



This MICCAI paper is the Open Access version, provided by the MICCAI Society. It is identical to the accepted version, except for the format and this watermark; the final published version is available on SpringerLink.

Lesion-Aware Hybrid State-Space and Residual-Encoder Fusion for Brain Metastasis Segmentation

Yutong Bian¹ and Jindong Sun^{1,2*}

¹ College of Intelligent Equipment, Shandong University of Science and Technology, Taian 271001, China

202583230013@sdust.edu.cn; jdsun@sdust.edu.cn

² Bioengineering Department and Imperial-X, Imperial College London, London W12 0BZ, U.K.

Abstract. Accurate brain metastasis segmentation must balance overlap, boundary fidelity, and lesion detection, especially for small multifocal lesions and postoperative resection cavities. We present METS-HybridFusion, the submitted frozen Docker inference pipeline for BraTS-METS 2026 Task 1. It combines a modality-aware state-space semantic expert and a residual-encoder region expert; their ET, RC, TC, and WT probabilities are fused by region, reconstructed hierarchically, and refined by RC gating and component verification. A conservative auxiliary lesion-recovery step is included only for rare case-level recovery rather than as a primary source of aggregate metric improvement.

In a development-set evaluator-aligned local analysis on the 202-scan, 122-patient split, BraTS-evaluation v0.0.8 yielded mean large-lesion DSC, large-lesion NSD, all-instance lesion-wise F1, and small-lesion F1 of 0.4704, 0.4908, 0.5841, and 0.2377, respectively. Patient-level cluster bootstrap showed wide RC uncertainty, including RC large-lesion DSC 0.1714 [0.0189, 0.3628]. These values are not Synapse validation-server results, because they were computed locally on our 202-scan development split. They should also not be interpreted as independent generalization estimates, because the same split was used for model selection, threshold tuning, OOF fitting of lightweight refinement modules, and diagnostic analysis.

Keywords: Brain metastasis segmentation · BraTS-METS · Mamba · residual encoder · lesion detection · ensemble learning

1 Introduction

Brain metastases may be spatially separated, small, heterogeneous, and affected by treatment. BraTS-METS extends the BraTS evaluation framework [13] to this setting, where small enhancing foci and postoperative resection cavities remain difficult [3,9,14,12]. Task 1 evaluates enhancing tumor (ET), tumor core (TC),

* Corresponding author.

resection cavity (RC), and whole tumor (WT) using large-lesion DSC, large-lesion normalized surface Dice (NSD), and lesion-level detection endpoints. We report all-instance lesion-wise F1 for overall lesion detection and small-lesion F1 for small-metastasis detection; RC adds severe class imbalance and visual heterogeneity.

We combine a Mamba-based semantic expert, which models long-range context efficiently [4], with a strong residual-encoder region expert inspired by nnU-Net and modern 3D ConvNets [6,18,7]. Their complementary predictions are fused per region and controlled at connected-component level. The resulting METS-HybridFusion pipeline integrates complementary state-space semantic and residual region experts, region-specific fusion with hierarchy reconstruction, and RC gating with learned component verification. We also include a conservative auxiliary lesion-recovery refinement that uses image and cross-model evidence for rare candidate recovery, while treating it as a precision-constrained safeguard rather than a central performance driver.

Unless otherwise stated, main results are development-set evaluator-aligned local metrics on the 202-scan, 122-patient split, computed with BraTS-evaluation v0.0.8. Because this split was also used for tuning and diagnostics, the values are not Synapse validation-server results. Related volumetric CNN, transformer, and state-space segmentation work motivates our pragmatic design: global context is modeled at low resolution, whereas overlap-sensitive regions and lesion-level decisions are handled by explicit region and component modules [17,1,8,15,6,19,5,18,7,4,10,20].

2 Dataset and Evaluation Protocol

Task and data. BraTS-METS 2026 Task 1 requires a single Dockerized method for four-modality mpMRI segmentation. Labels are background, non-enhancing tumor core (NETC), surrounding non-enhancing FLAIR hyperintensity (SNFH), enhancing tumor (ET), and resection cavity (RC). Composite regions are ET, RC, tumor core ($TC = \text{NETC} + \text{ET}$), and whole tumor ($WT = \text{NETC} + \text{SNFH} + \text{ET}$); RC is exclusive and not included in TC or WT.

The development set contained 1,296 scans from 810 patients, including 167 RC-positive scans. We used a patient-grouped split with 1,094 training scans and 202 validation scans from 122 patients. The 202-scan split contained 17 RC-positive scans and was used for model selection, threshold tuning, staged analysis, failure analysis, OOF fitting and threshold selection of lightweight refinement modules, and final internal evaluation. Independent Synapse validation-server scores for the 179-scan challenge validation cohort are not reported. Two development labels (BraTS-MET-01094-003 and BraTS-MET-01184-002) were replaced only with corrected annotations released by the organizers; no voxel labels were manually edited by the authors.

Evaluation metrics. We report evaluator-aligned large-lesion DSC [2], large-lesion NSD [16], all-instance lesion-wise F1, and small-lesion F1 for ET, RC,

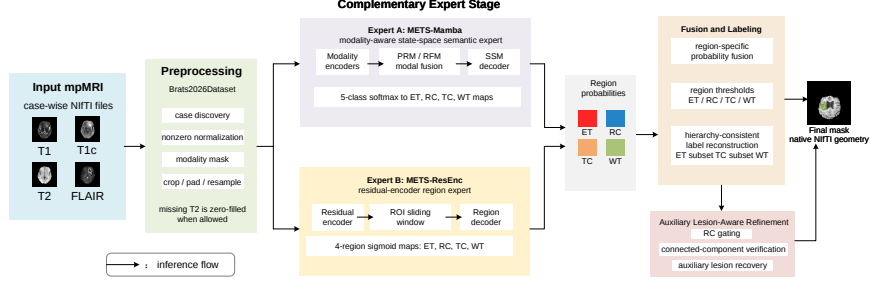


Fig. 1. Submitted frozen METS-HybridFusion Docker pipeline from four normalized mpMRI modalities to native-space segmentation.

TC, and WT. Lesion-wise F1 counts connected-component true positives, false positives, and false negatives; small-lesion F1 isolates the harder small-metastasis endpoint. Because challenge ranking is multi-metric and region-aware [11], we report means and weak regions, with explicit empty-mask handling for rare RC cases.

3 Method

Pipeline. Figure 1 shows the frozen pipeline: complementary experts, region-specific fusion, hierarchy reconstruction, RC gating, component verification, auxiliary lesion recovery, and native-space restoration. All stages and thresholds were fixed before Docker submission.

State-space semantic expert. The semantic expert uses modality-specific residual encoders, region-aware fusion blocks, and a selective state-space bottleneck to predict five exclusive semantic classes, which are converted into ET, RC, TC, and WT probabilities.

Residual-encoder region expert. The residual expert is a high-capacity residual-encoder 3D U-Net that directly predicts overlapping ET, RC, TC, and WT regions, with auxiliary boundary, center-detection, and hierarchy losses for boundary fidelity and lesion sensitivity.

Region-wise fusion and reconstruction. For region r , expert probabilities are combined as

$$p_r = \frac{\alpha_r p_r^M + \beta_r p_r^R}{\alpha_r + \beta_r}, \quad (1)$$

where p_r^M and p_r^R are the state-space and residual-encoder probabilities. Fusion weights and region thresholds were selected on the 202-scan, 122-patient internal validation split during development, frozen before Docker submission, and recorded in the released GitHub configuration files. The binary regions are then reconstructed into an exclusive label map with $ET \subseteq TC \subseteq WT$ and RC competition.

Component verification and auxiliary lesion recovery. The component verifier removes likely false-positive WT components using volume, confidence, cross-expert agreement, compactness, and core-support features. It was fitted as weighted logistic regression on 1,492 predicted components from the 202-scan development-validation split, with positives defined by $\text{IoU} \geq 0.01$ against a ground-truth WT component (990 positive, 502 negative). Twelve standardized features were used in patient-wise five-fold OOF training, and the threshold 0.0747 was selected only from OOF predictions by requiring recall ≥ 0.995 and maximizing a component-F1/false-removal objective. The OOF audit yielded precision 0.7273, recall 0.9970, and F1 0.8411, removing 132/502 false components while incorrectly removing 3/990 positives; mean DSC and F1 changed by +0.00166 and +0.01027. Because both fitting and threshold selection used the development-validation split, these quantities are reported as development diagnostics rather than independent validation evidence.

The auxiliary lesion-recovery step generated patch candidates from low-threshold WT residuals on the same 202-scan development-validation split: the union of fused, Mamba, and residual probabilities at 0.25 after excluding occupied tumor voxels and a one-iteration dilation. Components of 2–4,096 mm³ yielded 5,143 candidates (128 positive, 5,015 negative) and a 1,228-lesion bank. Patient-wise five-fold 32³ residual CNN classifiers used lesion-bank positives, candidate patches, 15 scalar descriptors, and hard negatives capped at 12 per positive; each fold excluded the held-out patients from candidate and lesion-bank training patches. The 0.98 threshold was selected only from OOF predictions to prioritize precision (precision ≥ 0.75 , recall ≥ 0.02 , at least three rescued true candidates, and no unacceptable metric degradation). The OOF audit rescued 7 true candidates and added 2 false candidates, with mean DSC/NSD/F1 changes of $-0.000005/-0.000002/+0.000351$ and small-lesion F1 change +0.001460; accordingly, this step is reported as an auxiliary case-level development diagnostic, not as a module that materially improves independent aggregate DSC, NSD, or F1.

4 Experiments

Implementation details. The Mamba semantic expert used 128³ patches, AdamW, and detection-aware supervision for 300 epochs. The residual-encoder expert used 128³ patches, effective batch size 2, and region, boundary, detection, focal-style, and Tversky losses for 250 epochs. Training-set diagnostics were recorded from the final training epochs: the Mamba expert ended with training loss 0.5963 and patch-level training Dice of 0.8294/0.1615/0.8595/0.8737 for ET/RC/TC/WT, whereas the residual-encoder expert ended with training loss 0.4448 and patch-level training Dice of 0.6657/0.0000/0.6952/0.7191. These training diagnostics were used only to document optimization behavior and are not directly comparable to the full-volume evaluator-aligned metrics. Docker inference used sliding-window prediction without test-time augmentation; all checkpoints, thresholds, fusion weights, gates, and auxiliary recovery settings

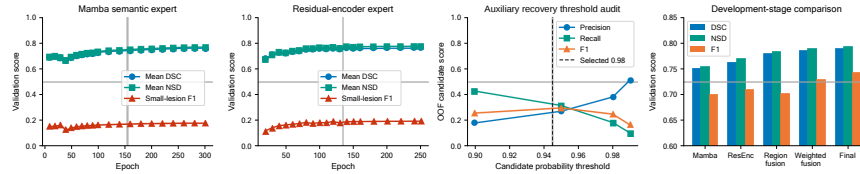


Fig. 2. Development diagnostics for model selection, not independent challenge-performance estimates.

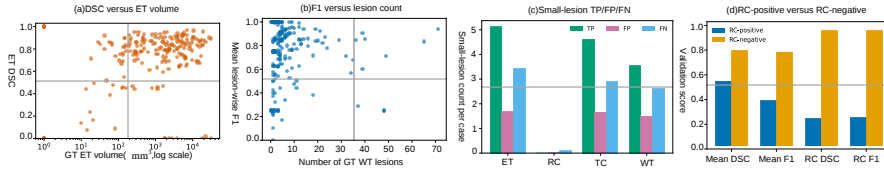


Fig. 3. Development-set failure analysis by lesion size, count, small-lesion status, and RC presence.

were fixed. All results come from the 202-scan, 122-patient internal validation split and are diagnostic internal estimates rather than Synapse validation-server results.

5 Results

Evaluator-aligned metrics. Table 1 reports evaluator-aligned local metrics for the submitted frozen Docker pipeline on the 202-scan, 122-patient internal validation split. These values are not Synapse validation-server results.

Table 1 reports both all-instance lesion-wise F1 and small-lesion F1 because they capture different detection endpoints. Patient-level cluster bootstrap resampled the 122 patient identifiers, thereby preserving longitudinal scans from the same patient within each resampled unit. RC intervals were widest, consistent with only 17 RC-positive validation scans. The small-lesion endpoint remained the most challenging detection criterion; the RC small-lesion F1 interval was not estimable under patient-level bootstrap because only one validation scan had a defined RC small-lesion endpoint, and the mean small-lesion F1 interval was therefore not reported.

Development variants. Table 2 summarizes representative configurations from more than ten internal variants using only evaluator-aligned endpoints. Internally implemented DSC, NSD, and connected-component lesion F1 were used during development to select checkpoints and thresholds, but they are not tabulated here to avoid mixing metric definitions with BraTS-evaluation endpoints. The F1 column in Table 2 reports small-lesion F1 computed with BraTS-evaluation on the same development split. Region-weighted fusion improved

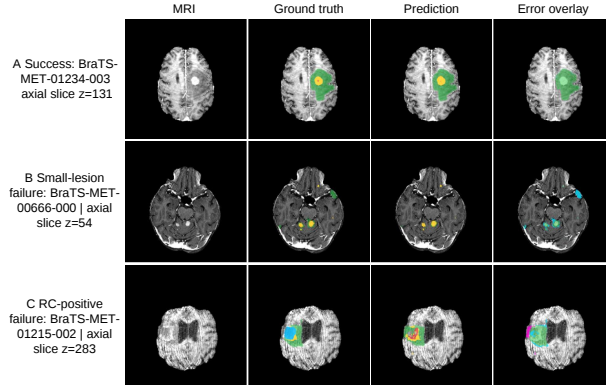


Fig. 4. Representative development-set examples with MRI, ground truth, prediction, and error overlay.

Table 1. Evaluator-aligned local metrics with patient-level cluster-bootstrap 95% CIs. Metrics were computed locally with BraTS-evaluation v0.0.8 on the 202-scan internal validation split; these are not Synapse validation-server results. N.E. denotes not estimable under patient-level bootstrap.

Region	Large DSC	Large NSD	All-inst. F1	Small F1
ET	0.5754 [0.5272, 0.6212] 0.1714	0.6488 [0.5941, 0.7010] 0.1027	0.7048 [0.6516, 0.7552] 0.2193	0.3129 [0.2400, 0.3833] 0.0000
RC	[0.0189, 0.3628] 0.5793	[0.0065, 0.2165] 0.6368	[0.0312, 0.4444] 0.7038	N.E. 0.3382
TC	[0.5296, 0.6278] 0.5554	[0.5819, 0.6908] 0.5750	[0.6496, 0.7558] 0.7087	[0.2622, 0.4109] 0.2997
WT	[0.5099, 0.6008]	[0.5257, 0.6241]	[0.6564, 0.7587]	[0.2238, 0.3747]
	0.4704	0.4908	0.5841	0.2377
Mean	[0.4198, 0.5311]	[0.4448, 0.5406]	[0.5233, 0.6533]	N.E.

mean large-lesion DSC, large-lesion NSD, and small-lesion F1 over uniform averaging by $+0.0153/ +0.0140/ +0.0075$, indicating selective complementary fusion rather than simple ensembling. These scores document model selection, not independent ablations.

Failure and staged analysis. Table 1 identifies RC and small disconnected lesions as the main failure modes, which are visualized in Figs. 3 and 4. Table 3 reports evaluator-aligned selected operating points rather than a strict factorial ablation. For the “+ component verifier” row, the verifier policy and threshold were selected from patient-wise five-fold out-of-fold (OOF) predictions on the 202-scan development split; each OOF prediction was produced by a verifier whose training fold excluded the corresponding held-out patient’s scans. After this OOF-based policy selection, the row was evaluated with BraTS-evaluation on the same 202-scan development split using the fixed verifier operating point. It should therefore be interpreted as an evaluator-aligned development-set result from an OOF-selected verifier policy, not as an independent held-out estimate or a Synapse validation-server result.

Table 2. Representative internal development variants evaluated with BraTS-evaluation on the same development split. Internally implemented metrics were used for model selection but are not tabulated as performance evidence.

Variant	Large DSC	Large NSD	Small F1
Mamba expert	0.4453	0.4591	0.2015
ResEnc expert	0.4498	0.4603	0.2133
Uniform avg.	0.4532	0.4582	0.2182
Region-weighted fusion	0.4685	0.4722	0.2257
Fusion + RC gate	0.4702	0.4700	0.2298
Complete pipeline	0.4704	0.4908	0.2377

Table 3. Evaluator-aligned selected operating points on the 202-scan development split. Rows are representative development configurations, not a strict factorial ablation. The “+ component verifier” row uses an OOF-selected verifier policy and is evaluated on the same development split with BraTS-evaluation. The final row is a jointly refrozen Docker operating point; its difference from the preceding row should not be read as the isolated marginal effect of lesion recovery.

Operating point	Changed settings	Large DSC	Large NSD	Small F1
Fusion only	Region probability fusion without post-fusion gates	0.4304	0.4408	0.2071
+ RC gate	Adds RC case gate and RC competition margin	0.4458	0.4631	0.2180
+ component verifier	Adds WT component verifier for false-positive control	0.4496	0.4684	0.2189
Jointly refrozen METS-HybridFusion operating point	Jointly selected final weights, thresholds, RC margin, verifier threshold, auxiliary recovery policy, and Docker realization	0.4704	0.4908	0.2377

Adding the RC gate changed mean large-lesion DSC, large-lesion NSD, and small-lesion F1 by +0.0154/ + 0.0223/ + 0.0109; the component verifier added +0.0038/ + 0.0053/ + 0.0009. The final row is labeled as a jointly refrozen operating point because it changes several deployment settings at once: final region-fusion weights, region thresholds, RC competition margin, RC gate, component-verifier threshold, conservative auxiliary recovery policy, and Docker realization. It should therefore not be interpreted as the isolated marginal effect of auxiliary lesion recovery. The auxiliary lesion-recovery OOF audit alone showed only small changes (−0.000005/ − 0.000002/ + 0.000351 in mean DSC/NSD/F1 and +0.001460 in small-lesion F1), so the final-row difference reflects the complete selected operating configuration rather than a single auxiliary-recovery gain.

6 Discussion

Evaluator-aligned metrics are stricter than locally computed development metrics but remain internal validation results; patient-level bootstrap quantifies uncertainty without making the 202-scan, 122-patient split independent. The pipeline shows moderate ET, TC, and WT performance, whereas wide RC intervals and low RC scores indicate instability from rare RC-positive cases. Small disconnected lesions strongly affect small-lesion F1; the auxiliary lesion-recovery step recovered a few true candidates but did not yield a meaningful

global DSC/NSD/F1 gain. The auxiliary recovery stage remained computationally bounded in the development audit, producing 5,143 candidates over 202 validation scans (25.5 candidates per scan on average). Because no matched recovery-disabled Docker timing was recorded, we do not report a separate inference-overhead estimate; the recorded labeled validation/evaluation run, which included metric computation rather than pure Docker inference only, required 1:29:38 for 202 scans on an RTX 4090D (26.6 s/scan).

The same split was used for checkpoint selection, threshold tuning, fusion-weight selection, and failure visualization. We therefore avoid presenting these numbers as independent Synapse validation-server results and interpret Tables 2 and 3 as development diagnostics. The final METS-HybridFusion row in Table 3 is the jointly frozen Docker operating point, including selected weights, thresholds, verification rules, auxiliary recovery settings, and Docker realization, not a separate rescue-only module or an estimate of lesion-recovery marginal benefit.

7 Conclusion

We presented METS-HybridFusion, a frozen hybrid Docker pipeline for BraTS-METS 2026 Task 1. On the 202-scan internal validation split, evaluator-aligned local metrics were 0.4704 large-lesion DSC, 0.4908 large-lesion NSD, 0.5841 all-instance lesion-wise F1, and 0.2377 small-lesion F1. These internal results identify small-lesion detection and RC-positive segmentation as the main bottlenecks and are reported cautiously because the same split supported development and evaluation.

CRedit. Yutong Bian: Conceptualization, Methodology, Software, Validation, Formal analysis, Data curation, Writing—original draft, Visualization. Jindong Sun: Conceptualization, Methodology, Supervision, Funding acquisition, Writing—review and editing.

Data availability. BraTS data are available under the challenge policy. Code, configuration files, and the Docker build context needed to inspect and rebuild the pipeline are publicly available at <https://github.com/takosainsbury59-hash/METS-HybridFusion-BraTS2026>. The frozen trained parameters used for challenge evaluation were packaged inside the submitted Synapse Docker image; standalone weights, challenge data, and intermediate caches are not redistributed in the public repository because of challenge-data and storage constraints.

Acknowledgments. This work was supported by the Shandong Provincial Natural Science Foundation (ZR2025QC700) and the Shandong Key Laboratory of Brain-Computer Interface and Intelligent Diagnosis, Treatment, and Rehabilitation System.

Disclosure of Interests. The authors declare no competing interests.

AI-assisted tools. AI-assisted tools were used for language editing only; all results and scientific decisions were verified by the authors.

References

1. Çiçek, Ö., Abdulkadir, A., Lienkamp, S.S., Brox, T., Ronneberger, O.: 3d U-Net: Learning dense volumetric segmentation from sparse annotation. In: Medical Image Computing and Computer-Assisted Intervention – MICCAI 2016. Lecture Notes in Computer Science, vol. 9901, pp. 424–432. Springer (2016). https://doi.org/10.1007/978-3-319-46723-8_49
2. Dice, L.R.: Measures of the amount of ecologic association between species. *Ecology* **26**(3), 297–302 (1945). <https://doi.org/10.2307/1932409>
3. Grøvik, E., Yi, D., Iv, M., Tong, E., Rubin, D., Zaharchuk, G.: Deep learning enables automatic detection and segmentation of brain metastases on multisequence MRI. *Journal of Magnetic Resonance Imaging* **51**(1), 175–182 (2020). <https://doi.org/10.1002/jmri.26766>
4. Gu, A., Dao, T.: Mamba: Linear-time sequence modeling with selective state spaces. In: International Conference on Learning Representations (2024), <https://openreview.net/forum?id=tEYskw1VY2>
5. Hatamizadeh, A., Nath, V., Tang, Y., Yang, D., Roth, H.R., Xu, D.: Swin UNETR: Swin transformers for semantic segmentation of brain tumors in MRI images. arXiv preprint arXiv:2201.01266 (2022). <https://doi.org/10.48550/arXiv.2201.01266>
6. Isensee, F., Jaeger, P.F., Kohl, S.A.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). <https://doi.org/10.1038/s41592-020-01008-z>
7. Isensee, F., Wald, T., Ulrich, C., Baumgartner, M., Roy, S., Maier-Hein, K.H., Jaeger, P.F.: nnU-Net revisited: A call for rigorous validation in 3d medical image segmentation. In: Medical Image Computing and Computer Assisted Intervention – MICCAI 2024. Lecture Notes in Computer Science, vol. 15010, pp. 488–498. Springer (2024). https://doi.org/10.1007/978-3-031-72114-4_47
8. Kamnitsas, K., Ledig, C., Newcombe, V.F.J., Simpson, J.P., Kane, A.D., Menon, D.K., Rueckert, D., Glocker, B.: Efficient multi-scale 3d CNN with fully connected CRF for accurate brain lesion segmentation. *Medical Image Analysis* **36**, 61–78 (2017). <https://doi.org/10.1016/j.media.2016.10.004>
9. Li, R., Guo, Y., Zhao, Z., Chen, M., Liu, X., Gong, G., Wang, L.: MRI-based two-stage deep learning model for automatic detection and segmentation of brain metastases. *European Radiology* **33**(5), 3521–3531 (2023). <https://doi.org/10.1007/s00330-023-09420-7>
10. Ma, J., Li, F., Wang, B.: U-Mamba: Enhancing long-range dependency for biomedical image segmentation. arXiv preprint arXiv:2401.04722 (2024). <https://doi.org/10.48550/arXiv.2401.04722>
11. Maier-Hein, L., Eisenmann, M., Reinke, A., Onogur, S., Stankovic, M., Scholz, P., Arbel, T., Bogunovic, H., Bradley, A.P., Carass, A., et al.: Why rankings of biomedical image analysis competitions should be interpreted with care. *Nature Communications* **9**(1), 5217 (2018). <https://doi.org/10.1038/s41467-018-07619-7>
12. Maleki, N., Amiruddin, R., Moawad, A.W., Yordanov, N., Gkampenis, A., Fehrer, P., Umeh, F., et al.: Analysis of the MICCAI brain tumor segmentation–metastases (BraTS-METS) 2025 lighthouse challenge: Brain metastasis segmenta-

- tion on pre- and post-treatment MRI. arXiv preprint arXiv:2504.12527 (2025), <https://arxiv.org/abs/2504.12527>
13. Menze, B.H., Jakab, A., Bauer, S., Kalpathy-Cramer, J., Farahani, K., Kirby, J., Burren, Y., Porz, N., Slotboom, J., Wiest, R., et al.: The multimodal brain tumor image segmentation benchmark (BRATS). *IEEE Transactions on Medical Imaging* **34**(10), 1993–2024 (2015). <https://doi.org/10.1109/TMI.2014.2377694>
 14. Moawad, A.W., et al.: The brain tumor segmentation (BraTS-METS) challenge 2023: Brain metastasis segmentation on pre-treatment MRI. arXiv preprint arXiv:2306.00838 (2023), <https://arxiv.org/abs/2306.00838>
 15. Myronenko, A.: 3d MRI brain tumor segmentation using autoencoder regularization. In: *Brainlesion: Glioma, Multiple Sclerosis, Stroke and Traumatic Brain Injuries. Lecture Notes in Computer Science*, vol. 11384, pp. 311–320. Springer (2019). https://doi.org/10.1007/978-3-030-11726-9_28
 16. Nikolov, S., Blackwell, S., Zverovitch, A., Mendes, R., Livne, M., De Fauw, J., Patel, Y., et al.: Clinically applicable segmentation of head and neck anatomy for radiotherapy: Deep learning algorithm development and validation study. *Journal of Medical Internet Research* **23**(7), e26151 (2021). <https://doi.org/10.2196/26151>
 17. Ronneberger, O., Fischer, P., Brox, T.: U-Net: Convolutional networks for biomedical image segmentation. In: *Medical Image Computing and Computer-Assisted Intervention – MICCAI 2015. Lecture Notes in Computer Science*, vol. 9351, pp. 234–241. Springer (2015). https://doi.org/10.1007/978-3-319-24574-4_28
 18. Roy, S., Koehler, G., Ulrich, C., Baumgartner, M., Petersen, J., Isensee, F., Jaeger, P.F., Maier-Hein, K.H.: MedNeXt: Transformer-driven scaling of ConvNets for medical image segmentation. In: *Medical Image Computing and Computer-Assisted Intervention – MICCAI 2023. Lecture Notes in Computer Science*, vol. 14224, pp. 405–415. Springer (2023). https://doi.org/10.1007/978-3-031-43901-8_39
 19. Wang, W., Chen, C., Ding, M., Li, J., Yu, H., Zha, S.: TransBTS: Multimodal brain tumor segmentation using transformer. In: *Medical Image Computing and Computer Assisted Intervention – MICCAI 2021. Lecture Notes in Computer Science*, vol. 12901, pp. 109–119. Springer (2021). https://doi.org/10.1007/978-3-030-87193-2_11
 20. Xing, Z., Ye, T., Yang, Y., Liu, G., Zhu, L.: SegMamba: Long-range sequential modeling mamba for 3d medical image segmentation. In: *Medical Image Computing and Computer Assisted Intervention – MICCAI 2024. Lecture Notes in Computer Science*, vol. 15008, pp. 578–588. Springer (2024). https://doi.org/10.1007/978-3-031-72111-3_54