al., 2008; Crawford et al., 2020), its application in distinguishing nuanced cerebrovascular structures remains largely unexplored.

2.4. Vectorization of Topological Features for ML Integration

Integrating topological priors into deep learning pipelines presents a fundamental computational challenge. The native output of PH is a Persistence Diagram (PD)—a variable-cardinality multi-set of birth-death coordinates that cannot be directly ingested by standard classifiers (Adams et al., 2017; Som et al., 2020). To bridge this gap, methods such as Persistence Images (PI) utilize lifetime-weighted Gaussian kernels to project PDs into stable Euclidean grids (Chepushtanova et al., 2015; Adams et al., 2017). Alternatively, Persistence Landscapes (PL) encode topological lifespans via layered, integrable functions mapped to Banach spaces (Bubenik, 2015). However, PI and PL often struggle to reliably capture the subtle H_2 topological void signatures necessary for identifying sub-3 mm

aneurysms, and standard PH calculations suffer from an $O(N^3)$ computational bottleneck on high-resolution 3D medical volumes. To circumvent this, recent advances introduced the Smooth Euler Characteristic Transform (SECT) (Turner et al., 2014; Crawford et al., 2020). SECT generates smooth, continuous vector curves based on directional sensitivity to asymmetric shape deformations, bypassing the traditional bottleneck of PH entirely.

Despite advancements in theoretical vectorization, existing methods have not systematically utilized explicit topological priors to resolve the geometric ambiguity between intracranial aneurysms and vessel bifurcations. By transforming complex vascular geometries into ML-compatible topological descriptors (PIs, PLs, and SECT vectors), our methodology explicitly leverages these representations to act as a definitive, shape-driven downstream filter for current CNN-based detection systems.

3. Methodology

3.1. Problem Formulation

We formalize intracranial aneurysm (IA) detection as a localized binary classification problem on 3D patches extracted from CT angiography (CTA) volumes, rather than a full-volume detection or segmentation task. Prior work and empirical observations indicate that aneurysms and vessel bifurcations exhibit similar intensity and texture characteristics in CTA, leading to high false positive rates in CNN-based systems.

The distinction between these structures is therefore hypothesized to be geometric rather than intensity-based, motivating representations that capture intrinsic shape properties. Thus, as illustrated in the classification stage of Figure 1, our problem revolves around discriminating aneurysm patches from anatomically similar vessel bifurcations at the patch level, allowing us to focus explicitly on clinically challenging small aneurysms (< 3 mm).

3.2. Dataset

All experiments are conducted on the RSNA Intracranial Aneurysm Detection Challenge 2025 dataset (Kaggle), a large-scale, multi-institutional collection of neuroimaging studies (Radiological Society of North America (RSNA), 2025). This dataset includes heterogeneous acquisitions across multiple scanner manufacturers and imaging protocols, providing a clinically realistic benchmark for evaluating model generalization.

For the scope of this work, we restrict our analysis exclusively to Computed Tomography Angiography (CTA) volumes. This restriction is motivated by CTA’s clinical prevalence as a primary, non-invasive screening modality. Furthermore, CTA presents a particularly challenging computational setting due to its lower spatial resolution and reduced sensitivity for small aneurysms compared to the Digital Subtraction Angiography (DSA) gold standard.

Leaving out the corrupted files, the accessible CTA cohort comprises ≈ 1808 unique scans. Within this subset, 973 patients have at least one expert-annotated intracranial aneurysm, localized via spatial coordinates and specific slice instance identifiers.

The final dataset consists of 1,201 aneurysm-positive patches and a heavily stratified negative class designed to test geometric discrimination. The negative class includes 3,576 Frangi-mined hard negatives (anatomically plausible vessel bifurcations), 1,787 random hard

negatives (other dense vascular structures), and 1,787 easy negatives (background tissue). The exact algorithmic extraction and class balancing protocols are detailed in Section B.

To ensure strict data hygiene and prevent data leakage, dataset splitting is strictly enforced at the patient level. Each aneurysm contributes at most one positive patch, and no patient appears in both the training and evaluation splits. Patients with multiple aneurysms may therefore contribute multiple positive patches, but all patches from a given patient are assigned to the same split.

3.3. Patch Construction and Pre-processing

To ensure geometric consistency and scanner invariance across multi-institutional acquisitions, raw DICOM series are first converted to volumetric NIfTI format while strictly preserving native scanner coordinate frames. Prior to patch extraction, volumes undergo strict intensity standardization: we apply anatomical window clipping and threshold-based masking to remove background air and bone artifacts, followed by per-volume Z-score normalization on the valid voxels.

Following volumetric standardization, we implement a Topology-Aware Extraction Sampling (TAXS) strategy, explicitly designed to emulate the false-positive distribution of a high-sensitivity candidate generation network. Positive patches are anchored to expert-annotated aneurysm centroids with a 15 mm physical radius and interpolated to a uniform $64 \times 64 \times 64$ voxel grid. To rigorously evaluate geometric discrimination, negative patches are stratified into three categories: (i) *Frangi-mined hard negatives*: Extracted using a multi-scale Frangi vesselness filter (Frangi et al., 1998) to specifically isolate anatomically plausible, complex vascular bifurcations, (ii) *Random hard negatives*: High-intensity local maxima geometrically distant from true lesions, (iii) *Easy tissue negatives*: Sampled from intermediate-intensity background regions. (See Appendix 1 for detailed algorithmic extraction)

Sampling is performed at the patient level to maintain a controlled negative-to-positive ratio while preventing data leakage. The dataset is explicitly constructed to target a key failure mode of existing systems—confusion between aneurysms and vascular bifurcations—by including anatomically plausible hard negatives. It is further enriched with small (<3 mm) aneurysms, enabling focused evaluation on clinically challenging cases where detection performance is most critical.

3.4. Topological Feature Extraction and Vectorization

Persistent Homology (PH) provides a principled framework for capturing multi-scale topological structure by tracking the birth and death of features across a continuous filtration (Zomorodian, 2005; Edelsbrunner, 2014). For a 3D scalar field, such as a CTA patch, PH summarizes the evolution of connected components (H_0), loops (H_1), and voids (H_2) as the intensity threshold varies. The native output of PH is a persistence diagram (PD), a variable-cardinality multi-set of birth–death coordinates, $PD = \{(b_i, d_i)\}_{i=1}^k$, where each point encodes the lifespan of a specific topological feature.

While PDs provide a complete topological summary, their non-Euclidean nature and varying set sizes render them incompatible with standard machine learning pipelines. Consequently, multiple vectorization methods have been proposed to embed PDs into fixed-

dimensional spaces while preserving stability under perturbations (Adams et al., 2017; Carrière et al., 2020). Among these, Persistence Images (PI) and Persistence Landscapes (PL) are widely adopted due to their theoretical guarantees and compatibility with conventional classifiers (Adams et al., 2017; Bubenik, 2020). However, both representations share a fundamental limitation: they produce direction-agnostic summaries. They encode *which* topological features exist and their relative persistence, but discard critical information regarding *how these features are spatially organized within the volume*.

This limitation is critical for intracranial aneurysm detection. Aneurysms and natural vessel bifurcations frequently exhibit similar local topological invariants, but differ fundamentally in their geometric configuration relative to the parent vessel. Capturing this asymmetry requires representations that preserve spatial context. To address this, we evaluate three complementary vectorization strategies: PI and PL as canonical baselines, and the Smooth Euler Characteristic Transform (SECT) (Crawford et al., 2020)—a directional topological representation explicitly designed to encode geometric asymmetry.

3.4.1. PERSISTENT HOMOLOGY SETUP

For each preprocessed patch $\tilde{P}_c$, we construct a cubical complex using the voxel grid as the underlying topological domain. Persistent homology is computed using the GUDHI library (v3.11.0) (Maria et al., 2014). To ensure that bright vascular structures enter the filtration first, we define a superlevel filtration over the normalized scalar intensity field $f : \tilde{P}_c \rightarrow [0, 1]$. This is implemented by evaluating the sublevel sets of the negated field:

$$\mathcal{F}_\alpha = \{x \in \tilde{P}_c \mid -f(x) \leq \alpha\}. \quad (1)$$

We specifically isolate the second homology group (H_2). Rather than representing literal enclosed cavities—as aneurysms are open structures— H_2 features in our superlevel filtration encode the 3D intensity gradient topology. They act as morphological signatures of localized convexity, distinguishing the sac-like geometry of aneurysm domes from the tubular branching of bifurcations. We empirically validate this representational capability using parameterized synthetic phantoms, demonstrating robust separation of saccular and branching structures across varying scales and noise regimes (see Appendix E.5).

For each clinical patch, we extract the H_2 persistence diagram, $PD_2(f)$. To suppress background noise and transient imaging artifacts, we apply a strict persistence threshold, retaining only finite features with a lifespan $(d_i - b_i) > \varepsilon$, where $\varepsilon = 0.15$ in normalized intensity units. This parameter aggressively filters noise while preserving the subtle, short-lived topological signatures characteristic of sub-3 mm aneurysms. The resulting sparse diagrams serve as direct inputs for the PI and PL vectorizations.

3.4.2. PERSISTENCE IMAGES

Persistence Images (PI) (Adams et al., 2017) provide a stable, finite-dimensional vector representation of persistence diagrams, enabling seamless integration with conventional learning algorithms. To construct a PI, each point in the persistence diagram is first transformed from birth–death coordinates to birth–persistence coordinates:

$$(b_i, d_i) \mapsto (b_i, p_i), \quad \text{where } p_i = d_i - b_i. \quad (2)$$

This transformation isolates the lifespan of a feature (p_i) from its spatial emergence (b_i). Each point then contributes a smooth Gaussian kernel centered at its new location, generating a continuous persistence surface:

$$\rho(x, y) = \sum_{i=1}^k w(b_i, p_i), \phi_\sigma(x - b_i, y - p_i), \quad (3)$$

where ϕ_σ is a Gaussian kernel with bandwidth σ , and $w(b_i, p_i)$ is a specialized weighting function. w is explicitly designed to reduce the effects of low-persistence features near the diagonal, effectively suppressing topological noise while modulating the contribution of each valid feature.

Finally, this continuous persistence surface $\rho(x, y)$ is discretized over a fixed, uniform grid $\Omega \subset \mathbb{R}^2$ to produce the persistence image. Each pixel value corresponds to the two-dimensional integral of ρ over that specific grid cell:

$$PI(u, v) = \iint_{\text{cell}_{u,v}} \rho(x, y), dx, dy. \quad (4)$$

The resulting image matrix is flattened into a fixed-length vector for downstream classification.

3.4.3. PERSISTENCE LANDSCAPES

Persistence Landscapes (PL) (Bubenik, 2015) address the non-Euclidean limitations of persistence diagrams by mapping them into a functional representation within a Banach space (completely normalized vector space), allowing the direct application of classical statistical tools.

Given a persistence diagram $PD = \{(b_i, d_i)\}_{i=1}^n$, each birth–death pair is first associated with a tent function:

$$f_{(b_i, d_i)}(t) = \begin{cases} 0, & t \notin (b_i, d_i), \\ t - b_i, & t \in (b_i, \frac{b_i + d_i}{2}], \\ d_i - t, & t \in (\frac{b_i + d_i}{2}, d_i). \end{cases} \quad (5)$$

Each tent function forms an isosceles triangle with a peak height proportional to the feature’s persistence, centered at its topological midpoint $(b_i + d_i)/2$.

The persistence landscape is then formally defined as a sequence of functions $\lambda_k : \mathbb{R} \rightarrow \mathbb{R}$, where $\lambda_k(t)$ is the k -th largest value of this set of tent functions:

$$\lambda_k(t) = k\text{-th largest value of } \{f_{(b_i, d_i)}(t)\}_{i=1}^n. \quad (6)$$

The resulting landscape can be interpreted as a stacked sequence of continuous curves derived from barcode intervals, where λ_1 captures the most dominant topological features and higher-order landscapes encode progressively weaker, noisier signals.

3.4.4. SMOOTH EULER CHARACTERISTIC TRANSFORM (SECT)

The Smooth Euler Characteristic Transform (SECT) (Crawford et al., 2020) provides a directional topological representation that maps each patch to a collection of smooth Euler

characteristic curves. Unlike PI and PL, which summarize topology without preserving directional organization, SECT encodes directional geometric structure of vascular shapes rather than relying on local intensity patterns. This is supported by controlled synthetic experiments (Appendix E.5), which demonstrate that SECT consistently distinguishes sacular and bifurcation geometries under varying noise, scale, and orientation.

The SECT builds upon the Euler characteristic (EC), a fundamental topological invariant that compactly encodes the number of connected components, loops, and higher-dimensional voids within a shape. To capture multiscale directional structure, SECT computes EC curves by tracking the topological evolution of the shape across sublevel set filtrations, defined by height functions along uniformly sampled directions on a sphere.

Because raw EC curves are piecewise constant, integer-valued functions containing sharp discontinuities, they must be smoothed. The curves are mean-centered and integrated to produce continuous, piecewise-linear functional representations known as Smooth Euler Characteristic (SEC) curves. This transformation ensures the representation resides in a well-behaved Hilbert space ($\mathbb{L}^2$) suitable for statistical learning. The formal mathematical derivations of the EC and the smoothing transformations are detailed in Appendix C.3.

Operationally, SECT is computed by first restricting the normalized patch to a vascular foreground mask, sweeping this mask along a uniformly sampled set of directions, and concatenating the discretized, smoothed EC curves into a fixed-dimensional feature vector. We utilize a vascular threshold $\tau_v = 0.40$ to define the foreground mask. This is conceptually distinct from the persistence threshold $\varepsilon = 0.15$ used in Section 3.4.1: τ_v determines the raw intensity support of the directional shape, whereas ε filters short-lived topological artifacts from persistence diagrams. The final SECT representation uses $D = 128$ directions, $T = 25$ filtration steps, Gaussian smoothing along the filtration axis, and ℓ_2 normalization. A procedural overview is provided in Appendix 2.

4. Experiments and Results

4.1. Experimental Setup and Cohort Characterization

All methods operate on the same set of CTA patches described in Section 3.2. Pre-processing is fixed across all experiments, including intensity clipping, Gaussian smoothing ($\sigma = 0.5$), and normalization to $[0, 1]$. Prior to downstream classification, we empirically verified that persistent homology natively captures a meaningful geometric signal; pairwise bottleneck distance analysis confirms a stable topological hierarchy distinguishing aneurysms from Frangi-mined bifurcations (see Appendix E.1).

Hyperparameters for each representation are selected via a preliminary grid search and then held fixed for all downstream evaluations to isolate representation quality from tuning effects. Sensitivity analysis (Appendix D.4) shows that performance is largely invariant to reasonable variations in key parameters, indicating that the chosen configurations lie in a stable regime.

4.2. Classification Benchmark Across Topological Representations

We evaluate the discriminative capacity, stability, and computational efficiency of three topological vectorization methods: Persistence Images (PI), Persistence Landscapes (PL),

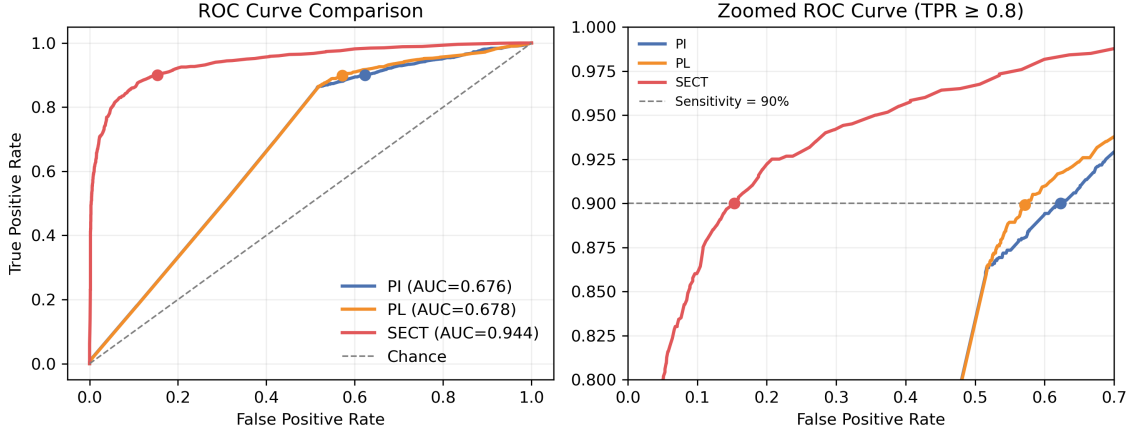


Figure 2: ROC comparison of PI, PL, and SECT using 5-fold out-of-fold predictions. Right: zoomed view of the high-sensitivity operating region. Markers indicate operating points at 90% sensitivity. SECT consistently achieves higher sensitivity at lower false positive rates, demonstrating improved discrimination of aneurysms from vascular structures.

and the Smooth Euler Characteristic Transform (SECT). Each method maps the 3D geometry of CTA patches into a fixed-dimensional feature vector, enabling evaluation under a shared classification pipeline (Section 3.4.4).

To assess whether performance gains are primarily *representation-driven* rather than classifier-dependent, we evaluate each method using two distinct classifiers: a non-linear Random Forest (RF) and an L1-regularized Linear Support Vector Classifier (L1-SVC). We report four complementary metrics: Area Under the ROC Curve (AUC) reported as mean $\pm$ standard deviation under 5-fold cross-validation, empirical Lipschitz constant L (measuring sensitivity to controlled Gaussian noise), feature dimensionality, and average end-to-end inference time per patch (including feature extraction and classification).

Both PI and PL achieve moderate discrimination performance, with AUC values clustered around 0.68, and exhibit sensitivity to classifier choice as relative rankings are not consistently preserved. In contrast, SECT achieves substantially higher performance under both Random Forest (0.9433 AUC) and L1-SVC (0.9312 AUC), indicating that its advantage is representation-driven rather than classifier-dependent.

Figure 2 presents ROC curves along with a zoomed view of the high-sensitivity operating region. SECT consistently dominates PI and PL across the ROC space, with particularly strong separation in the high-sensitivity regime. At clinically relevant operating points (e.g., 90% sensitivity), SECT achieves substantially lower false positive rates compared to persistence-based baselines. Detailed operating-point analysis, including FPR at fixed sensitivity and bifurcation-specific false positive rates, is provided in Appendix E.7.

Regarding robustness, PI exhibits high sensitivity to input perturbations (large empirical Lipschitz constant, L), whereas PL yields a smoother but less discriminative em-

Table 1: Comparison of topological representations across classifiers, stability, and computational cost. Performance is reported as mean AUC $\pm$ standard deviation under 5-fold cross-validation. SECT consistently outperforms persistence-based representations (PI, PL) across both classifiers, while exhibiting moderate stability and higher computational cost due to increased feature dimensionality. Computational efficiency is reported as average end-to-end inference time per patch, including both feature extraction and classification. (W_p stands for Diagram Space and σ for Image Space)

Method	Dim	L1-SVC AUC ($\uparrow$)	RF AUC ($\uparrow$)	Native Domain for L	Empirical L (95% CI)	Inference Time (s)
PI	900	0.6791 ± 0.0066	0.6745 ± 0.0077	W_p	430.5 [418.6, 444.6]	~ 1
PL	750	0.6515 ± 0.0086	0.6809 ± 0.0075	W_p	3.9 [1.9, 5.9]	~ 1
SECT	3200	0.9312 ± 0.0083	0.9433 ± 0.0076	σ	40.3 [28.9, 133.9]	~ 11

bedding. SECT occupies an optimal intermediate regime, maintaining bounded sensitivity without sacrificing discriminative power. Because SECT bypasses intermediate persistence diagrams, its stability is measured end-to-end directly against physical image noise (σ) rather than the Wasserstein diagram distance used for PI and PL. Although this difference in metric spaces prevents direct numerical comparison of L values (detailed in Appendix E.4), SECT’s formulation offers a more clinically relevant measure of robustness to actual scanner variance. Table 1 summarizes these findings, highlighting SECT’s consistent performance advantage.

4.3. Size-Stratified Evaluation of SECT

While overall classification metrics provide a global measure of discriminative capacity, clinical utility critically depends on performance across varying lesion sizes. Because small aneurysms (< 3 mm) constitute the primary failure mode of automated CNN pipelines (Bizjak and Špiclin, 2023), we evaluate whether SECT maintains its discriminative power across standard aneurysm size strata. This stratification isolates the true geometric properties of individual lesions and prevents overall performance metrics from being artificially skewed by easily detectable large aneurysms.

Positive patches are partitioned into three clinically standard categories (Small, Medium, and Large). Size annotations are extracted dynamically at the patch level to ensure spatial fidelity (detailed in Appendix E.3). For each stratum, we construct a balanced evaluation task by pairing the size-specific positive patches against the shared pool of Frangi-mined hard negatives. We report the Area Under the ROC Curve (AUC) alongside sensitivity measured at strict clinical operating points (95% and 99% specificity) using the Random Forest classifier.

As detailed in Table 2 and visualized in Figure 3, SECT exhibits exceptional performance on the sub-3 mm cohort, achieving an AUC of 0.9434 ± 0.0100 . Contrary to the characteristic sensitivity collapse observed in traditional intensity-based networks, SECT maintains a remarkable sensitivity of 0.7845 for small aneurysms even when forced to operate at a strict

Table 2: Size-stratified SECT performance. Performance is reported as mean AUC $\pm$ standard deviation. Sensitivity is evaluated at fixed specificity (Sp.) thresholds to simulate strict false-positive reduction constraints.

Size Stratum	n	AUC ($\uparrow$)	Sens @ 95% Sp.	Sens @ 99% Sp.
Small (< 3 mm)	942	0.9434 $\pm$ 0.0100	0.7845	0.6306
Medium (3–7 mm)	112	0.8634 $\pm$ 0.0209	0.5714	0.3036
Large (> 7 mm)	147	0.9178 $\pm$ 0.0147	0.6395	0.3946
Overall	1201	0.9405 $\pm$ 0.0095	0.7918	0.6003

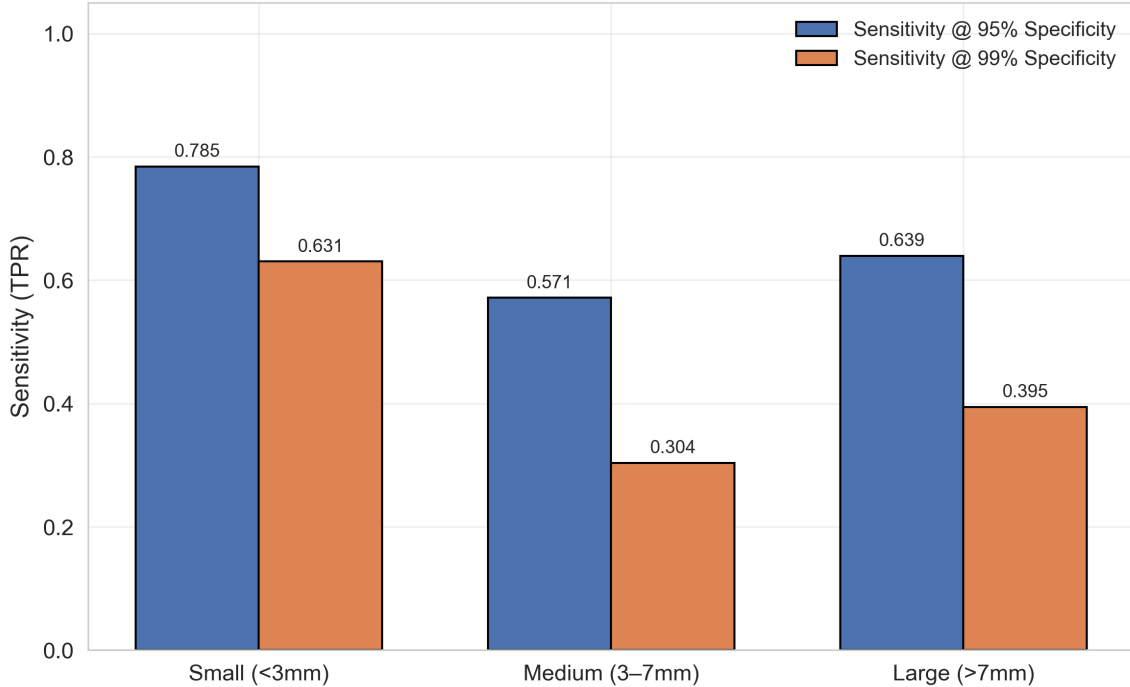


Figure 3: **Sensitivity across aneurysm size strata.** SECT maintains robust detection capabilities for clinically challenging small aneurysms, even under strict specificity constraints, bypassing the typical performance degradation observed in intensity-based CNNs.

95% specificity. Even at an extreme 99% specificity constraint, the method retains a 0.6306 detection rate.

Interestingly, while overall performance remains strong, sensitivity degrades for the Medium and Large strata at stricter operating points. We attribute this variance to two primary factors. First, the severe natural class imbalance in our dataset yields significantly

smaller sample sizes for these cohorts ($n = 112$ and $n = 147$, respectively), reducing statistical stability. Second, this degradation highlights a geometric constraint of our extraction pipeline: the fixed 15 mm physical patch radius may inadvertently truncate the topological boundaries of large aneurysms. Additionally, medium-sized aneurysms often exhibit transitional geometries that more closely mimic the smooth tubular branching of complex bifurcations.

4.4. Inter-Scanner Robustness of SECT

A primary barrier to the clinical translation of automated intracranial aneurysm detection is the severe performance degradation models experience across different scanner manufacturers and acquisition protocols. Traditional CNNs frequently overfit to scanner-specific intensity profiles and noise calibrations. We therefore evaluate whether the geometric signatures captured by SECT preserve discriminative performance across diverse acquisition domains.

Positive and negative patches are assigned scanner labels derived from native DICOM metadata (Siemens, GE, Toshiba/Canon, and Philips). We systematically evaluate scanner robustness under two distinct paradigms: within-scanner evaluation and leave-one-scanner-out (LOGO) validation. For the within-scanner baseline, we report stratified 5-fold cross-validation AUC using exclusively the patches originating from a single manufacturer. For the LOGO evaluation, we enforce strict domain separation: the model is trained on patches from all manufacturers *except* scanner S , and evaluated exclusively on the held-out scanner S . This prevents domain leakage and yields a definitive test of cross-scanner generalization.

Table 3: Cross-scanner robustness evaluation. The model demonstrates exceptional generalization, with LOGO performance matching or even exceeding within-scanner baselines.

Scanner	n	Within-Scanner	Leave-One-Scanner-Out (LOGO)	
		AUC ($\uparrow$)	AUC ($\uparrow$)	AP ($\uparrow$)
Siemens	377	0.9510 $\pm$ 0.0199	0.9369	0.9052
GE	333	0.9529 $\pm$ 0.0071	0.9585	0.8897
Toshiba/Canon	287	0.9011 $\pm$ 0.0178	0.8904	0.8240
Philips	204	0.8963 $\pm$ 0.0092	0.9233	0.8889

As detailed in Table 3, SECT achieves consistently high discriminative performance within individual acquisition domains, with within-scanner AUCs ranging from 0.8963 to 0.9529. More importantly, under the strict, leakage-free LOGO evaluation, SECT maintains highly robust performance across all held-out cohorts (mean AUC 0.9273 ± 0.0247).

Supporting Principal Component Analysis (Appendix E.6) reveals that while minor scanner-associated variance exists within the high-dimensional SECT feature space, the distributions across the four manufacturers remain substantially intertwined without forming isolated clusters. Consequently, while the representation is not mathematically perfectly scanner-invariant, the learned topological decision boundary transfers highly effectively

across acquisition domains, proving SECT’s resilience to the inter-scanner calibration shifts that typically confound intensity-based networks.

4.5. Generalizability Under Mixed Negative Distributions

We evaluate whether the topological representations remain discriminative under shifts in the negative class distribution, moving beyond the Frangi-only evaluation toward clinically realistic candidate mixtures. In clinical practice, an automated detection pipeline will generate a heterogeneous pool of candidates, encompassing vessel bifurcations, complex vascular tortuosity, and non-vascular background tissue artifacts.

We categorize negative patches into three groups: Frangi (F) bifurcation candidates, Hard (H) non-bifurcating vascular structures, and Easy (E) non-vascular background tissue. Our primary evaluation utilizes a clinically motivated mixture of 5:3:2 (F:H:E). To probe sensitivity to distributional variation, we evaluate additional configurations while keeping the positive cohort and the overall negative-to-positive ratio constant. Classification utilizes the Random Forest pipeline established in Section 4.2.

Table 4: SECT performance under varying negative class compositions.

Mixture Setting	Ratio (F:H:E)	AUC ($\uparrow$)
Frangi-only	1:0:0	0.940
Hard-dominant	3:5:2	0.926
Primary clinical	5:3:2	0.922
Uniform	1:1:1	0.914
Easy-dominant	4:2:4	0.905

As reported in Table 4, SECT maintains robust discriminative capacity across all mixture configurations. Under the primary clinical mixture (5:3:2), SECT achieves an AUC of AUC of 0.922 and an Average Precision of 0.844. This represents a minor, acceptable degradation from the idealized Frangi-only setting (AUC 0.940), indicating that the representation does not collapse when exposed to realistic background noise.

To understand the specific failure modes within this mixed distribution, we conducted a per-class false-positive analysis on the primary clinical mixture (Table 5).

Table 5: Per-class false-positive rates (FPR) under the primary 5:3:2 mixed setting.

Negative Type	n	FPR ($\downarrow$)	Mean Pred. Score
Hard Vascular	1073	0.003	0.078
Frangi Bifurcations	1788	0.024	0.113
Easy Non-Vascular	715	0.122	0.219

The breakdown reveals a highly illuminating, counter-intuitive result. Standard vascular structures (Hard) are almost perfectly rejected (FPR = 0.003), and Frangi-mined bifurcations are suppressed highly effectively (FPR = 0.024). However, the easy non-vascular negatives exhibit the highest false-positive rate (0.122) and the highest mean prediction score.

5. Discussion

The central clinical problem this work addresses is not detection sensitivity in the abstract, but the overwhelming false-positive burden that prevents automated IA detection from reaching the clinic. Current systems confuse aneurysmal sacs with normal vessel bifurcations because both structures share nearly identical local intensity distributions in CTA, forcing radiologists to manually review hundreds of spurious candidates per scan (Bizjak and Špiclin, 2023; Hsu et al., 2025). Our results show that SECT substantially reduces this burden at the patch level by encoding directional shape structure rather than intensity—giving the classifier access to information that CNN-based systems fundamentally lack.

Unlike PI and PL, which achieve only moderate discrimination (AUC ~ 0.68 , Table 1) due to severe H_2 diagram collapse (85% empty diagrams at $\varepsilon = 0.15$), SECT achieves robust classification (AUC 0.943 with RF, 0.931 with L1-SVC) by encoding the directional evolution of vascular geometry rather than individual feature lifespans. Crucially, SECT excels precisely where existing systems falter: it yields an AUC of 0.943 (78.4% sensitivity at 95% specificity) for sub-3 mm aneurysms, significantly outperforming typical CNN baselines. Furthermore, because SECT processes structural geometry rather than raw intensities, it demonstrates strong inter-scanner generalization, maintaining leave-one-scanner-out (LOGO) AUCs of 0.89–0.96 across four major manufacturers, albeit with minor residual scanner variance (Appendix E.6). Performance dips slightly for medium aneurysms (AUC 0.863), largely due to class imbalance and patch radius truncation for larger lesions.

SECT’s discriminative power is also resilient to diverse negative cohorts, dropping only ~ 0.015 AUC when evaluated on a primary clinical mixture. Notably, while it near-perfectly rejects challenging vascular bifurcations (FPR = 0.003), easy non-vascular negatives exhibit an anomalous FPR of 0.122, likely reflecting incidentally dome-like tissue cross-sections and RF miscalibration at the boundary. Although SECT’s current inference time (~ 11 s/patch) is non-trivial for clinical pipelines, reducing to 64 directions and GPU parallelization offer clear optimization pathways. Future work will leverage these optimizations to integrate SECT into hybrid CNN architectures, enabling direct benchmarking of these robust topological priors against learned intensity representations.

5.1. Limitations

Evaluation is patch-level within a fixed candidate-generation setting; full end-to-end pipeline validation remains future work. Size-strata imbalance limits conclusions for medium and large aneurysms. Residual scanner-associated variation in the feature space indicates partial rather than full scanner invariance. Computational cost at current settings requires further engineering before real-time deployment.

5.2. Conclusion

We present a systematic evaluation of topological representations for intracranial aneurysm detection. Persistence-based methods provide a meaningful but limited signal; SECT delivers substantially higher discrimination by encoding directional geometric structure that generalises across lesion size, scanner, and candidate distribution. The core clinical take-away is that shape-level representations can resolve the bifurcation ambiguity that defeats

intensity-based systems, and do so most reliably for the sub-3 mm lesions where the clinical need is greatest. Future work will integrate SECT into full detection pipelines, benchmark against CNN baselines, and pursue the computational optimisations needed for clinical deployment.

References

- Henry Adams, Tegan Emerson, Michael Kirby, Rachel Neville, Chris Peterson, Patrick Shipman, Sofya Chepushtanova, Eric Hanson, Francis Motta, and Lori Ziegelmeier. Persistence images: A stable vector representation of persistent homology. *Journal of Machine Learning Research*, 18(8):1–35, 2017.
- Žiga Bizjak and Žiga Špiclin. A systematic review of deep-learning methods for intracranial aneurysm detection in ct angiography. *Biomedicines*, 11(11):2921, 2023.
- Zi-Hao Bo, Hui Qiao, Chong Tian, Yuchen Guo, Wuchao Li, Tiantian Liang, Dongxue Li, Dan Liao, Xianchun Zeng, Leilei Mei, et al. Toward human intervention-free clinical diagnosis of intracranial aneurysm via deep neural network. *Patterns*, 2(2), 2021.
- H. F. Boehm, C. Fink, U. Attenberger, C. Becker, J. Behr, and M. Reiser. Automated classification of normal and pathologic pulmonary tissue by topological texture features extracted from multi-detector ct in 3d. *European Radiology*, 18(12):2745–2755, December 2008. ISSN 1432-1084. doi: 10.1007/s00330-008-1082-y.
- Peter Bubenik. Statistical topological data analysis using persistence landscapes. *Journal of Machine Learning Research*, 16(1):77–102, 2015.
- Peter Bubenik. The persistence landscape and some of its properties. In *Topological Data Analysis: The Abel Symposium 2018*, pages 97–117. Springer International Publishing, Cham, 2020.
- Mathieu Carrière, Frédéric Chazal, Yuichi Ike, Théo Lacombe, Martin Royer, and Yuhei Umeda. Perslay: A neural network layer for persistence diagrams and new graph topological signatures. In *Proceedings of the International Conference on Artificial Intelligence and Statistics*, pages 2786–2796. PMLR, 2020.
- Sofya Chepushtanova, Tegan Emerson, Eric Hanson, Michael Kirby, Francis Motta, Rachel Neville, Chris Peterson, Patrick Shipman, and Lori Ziegelmeier. Persistence images: An alternative persistent homology representation. *arXiv preprint arXiv:1507.06217*, 7, 2015.
- James R Clough, Ilkay Oksuz, Nicholas Byrne, Julia A Schnabel, and Andrew P King. Explicit topological priors for deep-learning based image segmentation using persistent homology. In *International Conference on Information Processing in Medical Imaging*, pages 16–28. Springer, 2019.
- Lorin Crawford, Anthea Monod, Andrew X. Chen, Sayan Mukherjee, and Raúl Rabadán. Predicting clinical outcomes in glioblastoma: An application of topological and functional data analysis. *Journal of the American Statistical Association*, 115(531):1139–1150, 2020.
- Herbert Edelsbrunner. *A Short Course in Computational Geometry and Topology*. Springer, Berlin, 2014.
- Alejandro F Frangi, Wiro J Niessen, Koen L Vincken, and Max A Viergever. Multiscale vessel enhancement filtering. In *International conference on medical image computing and computer-assisted intervention*, pages 130–137. Springer, 1998.

- Nils Hampe, Sanne GM van Velzen, Jelmer M Wolterink, Carlos Collet, José PS Henriques, Nils Planken, and Ivana Išgum. Graph neural networks for automatic extraction and labeling of the coronary artery tree in ct angiography. *Journal of Medical Imaging*, 11(3):034001–034001, 2024.
- Wei-Chan Hsu, Monique Meuschke, Alejandro F Frangi, Bernhard Preim, and Kai Lawonn. A survey of intracranial aneurysm detection and segmentation. *Medical Image Analysis*, 101:103493, 2025.
- Bio Joo. Methodological challenges in deep learning-based detection of intracranial aneurysms: a scoping review. *Neurointervention*, 20(2):52–65, 2025.
- Clément Maria, Jean-Daniel Boissonnat, Marc Glisse, and Mariette Yvinec. The gudhi library: Simplicial complexes and persistent homology. In *International congress on mathematical software*, pages 167–174. Springer, 2014.
- Nina Otter, Mason A Porter, Ulrike Tillmann, Peter Grindrod, and Heather A Harrington. A roadmap for the computation of persistent homology. *EPJ data science*, 6(1):17, 2017.
- Radiological Society of North America (RSNA). Rsnas 2025 intracranial aneurysm detection. <https://kaggle.com/competitions/rsna-intracranial-aneurysm-detection>, 2025. Kaggle competition.
- Arnaud Arindra Adiyoso Setio, Alberto Traverso, Thomas de Bel, Moira S. N. Berens, Cas van den Bogaard, Piergiorgio Cerello, Hao Chen, Qi Dou, Maria Evelina Fantacci, Bram Geurts, Robbert van der Gugten, Pheng Ann Heng, Bart Jansen, Michael M. J. de Kaste, Valentin Kotov, Jack Yu-Hung Lin, Jeroen T. M. C. Manders, Alexander Sónora-Mengana, Juan Carlos García-Naranjo, Evgenia Papavasileiou, Mathias Prokop, Marco Saletta, Cornelia M. Schaefer-Prokop, Ernst T. Scholten, Luuk Scholten, Miranda M. Snoeren, Ernesto Lopez Torres, Jef Vandemeulebroucke, Nicole Walasek, Guido C. A. Zuidhof, Bram van Ginneken, and Colin Jacobs. Validation, comparison, and combination of algorithms for automatic detection of pulmonary nodules in computed tomography images: The luna16 challenge. *Medical Image Analysis*, 42:1–13, 2017. ISSN 1361-8415. doi: <https://doi.org/10.1016/j.media.2017.06.015>.
- Yashbir Singh and Emilio Quaia. Unraveling the invisible: Topological data analysis as the new frontier in radiology’s diagnostic arsenal. *Tomography*, 11(1):6, 2025.
- Yashbir Singh, Colleen M Farrelly, Quincy A Hathaway, Tim Leiner, Jaidip Jagtap, Gunnar E Carlsson, and Bradley J Erickson. Topological data analysis in medical imaging: current state of the art. *Insights into Imaging*, 14(1):58, 2023.
- Anirudh Som, Hongjun Choi, Karthikeyan Natesan Ramamurthy, Matthew P Buman, and Pavan Turaga. Pi-net: A deep learning approach to extract topological persistence images. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition workshops*, pages 834–835, 2020.
- K. Turner, S. Mukherjee, and D. M. Boyer. Persistent homology transform for modeling shapes and surfaces. *Information and Inference*, 3(4):310–344, December 2014. ISSN 2049-8764, 2049-8772. doi: 10.1093/imaiai/iau011.

- Monique H. M. Vlak, Ale Algra, Raya Brandenburg, and Gabriel J. E. Rinkel. Prevalence of unruptured intracranial aneurysms, with emphasis on sex, age, comorbidity, country, and time period: a systematic review and meta-analysis. *The Lancet. Neurology*, 10 7: 626–36, 2011. URL <https://api.semanticscholar.org/CorpusID:19706811>.
- Iris N Vos, Ynte M Ruigrok, Ishaan R Bhat, Kimberley M Timmins, Birgitta K Velthuis, and Hugo J Kuijf. Graph convolutional networks for automated intracranial artery labeling. *Journal of Medical Imaging*, 11(1):014007–014007, 2024.
- Afra J. Zomorodian. *Topology for Computing*, volume 16. Cambridge University Press, 2005.

Appendix

Appendix A. Dataset Description

A.1. Scanner Distribution

The dataset exhibits realistic multi-scanner variability: Siemens: 377, GE: 333, Toshiba/Canon: 287, Phillips: 204. This diversity is critical for evaluating robustness to acquisition heterogeneity, a known limitation of CNN-based IA detection systems.

A.2. Aneurysm Size Distribution

The dataset is heavily skewed toward small aneurysms: Small (< 3 mm): 942 samples (78.4%), Medium (3–7 mm): 112 samples (9.3%), Large (> 7 mm): 147 samples (12.3%). This further strengthened our choice of the dataset as we can clearly perform size-stratified analysis; mainly focusing on small-aneurysms. See the method to find size estimate in [E.3](#).

Appendix B. Path Construction Details

Patch extraction is performed on NIfTI volumes using a unified pipeline for positives and negatives. In order to reduce the scanner-dependent intensity variability; all the volumes first undergo an intensity pre-processing where: (i) Intensity clipping to $[0, 839]$ HU, (ii) Removal of air/background via thresholding (< 10 HU) and (iii) Patch-level normalization using mean and standard deviation computed on non-air voxels.

Patches are extracted using voxel spacing-aware cropping. All boundary conditions are handled via padding with minimum intensity values and the final patches are resized to a fixed resolution of $64 \times 64 \times 64$. This ensures geometric consistency among patients with varying voxel spacings.

We construct a balanced dataset of positive and negative patches through a multi-stage procedure that enforces spatial, anatomical, and intensity constraints (see [Fig.4](#)).

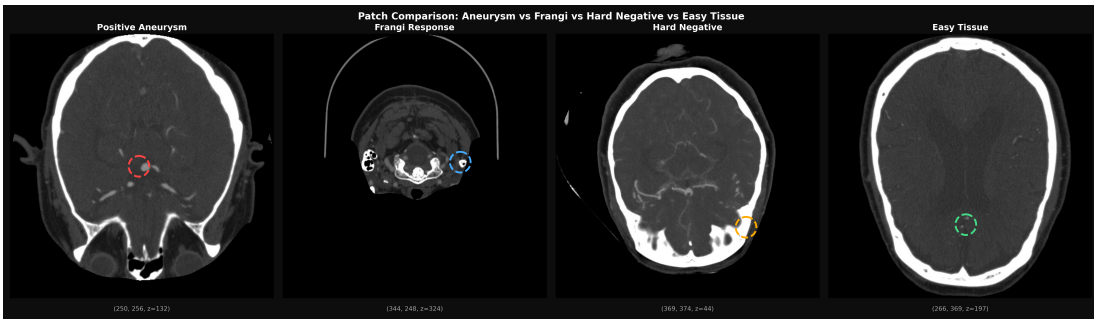


Figure 4: Representative examples of extracted CTA patches across classes. From left to right: (i) aneurysm-positive patch, (ii) Frangi-mined bifurcation candidate, (iii) hard vascular negative, and (iv) easy non-vascular tissue patch. Colored markers indicate patch centers used for extraction.

Algorithm 1: Topology-Aware Extraction Sampling (TAXS) Strategy**Input:** CTA volume V , annotation CSV $\mathcal{A}$, radius $r = 15$ mm**Output:** Set of labeled patches $\mathcal{P} = \mathcal{P}^+ \cup \mathcal{P}^-$ **Step 1: Positive Patch Extraction**

```

foreach annotation entry  $(x, y, SOPInstanceUID) \in \mathcal{A}$  do
  Reconstruct slice index  $z$  using DICOM ordering via ImagePositionPatient Define
  center  $c = (x, y, z)$  Extract patch  $P_c$  with physical radius  $r = 15$  mm Add  $P_c$  to  $\mathcal{P}^+$ 
end

```

Step 2: Random Negative Sampling

```

foreach patient volume  $V$  do
  Sample candidate locations  $c'$  such that:  $\|c' - c\|_2 > 40$  voxels for all positive centers  $c$ 
   $c'$  lies within axial bounds  $[0.15, 0.85]$  of volume depth
  foreach candidate  $c'$  do
    if  $intensity(c') > 150$  HU and local maximum then
      Add patch  $P_{c'}$  to hard negatives  $\mathcal{P}_{hard}^-$ 
    else
      if  $10 < intensity(c') < 150$  then
        Add patch  $P_{c'}$  to easy negatives  $\mathcal{P}_{easy}^-$ 
      end
    end
  end
end

```

Step 3: Frangi-Based Hard Negative Mining

```

foreach volume  $V$  do
  Downsample  $V$  to 25% resolution Apply multi-scale Frangi filter with  $\sigma \in \{1, 2\}$  Se-
  lect top 0.2% vesselness responses as candidate points Rescale coordinates to original
  resolution Extract corresponding patches and add to  $\mathcal{P}_{frangi}^-$ 
end

```

Step 4: Sampling Balance

```

foreach patient do
  Sample approximately  $2 \times |\mathcal{P}^+|$  hard negatives Sample approximately  $1 \times |\mathcal{P}^+|$  easy
  negatives
end

```

```

return  $\mathcal{P}^+ \cup \mathcal{P}_{hard}^- \cup \mathcal{P}_{easy}^- \cup \mathcal{P}_{frangi}^-$ 

```

To ensure that topological descriptors are evaluated under clinically realistic failure modes, the designed patch construction pipeline (see Algorithm:1) explicitly targets vascular structures prone to false positives, particularly bifurcations.

For the frangi-mined negatives, instead of full Hessian-based vesselness computation at native resolution, we use a down sampled approximation to reduce computational cost. Empirical validation on a held-out subset confirms that this approximation preserves candidate quality without loss of relevant geometric information (see AppendixD.2).

Appendix C. Topological Vectorization

C.1. Persistence Images

Persistence diagrams (PDs) provide a compact summary of topological features but do not naturally reside in a vector space, making them difficult to integrate with standard machine learning models. Persistence Images (PI) [Adams et al. \(2017\)](#), is a stable, finite-dimensional vector representation of PDs that enables the use of conventional learning algorithms.

In our pipeline, persistence images are computed from the H_2 persistence diagrams using:

- Gaussian kernel with bandwidth $\sigma = 1.0$,
- Linear weighting function $w(b, p) = p$,
- Fixed grid resolution of $R \times R$ pixels over a normalized domain.

All infinite intervals are discarded prior to transformation.

Despite their practicality, persistence images introduce several limitations— (i)Information loss due to discretization, (ii)Hyperparameter sensitivity and (iii)Loss of geometric structure. These limitations motivate the exploration of alternative representations that preserve more structural information while remaining stable and learnable.

C.2. Persistence Landscapes

Given a persistence diagram

$$PD = \{(b_i, d_i)\}_{i=1}^n, \quad (7)$$

each birth–death pair is associated with a **tent function**:

$$f_{(b_i, d_i)}(t) = \begin{cases} 0, & t \notin (b_i, d_i), \\ t - b_i, & t \in (b_i, \frac{b_i + d_i}{2}], \\ d_i - t, & t \in (\frac{b_i + d_i}{2}, d_i). \end{cases} \quad (8)$$

Each function forms an isosceles triangle with peak height proportional to persistence and centered at the midpoint $(b_i + d_i)/2$.

The Persistence Landscapes (PL), introduced by Bubenik [Bubenik \(2015\)](#) is defined as a sequence of functions

$$\lambda_k : \mathbb{R} \rightarrow \mathbb{R}, \quad k = 1, 2, \dots \quad (9)$$

where

$$\lambda_k(t) = k\text{-th largest value of } \{f_{(b_i, d_i)}(t)\}_{i=1}^n. \quad (10)$$

Equivalently, the persistence landscape can be viewed as a function

$$\lambda : \mathbb{N} \times \mathbb{R} \rightarrow \mathbb{R}, \quad \lambda(k, t) = \lambda_k(t), \quad (11)$$

forming a layered representation where — λ_1 captures the most persistent feature at each scale, and higher-order λ_k encode progressively weaker features.

In our pipeline, persistence landscapes are computed from H_2 persistence diagrams as follows:

- All intervals are converted into triangular functions,
- Landscapes are truncated to the first K levels,
- Functions are discretized over a fixed grid for numerical integration.

Despite their strong theoretical properties, PLs exhibit several limitations — (i) Dimensional truncation, (ii) Loss of pointwise correspondence, and (iii) Smoothing of features. These limitations highlight the trade-off between statistical tractability and structural fidelity in functional representations of persistent homology.

C.3. Smooth Euler Characteristic Transform (SECT) Formulation

The Smooth Euler Characteristic Transform (SECT), introduced by [Crawford et al. \(2020\)](#), provides a principled geometric alternative to persistence-based summaries by mapping shapes into a collection of smooth functions. Each shape is represented as a set of functions that lie in a Hilbert space $\mathbb{L}^2$, equipped with a well-defined inner product.

MATHEMATICAL FORMULATION

The SECT builds upon the Euler characteristic (EC), defined algebraically as the alternating sum of Betti numbers:

$$\chi(X) = \sum_{k=0}^{\infty} (-1)^k \beta_k, \quad (12)$$

which compactly encodes the number of connected components, loops, and higher-dimensional voids. For a discretized complex K , this admits an equivalent combinatorial form as an alternating sum over cell counts across dimensions:

$$\chi(K) = \sum_{k=0}^d (-1)^k n_k, \quad (13)$$

where n_k denotes the number of k -dimensional cells.

To capture multiscale directional structure, SECT considers a height function for a fixed direction $\nu \in S^{d-1}$:

$$r_\nu(x) = x \cdot \nu, \quad (14)$$

which induces a sublevel set filtration K_ν^x . Tracking the EC across this filtration yields the Euler characteristic curve $\chi_\nu(x) = \chi(K_\nu^x)$.

To ensure the representation is amenable to statistical learning, this piecewise constant curve must be smoothed. Following [Crawford et al. \(2020\)](#), we first compute the mean-centered EC curve over the domain $[a_\nu, b_\nu]$:

$$Z_\nu^K(x) = \chi_\nu^K(x) - \bar{\chi}_\nu^K. \quad (15)$$

Integrating this centered curve produces a continuous, piecewise-linear functional representation known as the Smooth Euler Characteristic (SEC) curve:

$$F\nu^K(y) = \int_{-\infty}^y Z\nu^K(x)dx. \quad (16)$$

The full SECT representation is formally defined as the collection of these SEC curves over all sampled directions: $\nu \mapsto F\nu^K \in \mathbb{L}^2$.

PROPERTIES AND LIMITATIONS

Compared to persistence landscapes and persistence diagrams, SECT offers several distinct representational advantages. It avoids the need for the arbitrary ordering or truncation of topological features, provides a globally consistent summary across spatial directions, and directly yields smooth functional representations without requiring auxiliary density transformations (such as Gaussian kerneling in Persistence Images).

Furthermore, SECT carries strong theoretical guarantees, including injectivity and a well-defined Hilbert space structure that naturally permits norms, inner products, and functional regression. This mathematical foundation natively supports advanced functional data analysis (FDA) methods, including Gaussian process models, while ensuring the EC-based summaries capture rich, global shape information.

However, the application of SECT is bounded by certain computational and theoretical limitations. The expressivity of the transform is highly dependent on the density of the directional sampling over the sphere and exhibits sensitivity to the granularity of the filtration steps. Additionally, because the Euler characteristic acts as a global summary scalar at each filtration step, highly localized geometric details can occasionally be lost in the aggregation. These limitations reflect the inherent trade-off between statistical tractability and granular structural expressivity common to all functional summaries in topological data analysis.

Appendix D. Hyper-parameter Selection

For hyper-parameter selection, and to test the stability of extracted features, from the full dataset we constructed a curated subset consisting of – 154 positive patient patches (each containing at least one annotated aneurysm), 394 Frangi-mined hard negatives, 221 random hard negatives, 145 easy tissue negatives.

D.1. Intensity Normalization and Clipping

To determine an appropriate intensity normalization range, we perform an exploratory analysis of voxel intensities across a representative subset of the dataset. For each volume, we compute summary statistics including minimum, maximum, mean, and upper percentiles of Hounsfield Units (HU).

We observe that while raw intensities vary across scans, the majority of informative signal is concentrated within a bounded range, with extreme values corresponding to noise or acquisition artifacts. In particular, the 99th percentile (P99) provides a stable upper bound across samples.

Algorithm 2: SECT Feature Extraction for a CTA Patch

Input: Preprocessed patch $\tilde{P}_c \in \mathbb{R}^{64 \times 64 \times 64}$, vascular threshold $\tau_v = 0.40$, directions $D = 128$, steps $T = 25$

Output: SECT feature vector $\xi \in \mathbb{R}^{3200}$

1. **Normalize patch:** Apply preprocessing (Sec. 3.2) to obtain $f : \tilde{P}_c \rightarrow [0, 1]$
2. **Vascular mask:** $M = \{x \in \tilde{P}_c : f(x) \geq \tau_v\}$ // $\tau_v \neq$ PH threshold ε
3. **Sample directions:** Generate $\{v_i\}_{i=0}^{D-1} \subset S^2$ via Fibonacci sphere:

$$y_i = 1 - \frac{2i}{D-1}, \quad r_i = \sqrt{1 - y_i^2}, \quad \theta_i = \phi i, \quad \phi = \pi(3 - \sqrt{5})$$

$$v_i = (r_i \cos \theta_i, y_i, r_i \sin \theta_i)$$

for $k \leftarrow 1$ **to** D **do**

 4. **Height function:** $h_{v_k}(x) = v_k \cdot x$, normalized to $[-1, 1]$

 5. **Sublevel sweep:** $\{t_j\}_{j=1}^T = \text{linspace}(-1, 1, T)$

for $j \leftarrow 1$ **to** T **do**

 | $A_{v_k}(t_j) = M \cap \{x : h_{v_k}(x) \leq t_j\}$ $\chi_{k,j} = \chi(A_{v_k}(t_j))$ using 26-connectivity

end

 6. **Smooth:** $\tilde{\chi}_{k,:} = \text{GaussianSmooth}(\chi_{k,:}, \sigma = 1.0)$

end

7. **Output:** $\Xi \in \mathbb{R}^{D \times T}$, $\xi = \text{vec}(\Xi) \in \mathbb{R}^{3200}$, ℓ_2 -normalized

return ξ

Based on this analysis, we adopt a clipping range of $[0, 839]$ HU, which retains the relevant vascular and soft-tissue contrast while suppressing extreme outliers. This choice is consistent with the empirically observed intensity distribution and is used uniformly across all experiments.

D.2. Validation of Fast Frangi Approximation

To enable large-scale hard negative mining, we evaluate a computationally efficient approximation of the Frangi vesselness filter. The reference implementation applies the filter at $0.5 \times$ resolution with $\sigma \in \{1, 2, 3\}$, while the proposed approximation operates at $0.25 \times$ resolution with $\sigma \in \{1, 2\}$, resulting in a $\sim 4 \times$ reduction in processing time.

Since the two methods produce spatially independent candidate patches, we assess equivalence by analyzing the structural properties of the extracted patches themselves. For each method, a subset of patches is evaluated using in-patch vesselness and intensity-based metrics, including the 95th percentile Frangi response, mean vesselness, vessel occupancy (fraction of voxels > 150 HU), mean HU, and intensity standard deviation.

Statistical comparison using a two-sample Kolmogorov–Smirnov test shows no significant difference between the two distributions across all metrics ($p > 0.05$). The primary hardness measure, in-patch Frangi P95, yields $p = 0.068$ with negligible Wasserstein distance, indicating near-identical structural content.

D.3. Persistence Diagram Preprocessing Selection

To determine appropriate preprocessing for persistence diagram (PD) computation, we perform a grid search over key parameters controlling intensity normalization and smoothing. The evaluated grid includes Gaussian smoothing ($\sigma \in \{0.0, 0.5, 1.0\}$), intensity clipping ($\max \in \{800, 1000, 1200\}$), and normalization method (z-score vs. min-max).

Each parameter combination is evaluated on a balanced subset of positive, Frangi, hard, and easy patches. Performance is assessed using a topology-based separability ratio (TSR), defined as the ratio of inter-class distance (positive vs. Frangi) to intra-class distance (positive vs. positive), computed using bottleneck distance.

The optimal configuration corresponds to Gaussian smoothing $\sigma = 0.5$, intensity clipping at 800 HU, and min-max normalization, which maximizes the TSR while maintaining stable intra-class structure. These parameters are used for all subsequent persistence-based analyses.

D.4. SECT Parameter Sensitivity and Robustness

SECT hyperparameters are selected via a grid search over the number of projection directions ($n_{\text{dirs}} \in \{32, 64, 128, 256\}$), filtration resolution ($t_{\text{steps}} \in \{25, 50, 100\}$), and smoothing strategy. The optimal configuration (128 directions, 25 steps, Gaussian smoothing) provides the best trade-off between performance and computational cost.

To assess robustness, we analyze performance across a broad range of parameter settings. As shown in Fig. 5, SECT-derived representations exhibit minimal variation across resolution and landscape depth, with AUC remaining effectively constant (≈ 0.675 – 0.677). Similarly, Fig. 6 shows that persistence image performance remains stable across grid resolutions and Gaussian variances, with AUC variation below 0.01.

We additionally evaluate sensitivity to the vascular mask threshold $\tau_v \in \{0.30, 0.40, 0.50, 0.60\}$. Across all settings, performance remains unchanged (AUC = 0.937), indicating that SECT features are insensitive to moderate variations in thresholding.

Summary. These results demonstrate that SECT operates in a stable regime, where performance is largely invariant to reasonable choices of hyperparameters. This robustness suggests that the method does not rely on fine-tuning and generalizes consistently across parameter configurations.

Appendix E. Ablation Studies

E.1. Topological Separability via Persistence Diagrams

Before evaluating downstream classification, we first investigate whether persistent homology natively captures a meaningful geometric signal capable of distinguishing aneurysm patches from vascular structures. For each patch, we compute H_2 persistence diagrams using a cubical complex under a superlevel filtration. We evaluate topological similarity using the pairwise bottleneck distance, $d_B(\cdot, \cdot)$, across three distinct comparison groups: (i) intra-class (Aneurysm vs. Aneurysm), (ii) inter-class hard-negatives (Aneurysm vs. Frangi-mined bifurcations), and (iii) non-aneurysmal vascular inter-class (Aneurysm vs. Random Hard

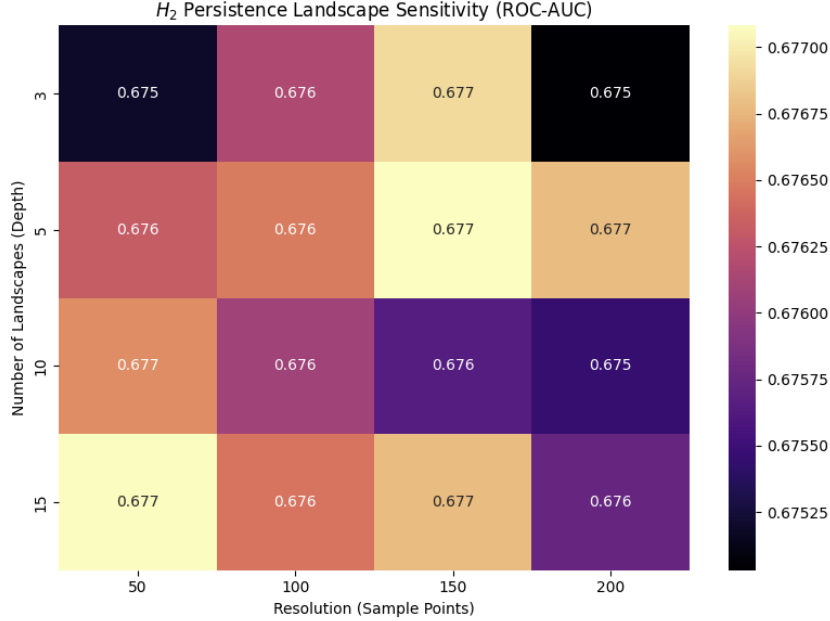


Figure 5: Sensitivity of persistence landscape features to resolution and number of landscapes. Performance remains stable across all configurations, indicating low dependence on discretization choices.

Negatives). The table in Figure 7 reports the mean bottleneck distances across $N = 1201$ samples per class at persistence threshold $\varepsilon = 0.15$.

We observe a consistent distance hierarchy:

$$d_B(\text{Pos}, \text{Pos}) < d_B(\text{Pos}, \text{Frangi}) < d_B(\text{Pos}, \text{Hard}),$$

This hierarchy, which remains stable across persistence thresholds $\varepsilon \in [0.10, 0.20]$, confirms that aneurysm patches are topologically more similar to each other than to vascular structures. Furthermore, Frangi-mined bifurcations occupy an intermediate topological class.

These results provide empirical evidence that persistent homology captures a non-trivial geometric signal aligned with anatomical structure. Importantly, this signal does not correspond to persistently enclosed cavities, but rather reflects the differences in local geometric structure under the filtration. While global separability is observed at the pairwise distance level, the underlying H_2 topological signal is inherently sparse and highly sensitive to noise thresholding. A detailed diagnostic audit of threshold selection and H_2 feature degeneracy is provided in Appendix E.2.

E.2. Details about Topological Separability and Degeneracy

While we try to answer the question (i)do aneurysm-positive patches induce a coherent geometry in PD space, i.e., are positives closer to each other than to hard negatives?; it is

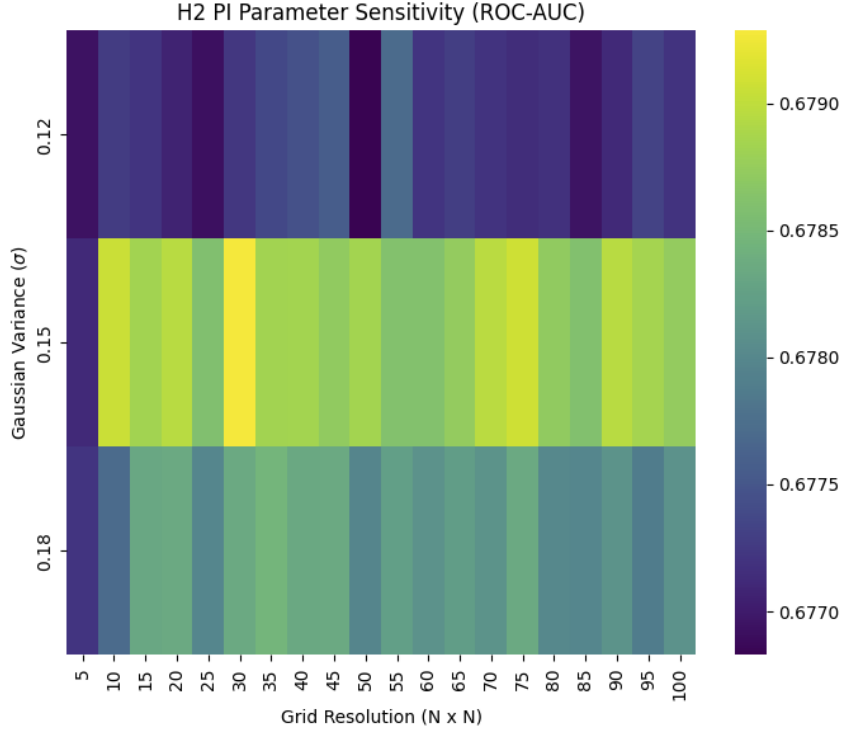


Figure 6: Sensitivity of persistence image features to grid resolution and Gaussian variance. AUC varies minimally across configurations, demonstrating robustness to preprocessing parameters.

also necessary to answer (ii) how often do persistence channels collapse to empty diagrams under denoising thresholds, and how is this affected by homology dimension and negative type?

We evaluate: (i) aneurysm-positive (**positive**), Frangi-based candidate negatives (**fast_frangi**), and a negative baseline (**negative_hard**).

To quantify class geometry directly in diagram space, we use the bottleneck distance $d_B(\cdot, \cdot)$ between H_2 persistence diagrams.

For each class and threshold, we report the *degeneracy rate* in each homology dimension, defined as the fraction of patches whose thresholded diagram is empty: $\delta_k(\varepsilon) = \Pr(|D_k| = 0)$. We additionally report the mean surviving feature count $\mathbb{E}[|D_k|]$. This audit is essential, since high degeneracy implies that any PD-vectorisation (including PI/PL) may be dominated by collapse patterns rather than rich persistence geometry. We also report the Topological Separability Ratio (TSR),

$$\text{TSR}(\varepsilon) = \frac{\mathbb{E}[d_B(\text{Pos}, \text{Frangi})]}{\mathbb{E}[d_B(\text{Pos}, \text{Pos})]},$$

which remains > 1 for both thresholds ($\text{TSR}_{0.15} \approx 1.79$, $\text{TSR}_{0.20} \approx 1.93$), confirming that inter-class separation is not merely intra-class noise.

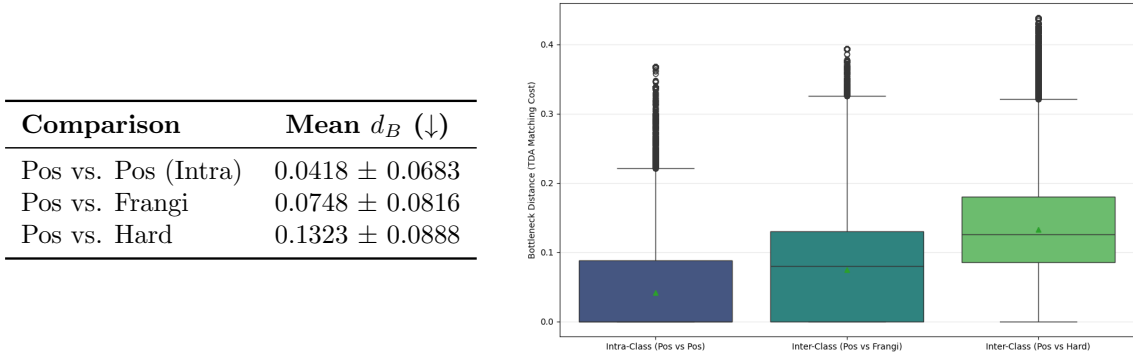


Figure 7: **Empirical Topological Separability via H_2 Bottleneck Distance.** The table (left) reports mean distances $\pm$ standard deviation, while the box plot (right) illustrates the full distributions across comparison groups (green triangles denote distribution means; horizontal lines denote medians). Aneurysm patches exhibit the highest topological similarity to each other (lowest d_B), while Frangi-mined bifurcations occupy an intermediate regime, mathematically validating their role as the primary geometric confounder in CNN-based detection.

Table 6: Empirical H_2 bottleneck separability across persistence thresholds ε .

ε	$\mathbb{E}[d_B(\text{Pos}, \text{Pos})]$	$\mathbb{E}[d_B(\text{Pos}, \text{Frangi})]$	$\mathbb{E}[d_B(\text{Pos}, \text{Hard})]$	$\text{TSR}(\varepsilon)$
0.15	0.0418	0.0748	0.1323	1.79
0.20	0.0311	0.0600	0.1168	1.93

For positives, H_2 degeneracy rises from 77.10% at $\varepsilon = 0.10$ to 85.01% at $\varepsilon = 0.15$ and 90.34% at $\varepsilon = 0.20$, while the mean surviving H_2 counts drop from 2.85 to 0.58 to 0.23 respectively. In contrast, negative classes retain more H_2 structure under the same thresholds. Notably, H_1 degeneracy is consistently lower than H_2 degeneracy in positives, suggesting that if persistence is to be used, loop-like structure (H_1) may carry more stable signal than volumetric voids (H_2) under thresholding.

E.3. Patch-Level Aneurysm Size Verification

Aneurysm size annotations are derived using a dedicated patch-driven pipeline that estimates the physical size of each aneurysm directly from imaging metadata. Unlike prior approaches that aggregate size at the patient or series level, this pipeline produces one size estimate per patch, ensuring correct handling of multi-aneurysm cases. As a result, performance differences observed across size strata can be attributed to representation quality rather than annotation artifacts.

Each patch is assigned to a size category using clinically standard thresholds– (i)**Small**: < 3 mm, (ii)**Medium**: 3 – 7mm and (iii)**Large**: > 7 mm. These thresholds align with prior clinical literature and reflect known differences in detection difficulty across aneurysm sizes.

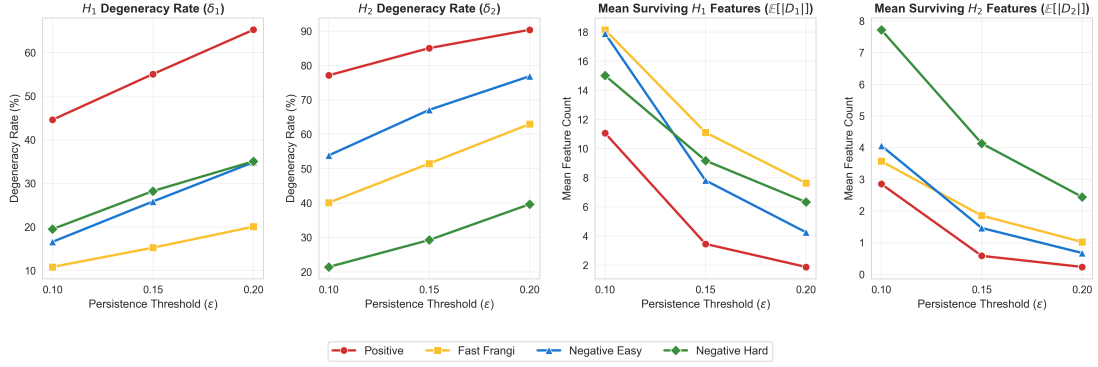


Figure 8: Topological degeneracy and feature survival across persistence thresholds (ϵ): From left to right, the panels illustrate the degeneracy rates—the percentage of patches yielding empty persistence diagrams—for 1-dimensional (H_1) and 2-dimensional (H_2) topological features, followed by the mean number of surviving features per diagram for H_1 and H_2 , respectively. Notably, true Positive aneurysm patches exhibit significantly higher degeneracy rates and lower surviving feature counts compared to Frangi, Hard, and Easy negative patches. This divergence underscores the morphological simplicity of true saccular aneurysms, which consist of isolated, prominent structures, whereas false positive mimics contain complex, persistent structural noise.

Each positive patch is identified by a unique `patch_id`, derived from its filename of the form: `<SeriesInstanceUID>_positive-<x>_<y>_<z>`. The spatial coordinates (x, y, z) are used to match each patch to its corresponding aneurysm annotation in the source dataset. Matching is performed in two stages:

- **Series-level filtering:** Candidate annotations are first restricted to the same `SeriesInstanceUID`.
- **Spatial association:** The patch location is matched to the closest annotation centroid in the (x, y) plane, optionally using aneurysm labels when available, or otherwise via proximity within a fixed tolerance.

This procedure ensures that each patch is associated with the correct aneurysm, even in the presence of multiple lesions within the same scan. :contentReference[oaicite:0]index=0

Aneurysm size is estimated in millimeters using a 3D bounding-box proxy derived from both in-plane and through-plane measurements:

- **XY diameter:** Computed as the spatial extent of annotation coordinates in the image plane, scaled by voxel spacing obtained from NIfTI headers.
- **Z height:** Computed from slice positions using DICOM `ImagePositionPatient` metadata. When full slice coverage is unavailable, fallback strategies based on slice thickness are used.

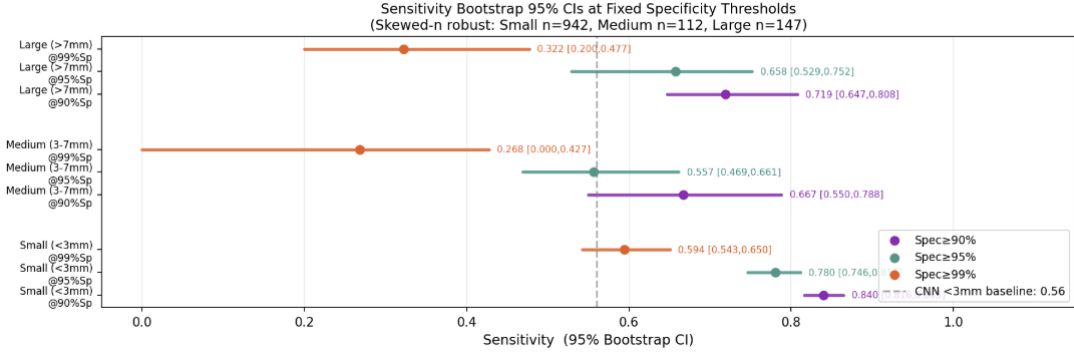


Figure 9: Size-stratified SECT performance with 95% bootstrap confidence intervals. Confidence intervals widen for medium and large strata due to limited sample size, while small aneurysms exhibit stable estimates despite class imbalance.

To provide a robust upper-bound estimate of aneurysm extent in 3D; the final size estimate (dome diameter proxy) is defined as:

$$\text{Dome}_{\text{mm}} = \max(\text{XY diameter}, \text{Z height}),$$

To ensure robustness across heterogeneous imaging data, the pipeline incorporates multiple fallback strategies — (i) If NIfTI spacing is unavailable, pixel-based XY measurements are used, (ii) If DICOM slice positions are incomplete, Z-extent is approximated using slice count and thickness, and (iii) If slice thickness is missing, a default physical thickness is used.

Patches for which no valid annotation can be matched are excluded from stratified.

To assess statistical differences across strata, we perform a Kruskal–Wallis test on predicted scores, which indicates a significant difference across size groups ($H = 37.23$, $p < 0.001$). Pairwise DeLong tests show that performance for small aneurysms is significantly higher than medium ($\Delta\text{AUC} = 0.0763$, $p = 0.0009$), while differences between medium and large are not statistically significant.

Sensitivity at fixed specificity further reflects the effect of sample size (see Fig.9) – confidence intervals remain tight for small aneurysms but widen substantially for medium and large strata, consistent with reduced statistical stability under limited sample sizes.. Diagnostic logging confirms near-complete correspondence between patch files and size entries.

E.4. Methodology and Domain Definitions for Empirical Lipschitz Analysis

To evaluate the robustness of our topological representations against scanner noise, we conducted an empirical Lipschitz stability analysis. We applied uniform physical perturbations to the source data by injecting Gaussian noise $\mathcal{N}(0, \sigma^2)$ directly into the normalized 3D CTA scalar fields, sweeping $\sigma \in \{0.01, 0.05, 0.10, 0.15, 0.20\}$.

It is critical to note that the empirical Lipschitz constant L is calculated across different input metric spaces depending on the representation. For Persistence Images (PI) and

Persistence Landscapes (PL), stability is traditionally bounded relative to the intermediate Persistence Diagram (PD). Let $W_p(D_1, D_2)$ denote the p -Wasserstein distance between diagrams. The empirical Lipschitz constant for PI and PL is calculated as:

$$L_\Phi \approx \max \frac{\|\Phi(D_{base}) - \Phi(D_{noisy})\|_2}{W_p(D_{base}, D_{noisy})} \quad \text{for } \Phi \in \{\text{PI}, \text{PL}\}$$

Conversely, the Smooth Euler Characteristic Transform (SECT) operates directly on the scalar field and does not compute a persistence diagram. Therefore, its stability is measured end-to-end against the intensity perturbation domain:

$$L_{SECT} \approx \max \frac{\|\text{SECT}(X_{base}) - \text{SECT}(X_{noisy})\|_2}{\sigma}$$

Because the denominators occupy different metric spaces (Wasserstein distance versus image intensity variation), the raw numerical L values cannot be compared directly. However, analyzing them within their respective domains reveals important behavioral differences: PI exhibits extreme sensitivity (high amplification) to minor topological shifts $L \approx 430.5$, whereas SECT demonstrates bounded, moderate variation $L \approx 40.3$ directly with respect to physical image noise.

E.5. Synthetic Validation of SECT under Controlled Geometric Perturbations

To validate that SECT captures intrinsic geometric differences between aneurysm-like and bifurcation-like structures, we conduct controlled synthetic experiments using parameterized 3D phantoms.

We generate two classes of volumetric phantoms: (i) **saccular structures**, modeled as a spherical outpouching attached to a parent vessel, and (ii) **bifurcation structures**, modeled as a branching tubular configuration. All phantoms are embedded in a fixed 32^3 grid, with geometry controlled via scale, orientation, and noise.

We systematically vary three factors—(i) **Noise**: $\sigma \in \{0.0, 0.15, 0.30\}$ (clean, moderate, heavy), (ii) **Scale**: $s \in \{0.6, 1.0, 1.5\}$ (small, baseline, large), and (iii) **Orientation**: fixed, limited ($\pm 30^\circ$), and random ($[0^\circ, 360^\circ)$). This results in $3 \times 3 \times 3 = 27$ conditions. For each condition, we generate $n = 60$ samples per class and compute SECT features using the same configuration as the main experiments. Classification is performed using a Random Forest with 5-fold cross-validation, and performance is reported as mean AUC with bootstrap confidence intervals.

SECT consistently separates saccular and bifurcation phantoms across most conditions. Under clean and moderate noise, performance is near-perfect ($\text{AUC} \approx 0.99$ across all scales and orientations). Under heavy noise, performance degrades but remains above chance in all cases, with mean $\text{AUC} 0.798 \pm 0.117$.

Across all 27 conditions, the overall mean AUC is 0.9295 ± 0.1152 . Performance decreases with increasing noise and random orientation, particularly for small-scale structures, but remains strongly discriminative in most regimes.

These results demonstrate that SECT captures stable geometric signatures distinguishing saccular and bifurcation structures, and that this signal persists under realistic perturbations in scale, orientation, and noise. Performance degradation under extreme noise conditions is expected, reflecting loss of topological structure in the underlying signal rather than failure of the representation.

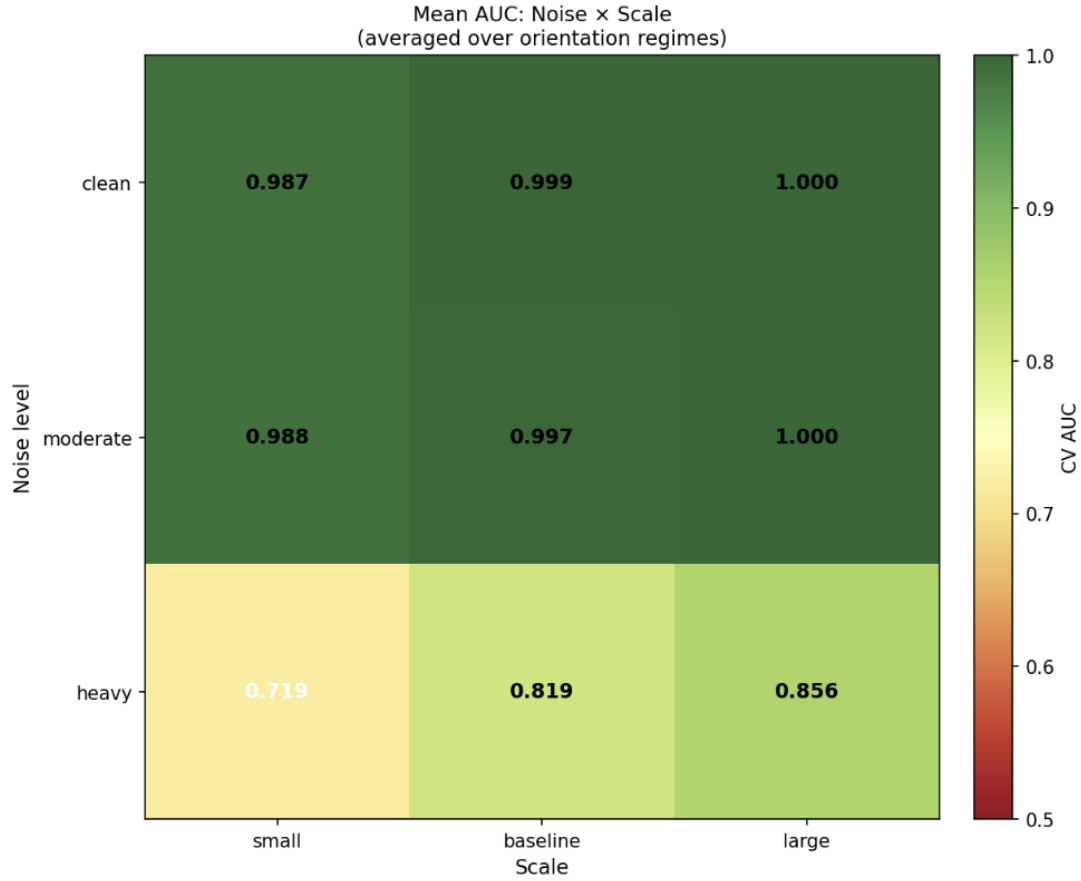


Figure 10: Synthetic validation of SECT under controlled perturbations. Heatmap shows mean AUC for distinguishing saccular and bifurcation phantoms across noise and scale conditions (averaged over orientations). SECT achieves near-perfect separation under clean and moderate noise, with degradation under heavy noise, particularly for small-scale structures.

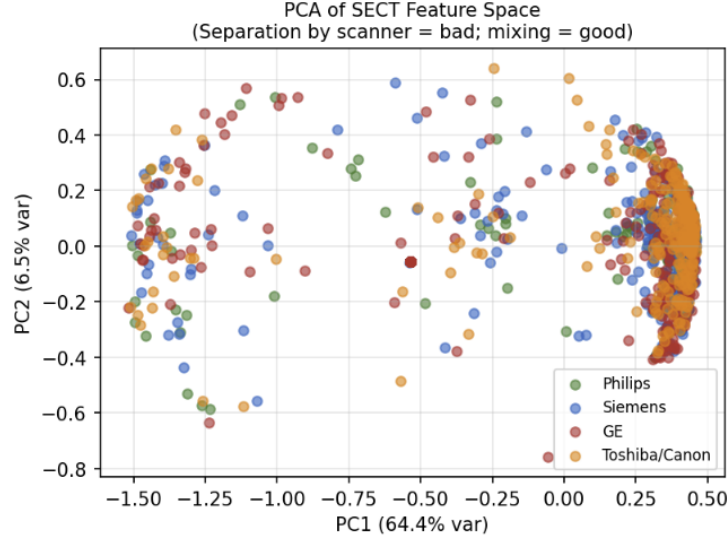


Figure 11: The scatter plot visualizes the first two principal components (accounting for 64.4% and 6.5% of the variance, respectively) for positive aneurysm patches sourced from a multi-institutional dataset.

E.6. Inter-Scanner Metadata and Feature-Space Analysis

For each unique imaging series, a single DICOM file is accessed to extract acquisition metadata, including manufacturer and model information. Manufacturer names are normalized into canonical groups (Siemens, GE, Philips, Toshiba/Canon) to account for variations in raw DICOM strings.

Each positive and negative patch, inherits the scanner label of its source series, resulting in a consistent patch-level assignment without requiring full-series reconstruction. This strategy ensures efficient metadata extraction while preserving correct scanner attribution across all samples.

The dataset exhibits substantial diversity across scanner manufacturer (see Appendix A.1) which reflects a multi-institutional acquisition setting and provides a suitable basis for evaluating scanner robustness.

E.6.1. FEATURE-SPACE ANALYSIS.

To assess whether SECT embeddings encode scanner-specific structure, we analyze the feature space using Principal Component Analysis (PCA). High-dimensional SECT vectors are projected onto a lower-dimensional subspace, and scanner labels are examined for clustering behavior. If SECT representations are scanner-agnostic, samples from different scanners should exhibit substantial overlap in the projected feature space rather than forming distinct clusters. Fig. 11 shows a heavy intertwining of distributions across Siemens, GE, Philips, and Toshiba/Canon scanners indicating strong feature-level mixing; although mild scanner-associated shifts are observable in lower-dimensional projections.

E.6.2. ANOVA-BASED SCANNER EFFECT QUANTIFICATION.

To quantify scanner-associated variation, we perform one-way ANOVA across scanner groups on the leading principal components. Specifically, we evaluate whether the distribution of feature values differs significantly across scanners for each principal component. The analysis yields the following results:

- Mean F-statistic (first 10 PCs): 5.42
- Number of significant PCs ($p < 0.05$): 7/10

The observed F-statistics indicate the presence of moderate scanner-associated variation in the SECT feature space. However, qualitative inspection of PCA projections shows substantial mixing of samples across scanner groups, with no clear separation by manufacturer. This suggests that while scanner identity contributes to feature variance, it does not dominate the representation. In particular, scanner-specific variation appears secondary to the underlying class-discriminative structure captured by SECT.

E.7. Detailed Operating-Point and False Positive Analysis

To assess the clinical viability of the topological representations, we conducted a granular operating-point analysis focusing on the high-sensitivity regime strictly required for intracranial aneurysm screening. Using out-of-fold predictions from a 5-fold cross-validation setup, we evaluated the False Positive Rate (FPR) at a fixed 90% sensitivity threshold. Additionally, we computed a bifurcation-specific FPR (utilizing the Youden-optimal threshold) to isolate the model’s performance on the most challenging morphological mimics: healthy vessel bifurcations.

As illustrated in the zoomed ROC analysis (Figure 12, Panel 2), SECT maintains robust discriminative power even in the aggressive True Positive Rate (TPR) $\geq 80\%$ regime. At a fixed 90% sensitivity, SECT achieves an FPR of 0.153 (Panel 3). In stark contrast, the direction-agnostic baselines suffer from severe specificity collapse, with PL and PI yielding FPRs of 0.572 and 0.623, respectively. This demonstrates that SECT can operate safely at clinical thresholds without inundating radiologists with false positives.

To confirm that SECT resolves the primary failure mode of standard CNNs, we isolated the false positive evaluation strictly to bifurcation mimics (Frangi negatives). Under the optimal classification threshold, SECT successfully suppressed the bifurcation-specific FPR to 0.108 (Figure 12, Panel 4). PI and PL, which lack the directional encoding necessary to capture asymmetric dome geometry, exhibited error rates nearly five times higher (0.516 and 0.517, respectively).

These results provide conclusive evidence that the performance gains of SECT are not merely aggregate improvements, but stem directly from its ability to resolve the specific geometric ambiguity between saccular aneurysms and branching vessels.

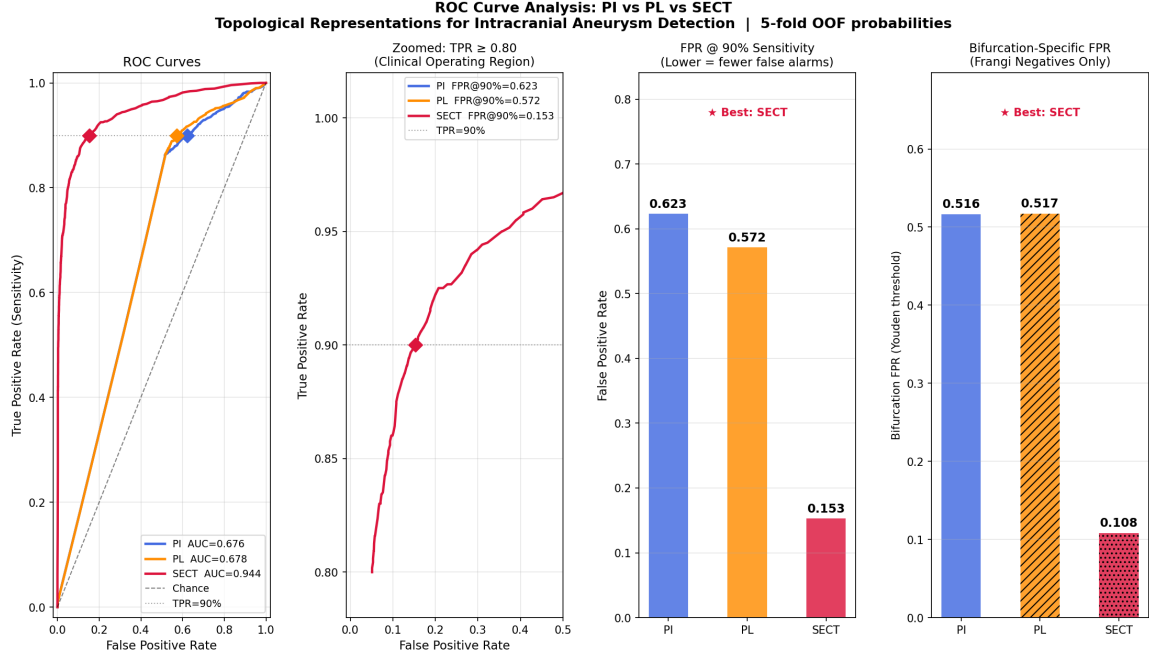


Figure 12: **Detailed Operating-Point and False Positive Analysis.** (Left to Right)
Panel 1: Standard ROC curves across the three topological representations.
Panel 2: Zoomed ROC focusing on the clinically critical high-sensitivity regime (TPR ≥ 0.80). **Panel 3:** False Positive Rate (FPR) forced at a strict 90% sensitivity. SECT reduces the false alarm rate to 15.3%, avoiding the specificity collapse seen in PI and PL. **Panel 4:** Bifurcation-specific FPR measured exclusively on Frangi-mined hard negatives, confirming SECT’s superior capacity to distinguish true saccular geometry from confounding vascular branches.