



ARTICLE

# LiteDKT-Net: A Lightweight Diverse Kernel Transformer Network for Brain Tumor Segmentation

Ronak Patel<sup>1</sup>, Miral Patel<sup>2</sup>, Deep Kothadiya<sup>3</sup>, Bayan AlGhofaily<sup>4</sup>, Faten S. Alamri<sup>5,\*</sup>, Awad Alyousef<sup>4</sup> and Amjad R Khan<sup>4</sup>

<sup>1</sup>U & P U Patel Department of Computer Engineering, Chandubhai S. Patel Institute of Technology (CSPIT), Faculty of Technology (FTE), Charotar University of Science and Technology (CHARUSAT), Changa, Anand, India

<sup>2</sup>G H Patel College of Engineering and Technology, CVM University, V V Nagar, Anand, India

<sup>3</sup>Symbiosis Center for Information Technology, Symbiosis International (Deemed University), Pune, India

<sup>4</sup>AIDA Lab. CCIS, Prince Sultan University, Riyadh, Saudi Arabia

<sup>5</sup>Department of Mathematical Sciences, College of Science, Princess Nourah bint Abdulrahman University, Riyadh, Saudi Arabia

\*Corresponding Author: Faten S. Alamri. Email: fsalamri@pnu.edu.sa

Received: 16 May 2026; Accepted: 13 July 2026; Published: 15 September 2026

**ABSTRACT:** Growth of cancerous cells is unpredictable, and their effects vary across organs and levels of aggression. Identification of the pattern, size, and shape of the growth helps assess severity for better treatment. The proposed LiteDKT-Net combines the DK-IRB (Diverse Kernel Inverted Residual Block) block and Transformer to target conceptual information about shape and location. For better edge detection, LiteDKT-Net uses GAG (Group Attention Gate) followed by CBAM (Convolutional Block Attention Module). LiteDKT-Net is a lightweight encoder-decoder-based network optimized for accurate brain tumor segmentation. The network parameter optimization and reduced computational complexity in LiteDKT-Net enable high segmentation accuracy while maintaining a lightweight model size for deployment in clinical environments with limited resources. The proposed architecture achieves Dice score similarities of 0.80, 0.82, and 0.87 for ET (Enhancing Tumor), TC (Tumor Core), and Whole Tumor (WT), respectively. To check the difference between the predicted and ground truth with HD95, the results are 5.24, 8.12, and 8.01 mm for ET, TC, and WT. The proposed architecture evaluates 0.869M parameters and uses 0.856 Giga Floating-Point Operations Per Second (GFLOPs) of computation. Ablation analysis is carried out based on the channel-wise and kernel-wise parameters, which helps to understand model efficiency and lightweight computational cost.

**KEYWORDS:** UNET; transformer; group attention gate; CBAM; 3D MRI segmentation; FLOPs

## 1 Introduction

Brain tumors are abnormal cell growths that impair critical brain functions and may be benign or malignant. Their complex structures, indistinct boundaries, and diverse textures make early diagnosis challenging, emphasizing the need for accurate and reliable automated detection and segmentation methods [1]. It is extremely difficult to diagnose an early abnormality like a tumor because of its intricate structure, ambiguous cell boundaries, and a wide variety of textures.

Deep learning has become essential in medical image analysis, supporting feature extraction, diagnosis, and classification [2], with Convolutional Neural Networks (CNNs) forming the foundation of early semantic segmentation [3]. Later in 2015, Ronneberger et al. designed the U-Net architecture that is based on

an encoder-decoder structure for biomedical image segmentation [4]. U-Net is regarded as the original architecture that is the basis of other similar architectures on 2D/3D, including U-Net++ [5], V-Net [6], DeepMedic [7], and nnUNET [8]. While effective, these approaches often come with high computational overhead due to large image sizes [9]. Real-world applications demand computationally efficient solutions that use fewer training parameters and floating-point operations (FLOPs) while achieving accurate segmentation. CNN [10]-based approaches are more focused on local feature extraction to identify strong relationships between spatial features, whereas transformer-based models [11] are more useful for long-range dependencies and global contextual information. The concept of multi-head attention helps identify and target relevant spatial patterns in the scans.

Based on the observation, the proposed approach deals with accurate segmentation through lower computation cost, and the key contributions are:

**Key contribution summarized as follows:**

- Proposed a multi-kernel depth-wise block called DK-IRB for efficient extraction of spatial information from the MRI scans. This block enhances tumor boundary granularity while reducing computational costs.
- Long-range dependencies are captured by an embedded transformer block in the bottleneck of the architecture in order to extract the global features in tumor size, shape, and location.
- Evaluate the LiteDKT-Net architecture on multi-class segmentation tasks, including ET, TC, and WT. On the expansion path, GAG and CBAM attention modules help to highlight different tumor regions.

The remaining article is organized as follows: A complete literature review of the various segmentation techniques and analyses in [Section 2](#). A detailed explanation of the proposed lightweight architecture and all of its components will be provided in the methodology portion of [Section 3](#). The results of the comparison between the proposed segmenting method and the currently recognized “state-of-the-art” segmenting methods will be given in [Sections 4 and 5](#). The ablation study is reported in [Section 6](#). [Section 7](#) represents the discussion section of the overall architecture impact on tumor segmentation.

## 2 Literature Analysis

Brain tumor segmentation using a deep learning approach is classified in the literature survey into three major bifurcations: U-Net and Its Derivative Architectures, existing lightweight approaches, and attention-based lightweight architectures.

### 2.1 U-Net and its Derivative Architectures

Ronneberger et al. introduce an encoder-decoder architecture, U-Net, with skip connections for medical image segmentation [4]. Simple convolution layers face the challenge of identifying fine-grained features. In brain tumor segmentation on the BraTS MRI dataset, Zhang et al. recently proposed CU-Net [12], which represents a significant advancement. CU-Net focuses more on deeper feature extraction and also changes its upsampling strategies. On the BraTS2019 dataset, CU-Net achieved a Dice score of 0.82, which is more than Swin-Unet (0.81) [13] and TransUnet, which also get the same Dice score of 0.82 [14]. Walsh et al. [15] proposed a lightweight UNet that processes MRI scans in three planes (Axial, Sagittal, and Coronal).

### 2.2 Existing Lightweight Approach

The limitation of major U-Net-based architectures is that an MRI scan for segmentation; therefore, they require more computational power. Zhang et al. proposed EHFF (Enhanced by Hierarchical Feature Fusion) [16], which used around 2M parameters and 6 GFLOPs for processing the BraTS2021 dataset. Zhong

et al. proposed the Wavelet-Guided Iterative Axial Factorization (WIAF) [17] approach to maximize accuracy and minimize computational cost. WIAF basically focuses on wavelet decomposition and iterative axial attention. Overall, WIAF used 5.23M parameters for the process and 9.75 GFlops as computational cost to achieve 0.85 average Dice score. Luo et al. proposed MPEDA-Net, which extracted features from multiple perspective directions, capturing both global context and fine-grained local details. MPEDA-Net [18] uses dense attention to focus on the most relevant features using CBAM attention. Evaluation was done based on 0.91M parameters and 4.62 GFlops for achieving results of 0.82, 0.93, and 0.87 for ET, WT, and TC on the BraTS2021 dataset. Shahid et al. proposed inception-based blocks for the classical U-Net. LIU-Net [19] employs parallel convolutional pathways for global and local feature extraction. LIU-Net uses 1.32 M parameters and 5.41 GFlops to achieve a mean DSC of 0.85 on the BraTS2021 Dataset.

### 2.3 Attention-Based Lightweight Architectures

Lyu and Tian proposed a Generative Adversarial Network (GAN)-based approach with attention gates for improving segmentation accuracy. MWG-UNet++ [20] architecture holds a transformer with UNet followed by attention, making the model heavyweight; the study achieves around 0.89 accuracy on multi-model brain images. Cao et al. proposed a spatial attention-based UNet for edge feature extraction. But due to complex computation, BSAU-Net [21] is heavyweight. Evaluations of BSAU-Net are carried out on the BraTS 2018 and BraTS 2021 datasets, and the Dice score achieved 0.75 in both cases. Afridi et al. proposed 3D-ViT-UNet [22], which combined 3D and dilated windows with self-attention on volumetric brain tumor MRI scans to achieve a mean Dice Score Coefficient (DSC) of 0.84 and an HD95 of 4.87 mm on BraTS2020. Rahim et al. proposed a transformer and state-space model [23] for 3D segmentation, which focuses on global context modelling with sequential feature extraction. Comparative analysis with a recent study is demonstrated in Table 1.

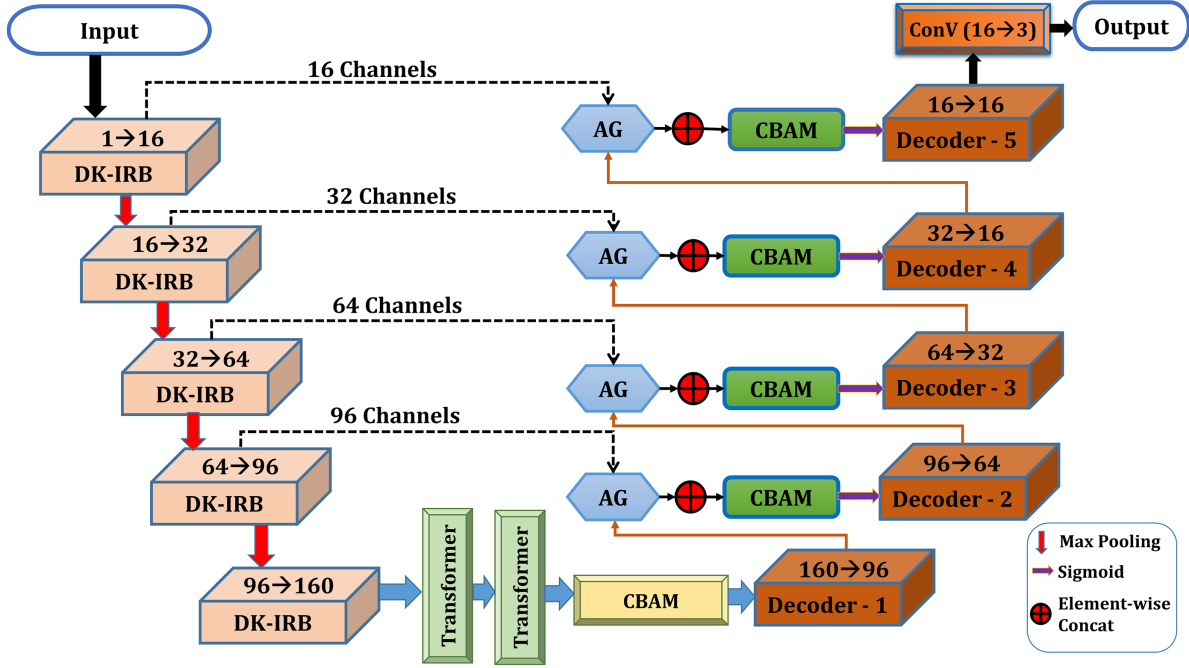
**Table 1:** Encoder-based approach comparison on BraTS2020.

| Model                 | Encoder                       | Attention | DSC [ET/TC/WT] |
|-----------------------|-------------------------------|-----------|----------------|
| Eff U-Net [24]        | VGG-19, ResNet50, MobileNetV2 | Yes       | 0.70/0.80/0.87 |
| TransDoubleU-Net [25] | Dual-scale Swin Transformers  | Yes       | 0.71/0.80/0.85 |
| TransUNet [14]        | Transformer + CNN             | Yes       | 0.78/0.83/0.91 |
| Swin-UNet [13]        | Swin Transformer              | Yes       | 0.78/0.82/0.91 |
| dResU-Net [26]        | ResNet                        | No        | 0.72/0.79/0.84 |
| BT-UNet [27]          | Standard U-Net                | No        | 0.90 Mean DSC  |

From reviewing recent literature analysis represented in Table 1, it is noted that numerous critical aspects need to be addressed. U-Net and similar architectures (CU-NET [12] and TransUNet [14]) produce a very high Dice score; however, these same models struggle when attempting to calculate a high computational cost. Existing lightweight methodologies (EHFF [16], WIAF [17], MPEDA-Net [18], LIU-NET [19]) possess fewer parameters than U-NET models, but can still find it difficult to compensate for issues presented by single-scale convolutional operations and restricted kernel execution, resulting in poor performance in capturing heterogeneous spatial data. Attention mechanisms can improve boundary detection but often increase memory usage and computational cost. The proposed lightweight architecture overcomes these limitations by combining a DK-IRB encoder, lightweight transformer bottleneck, and GAG+CBAM dual-attention decoder to effectively capture multi-scale features, model global dependencies, and achieve precise boundary segmentation with low computational overhead.

### 3 Methodology

The proposed encoder-decoder-based architecture makes the architecture lighter in terms of computational cost. On the encoder side, a Diverse Kernel Inverted Residual Block (DK-IRB) is used. On the decoder side, a Group Attention Gate (GAG) is used, followed by the Convolution Block Attention Module (CBAM) [28]. The contribution of all the components and the effective use of the convert model as lighter and more effective with a lower parameter count is represented in Fig. 1.



**Figure 1:** Architecture of LiteDKT-Net: A five-stage DK-IRB encoder with dual transformers captures multi-scale and global features, while a GAG- and CBAM-assisted decoder reconstructs accurate segmentation masks through skip connections.

#### 3.1 Diverse Kernel Inverted Residual Block (DK-IRB)

The BraTS [29] dataset published as BraTS2020 has heterogeneous features for tumor identification. To address this, the DK-IRB is introduced (illustrated in Fig. 2). To maintain a lightweight architecture and focus on multi-scale conceptual representation, the DK-IRB is well-suited for MRI scans. Enhancement of multi-scale features is performed internally by the DKDC (Diverse Kernel Depth-wise Convolutions) block. The DKDC block is specifically used for capturing volumetric information and sharpening the borders of tumor tissues from MRI scans. At the initial level, pointwise expansion convolution ( $1 \times 1 \times 1$ ) is used to target input features through an expansion factor, as shown in Eq. (1).

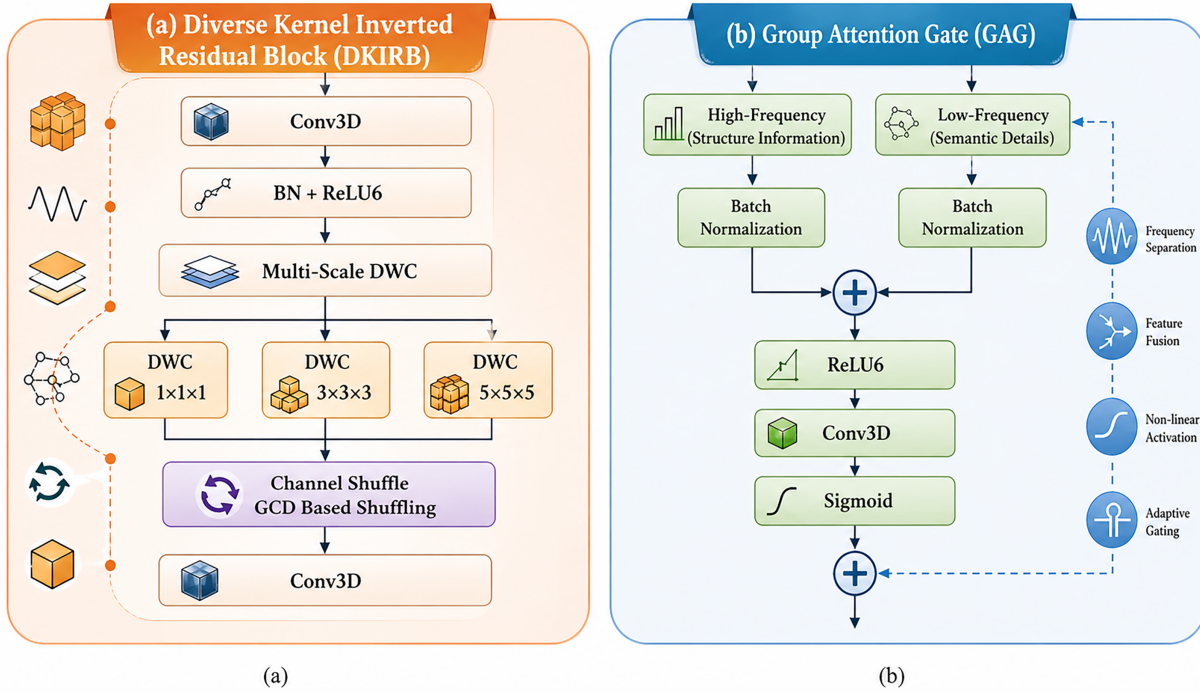
$$Z_{exp} = \delta \left( BN \left( WE_{pconv} * Z_{in} \right) \right) \quad (1)$$

where  $Z_{in}$  is the input feature map,  $WE_{pconv}$  represents the learnable weights of the pointwise expansion convolution, and  $BN$  denotes batch normalization.  $\delta$  represents the non-linear activation function ReLU6 [30]. This expansion helps identify higher-dimensional feature maps for tumor regions and enhances computational efficiency. This block implements depth-wise convolution with three kernel sizes (1, 3, 5), appropriate padding values of 0, 1, and 2 are applied to preserve the spatial dimensions of the corresponding feature maps. Depth-wise convolution is defined by the Eq. (2).

$$Z_{dw}^k = \delta(BN(W_{dw}^k * Z_{exp})) \quad (2)$$

where  $Z_{dw}^k$  represents features identified from the kernel,  $W_{dw}^k$  represents the weight of  $k \times k \times k$  kernel.  $Z_{exp}$  is the expanded features after BN and ReLU [30]. Outputs of each channel are concatenated with channel dimension for unique tensors, which is represented by Eq. (3).

$$DKDC = Concat(Z_{dw}^{1 \times 1}, Z_{dw}^{3 \times 3}, Z_{dw}^{5 \times 5}) \quad (3)$$



**Figure 2:** Internal Structure of (a) DK-IRB integrates multi-scale depth-wise convolutions with channel shuffling based on GCD, (b) GAG fuses high/low-frequency features through BN, ReLU6, Conv3D, and Sigmoid to enable adaptive gating.

Greatest Common Divisor (GCD)-based channel shuffling evenly redistributes concatenated features across groups, promoting inter-group feature interaction and efficient information exchange without introducing additional parameters or increasing model complexity. The channel shuffling operation is related to the channel shuffle mechanism of ShuffleNet [31]. This operation provides a mechanism to interleave the feature maps extracted from each kernel shape ( $1 \times 1 \times 1$ ,  $3 \times 3 \times 3$ , and  $5 \times 5 \times 5$ ) together so that subsequent convolutions are able to combine local and contextually complementary features from each of the feature maps. To preserve original encoder features, a residual connection is applied using a  $1 \times 1 \times 1$  convolution. To reduce the channel dimension back to the output size while maintaining feature consistency, a  $3 \times 3 \times 3$  convolution and normalization with ReLU6 [30] are applied in Eq. (4).

$$DKIR = \delta(BN(Conv_{3 \times 3}(DKDC + Conv_{1 \times 1}(Z)))) \quad (4)$$

where  $Z$  represents the input feature map,  $\delta$  represents the ReLU6 [30]. This will help to capture fine-grained tumor boundary, local spatial context and border contextual information for ET, TC and WT.

### 3.2 Transformer Block

The DK-IRB captures fine-grained structural details and short-range spatial information but struggles with long-range relationships across MRI scans. To overcome this, a transformer block is used in the bottleneck for better enhancement of size, shape, and regions in MRI scans.

The DK-IRB encoder outputs volumetric features denoted as  $DKIR_{exc} \in R^{C \times D \times H \times B}$ , where  $C$  represents the number of feature channels and  $D$ ,  $H$ , and  $W$  denote the depth, height, and width of the feature volume, respectively. The spatial dimensions are then flattened into a sequence of  $N = D \times H \times W$  tokens, where each token has dimension of  $C$ . The feature sequence is embedded with positional encoding and represented as,  $T_{pos} \in R^{N \times C}$  to preserve the spatial arrangement of tokens before being processed by the Multi-Head Self-Attention (MHSA) module, which models pairwise correlations among all tokens as defined in Eq. (5).

$$Attention(Q, K, V) = Softmax\left(\frac{QK^T}{\sqrt{d_k}}\right)V \quad (5)$$

where  $Q$ ,  $K$  and  $V$  represent linear projections of input tokens, and  $d_k$  identifies each vector dimension. After attaching the transformer at the bottleneck, it provides a more representative view of local and global tumor features. Combining the use of both DK-IRB and the transformer helps to identify the complex boundaries of the tumor.

### 3.3 Group Attention Gate (GAG)

To enhance contextual information between the encoder and decoder throughout the segmentation process, the Group Attention Gate (GAG) was incorporated. This component focuses on tumor-related information and suppresses irrelevant or background noise, as the main focus is on multi-class segmentation in BraTS2020 [29], such as ET, TC, and WT. To identify diverse characteristics, GAG divides the input feature map into channel-wise groups, allowing finer granularity in feature modulation.

Encoder feature maps  $f \in R^{B \times C \times D \times H \times W}$ , transmitted through the skip connection, and decoder gating features  $gf \in R^{B \times C' \times D \times H \times W}$  are first independently projected using separate  $1 \times 1 \times 1$  convolution layers with learnable weights  $W_f$  and  $W_{gf}$ , respectively, to align their channel representations. The transformed feature maps are then combined through element-wise addition, followed by a ReLU6 activation function. Subsequently, another  $1 \times 1 \times 1$  convolution with learnable weights  $W_l$  is applied, and a sigmoid activation generates the Attention Coefficient Map (ACM), as defined in Eq. (6).

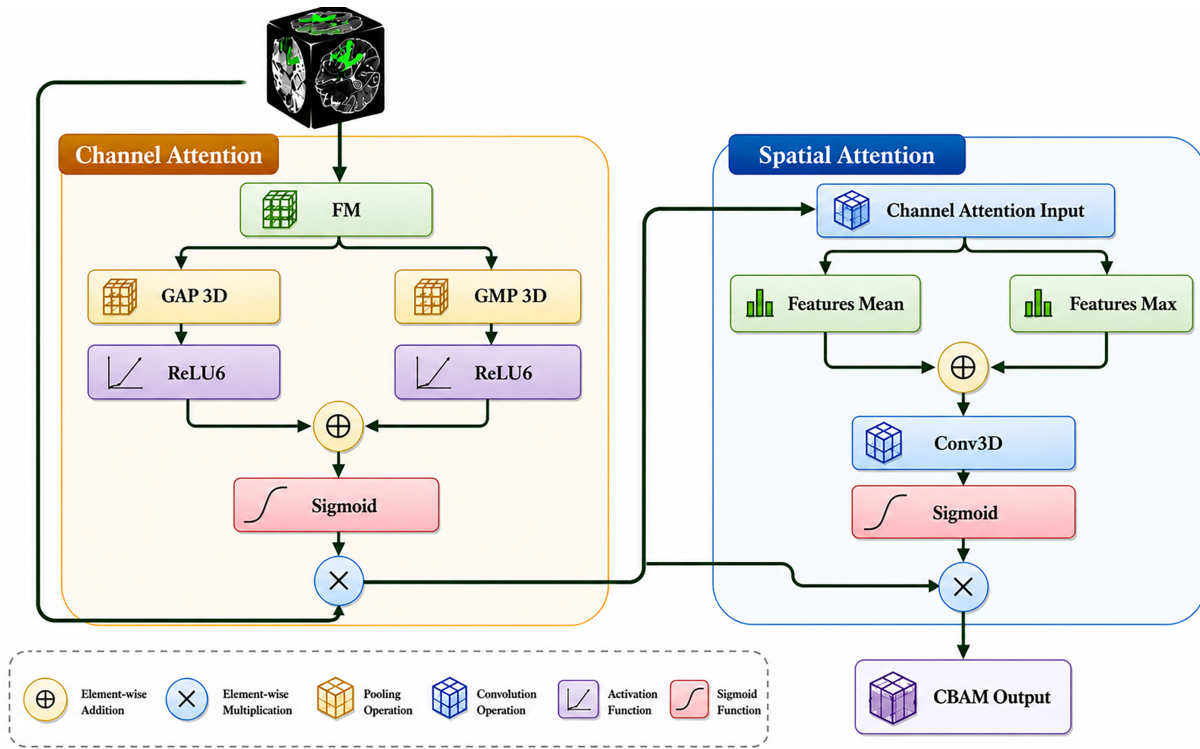
$$ACM = \sigma(W_l \cdot (\delta(W_f f + W_{gf} gf))) \quad (6)$$

where  $f$  represents encoder feature maps and  $gf$  represents the gating features from the decoder.  $W_l$  represent the learnable weights for the attention gate.  $W_f$  and  $W_{gf}$  are learnable weight representations based on the  $f$  and  $gf$ .  $\delta$  and  $\sigma$  represent the ReLU6 and sigmoid. The final output is element-wise multiplied with the encoder features, represented as below Eq. (7).

$$f' = ACM \odot f \quad (7)$$

### 3.4 Convolutional Block Attention Module (CBAM)

A dual-attention mechanism helps enhance channel-wise and spatial-wise features to optimize tumor segmentation after GAG. A difficulty in this dataset is the inconsistency in tumor sizes, shapes, and positions. The attention mechanism helps the network focus on semantic features and remove disturbing background information in Fig. 3.



**Figure 3:** Architecture of 3D CBAM with channel and spatial attention mechanisms.

The first phase is channel attention, which exploits global context information to select informative feature maps by both Global Average Pooling (GAP) and Global Max Pooling (GMP). These preserve essential information and structural boundaries from each slice. The result of this pooling pass is fed through a two-layer Multilayer Perceptron (MLP) with a reduction ratio  $\gamma = 16$ , ReLU, and then sigmoid activation to produce normalized attention weights. Channel attention Eq. (8).

$$CA(F) = \sigma(W_2 \cdot \delta(W_1 \cdot GAP(F)) + W_2 \cdot \delta(W_1 \cdot GMP(F))) \odot F \quad (8)$$

where  $\delta(\cdot)$  represent ReLU6,  $\sigma(\cdot)$  represent the sigmoid function, and  $\odot$  denotes channel-wise multiplication. This attention focuses on high-response channels associated with enhancing tumor boundaries and infiltrated regions.

Spatial attention follows, focusing on voxel-level dependencies crucial for representing irregular contours and small-scale features like the tumor core. Computation of channel-wise average and max projections, concatenate them, and pass through Conv3D followed by a sigmoid defined in Eq. (9).

$$SA(F') = \sigma(Conv3D \oplus [AvgPool_c(F'), MaxPool_c(F')]) \odot F' \quad (9)$$

where  $F' = CA(F)$  and  $AvgPool_c(\cdot)$  and  $MaxPool_c(\cdot)$  denotes average and maximum pooling operations performed along the channel dimension, respectively, while preserving the spatial dimension. This operation helps to identify spatial zones with higher tumor tissues. The combination of these two attention branches balances high-intensity feature localization with consistent semantic feature propagation, enabling the decoder to reconstruct complex spatial patterns.

### 3.5 Significance of the LiteDKT-Net

The proposed LiteDKT-Net employs a five-stage encoder–decoder architecture using Diverse Kernel Inverted Residual Blocks (DK-IRB) to efficiently capture contextual features through multi-kernel depth-wise convolutions, achieving high segmentation accuracy with minimal parameters and computational cost on the BraTS2020 dataset.

During decoding, the GAG module leverages skip connections to selectively refine encoder features using decoder guidance. The refined features are subsequently processed by the Convolutional Block Attention Module (CBAM), which sequentially applies channel and spatial attention to enhance feature representation and suppress irrelevant responses. The attended features are then fused within the decoder to progressively reconstruct high-resolution feature maps. Finally, a sigmoid-based segmentation layer generates multi-class outputs, while Dice loss optimizes segmentation performance during supervised training.

## 4 Experimental Configuration

### 4.1 Training Setup

The LiteDKT-Net architecture is designed and experimented with PyTorch 2.4.0 and MONAI 0.6 on 2 NVIDIA Tesla T4 GPUs with 16 GB of RAM, which effectively utilizes the resources to efficiently preprocess the high-resolution 3D MRI scans.

The learning rate was selected based on the set  $\{0.0001, 0.001, 0.01\}$  and the AdamW [32] optimizer was employed with weight decay values selected from the set  $\{0.00001, 0.0001, 0.001\}$ . The highest average mean Dice found on validation data was with  $lr = 0.01$  and  $weight\ decay = 0.0001$ . Only batch size 1 was possible due to memory limitations on the GPU, which resulted from using full 3D volumes. The entire training process was performed using FP16 mixed precision to store less data than required for training with FP32. The model size corresponds to the FP16 PyTorch checkpoint, whereas the reported parameter count refers only to the trainable network parameters.

In this architecture, a diverse kernel with depth convolution (1, 3, 5) has been incorporated to implement the DK-IRB block. For validating the proposed architecture, standardizing channel configuration [16, 32, 64, 96, 160] was kept constant across all tests. The remaining training hyperparameters are given in Table 2.

**Table 2:** Configuration used to train the proposed approach.

| Parameters          | Values                      |
|---------------------|-----------------------------|
| Framework           | PyTorch 2.4.0/MONAI 0.6     |
| GPU                 | NVIDIA Tesla T4 $\times$ 2  |
| Batch Size          | 1                           |
| Epochs              | 50                          |
| Optimizer           | AdamW                       |
| Learning Rate       | 0.01                        |
| Weight Decay        | 0.0001                      |
| Loss Function       | Dice Loss                   |
| Precision Mode      | FP16                        |
| Model Size          | 5.6 MB (FP16)               |
| Average Epoch Time  | $\approx$ 154.8 s per epoch |
| Total Training Time | $\approx$ 2.09 h            |
| Total Parameter     | 0.869M                      |

Since the experimentation is conducted on a 3D MRI image, to decrease the training time and utilize the memory efficiently, FP16 precision mode was used. Also, there was no provision of early stopping, and each epoch was stored for validation performance calculation.

#### 4.2 Evaluation Metrics

DSC [33] and HD95 [34] are used for the evaluation of the LiteDKT-Net architecture. Dice score is computed from the similarity between the ground segmentation area and the generated segmented output [33]. Dice Score Coefficient (DSC) can be computed from Eq. (10).

$$DICE = \frac{2|GS \cap PS|}{|GS| + |PS|} \quad (10)$$

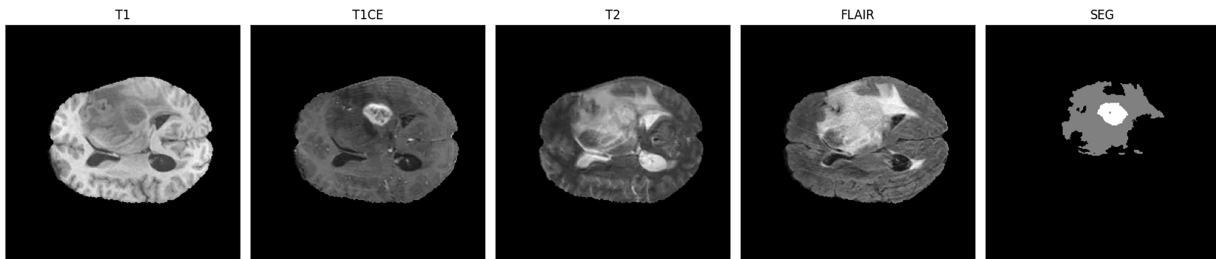
where GS represents the ground segmentation part annotated manually, PS represents predicted segmentation part by the proposed architecture.

The Hausdorff distance measures the maximum boundary error between the ground truth and predicted segmentation [34]. To reduce sensitivity to outliers, the 95th percentile Hausdorff distance (HD95) is measured based on the Eq. (11).

$$HD_{95}(GS, PS) = \max \left\{ percentile_{95} \left( \min_{ps \in PS} \|gs - ps\| \right), percentile_{95} \left( \min_{gs \in GS} \|ps - gs\| \right) \right\} \quad (11)$$

#### 4.3 BraTS Dataset

This work relies on the BraTS [29] dataset published as BraTS2020 released by the MICCAI in 2020 for Brain Tumor Segmentation Challenge. This dataset was selected to facilitate the development of an automatic brain tumor segmentation algorithm using 3D multimodal MRI volumes and provides detailed segmentation of three different tumor regions: Whole Tumor (WT), Tumor Core (TC), Enhancing Tumor (ET). An example is illustrated in Fig. 4 with the 4 modalities for the same location.



**Figure 4:** MRI scans of BraTS 2020 dataset.

There are 280 MRI volumes available for the training split, 50 for the validation split, and 39 for the test split. Each subject scan includes four co-registered 3D MRI modalities: T1: good for describing the shape of the tissues and boundaries. T1CE: sensitive to tumor infiltration. T2: used for segmenting brain tumors; Fluid-Attenuated Inversion Recovery (FLAIR): is more sensitive to tumor-induced edema. Ground Truth (SEG): tumor segments labeled by a neuroradiologist. Ground-truth volumes were hand-segmented by expert neuroradiologists, and a consensus was reached.

The BraTS 2020 dataset provides skull-stripped, co-registered MRI volumes resampled to  $240 \times 240 \times 155$  resolution in NIfTI format. Images are loaded using MONAI, requiring no additional preprocessing or augmentation due to prior standardization.

## 5 Analysis of Outcome

The experimental result of the LiteDKT-Net model is to integrate the channel-spatial attention mechanism, group attention gate, and DK-IRB blocks to enhance segmentation performance at a minimum computational cost. To ensure a comprehensive assessment of segmentation accuracy and boundary accuracy, the DSC and HD95 are employed as primary performance metrics.

### 5.1 Quantitative Comparison with Existing Methods

Table 3 shows performance on BraTS2020, as measured by the DSC and the 95th-percentile HD95, for ET, TC, and WT. The proposed LiteDKT-Net achieves DSC = 0.80 (ET), 0.82 (TC), and 0.87 (WT), competitive with heavier baselines. While delivering substantially better boundary quality of HD95 = 5.24 (ET), 8.12 (TC), and 8.01 (WT). In contrast, certain baselines exhibit ET failures (HD95 = INF), and even strong DSC models (nnUNet) record a high ET HD95 of 26.14, underscoring that overlap alone can mask clinically relevant boundary errors. Overall, LiteDKT-Net offers a favorable accuracy–boundary trade-off with consistent sub-region behavior (ET/TC/WT).

**Table 3:** Comparative analysis on BraTS2020 dataset for ET/TC/WT. LiteDKT-Net attains competitive DSC with markedly lower HD95, avoiding ET failures (INF) observed for some baselines.

| Model            | DSC   |       |       | HD95  |       |       |
|------------------|-------|-------|-------|-------|-------|-------|
|                  | ET    | TC    | WT    | ET    | TC    | WT    |
| nnUNET [35]      | 0.798 | 0.857 | 0.911 | 26.14 | 3.73  | 5.64  |
| dResU-Net [26]   | 0.72  | 0.79  | 0.84  | 37.67 | 12.24 | 10.93 |
| UNETR [36]       | 0.75  | 0.71  | 0.86  | INF   | 11.06 | 11.01 |
| FPN-UNET [37]    | 0.77  | 0.86  | 0.85  | 11.96 | 20.06 | 20.88 |
| SegNet 3D [38]   | 0.67  | 0.76  | 0.85  | INF   | 14.82 | 14.81 |
| U-SegNet [39]    | 0.69  | 0.64  | 0.83  | 9.44  | 10.01 | 10.32 |
| 3D-ViT-UNet [22] | 0.79  | 81.40 | 90.43 | 5.74  | 4.65  | 4.23  |
| SwinMamba [23]   | 0.86  | 0.90  | 0.91  | –     | –     | –     |
| LiteDKT-Net      | 0.80  | 0.82  | 0.87  | 5.24  | 8.12  | 8.01  |

LiteDKT-Net achieves DSC scores of 0.80 (ET), 0.82 (TC), and 0.87 (WT), approximately 4%–5% lower than nnUNet (0.798/0.857/0.911) in WT and TC. However, nnUNet contains 34.33M parameters and requires 1405.78 GFLOPs, more than 1600 times as much computation cost of LiteDKT-Net with 0.856 GFLOPs, as shown in Table 4.

**Table 4:** Comparison of parameters and FLOPs for different models.

| Model          | Params (M) | FLOPs (G) |
|----------------|------------|-----------|
| nnUNET [35]    | 34.33      | 1405.7787 |
| dResU-Net [26] | 15.83      | 56.84     |
| UNETR [36]     | 92.78      | 82.6      |
| FPN-UNET [37]  | 1.4339     | 102.24    |
| SegNet 3D [38] | 31.5       | 178.8     |
| U-SegNet [39]  | 28.5       | 35.8      |
| LiteDKT-Net    | 0.869      | 0.856     |

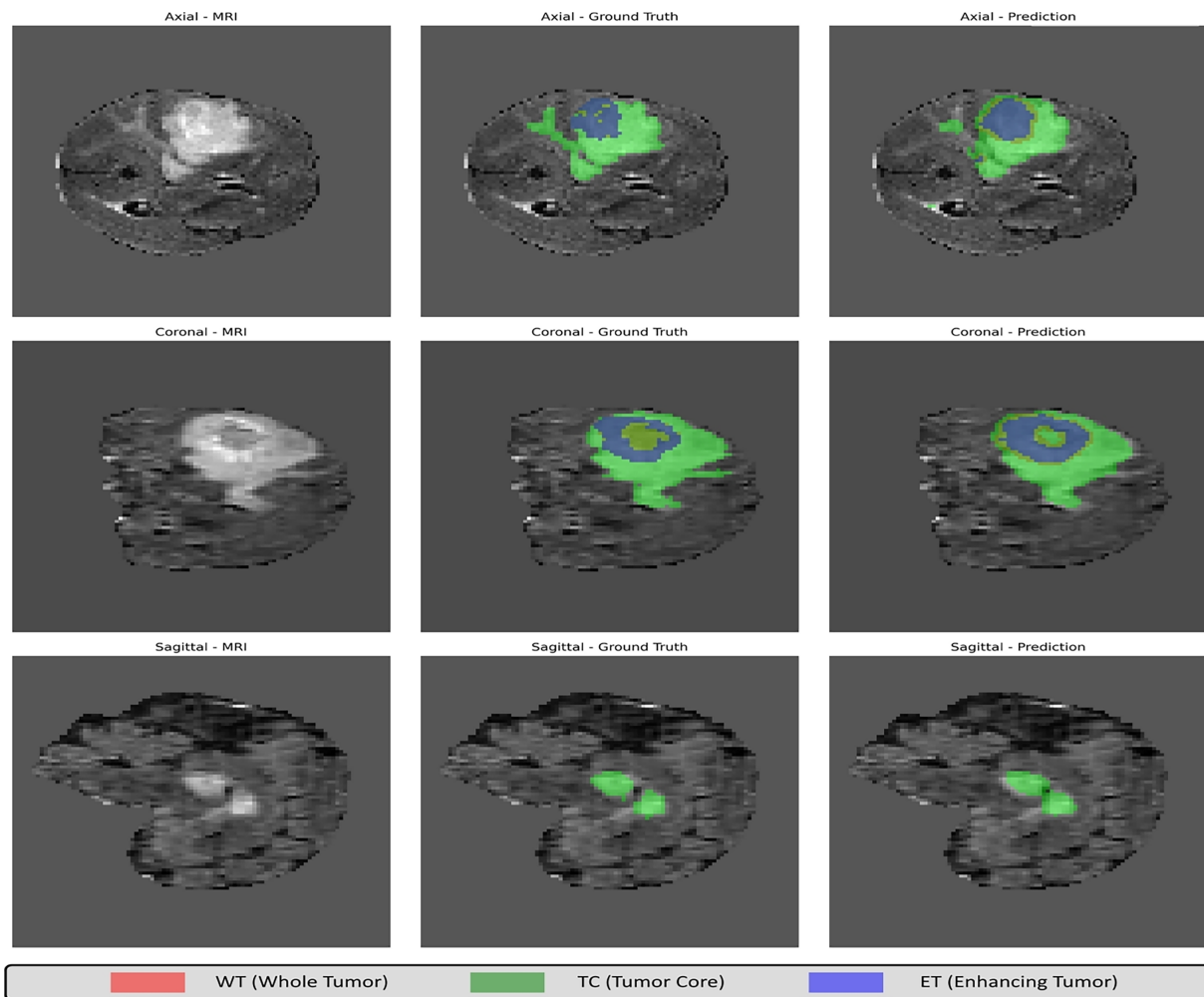
## 5.2 Computational Efficiency Comparison

LiteDKT-Net achieves high deployment efficiency with only 0.869M parameters and 0.856 GFLOPs. Its lightweight design significantly reduces memory usage and inference latency while maintaining competitive segmentation performance compared to larger CNN and transformer-based architectures.

## 5.3 Volumetric & Multi-View Qualitative Analysis

### 5.3.1 Tri-Planar Views

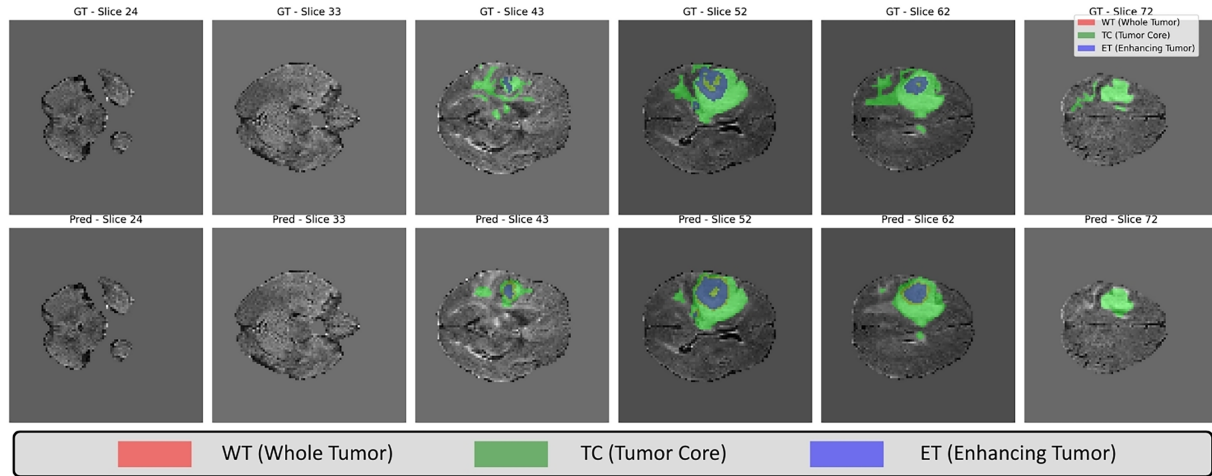
Axial, coronal, and sagittal overlays shown in Fig. 5 juxtapose ground truth (GT) and predictions for WT/TC/ET, enabling cross-sectional assessment of spatial consistency across orthogonal planes. The tri-planar layout reveals how LiteDKT-Net preserves internal structure while adhering closely to GT boundaries, aspects not fully conveyed by single-slice displays.



**Figure 5:** Tri-planar views. Axial, coronal, and sagittal overlays with GT and predictions for WT/TC/ET, illustrating cross-plane spatial consistency, internal structure preservation, and boundary adherence.

### 5.3.2 Slice Progression

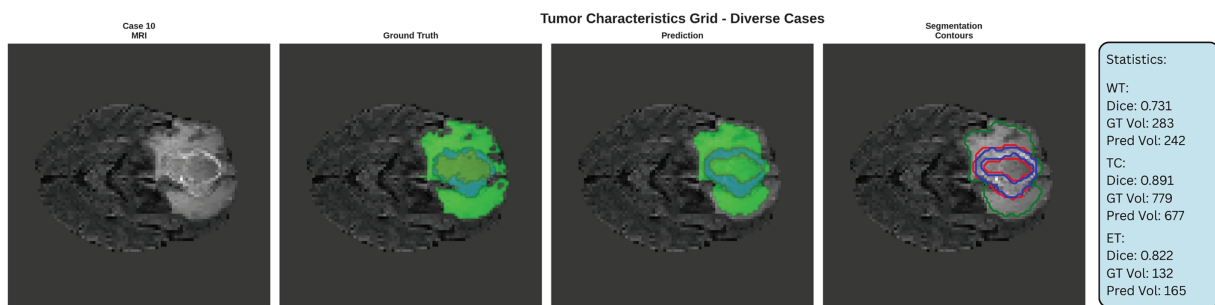
A slice-wise GT vs. Prediction progression across contiguous axial slices tracks the evolution of tumor morphology from WT to TC to ET, as shown in Fig. 6. Paired rows (GT above, prediction below) facilitate rapid inspection of boundary fidelity, shape continuity, and sub-region coherence as the anatomy varies with slice index.



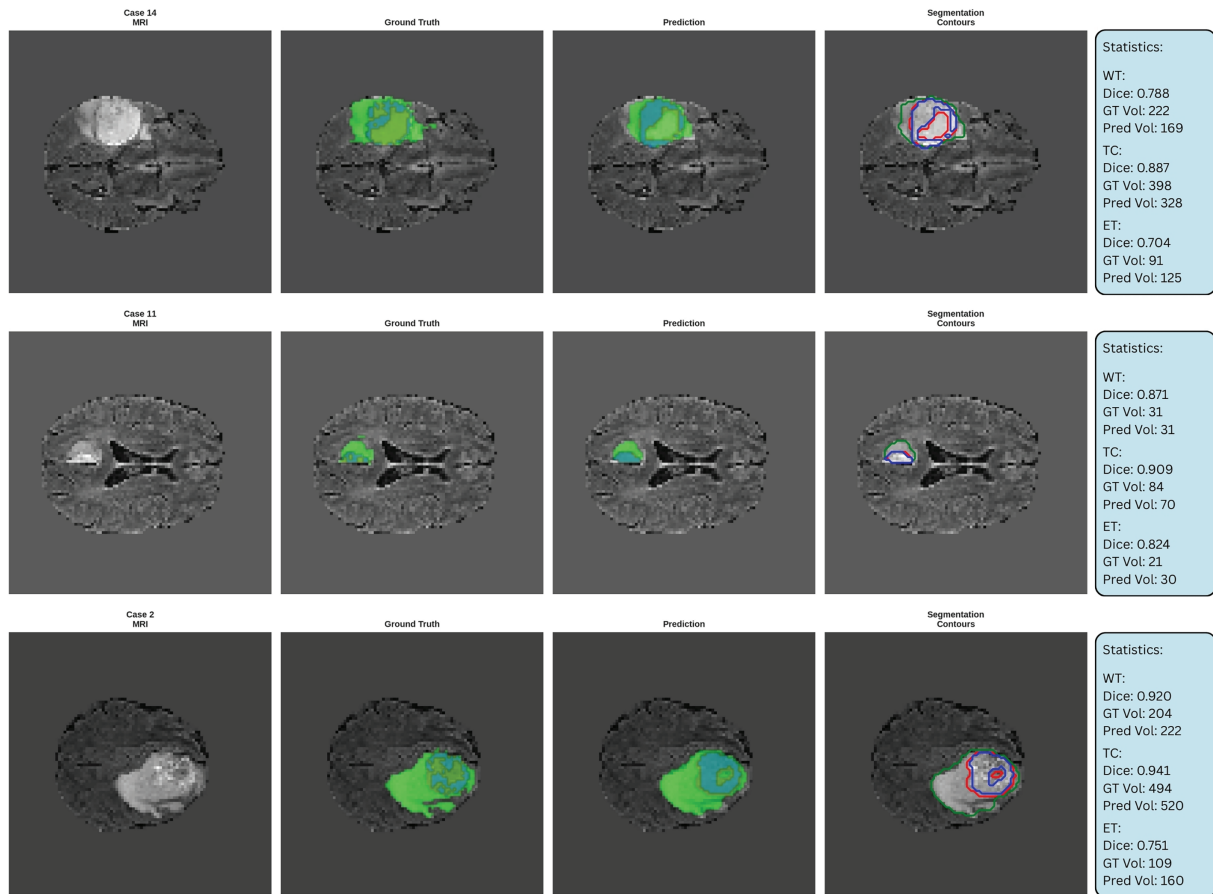
**Figure 6:** Slice progression. Contiguous axial slices with GT (top row) and predictions (bottom row). The sequence demonstrates morphological continuity and boundary fidelity for WT (red), TC (green), and ET (blue).

### 5.4 Error and Boundary Analysis

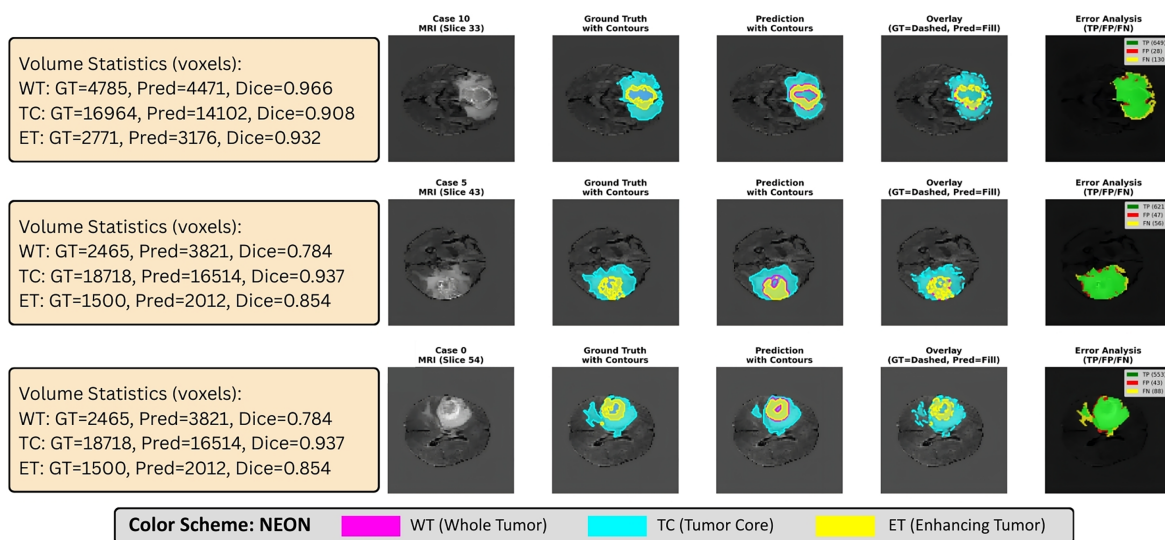
Figs. 7 and 8 present multi-case visualizations comprising MRI slices, ground truth masks, predictions, contour overlays, and TP/FP/FN maps. The results highlight improved boundary localization, reduced over- and under-segmentation, and better contour adherence. The integration of DK-IRB, lightweight transformers, and dual-attention decoding enhances segmentation accuracy with low computational cost.



**Figure 7:** (Continued)



**Figure 7:** Tumor characteristic grid. Multi-case grid with MRI, GT, prediction, and contour overlays; per-case dice and voxel counts reported at right.



**Figure 8:** Error analysis (TP/FP/FN Maps). Voxel-wise TP/FP/FN maps with contour overlays localize over- and under-segmentation near ET margins; paired per-case metrics support the HD95 improvements observed quantitatively.

### 5.5 Evaluation on BraTS2021

LiteDKT-Net was evaluated on the BraTS2021 dataset, which was used exclusively for external evaluation without any retraining, fine-tuning, parameter updating, or model selection being performed to assess cross-dataset generalization. On 1251 (1000 Training, 125 Testing, 126 Validation) multimodal MRI scans, it achieved DSC scores of 0.81, 0.84, and 0.89 and HD95 values of 5.12, 7.98, and 7.68 mm for ET, TC, and WT, respectively. These results demonstrate that features learned from BraTS 2020 generalize effectively to heterogeneous clinical datasets.

## 6 Ablation to Understand Lightweight

To compare proposed lightweight segmentation model, a two-stage ablation study was conducted on BraTS2020, evaluating transformer bottleneck channel scaling and convolutional kernel sizes under strict parameter and computational efficiency constraints.

### 6.1 Effect of Channel Scaling in Transformer Bottleneck

The ablation study presented in this section investigates the effect of transformer bottleneck scaling while employing a fixed single  $3 \times 3$  convolutional configurations for all evaluated models. Among the evaluated configurations, the TB\_Tiny model achieves the highest mean Dice score (0.8718) with only 0.079M parameters, whereas the TB\_Base model attains comparable segmentation performance with 0.869M parameters and 0.856 GFLOPs. Although TB\_Tiny achieves the highest Dice score, the TB\_Base configuration was further experimented with because it provides a more balanced compromise between segmentation accuracy, representational capacity, and computational efficiency. Additionally, TB\_B was evaluated using the most commonly used channel capacities of 16 to 160.

To ensure a fair comparison, the experiments reported in Table 5 investigate only the influence of the transformer bottleneck dimension on network performance. Throughout this study, all other architectural components were kept unchanged. Specifically, every evaluated model employs the same encoder-decoder architecture, identical convolutional layers using a single  $3 \times 3$  convolution in each block, the same number of transformer layers, identical skip connections, training strategy, optimizer, learning rate schedule, batch size, and loss function. The only modified parameter is the transformer channel dimension, which is varied to analyze its effect on segmentation accuracy and computational complexity. Consequently, the observed performance differences in Table 5 can be attributed solely to the change in transformer bottleneck capacity rather than to any other architectural modification.

**Table 5:** Comparison of transformer block (TB) variants with parameters, FLOPs, and dice scores.

| Model | Channel                | Param (M) | FLOPs (G) | ET   | TC   | WT   | Mean DSC | ET   | TC   | WT   | Mean HD95 |
|-------|------------------------|-----------|-----------|------|------|------|----------|------|------|------|-----------|
| TB_T  | 4, 8, 16, 24, 32       | 0.079     | 0.266     | 0.86 | 0.84 | 0.92 | 0.87     | 5.12 | 7.93 | 7.83 | 6.96      |
| TB_S  | 8, 16, 32, 48, 80      | 0.268     | 0.441     | 0.81 | 0.81 | 0.89 | 0.84     | 5.17 | 8.01 | 7.90 | 7.03      |
| TB_B  | 16, 32, 64, 96, 160    | 0.869     | 0.856     | 0.80 | 0.82 | 0.87 | 0.83     | 5.24 | 8.12 | 8.01 | 7.12      |
| TB_M  | 32, 64, 128, 192, 320  | 3.069     | 1.959     | 0.80 | 0.81 | 0.84 | 0.82     | 5.52 | 8.53 | 8.42 | 7.49      |
| TB_L  | 64, 128, 256, 384, 512 | 8.472     | 5.213     | 0.79 | 0.80 | 0.81 | 0.80     | 6.20 | 9.62 | 9.48 | 8.43      |

Note: T = Tiny, S = Small, B = Base, M = Medium, L = Large.

As shown in Table 5, the transformer bottleneck dimension has a direct influence on segmentation performance. Since all remaining architectural components are fixed and unchanged, the observed performance differences are solely attributable to variations in transformer bottleneck capacity. Increasing the transformer bottleneck dimension results in a larger number of trainable parameters and higher

computational complexity without providing a consistent improvement in segmentation accuracy. Therefore, the TB\_Base configuration was adopted for all subsequent experiments, as it provides an effective balance between segmentation accuracy, computational complexity, and model scalability for lightweight medical image segmentation.

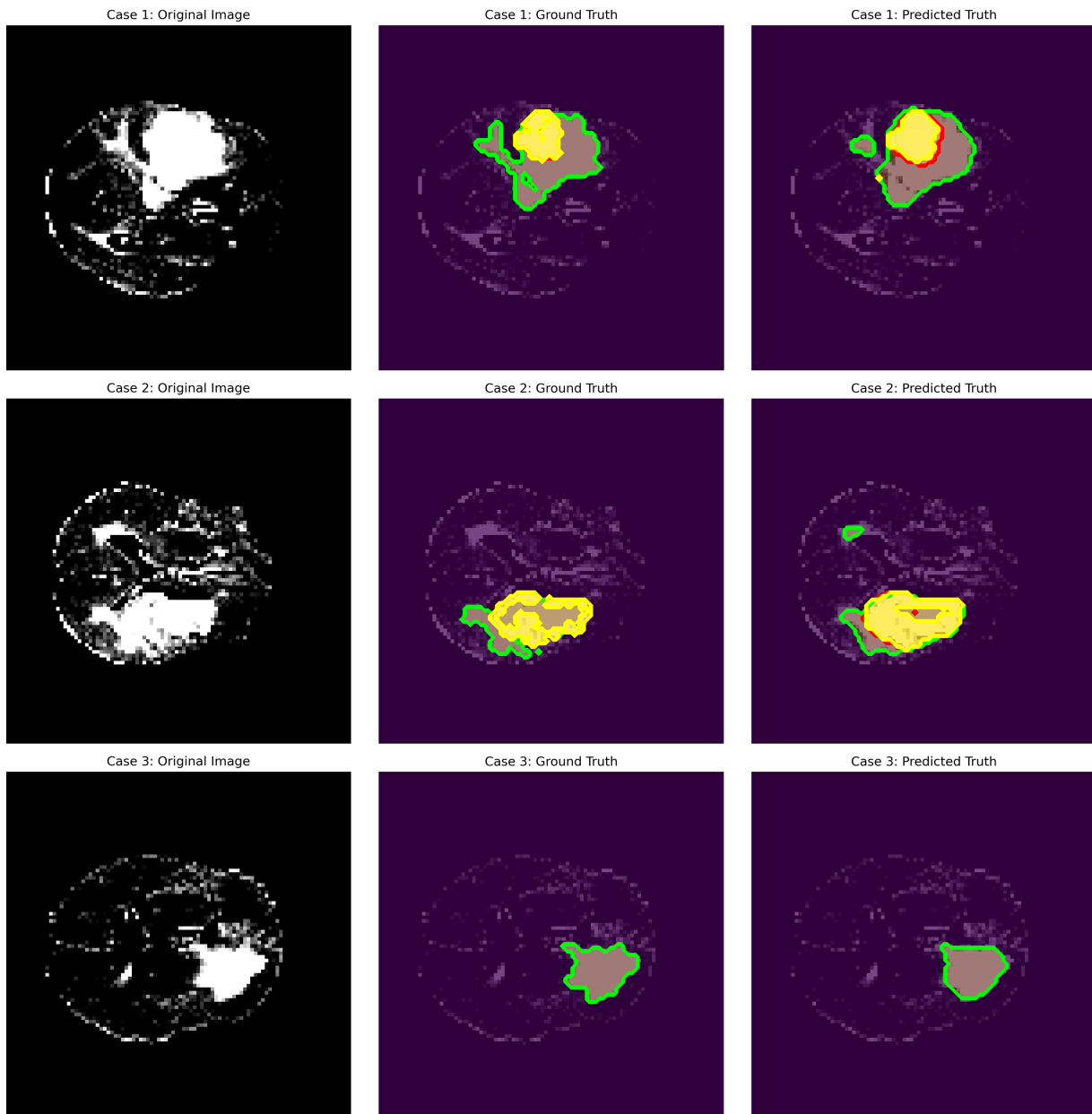
## 6.2 Effect of Convolutional Kernel Size with Fixed Channels

With the TB\_Base transformer bottleneck configuration established from the transformer channel scaling study, the second phase tested the effects of different kernel configurations within the convolutional feature extractor block. During this experiment, the transformer bottleneck remained fixed at TB\_Base, while only the convolutional kernel configurations were varied. Testing started with the single-scale kernels ( $1 \times 1$ ,  $3 \times 3$ ,  $5 \times 5$ ) against different multi-scale configurations to maintain both fine details and larger contextual features, as represented in Table 6. The results show that all single-scale kernels produced lower segmentation performance than the multi-scale configurations expect ( $5 \times 5$ ) kernel outperform. Among the evaluated settings, the combination of ( $1 \times 1$ ,  $3 \times 3$ ,  $5 \times 5$ ) kernels achieved the best overall performance, yielding a Mean DSC of 0.8300 and a Mean HD95 of 7.12. These results indicate that integrating multiple receptive fields enables more effective feature extraction than using a single kernel size alone. Consequently, the LiteDKT-Net ( $1 \times 1$ ,  $3 \times 3$ ,  $5 \times 5$ ) variant with a multi-scale convolutional module perform stable over multiple trials and is applicable for multiscale variations.

**Table 6:** Comparison of different kernel sizes with FLOPs, params, and mean DSC.

| Kernel Size                          | FLOPs (G) | Params (M) | Mean DSC | Mean HD95 |
|--------------------------------------|-----------|------------|----------|-----------|
| $1 \times 1$                         | 0.328     | 0.689      | 0.7605   | 9.36      |
| $1 \times 1, 1 \times 1$             | 0.344     | 0.692      | 0.7702   | 9.26      |
| $3 \times 3$                         | 0.414     | 0.719      | 0.7605   | 9.30      |
| $1 \times 1, 3 \times 3$             | 0.430     | 0.722      | 0.7905   | 8.87      |
| $1 \times 1, 1 \times 1, 3 \times 3$ | 0.447     | 0.726      | 0.8050   | 8.01      |
| $1 \times 1, 3 \times 3, 5 \times 5$ | 0.856     | 0.869      | 0.8300   | 7.12      |
| $3 \times 3, 3 \times 3$             | 0.516     | 0.752      | 0.7982   | 8.65      |
| $1 \times 1, 3 \times 3, 3 \times 3$ | 0.533     | 0.756      | 0.8101   | 7.95      |
| $5 \times 5$                         | 0.737     | 0.832      | 0.7750   | 8.95      |
| $1 \times 1, 5 \times 5$             | 0.754     | 0.835      | 0.7630   | 9.10      |
| $3 \times 3, 3 \times 3, 3 \times 3$ | 0.618     | 0.786      | 0.8040   | 8.03      |
| $3 \times 3, 5 \times 5$             | 0.840     | 0.865      | 0.8104   | 7.85      |
| $5 \times 5, 5 \times 5$             | 1.164     | 0.978      | 0.8010   | 7.95      |
| $5 \times 5, 5 \times 5, 5 \times 5$ | 1.590     | 1.125      | 0.7814   | 8.76      |

Using multi-scale convolutional kernels ( $1 \times 1$ ,  $3 \times 3$ , and  $5 \times 5$ ) enhances local feature extraction and global context modeling, achieving superior segmentation accuracy with minimal computational cost. Combined with the Base channel configuration, the proposed architecture attains robust tumor segmentation using only 0.869M parameters and 0.856 GFLOPs, substantially lower than conventional transformer-based models while maintaining excellent performance, as illustrated in Fig. 9.



**Figure 9:** Analysis of the BraTS2020 dataset in which the first column represents the original brain tumor MRI, the second column represents the ground truth, and the last column represents the predicted truth by the proposed approach.

## 7 Discussion

Based on the ablation study, TB\_Tiny achieved an average Dice score of 0.87, which is higher than TB\_Base's (0.869M, Dice 0.8293). Although the TB\_Tiny configuration achieved the highest Dice score during the isolated channel-scaling ablation, the TB\_Base configuration was selected for the final LiteDKT-Net architecture because it provides a more suitable balance between segmentation accuracy, model capacity, computational efficiency, and effective integration with the proposed multi-scale ( $1 \times 1$ ,  $3 \times 3$ ,

$5 \times 5$ ) convolutional module. This requires greater channel capacity to optimally leverage multi-scale feature diversity.

LiteDKT-Net is about  $2.3\times$  less weighty when compared with EHFF [16] (2M parameters and 6 GFLOPs, WT Dice 89.3) while only seeing a reduction of  $\sim 2\%$  in WT Dice and  $5.4\times$  fewer FLOPs than with MPEDA-Net [18] given similar parameter counts. Compared to WIAF [17] (5.23M parameters and 9.75 GFLOPs with an 85% mean Dice), LiteDKT-Net achieves competitive accuracy at far lower computational cost. The aforementioned comparisons provide sufficient evidence that LiteDKT-Net is the best among lightweight models for brain tumor segmentation in terms of its overall balance between accuracy and resource allocation. With an approximate size of 0.869M parameters and 0.856 GFLOPs, LiteDKT-Net is highly practical for clinical environments that lack sufficient resources for larger models. Future work will integrate explainable AI techniques, such as Grad-CAM, to interpret model predictions and enhance clinical validation and transparency.

## 8 Conclusion

LiteDKT-Net, a sparse and lightweight deep learning architecture tailored for brain tumor segmentation. Extensive experimentation demonstrates that LiteDKT-Net attains segmentation accuracy comparable to leading architectures while substantially reducing computational cost. The model incorporates only 0.869 million parameters and requires 0.856 GFLOPs, making it well-suited for resource-constrained. The architectural design enables accurate tumor annotation with rapid inference and efficient memory utilization. With the optimal kernel configuration of (1, 3, 5) and channel dimensions of (16, 32, 64, 96, 160), the model achieves Dice scores of 0.80, 0.82, and 0.87 for ET, TC, and WT, respectively. Future research directions include enhancing the model's generalization capability and validating its robustness on larger and more diverse datasets to ensure adaptability in complex segmentation.

**Acknowledgement:** The authors want to acknowledge the fund by Princess Nourah bint Abdulrahman University Researchers Supporting Project number (PNURSP2026R346) Princess Nourah bint Abdulrahman University, Riyadh, Saudi Arabia. The authors would also like to acknowledge support of Prince Sultan University, Riyadh Saudi Arabia for APC of this publication.

**Funding Statement:** This research was funded by Princess Nourah bint Abdulrahman University Researchers Supporting Project number (PNURSP2026R346), Princess Nourah bint Abdulrahman University, Riyadh, Saudi Arabia.

**Author Contributions:** Ronak Patel, Amjad R Khan and Faten S. Alamri contributed to Methodology, Implementation, Investigation and Writing—Original Draft; Miral Patel, Awad Alyousef, Bayan AlGhofaily and Deep Kothadiya contributed to Visualization and Supervision and Writing—Review & Editing. All authors reviewed and approved the final version of the manuscript.

**Availability of Data and Materials:** The BraTS 2020 and BraTS2021 dataset used in this study are publicly available and can be accessed through the official Brain Tumor Segmentation (BraTS) Challenge repository at <https://www.med.upenn.edu/cbica/brats2020/data.html>. No datasets were generated.

**Ethics Approval:** This article does not contain any studies with human participants or animals performed by any of the authors.

**Conflicts of Interest:** The authors declare no conflicts of interest.

## References

1. Neamah K, Mohamed F, Adnan MM, Saba T, Ali Bahaj S, Kadhim KA, et al. Brain tumor classification and detection based DL models: a systematic review. *IEEE Access*. 2024;12:2517–42. doi:10.1109/ACCESS.2023.3347545.
2. Naheed N, Shaheen M, Khan SA, Alawairdhi M, Khan MA. Importance of features selection, attributes selection, challenges and future directions for medical imaging data: a review. *Comput Model Eng Sci*. 2020;125(1):314–44. doi:10.32604/cmes.2020.011380.
3. Huang SY, Hsu WL, Hsu RJ, Liu DW. Fully convolutional network for the semantic segmentation of medical images: a survey. *Diagnostics*. 2022;12(11):2765. doi:10.3390/diagnostics12112765.
4. Ronneberger O, Fischer P, Brox T. U-Net: convolutional networks for biomedical image segmentation. In: *Proceedings of the 18th International Conference on Medical Image Computing and Computer-Assisted Intervention—MICCAI 2015*; 2015 Oct 5–9; Munich, Germany. p. 234–41.
5. Zhou Z, Rahman Siddiquee MM, Tajbakhsh N, Liang J. UNet++: a nested U-Net architecture for medical image segmentation. In: *Deep learning in medical image analysis and multimodal learning for clinical decision support*. Cham, Switzerland: Springer; 2018. p. 3–11.
6. Milletari F, Navab N, Ahmadi SA. V-Net: fully convolutional neural networks for volumetric medical image segmentation. In: *Proceedings of the 2016 Fourth International Conference on 3D Vision (3DV)*; 2016 Oct 25–28; Stanford, CA, USA. p. 565–71.
7. Kamnitsas K, Ferrante E, Parisot S, Ledig C, Nori AV, Criminisi A, et al. DeepMedic for brain tumor segmentation. In: *Brainlesion: glioma, multiple sclerosis, stroke and traumatic brain injuries*. Cham, Switzerland: Springer; 2016. p. 138–49.
8. Isensee F, Jaeger PF, Kohl SAA, Petersen J, Maier-Hein KH. nnU-Net: a self-configuring method for deep learning-based biomedical image segmentation. *Nat Meth*. 2021;18(2):203–11. doi:10.1038/s41592-020-01008-z.
9. Bakhtiarnia A, Zhang Q, Iosifidis A. Efficient high-resolution deep learning: a survey. *ACM Comput Surv*. 2024;56(7):1–35.
10. Ramzan F, Khan MUG, Iqbal S, Saba T, Rehman A. Volumetric segmentation of brain regions from MRI scans using 3D convolutional neural networks. *IEEE Access*. 2020;8:103697–709. doi:10.1109/ACCESS.2020.2998901.
11. Vaswani A, Shazeer N, Parmar N, Uszkoreit J, Jones L, Gomez AN, et al. Attention is all you need. In: *Proceedings of the 31st Conference on Neural Information Processing Systems (NIPS 2017)*; 2017 Dec 4–9; Long Beach, CA, USA.
12. Zhang Q, Qi W, Zheng H, Shen X. CU-Net: a U-Net architecture for efficient brain-tumor segmentation on BraTS 2019 dataset. In: *Proceedings of the 2024 4th International Conference on Machine Learning and Intelligent Systems Engineering (MLISE)*; 2024 Jun 28–30; Zhuhai, China. p. 255–8.
13. Cao H, Wang Y, Chen J, Jiang D, Zhang X, Tian Q, et al. Swin-Unet: Unet-like pure transformer for medical image segmentation. In: *Proceedings of the 17th European Conference on Computer Vision—ECCV 2022*; 2022 Oct 23–27; Tel Aviv, Israel. p. 205–18.
14. Chen J, Mei J, Li X, Lu Y, Yu Q, Wei Q, et al. TransUNet: rethinking the U-Net architecture design for medical image segmentation through the lens of transformers. *Med Image Anal*. 2024;97:103280. doi:10.1016/j.media.2024.103280.
15. Walsh J, Othmani A, Jain M, Dev S. Using U-Net network for efficient brain tumor segmentation in MRI images. *Healthc Anal*. 2022;2:100098. doi:10.1016/j.health.2022.100098.
16. Zhang L, Zhang R, Zhu Z, Li P, Bai Y, Wang M. Lightweight MRI brain tumor segmentation enhanced by hierarchical feature fusion. *Tomography*. 2024;10(10):1577–90. doi:10.3390/tomography10100116.
17. Zhong Y, Wang S, Miao Y, Zhang T, Li H. Lightweight brain tumor segmentation through wavelet-guided iterative axial factorization attention. *Brain Sci*. 2025;15(6):613. doi:10.3390/brainsci15060613.
18. Luo H, Zhou D, Cheng Y, Wang S. MPEDA-Net: a lightweight brain tumor segmentation network using multi-perspective extraction and dense attention. *Biomed Signal Process Control*. 2024;91:106054. doi:10.1016/j.bspc.2024.106054.

19. Shahid GES, Ahmad J, Warraich CAR, Ksibi A, Alsenan S, Arshad A, et al. LIU-NET: lightweight inception U-Net for efficient brain tumor segmentation from multimodal 3D MRI images. *PeerJ Comput Sci.* 2025;11:e2787. doi:10.7717/peerj-cs.2787.
20. Lyu Y, Tian X. MWG-UNet++: hybrid transformer U-Net model for brain tumor segmentation in MRI scans. *Bioengineering.* 2025;12(2):140. doi:10.3390/bioengineering12020140.
21. Cao J, Liu J, Chen J. A brain tumor segmentation method based on attention mechanism. *Sci Rep.* 2025;15:15229. doi:10.1038/s41598-025-98355-8.
22. Afridi S, Jan A, Irfan MA, Khattak MI, Khan TA. 3D-ViT-UNet: 3D vision transformer based Unet-like model for volumetric brain tumor segmentation. *PLoS Digit Health.* 2026;5(3):e0001323. doi:10.1371/journal.pdig.0001323.
23. Rahim I, Arpa NR, Sarkar S, Nibir MMI, Khan R. Bridging transformers and state-space models: an efficient hybrid framework for 3D brain tumor segmentation. *Biomed Signal Process Control.* 2026;122:110406. doi:10.1016/j.bspc.2026.110406.
24. Aboussaleh I, Riffi J, Fazazy KE, Mahraz MA, Tairi H. Efficient U-Net architecture with multiple encoders and attention mechanism decoders for brain tumor segmentation. *Diagnostics.* 2023;13(5):872. doi:10.3390/diagnostics13050872.
25. Vatanpour M, Haddadnia J. TransDoubleU-Net: dual scale swin transformer with dual level decoder for 3D multimodal brain tumor segmentation. *IEEE Access.* 2023;11:125511–8. doi:10.1109/ACCESS.2023.3330958.
26. Raza R, Ijaz Bajwa U, Mehmood Y, Waqas Anwar M, Hassan Jamal M. dResU-Net: 3D deep residual U-Net based brain tumor segmentation from multimodal MRI. *Biomed Signal Process Control.* 2023;79:103861. doi:10.1016/j.bspc.2022.103861.
27. Punns NS, Agarwal S. BT-Unet: a self-supervised learning framework for biomedical image segmentation using barlow twins with U-Net models. *Mach Learn.* 2022;111(12):4585–600. doi:10.1007/s10994-022-06219-3.
28. Woo S, Park J, Lee JY, Kweon IS. CBAM: convolutional block attention module. In: *Proceedings of the European Conference on Computer Vision—ECCV 2018*; 2018 Sep 8–14; Munich, Germany. p. 3–19.
29. Mehta R, Filos A, Baid U, Sako C, McKinley R, Rebsamen M, et al. QU-BraTS: MICCAI BraTS 2020 challenge on quantifying uncertainty in brain tumor segmentation-analysis of ranking scores and benchmarking results. *J Mach Learn Biomed Imaging.* 2022;1:1–60. doi:10.59275/j.melba.2022-354b.
30. Sandler M, Howard A, Zhu M, Zhmoginov A, Chen LC. Mobilenetv2: inverted residuals and linear bottlenecks. In: *Proceedings of the 2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition*; 2018 Jun 18–23; Salt Lake City, UT, USA. p. 4510–20.
31. Zhang X, Zhou X, Lin M, Sun J. ShuffleNet: an extremely efficient convolutional neural network for mobile devices. In: *Proceedings of the 2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition*; 2018 Jun 18–23; Salt Lake City, UT, USA. p. 6848–56.
32. Loshchilov I, Hutter F. Decoupled weight decay regularization. *arXiv:1711.05101.* 2017.
33. Zou KH, Warfield SK, Bharatha A, Tempany CMC, Kaus MR, Haker SJ, et al. Statistical validation of image segmentation quality based on a spatial overlap index1 scientific reports. *Acad Radiol.* 2004;11(2):178–89. doi:10.1016/S1076-6332(03)00671-8.
34. Huttenlocher DP, Klanderman GA, Rucklidge WJ. Comparing images using the Hausdorff distance. *IEEE Trans Pattern Anal Mach Intell.* 1993;15(9):850–63. doi:10.1109/34.232073.
35. Isensee F, Jäger PF, Full PM, Vollmuth P, Maier-Hein KH. nnU-net for brain tumor segmentation. In: *Brainlesion: glioma, multiple sclerosis, stroke and traumatic brain injuries.* Cham, Switzerland: Springer; 2021. p. 118–32.
36. Hatamizadeh A, Tang Y, Nath V, Yang D, Myronenko A, Landman B, et al. UNETR: transformers for 3D medical image segmentation. In: *Proceedings of the 2022 IEEE/CVF Winter Conference on Applications of Computer Vision (WACV)*; 2022 Jan 3–8; Waikoloa, HI, USA.
37. Sun H, Yang S, Chen L, Liao P, Liu X, Liu Y, et al. Brain tumor image segmentation based on improved FPN. *BMC Med Imaging.* 2023;23(1):172. doi:10.1186/s12880-023-01131-1.

38. Zhang C, Lu W, Wu J, Ni C, Wang H. SegNet network architecture for deep learning image segmentation and its integrated applications and prospects. *Acad J Sci Technol.* 2024;9(2):224–9. doi:10.54097/rfa5x119.
39. Kumar P, Nagar P, Arora C, Gupta A. U-segnet: fully convolutional neural network based automated brain tissue segmentation tool. In: *Proceedings of the 2018 25th IEEE International Conference on Image Processing (ICIP)*; 2018 Oct 7–10; Athens, Greece. p. 3503–7.