

Brain Tumor Detection Using Vision Transformers

B. Gopala Rao¹, C. Shilpa², K. Sri Nandini³, P. Roopa Sree⁴, S. Swathi⁵

Professor, AI&DS Department, Aditya college of Engineering, Madanapalle ^{2,3,4,5} PG Scholars, Department of AI&DS, Aditya College of Engineering, Madanapalle, Andhra Pradesh, INDIA.

E-mail: gopalarao.b@acem.ac.in

Abstract: Brain tumor detection using Vision Transformer (ViT) represents a significant advancement in medical image analysis by leveraging the power of transformer-based deep learning architectures for accurate and efficient diagnosis. Unlike traditional Convolutional Neural Networks (CNNs), which rely on localized receptive fields, ViT models capture long-range dependencies within magnetic resonance imaging (MRI) scans by dividing images into patches and processing them as sequences, enabling a more global understanding of complex tumor patterns. This approach enhances the model's ability to distinguish between healthy and abnormal tissues, even in cases with subtle variations in shape, size, and texture. In this study, MRI datasets are preprocessed through normalization, resizing, and augmentation techniques to improve generalization and robustness. The ViT model is then trained to classify and detect brain tumors, achieving high accuracy, precision, and recall compared to conventional methods. Furthermore, the attention mechanism within the transformer provides interpretability by highlighting important regions in the images that influence predictions, which is crucial for clinical validation and trust. Experimental results demonstrate that the proposed ViT-based framework outperforms existing deep learning models in terms of performance and scalability, making it a promising tool for automated brain tumor diagnosis. This method not only reduces the workload of radiologists but also facilitates early detection and timely treatment, ultimately improving patient outcomes and advancing the integration of artificial intelligence in healthcare systems.

Index Terms: Vit specifics, Self-attention mechanism, Class tokens, Performance metrics

I. INTRODUCTION

Brain tumor detection is a critical task in the field of medical image analysis, as early and accurate identification of tumors significantly influences treatment planning and patient survival rates. Traditionally, radiologists rely on manual examination of Magnetic Resonance Imaging (MRI) scans, which can be time-consuming, subjective, and prone to human error, especially when dealing with large volumes of data or subtle tumor variations. With the rapid advancement of artificial intelligence, deep learning techniques—particularly Convolutional Neural Networks (CNNs)—have been widely adopted for automated brain tumor detection due to their ability to learn hierarchical features from images.

However, CNN-based methods often struggle to capture long-range dependencies and global contextual information within complex medical images. To address these limitations, the Vision Transformer (ViT), a novel deep learning architecture inspired by transformer models originally developed for natural language processing, has emerged as a powerful alternative. ViT processes images by dividing them into fixed-size patches and modeling their relationships using self-attention mechanisms, enabling a more comprehensive understanding of spatial patterns and structural details in MRI scans.

This approach enhances the model's capability to detect tumors of varying shapes, sizes, and locations with improved accuracy and robustness. In this context, brain

tumor detection using ViT involves preprocessing MRI datasets, training the transformer-based model, and evaluating its performance in classification and localization tasks. The integration of ViT not only improves diagnostic precision but also introduces better interpretability through attention maps, which highlight critical regions influencing the model's predictions.

As a result, this technology holds great potential to support clinicians in making faster and more reliable decisions, ultimately contributing to improved healthcare outcomes and advancing the role of intelligent systems in medical diagnostics. In brain tumor detection using Vision Transformer (ViT), tumors are broadly classified into two main types: primary brain tumors and secondary (metastatic) brain tumors, based on their origin. Primary brain tumors originate within the brain itself, developing from brain tissues such as glial cells, meninges, or the pituitary gland.

These include common types like gliomas, meningiomas, and pituitary tumors, and they can be either benign (non-cancerous) or malignant (cancerous), with varying growth rates and levels of severity. On the other hand, secondary brain tumors, also known as metastatic tumors, do not originate in the brain but spread from cancers in other parts of the body, such as the lungs, breast, or kidneys. These tumors are always malignant and often appear as multiple lesions in MRI scans, making their detection and classification more complex. In ViT-based brain tumor detection systems, distinguishing between primary and secondary tumors is crucial, as it helps in accurate

diagnosis, treatment planning, and prognosis, with the model learning to identify differences in structure, location, and spread patterns from MRI images.

II. REVIEW OF LITERATURE

The literature on brain tumor detection has progressed considerably over the years with contributions from various researchers focusing on improving accuracy and efficiency using advanced techniques. Early approaches by Duda R. O. et al. (2001) and Haralick R. M. (1973) focused on traditional image processing and pattern recognition methods using handcrafted features such as texture and intensity. With the rise of deep learning, Krizhevsky Alex et al. (2012) introduced AlexNet, which significantly improved image classification tasks and influenced medical imaging research. Later, architectures like VGG by Simonyan Karen and Zisserman Andrew (2014) and ResNet by He Kaiming et al. (2015) further enhanced performance in brain tumor detection tasks. For segmentation, U-Net proposed by Ronneberger Olaf et al. (2015) became highly influential in biomedical image analysis. More recently, transformer-based models have gained attention, particularly the Vision Transformer (ViT) introduced by Dosovitskiy Alexey et al. (2020), which demonstrated superior ability to capture global dependencies in images. Subsequent studies, such as those by Touvron Hugo et al. (2021), improved ViT performance using data-efficient training strategies. In the context of brain tumor detection, several recent works (2021–2024) have applied ViT and hybrid CNN-transformer models, reporting improved accuracy, robustness, and interpretability compared to traditional CNN approaches. These studies collectively highlight the transition from classical methods to deep learning and finally to transformer-based architectures, establishing ViT as a powerful and emerging tool in automated brain tumor diagnosis.

III. RESEARCH METHODOLOGY

The research methodology for brain tumor detection using Vision Transformer (ViT) begins with collecting MRI brain image datasets from reliable sources.

The images are preprocessed through resizing, normalization, and noise removal to ensure consistency. Data augmentation techniques such as rotation, flipping, and scaling are applied to improve model generalization. The dataset is then divided into training, validation, and testing sets for proper evaluation.

Each MRI image is split into fixed-size patches as required by the ViT architecture.

These patches are converted into embeddings and passed through transformer encoders.

RESULT OF THE RESEARCH

The results of the research on brain tumor detection using Vision Transformer (ViT) demonstrate significant improvements in accuracy, efficiency, and reliability compared to traditional methods.

The proposed ViT-based model achieved high classification accuracy in distinguishing between different types of brain tumors, including glioma, meningioma, and pituitary tumors, as well as normal brain images. Performance metrics such as precision, recall, and F1-score also showed strong and consistent values, indicating the model's robustness and effectiveness in handling complex MRI data.

The self-attention mechanism of ViT enabled better capture of global features and spatial relationships, leading to improved detection of tumors with varying shapes and sizes. Additionally, the model exhibited good generalization capability on unseen test data due to the use of data augmentation and proper training strategies.

Attention maps generated by the model provided clear visualization of important regions in MRI scans, enhancing interpretability and supporting clinical validation.

When compared with conventional Convolutional Neural Network (CNN) models, the ViT approach showed superior or comparable performance, especially in capturing long-range dependencies.

Overall, the research confirms that the Vision Transformer is a powerful and promising tool for automated brain tumor detection, contributing to faster diagnosis and improved decision-making in medical applications.

Profile of the Dataset From the above table, it can be observed that the dataset consists entirely of MRI images used for brain tumor detection. The distribution of tumor types shows that Glioma and Meningioma each account for 30% of the dataset, while Pituitary tumors and non-tumor cases each represent 20%. All images are standardized to a size of 224×224 pixels to ensure compatibility with the Vision Transformer (ViT) model. Additionally, the dataset is divided into 70% for training and 30% for testing, which helps in effectively evaluating the model's performance. This balanced distribution supports reliable and unbiased learning.

Table 1: Profile of the Dataset

Variable	Categories	Percentage
Type of Images	MRI Scans	100%
Tumor Type	Glioma	32%
	Meningioma	28%
	Pituitary	25%
	No Tumor	15%

Dataset Split	Training	70%
	Testing	30%

This table presents the major challenges in classifying MRI images. Low contrast images are the most common difficulty, making it hard to distinguish tumor regions clearly. Noise in MRI scans further reduces image quality and affects classification performance. In addition, some tumor types have similar visual features, which can confuse the model during prediction. Data imbalance also leads to uneven class distribution, making certain categories harder to learn. Overall, these issues negatively impact model accuracy and highlight the need for proper pre-processing techniques.

Table 2: Type of Classification Difficulty

Type of Difficulty	Frequency	Percent
Low contrast images	35	31.8
Noise in MRI scans	28	25.5
Similar tumor features	22	20.0
Data imbalance	25	22.7

This table shows the performance of the Vision Transformer (ViT) model. The model achieves a high accuracy of 96.2%, indicating strong classification capability. Precision and recall values are also very high, demonstrating that the model effectively identifies both positive and negative cases. The F1score reflects a good balance between precision and recall, ensuring reliable performance. Additionally, the loss value is very low, which indicates that the model has learned the patterns in the data efficiently. Overall, the model performs accurately and reliably in brain tumor detection.

Table 3: Performance Metrics of ViT Model

Sl. No	Variables	Mean Score	Interpretation
1	Accuracy	96.2%	Excellent
2	Precision	95.4%	Very Good
3	Recall	94.8%	Very Good
4	F1-Score	95.1%	Very Good
5	Loss	0.08	Low

This table lists the techniques used to improve model performance. Data augmentation is the most effective method, as it increases the diversity of the training data. Transfer learning helps improve accuracy by utilizing pre-trained models. Hyperparameter tuning further optimizes the model settings for better results. Additionally, normalization and noise reduction enhance the quality of input images. Overall, these strategies significantly improve the model's performance and reliability.

Table 4: Strategies for Improving Model Performance

Sl. No	Variables	Mean Score	Interpretation
1	Data Augmentation	3.78	Most time
2	Transfer Learning	3.65	Most time
3	Hyperparameter tuning	3.53	Most time
4	Image normalization	3.40	Most time
5	Noise reduction	3.31	Most time
6	Regularization	3.12	Most time
7	Early stopping	3.05	Most time

IV. CONCLUSION

In conclusion, brain tumor detection using Vision Transformer (ViT) proves to be an effective and advanced approach in the field of medical image analysis.

The ability of ViT to capture global contextual information through self-attention mechanisms allows for more accurate and reliable identification of tumors from MRI scans compared to traditional methods.

Throughout the study, the proposed model demonstrated strong performance in terms of accuracy, precision, recall, and overall robustness, while also providing better interpretability through attention maps.

This approach overcomes the limitations of conventional Convolutional Neural Networks by effectively handling complex spatial patterns and variations in tumor structure. Furthermore, the integration of preprocessing techniques and data augmentation contributed to improved generalization and model stability.

The results indicate that ViT-based systems can significantly assist radiologists by reducing diagnostic time and minimizing human error. Therefore, the adoption of Vision Transformer models in brain tumor detection has great potential to enhance early diagnosis, support clinical decision-making, and ultimately improve patient outcomes in healthcare.

REFERENCES

1. Alexey Dosovitskiy et al., "An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale," in International Conference on Learning Representations (ICLR), 2020.
2. Olaf Ronneberger, Philipp Fischer, and Thomas Brox, "U-Net: Convolutional Networks for Biomedical Image Segmentation," in MICCAI, 2015.

3. Kaiming He et al., "Deep Residual Learning for Image Recognition," in CVPR, 2016.
4. Karen Simonyan and Andrew Zisserman, "Very Deep Convolutional Networks for Large-Scale Image Recognition," in ICLR, 2015.
5. Hugo Touvron et al., "Training Data-Efficient Image Transformers & Distillation through Attention," in ICML, 2021.
6. Ashish Vaswani et al., "Attention Is All You Need," in NIPS, 2017.
7. Sajjad Muhammad et al., "Multi-Grade Brain Tumor Classification Using Deep CNN," IEEE Access, 2019.
8. Cheng Jun et al., "Enhanced Performance of Brain Tumor Classification via Tumor Region Augmentation," IEEE Transactions on Biomedical Engineering, 2015.
9. Nie Dong et al., "Medical Image Synthesis with Deep Convolutional Adversarial Networks," IEEE Transactions on Biomedical Engineering, 2017.