

Anatomy-Conditioned Domain Randomization for Zero-Shot 3D Cerebrovascular Segmentation

Matteo Pentassuglia¹[0009–0004–1132–4745], Xiaoming Zhang¹[0009–0000–3986–885X], Benjamin Billot³[0000–0002–3018–1282], and Maria A. Zuluaga^{1,2}[0000–0002–1147–766X]

¹ EURECOM, Biot, France

² School of Biomedical Engineering & Imaging Sciences, King’s College London, UK

³ Université Côte d’Azur, Inria, Epione team, France

maria.zuluaga@eurecom.fr

Abstract. Domain randomization methods have shown strong cross-domain robustness in neuroimaging, yet their use for 3D cerebrovascular segmentation remains limited. Existing vascular synthetic pipelines largely model vessel morphology and appearance in isolation, with limited explicit control of organ-level anatomical context during image generation. We propose an anatomy-conditioned DR framework for zero-shot 3D cerebrovascular segmentation using only synthetic training images that addresses these challenges through three task-specific components: anatomically consistent background context, morphology-aware vessel enrichment, and contrast-aware randomized label-to-image synthesis. A 3D CNN is trained exclusively on synthetically generated images and corresponding vessel masks, with no target-domain images or masks used for training or model selection. Under strict zero-shot evaluation on unseen TOF-MRA and CTA datasets, our method improves over a vessel-centric synthetic baseline and matches or exceeds a vessel-specific foundation model on CTA, while remaining competitive on TOF-MRA. These results show that anatomy-conditioned DR can provide a label-efficient alternative to large foundation-model pipelines for robust zero-shot 3D cerebrovascular segmentation. Our code is available at https://github.com/erc-caravel/anatomical_vessel_dr

Keywords: Cerebrovascular segmentation · Domain randomization · Anatomy-conditioned synthesis.

1 Introduction

3D brain vessel segmentation remains a challenging task due to the extreme sparsity and complex topology of vascular structures, as well as the high variability across imaging modalities and acquisition protocols. Although supervised deep learning approaches, typically 3D convolutional neural networks [13] achieve strong performance within individual datasets, they often fail to generalize to unseen domains, requiring new manual annotations for each site or modality.

Cross-domain robustness is typically pursued in two ways. Domain adaptation aligns source and target distributions using feature normalization, adversarial learning, or reconstruction objectives [6,14,16], but requires target-domain data and retraining for each deployment [5]. Domain generalization instead seeks robustness to unseen domains without target-domain data; a prominent strategy is large-scale pretraining of generic medical foundation models (FMs), such as VISTA3D and SAM-Med3D [8,17]. However, these generic FMs underperform on vascular segmentation in zero-/few-shot transfer [19], motivating vessel-specific FMs [10,19] that improve transfer but still require extensive expert-labelled data.

Domain randomization (DR) is attractive for cerebrovascular segmentation because it supports synthetic-image-only training with exact image-mask alignment by construction and robustness to modality and acquisition shifts without target-domain retraining [9]. In neuroimaging, SynthSeg [1] and SynthStrip [11] have shown that anatomically consistent synthetic images can enable cross-domain generalization without retraining. However, applying DR to cerebrovascular segmentation remains challenging for two reasons: vessels are extremely sparse ($\sim <1\%$ of intracranial volume), making anatomically consistent background modelling critical, and high-resolution vessel labels are scarce, limiting morphological priors for DR synthesis. Existing vascular DR pipelines [3,19] largely model vessel morphology and appearance in isolation, and do not explicitly account for organ-level anatomical context.

In this work, we propose an anatomy-conditioned domain randomization framework for strict zero-shot 3D brain vessel segmentation across angiographic modalities. Our contributions are threefold: (1) a synthetic-image training pipeline that conditions generation on whole-brain anatomical label maps to preserve global structure; (2) a vascular-specific synthesis strategy that combines anatomy-constrained vessel enrichment with contrast-aware label grouping in randomized label-to-image synthesis; and (3) a strict zero-shot evaluation on unseen TOF-MRA and CTA datasets showing strong gains over a synthetic vascular baseline and competitive performance with a vessel-specific foundation model despite training on synthetic images only.

2 Related Work

Vascular Segmentation Generalization. Beyond domain adaptation, vascular segmentation generalization has also been addressed, especially in retinal imaging, through domain-generalization methods based on augmentation-policy search and vessel-structure invariants [15,12]. A prominent alternative is large-scale supervised pretraining on heterogeneous annotated vascular datasets. Vessel-specific FMs follow this direction: tUbeNet [10] emphasizes few-shot specialization via fine-tuning after broad pretraining, whereas vesselFM [19] combines curated real data with synthetic data from domain randomization and flow matching for zero-/few-shot transfer. While effective, these methods rely heavily on extensive expert-labeled datasets.

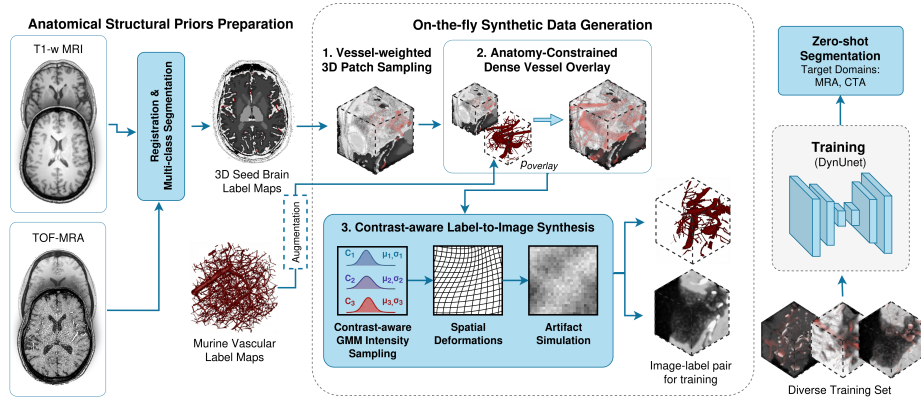


Fig. 1. Anatomy-conditioned DR framework. Whole-brain seed label maps provide anatomical priors. During training, vessel-weighted patches are sampled, optionally enriched with donor-vessel overlays, and synthesized into images via randomized label-to-image generation. A 3D segmentation network is trained on these synthetic samples and applied zero-shot to unseen MRA/CTA domains.

Synthetic Data Generation with Domain Randomization. Domain randomization (DR) uses a controllable parametric generator and randomizes parameters governing domain-varying factors such as appearance, artifacts, geometry, and background texture. By varying these factors across synthetic samples, DR encourages learning features that remain stable and transfer to unseen images without target-domain retraining [9]. In neuroimaging, anatomy-driven label-to-image DR frameworks [1,11] successfully used structured whole-brain label maps to preserve realistic spatial organization during randomization. In vascular segmentation, prior DR pipelines are mostly vessel-centric: a serial-section OCT pipeline synthesizes bifurcating spline-based vascular trees together with randomized geometric and noise-based backgrounds and simulated sOCT effects [3], while vesselFM’s DR framework combines transformed murine corrosion-cast vessel foreground patches with algorithmically generated randomized backgrounds built from geometric primitives and noise fields [18,19]. Neither explicitly conditions synthesis on dense organ-level anatomical segmentations, leaving global anatomical context unconstrained. In contrast, our approach conditions synthesis on dense whole-brain label maps to preserve anatomical context while enriching vessel morphology.

3 Method

We train a 3D vessel segmentation network using anatomy-conditioned domain randomization driven by dense whole-brain seed maps, targeting generalization across modalities and acquisition protocols.

We use a strict zero-shot protocol: no target-domain images or labels are used for training, validation, or model selection, and no test-time adaptation or fine-tuning is performed. Training samples are synthesized on the fly from dense whole-brain label maps using a randomized label-to-image generator that samples intensities and acquisition effects, yielding synthetic images with voxel-aligned vessel supervision by construction (Fig. 1). Our pipeline combines (i) anatomy-consistent background context from whole-brain seed maps, (ii) anatomy-constrained vessel enrichment to broaden morphological support, and (iii) contrast-aware label grouping within SynthSeg-style label-to-image synthesis.

3.1 Anatomical Structural Priors

Because cerebral vessels are sparse and embedded in complex brain anatomy, synthetic samples that do not preserve whole-brain organization can introduce implausible spatial relationships and bias learning. To enforce anatomically consistent structural context, we condition our DR framework on dense whole-brain label maps.

We construct multi-class 3D seed maps from paired T1-weighted (T1w) and TOF-MRA scans from the IXI dataset⁴. Whole-brain tissue segmentations are obtained from the T1w images using FastSurfer and registered to the corresponding TOF-MRA volumes. The registered tissue labels are aggregated into nine anatomical classes: cerebrospinal fluid (CSF), cortical gray matter, white matter, basal ganglia, ventricles, cerebellum, deep gray matter, brainstem, and hippocampus/amygdala. They are then combined with expert-reviewed intracranial vessel annotations from VesselVerse [4]. Because both labels are defined in TOF-MRA space after registration, they are voxel-aligned. We model non-brain and acquisition-dependent head/background structure by Gaussian-mixture clustering ($k = 5$) of non-brain voxels in the TOF-MRA volumes, yielding five coarse classes that capture heterogeneous surrounding tissues. Clustering is restricted to non-brain voxels, as intensity-based clustering in the brain would replace anatomical labels with contrast-specific partitions and weaken the synthesis prior. All seed maps are resampled to 0.5 mm isotropic resolution and used as structural priors for subsequent synthetic training sample generation. Although we use paired T1w scans to maximize seed-map quality, the framework only requires vessel-aligned anatomical labels; for MRI without paired structural scans, approximate labels from contrast-robust or modality-appropriate parcellation can still condition synthesis.

3.2 Sparse-Structure Learning

Vessel-weighted Patch Sampling. Cerebrovascular segmentation is highly class-imbalanced, with vessels occupying $< 1\%$ of intracranial volume; uniform patch sampling would therefore produce mainly background-dominated samples. We mitigate this by vessel-weighted sampling of 128^3 patches: 80% are centered

⁴ <https://brain-development.org/ixi-dataset/>

on vessel voxels, while 20% are sampled uniformly from non-vessel locations. Sampling is applied on the seed maps before vessel enrichment and label-to-image synthesis, balancing the foreground signal prior to appearance randomization.

Morphology-aware Vessel Enrichment. Available human intracranial vessel annotations [4], used in our seed maps, predominantly capture major arterial structures and under-represent the dense small-vessel patterns visible in higher-resolution angiographic acquisitions. To expand the training morphological variety without additional human annotations, we enrich the seed-map vessel labels with dense donor vessel foreground patches derived from graph-based murine corrosion-cast data [18].

Rather than modelling vessels independently of anatomy, we merge the murine donor vessel patch with the human vessel label within the sampled whole-brain seed map patch. During training, with probability $p_{overlay}$, a donor vessel patch is sampled, spatially transformed (flips and 90° rotations), cropped to 128^3 , and merged with the human vessel label of the sampled seed patch. To preserve anatomical plausibility, we apply an intracranial masking constraint: donor vessel voxels are retained only within valid intracranial regions defined by brain tissue classes, and voxels outside these regions are discarded. Thus, the murine donor patch modifies local vessel morphology, while global human anatomical context remains determined by the seed map.

3.3 Contrast-aware Label-to-image Generator

Given the sampled seed-map patch after vessel enrichment, we synthesize image patches on the fly via label-to-image domain randomization, keeping anatomy fixed while varying appearance across training iterations. To match angiographic imaging, we use a contrast-aware label grouping strategy for intensity sampling. In TOF-MRA and CTA, vessel visibility is driven primarily by flow or contrast enhancement rather than detailed T1/T2-like tissue contrast; assigning distinct intensity distributions to all brain tissues could encourage reliance on tissue boundaries whose contrast is not consistently informative across domains. Specifically, the nine brain tissue classes are merged into two groups (CSF+ventricles, and all solid brain tissues), while vessels and a true-background class are modelled separately. Of the five clustered non-brain classes, the remaining four are grouped and, with 50% probability, mapped to true background to improve robustness to skull-stripped acquisitions. Synthesis then applies randomized class-conditional Gaussian mixture model (GMM) intensity sampling, spatial deformations, intensity perturbations, resolution simulation, and Gaussian noise. The output is a synthetic 128^3 image patch paired with a binary vessel target derived from the same spatially transformed label patch (vessel / non-vessel).

4 Experimental Setup

Evaluation Datasets. We evaluate strict zero-shot generalization on three cerebrovascular datasets spanning different modalities and cohorts: (i) **CereVess-MRA** [7] ($N = 121$): TOF-MRA scans with vessel annotations. To preserve



Fig. 2. 3D renderings of predicted vessel masks on a TopCow24 CT case obtained using vesselFM-DR, vesselFM-full, and our method, alongside the ground truth (GT).

strict zero-shot conditions, we use only the 121 pathological cases and exclude healthy subjects overlapping with IXI. (ii) **SMILE-UHURA** [2] ($N = 14$): High-resolution 7T TOF-MRA angiographies with detailed expert vessel annotations. (iii) **TopCow24 CT** [20] ($N = 35$): CTA scans with manually refined binary vessel labels sourced from VesselVerse [4]. We retain only the 35 cases added in 2024, as the 2023 cases are part of vesselFM’s training data [19].

Baselines. We compare our proposed framework against three baselines: (i) **Supervised IXI baseline.** A 3D CNN trained directly on the real IXI TOF-MRA scans and corresponding VesselVerse vessel annotations from the same 112 subjects used to construct anatomical seed maps. Images and labels are resampled to 0.5 mm isotropic resolution. Training uses vessel-weighted patch sampling and standard augmentation consistent with vesselFM [19]. (ii) **vesselFM-DR.** We reproduce the vesselFM DR pipeline and train on 50k synthetic patches using the 25 publicly available graph-derived murine foreground volumes from [18]. This provides a reproducible vessel-centric DR reference (no anatomical conditioning) and is conservative relative to the original vesselFM study, which trains on 500k synthetic patches with full foreground access. (iii) **vesselFM-full.** The publicly released vesselFM checkpoint trained on the full real and synthetic dataset [19].

Evaluation Protocol. All models are evaluated using identical preprocessing, inference, and post-processing settings, using 3D sliding-window inference with 128^3 patches and 50% overlap. Logits from overlapping windows are averaged prior to thresholding. For SMILE-UHURA and TopCow24, intensities are normalized using percentile-based clipping (1st–99th percentile) followed by linear scaling to $[0, 1]$. CereVessMRA volumes are resampled to an isotropic resolution of 0.5 mm. Performance is reported as mean $\pm$ standard deviation Dice and cDice scores across subjects. Post-hoc paired two-sided Wilcoxon tests (Bonferroni-corrected, $\alpha = 0.05/3$) were used for pairwise comparisons per dataset and metric. No images or labels from the evaluation datasets are used for training, validation, or hyperparameter tuning at any stage.

Implementation Details. All models use MONAI’s DynUNet, adopting vesselFM’s backbone [19], and are trained with batch size 8, Dice + Cross-Entropy loss (0.9/0.1), Adam, and linear warmup for 46,000 iterations, matching ves-

Table 1. Zero-shot segmentation performance. We report mean (standard deviation); best results are in **bold**. * marks the best method when it significantly outperforms all others in the same dataset/metric column (paired two-sided Wilcoxon, Bonferroni-corrected $p_{\text{corr}} < 10^{-2}$).

	CereVessMRA		SMILE-UHURA		TopCow24 CT	
	Dice $\uparrow$	clDice $\uparrow$	Dice $\uparrow$	clDice $\uparrow$	Dice $\uparrow$	clDice $\uparrow$
Supervised IXI	64.2* (3.0)	67.2 (3.3)	59.2 (5.2)	71.0 (4.8)	1.7 (2.0)	3.7 (4.1)
vesselFM-DR	22.1 (4.4)	31.1 (5.2)	43.2 (6.9)	53.6 (4.0)	4.8 (2.3)	9.7 (3.9)
vesselFM-full	53.5 (8.7)	60.5 (8.5)	70.8 (6.9)	71.0 (5.0)	27.4 (9.9)	33.4 (9.7)
Our method	58.7 (4.5)	67.0 (4.7)	68.2 (6.9)	72.7 (5.2)	41.3*(8.0)	49.6*(8.2)

selfFM’s training budget. For the Supervised IXI and vesselFM-DR baselines, we adopt the hyperparameters reported in [19]: initial learning rate 10^{-5} , weight decay 0.1, and cosine annealing with $T_{\text{max}} = 5000$. For the proposed model, synthetic samples are generated on-the-fly from 112 IXI-derived seed label maps at each iteration; no synthetic dataset is stored offline. Sample generation averages 0.33 s per 128^3 patch, ranging from 0.24 to 0.65 s. Label-to-image synthesis is implemented using the `cornucopia` library with default transform settings and no dataset-specific tuning. To accommodate higher variance introduced by continuous domain randomization, we use an initial learning rate of 5×10^{-4} , weight decay 10^{-2} , and cosine annealing with $T_{\text{max}} = 20,000$. Unless otherwise stated, $p_{\text{overlay}} = 0.4$. Checkpoint selection uses no target-domain data. Validation relies only on held-out IXI and/or fixed synthetic samples, depending on the model. All experiments are conducted on a single NVIDIA A100 GPU.

5 Results

Benchmark Table 1 summarizes the zero-shot performance of baselines and our model on the three test sets. On CereVessMRA, the supervised IXI baseline achieves the highest Dice, reflecting the structural similarity between this dataset and the IXI source domain. Our synthetic-only DR model ranks second and improves over vesselFM-full, particularly in clDice; paired post-hoc tests confirm significant gains over vesselFM-full for both Dice and clDice. While Supervised IXI remains superior in Dice, clDice shows no significant difference between the two models, indicating comparable topological consistency. On SMILE-UHURA, our method is comparable to vesselFM-full, with slightly lower Dice but higher clDice; the clDice improvement over vesselFM-full is significant, whereas the Dice difference is not. Crucially, under the strongest domain shift (TopCow24 CT), our method achieves the best results (41.3 Dice), outperforming vesselFM-full (27.4 Dice); these gains are significant for both metrics. This improvement is notable given that vesselFM-full includes TopCow23 CTA data in its real training set, albeit with CoW-focused labeling [20]. Finally, vesselFM-DR performs poorly across targets and tends to over-segment non-vascular structures; our method is significantly better on all datasets and metrics, consistent with the benefit of enforcing anatomically structured context during synthesis (Fig. 2).

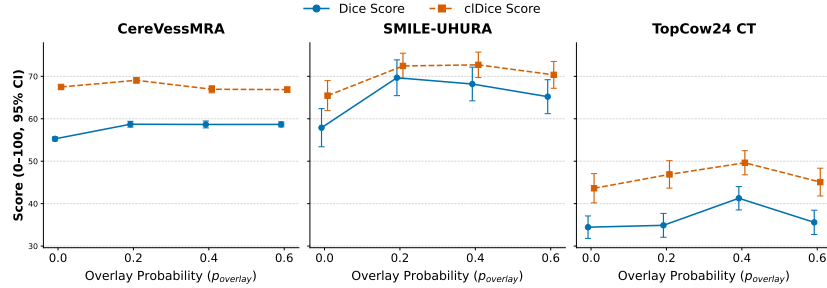


Fig. 3. Effect of p_{overlay} on zero-shot performance. Curves report mean Dice and cDice across volumes. Error bars show 95% confidence intervals over test subjects.

Ablation: Dense vessel overlay probability. We quantify the effect of dense vessel enrichment by training variants with $p_{\text{overlay}} \in \{0, 0.2, 0.4, 0.6\}$ (Figure 3). On SMILE-UHURA, introducing overlays already at $p_{\text{overlay}} = 0.2$ substantially improves performance (Dice increases from 57.9 to 69.7, cDice from 65.5 to 72.5), indicating that donor-vessel enrichment helps bridge the morphological gap to higher-resolution angiography. In contrast, gains are less pronounced on the other target datasets, whose annotations contain less fine vasculature. TopCow24 CT performs best at a moderate overlay probability ($p_{\text{overlay}} = 0.4$). Increasing it to $p_{\text{overlay}} = 0.6$ degrades performance on both SMILE-UHURA and TopCow24 CT. Overall, the non-monotonic trend supports dense overlay as a controlled morphology-enrichment mechanism, with the largest benefit when the target domain exhibits richer small-vessel detail than the base human vessel priors.

Ablation: Anatomy-constrained vessel enrichment. We ablate the intracranial masking constraint in the vessel-enrichment step by merging donor vessel voxels without restricting overlays to intracranial regions defined by the seed map, keeping all other settings fixed ($p_{\text{overlay}}=0.4$). The effect is modality-dependent. On TopCow24 CT, the constraint is essential, improving Dice from 28.2 to 41.3 ($p < 10^{-6}$). On TOF-MRA, Dice decreases slightly: SMILE-UHURA 69.4 vs 68.2 ($p=0.0676$) and CereVessMRA 60.5 vs 58.7 ($p=10^{-6}$). All p -values are from paired Wilcoxon tests. cDice exhibits the same modality-dependent trend. Overall, intracranial masking is essential for CTA robustness, preventing off-target insertions that can mimic vessel-like CTA appearance.

6 Conclusion

We proposed an anatomy-conditioned domain-randomization pipeline for strict zero-shot 3D cerebrovascular segmentation that generates training samples on the fly by combining whole-brain label maps, anatomy-constrained vessel enrichment, and contrast-aware randomized label-to-image synthesis. Training relies exclusively on synthetic data, isolating the contribution of anatomy-conditioned DR without access to real target-domain images or annotations. This setting

targets new domains where task-specific annotations are unavailable, costly, or inconsistent across labelling protocols; when in-domain accuracy is the primary goal, hybrid real-synthetic or target-supervised training may be preferable. When evaluated on unseen TOF-MRA and CTA datasets, our method outperformed a synthetic vascular baseline and achieved competitive performance with a vessel-specific foundation model, with the strongest gains on CTA. These findings indicate that enforcing organ-level anatomical context during domain randomization enhances robustness under pronounced modality shift, supporting anatomy-aware DR as a label-efficient and reproducible alternative for cross-domain cerebrovascular segmentation. Future work will study adaptive or learned synthesis priors, broader anatomical source datasets, and extension to additional vascular territories and pathologies.

Acknowledgments. This work is funded by the European Union (ERC CARAVEL, 101171357). We acknowledge computational resources provided by GENCI-IDRIS on Jean Zay A100 partition (grant 2026-A0200617550) and by the EuroHPC JU on LEO-NARDO, hosted by CINECA, Italy (EHPC-DEV-2026D01-008).

Disclosure of Interests. The authors have no competing interests to declare that are relevant to the content of this article. Views and opinions expressed are those of the authors only and do not necessarily reflect those of the EU or the European Research Council. Neither the EU nor the granting authority can be held responsible for them.

References

1. Billot, B., Greve, D.N., Puonti, O., Thielscher, A., Van Leemput, K., Fischl, B., Dalca, A.V., Iglesias, J.E.: SynthSeg: Segmentation of brain MRI scans of any contrast and resolution without retraining. *Medical Image Analysis* **86**, 102789 (2023)
2. Chatterjee, S., Mattern, H., Dörner, M., Sciarra, A., Dubost, F., Schnurre, H., Khatun, R., Yu, C.C., Hsieh, T.L., Tsai, Y.S., Fang, Y.Z., Yang, Y.C., Huang, J.D., Xu, M., Liu, S., Ribeiro, F.L., Bollmann, S., Chintalapati, K.V., Radhakrishna, C.M., Kumara, S.C.H.R., Sutrave, R., Qayyum, A., Mazher, M., Razzak, I., Rodero, C., Niederren, S., Lin, F., Xia, Y., Wang, J., Qiu, R., Wang, L., Panah, A.Y., Jurdi, R.E., Fu, G., Arslan, J., Vaillant, G., Valabregue, R., Dormont, D., Stankoff, B., Colliot, O., Vargas, L., Chacón, I.D., Pitsiorlas, I., Arbeláez, P., Zuluaga, M.A., Schreiber, S., Speck, O., Nürnberger, A.: SMILE-UHURA Challenge – Small Vessel Segmentation at Mesoscopic Scale from Ultra-High Resolution 7T Magnetic Resonance Angiograms (Nov 2024). <https://doi.org/10.48550/arXiv.2411.09593>
3. Chollet, E., Balbastre, Y., Mauri, C., Magnain, C., Fischl, B., Wang, H.: Neurovascular Segmentation in sOCT with Deep Learning and Synthetic Training Data (Jul 2024). <https://doi.org/10.48550/arXiv.2407.01419>
4. Falcetta, D., Marciano, V., Yang, K., Cleary, J., Legris, L., Rizzaro, M.D., Pitsiorlas, I., Chaptoukaev, H., Lemasson, B., Menze, B., Zuluaga, M.A.: VesselVerse: A Dataset and Collaborative Framework for Vessel Annotation. In: Gee, J.C., Alexander, D.C., Hong, J., Iglesias, J.E., Sudre, C.H., Venkataraman, A., Golland,

- P., Kim, J.H., Park, J. (eds.) Medical Image Computing and Computer Assisted Intervention – MICCAI 2025. pp. 655–665 (2026)
5. Galati, F., Falcetta, D., Cortese, R., Prados, F., Burgos, N., Zuluaga, M.A.: Multi-Domain Brain Vessel Segmentation Through Feature Disentanglement. *Machine Learning for Biomedical Imaging* **3**(September 2025 issue), 477–495 (Sep 2025)
6. Gu, R., Zhang, J., Wang, G., Lei, W., Song, T., Zhang, X., Li, K., Zhang, S.: Contrastive Semi-Supervised Learning for Domain Adaptive Segmentation Across Similar Anatomical Structures. *IEEE Transactions on Medical Imaging* **42**(1), 245–256 (Jan 2023)
7. Guo, B., Chen, Y., Lin, J., Huang, B., Bai, X., Guo, C., Gao, B., Gong, Q., Bai, X.: Self-supervised learning for accurately modelling hierarchical evolutionary patterns of cerebrovasculature. *Nature Communications* **15**(1), 9235 (Oct 2024)
8. He, Y., Guo, P., Tang, Y., Myronenko, A., Nath, V., Xu, Z., Yang, D., Zhao, C., Simon, B., Belue, M., Harmon, S., Turkbey, B., Xu, D., Li, W.: VISTA3D: A Unified Segmentation Foundation Model For 3D Medical Imaging. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 20863–20873 (2025)
9. Hoffmann, M.: Domain-Randomized Deep Learning for Neuroimage Analysis: Selecting Training Strategies, Navigating Challenges, and Maximizing Benefits. *IEEE signal processing magazine* **42**(4), 78–90 (Jul 2025)
10. Holroyd, N.A., Li, Z., Walsh, C., Brown, E., Shipley, R.J., Walker-Samuel, S.: tUbeNet: A generalizable deep learning tool for 3D vessel segmentation. *Biology Methods and Protocols* **10**(1), bpaf087 (Jan 2025)
11. Hoopes, A., Mora, J.S., Dalca, A.V., Fischl, B., Hoffmann, M.: SynthStrip: Skull-stripping for any brain image. *NeuroImage* **260**, 119474 (Oct 2022)
12. Hu, D., Li, H., Liu, H., Oguz, I.: Domain generalization for retinal vessel segmentation via Hessian-based vector field. *Medical Image Analysis* **95**, 103164 (Jul 2024)
13. Isensee, F., Jaeger, P.F., Kohl, S.A., Petersen, J., Maier-Hein, K.H.: nnu-net: a self-configuring method for deep learning-based biomedical image segmentation. *Nature methods* **18**(2), 203–211 (2021)
14. Kumari, S., Singh, P.: Deep learning for unsupervised domain adaptation in medical imaging: Recent advancements and future perspectives. *Computers in Biology and Medicine* **170**, 107912 (Mar 2024)
15. Lyu, J., Zhang, Y., Huang, Y., Lin, L., Cheng, P., Tang, X.: AADG: Automatic Augmentation for Domain Generalization on Retinal Image Segmentation. *IEEE Transactions on Medical Imaging* **41**(12), 3699–3711 (Dec 2022)
16. Peng, L., Lin, L., Cheng, P., Huang, Z., Tang, X.: Unsupervised Domain Adaptation for Cross-Modality Retinal Vessel Segmentation via Disentangling Representation Style Transfer and Collaborative Consistency Learning. In: *2022 IEEE 19th International Symposium on Biomedical Imaging (ISBI)*. pp. 1–5 (Mar 2022)
17. Wang, H., Guo, S., Ye, J., Deng, Z., Cheng, J., Li, T., Chen, J., Su, Y., Huang, Z., Shen, Y., Fu, B., Zhang, S., He, J.: SAM-Med3D: A Vision Foundation Model for General-Purpose Segmentation on Volumetric Medical Images. *IEEE Transactions on Neural Networks and Learning Systems* **36**(10), 17599–17612 (Oct 2025)
18. Wittmann, B., Glandorf, L., Paetzold, J.C., Amiranashvili, T., Wälchli, T., Razansky, D., Menze, B.: Simulation-Based Segmentation of Blood Vessels in Cerebral 3D OCTA Images. In: *Linguraru, M.G., Dou, Q., Feragen, A., Giannarou, S., Glocker, B., Lekadir, K., Schnabel, J.A. (eds.) Medical Image Computing and Computer Assisted Intervention – MICCAI 2024*. pp. 645–655 (2024)

19. Wittmann, B., Wattenberg, Y., Amiranashvili, T., Shit, S., Menze, B.: vesselFM: A Foundation Model for Universal 3D Blood Vessel Segmentation. In: 2025 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 20874–20884 (Jun 2025)
20. Yang, K., Musio, F., Ma, Y., Juchler, N., Paetzold, J.C., Al-Maskari, R., Höher, L., Li, H.B., Hamamci, I.E., Sekuboyina, A., Shit, S., Huang, H., Prabhakar, C., de la Rosa, E., Wittmann, B., Waldmannstetter, D., Kofler, F., Navarro, F., Menten, M.J., Ezhov, I., Rueckert, D., Vos, I.N., Ruigrok, Y.M., Velthuis, B.K., Kuijff, H.J., Shi, P., Liu, W., Ma, T., Rokuss, M.R., Kirchhoff, Y., Isensee, F., Maier-Hein, K., Zhu, C., Zhao, H., Bijlenga, P., Hämmerli, J., Wurster, C., Westphal, L., Bisschop, J., Colombo, E., Baazaoui, H., Handelsmann, H.L., Makmur, A., Hallinan, J., Soundararajan, A., Wiestler, B., Kirschke, J.S., Wiest, R., Montagnon, E., Letourneau-Guillon, L., Oh, K., Lee, D., Aydin, O.U., Hilbert, A., Rieger, J., Rallios, D., Tanioka, S., Koch, A., Frey, D., Qayyum, A., Mazher, M., Niederer, S., Disch, N., Holzschuh, J., LaBella, D., Galati, F., Falcetta, D., Zuluaga, M.A., Lin, C., Zhao, H., Zhang, Z., Zhang, M., You, X., Zhang, H., Yang, G.Z., Gu, Y., Ra, S., Hwang, J., Park, H., Chen, J., Wodzinski, M., Müller, H., Mansouri, N., Autrusseau, F., Yalçin, C., Hamadache, R.E., Lisazo, C., Salvi, J., Casamitjana, A., Lladó, X., Estrada, U.M.L.T., Abramova, V., Giancardo, L., Oliver, A., Casademunt, P., Galdran, A., Delucchi, M., Liu, J., Huang, H., Cui, Y., Lin, Z., Liu, Y., Zhu, S., Patel, T.R., Siddiqui, A.H., Tutino, V.M., Orouskhani, M., Wang, H., Mossa-Basha, M., Sato, Y., Hirsch, S., Wegener, S., Menze, B.: Benchmarking the CoW with the TopCoW Challenge: Topology-Aware Anatomical Segmentation of the Circle of Willis for CTA and MRA. ArXiv p. arXiv:2312.17670v4 (Jul 2025). <https://doi.org/10.48550/arXiv.2312.17670>