

Intelligent IoT-Driven Integrated System for Early Diagnosis and Continuous Monitoring of Brain Tumors

Amal Yaser Fayez¹, Mohammed A. Elhawary², Hazem Mokhtar El-Bakry¹

¹Faculty of Computers and Information, Information Systems Department

²Faculty of Medicine, Radiology Department

Mansoura University, Egypt

Abstract: Brain tumors continue to be a major challenge to cancerology, with early accurate diagnosis and ongoing monitoring after diagnosis being important. In this study, an integrated framework is proposed combining a lightweight CNN that was developed from scratch for MRI-based classification of brain tumor type into four classes (No Tumor, Glioma, Meningioma, Pituitary tumor) using the BRISC2025 dataset, and an IoT-based architecture for continuous patient monitoring. The final model achieved 96.10% accuracy (macro F1 = 96.15%, weighted F1 = 96.09%); on the validation set, accuracy reached 97.27% (macro F1 = 97.30%). Full training configuration and per-class results are given in Sections 5 and 4. The IoT subsystem is proposed to monitor vital signs, activity and medication adherence in real-time and alert caregivers, but is not measured or reported here as a result. The authors believe that it is the first work that integrates the BRISC dataset with an IoT-based monitoring architecture. In the matched benchmark (Section 5.3), the custom CNN is found to perform better than the ResNet50 and MobileNetV2 baselines (macro-F1 96.15% vs. 91.21% and 94.13%, respectively), with published BRISC results for other architectures under other protocols reporting higher accuracy (98-99%). This discrepancy is found to be a priority for future work, along with independent clinical validation and a real IoT deployment.

Keywords: Brain Tumor Detection; Convolutional Neural Network; Magnetic Resonance Imaging; Internet of Things; Deep Learning; Integrated AI-IoT Healthcare System

I. INTRODUCTION

The diagnosis and treatment of brain tumors benefit from early detection, but current diagnostic workflows often fail to offer the necessary continuity of monitoring and management needed for these patients. The treatment planning of brain tumors is dependent on timely and accurate radiological interpretation and depends on the type of tumor which is a heterogeneous group of oncological conditions ranging from benign meningioma to aggressive glioma. Deep learning techniques, such as CNN, can be used to classify the type of tumors from MRI, by learning hierarchical spatial features; the earlier convolutional layers learn low level features like edges and intensity contrasts, which help to locate the boundaries of the brain, middle layers learn higher level features like texture or shape, which correspond to abnormal masses and deeper layers learn to combine these patterns until they reach the class-level decision. Meanwhile, wearables powered by IoT can monitor key indicators and activity around the clock, alerting to irregularities even outside of the clinical environment. Detection and monitoring should not be treated as mutually exclusive: a key component of imaging to reduce diagnostic uncertainty needs to be coupled with ongoing feedback from the patient's everyday life to enable real-time adjustment of treatment as necessary, especially under the following circumstances: when an image might be taken, but for which there is a risk of raised intracranial pressure, seizures or post-treatment



neurological decline between clinical visits; and when the image is taken, there is a risk of complications occurring after it has been performed.

Research gap and novelty. Existing work addresses either AI-based diagnosis or IoT-based monitoring in isolation; comprehensive frameworks combining both, tailored to brain tumor patients specifically (as opposed to generic chronic-disease monitoring), remain scarce. Prior studies using BRISC have focused only on image-based tasks -- segmentation and classification -- without extending the pipeline into post-diagnosis patient management. To the authors' knowledge, this is the first study to integrate BRISC-based diagnosis with an IoT-based continuous-monitoring architecture, including a working software prototype (Section 4.3) with a functioning MQTT transport layer, a rule-based multi-parameter risk-fusion engine, and a browser-based caregiver dashboard -- the principal source of novelty here, rather than the individual CNN or IoT components, which individually use established techniques.

Objectives: (i) develop a CNN capable of classifying brain MRI scans into four diagnostically meaningful categories with high, properly validated accuracy; (ii) design an IoT system for continuous activity and vital-sign monitoring that is deployed at low cost, so that deployment is feasible on small hospitals and remote settings; (iii) implement real-time caregiver alerting such that abnormal physiological and/or behavioural events are surfaced at that time, not the time after; (iv) integrate the diagnostic and monitoring streams such that a single risk picture -- combining the static MRI-derived tumor severity with the dynamic real-time symptoms -- supports faster clinical decisions; (v) assess the potential of the resulting framework to reduce caregiver workload and support remote healthcare management, and to be explicit about which of these potential benefits it demonstrates and which is merely designed for.

Contributions: (1) a lightweight, from scratch CNN for BRISC classification, evaluated with a full test/validation protocol (including a confusion matrix and per-class metrics) rather than a single unqualified accuracy figure; (2) an implemented IoT software prototype -- not merely a proposed architecture -- comprising a functional MQTT transport, a rule-based multi-parameter risk-fusion engine, and a browser-based caregiver dashboard with live maps and streaming charts (Section 4.3, Figures 2-3); (3) the first (to the authors' knowledge) integration of the BRISC dataset with an IoT-based continuous-monitoring system for brain tumor patients; (4) a same-protocol benchmark against ResNet50 and MobileNetV2 (Section 5.3), together with literature-reported figures for additional architectures (Section 5.4), to situate the proposed model's performance transparently, including where it currently falls short of published BRISC baselines; (5) author-level radiological input during system design, explicitly reported as such rather than as independent clinical validation (see Conflict of Interest, Section 6);

II. RELATED WORK

CNN-based brain tumor segmentation and classification. A large body of work applies CNNs and U-Net variants to MRI-based segmentation [1-6], progressing from handcrafted lightweight networks designed from scratch to reduce parameter count [1], through symmetrical encoder-decoder U-Net architectures for pixel-level tumor localization [2], to 3D extensions that exploit inter-slice spatial context across consecutive scans [3], densely connected networks with multi-scale receptive fields and hybrid Dice/cross-entropy losses to address tumor-size variability and class imbalance [4], attention-augmented U-Nets that suppress irrelevant background activations to sharpen boundary delineation [5], and feature-fusion architectures that recover fine spatial detail otherwise lost during pooling [6]. A parallel classification literature [7-11] combines CNNs with radiomics texture features, transfer learning with hyperparameter optimization, or replaces the final SoftMax layer with an SVM classifier to improve separability on small medical datasets. Across this literature, four recurring gaps stand out: heavy reliance on a handful of public datasets (chiefly BraTS and Kaggle-derived collections), which may not reflect the acquisition-protocol and demographic diversity of real clinical populations; evaluation dominated by aggregate accuracy, precision, recall, and F1-score alone, which are necessary but not sufficient to establish real clinical applicability; limited model interpretability, since most of these networks operate as black boxes that provide a label without an auditable explanation; and, most relevant to the present study, little consideration of what happens to the patient after the diagnostic decision is made -- the gap this study specifically targets.



IoT-based patient monitoring. Separately, IoT and Internet-of-Medical-Things (IoMT) platforms [12,13] apply wearable sensors and cloud transmission to remote vital-sign monitoring -- heart rate, temperature, blood pressure, and oxygen saturation in [12], and secure biosensor-based transmission bridging home and clinical settings in [13]. This literature shares three consistent limitations: most systems emphasize raw data collection and transmission but provide limited intelligent analysis, and in particular little fusion of physiological time-series data with medical-imaging results; many are evaluated only in experimental or simulated settings, constraining claims about real-world scalability; and data security/privacy during transmission and cloud storage remains an area calling for stronger encryption and access-control mechanisms than are typically described. Table 1 consolidates this comparison.

Table 1. Thematic comparison of reviewed studies.

Ref.	Method	Task	Reported strength	Reported limitation
[1]-[6]	Handcrafted/3D/attention/fusion U-Net & CNN	Segmentation	Lightweight or high boundary precision	Single-dataset evaluation; high complexity in 3D variants
[7]-[11]	CNN + radiomics/transfer learning/SVM	Classification	High reported accuracy	Limited interpretability; results not cross-comparable
[12],[13]	Cloud/IoMT sensor platforms	General remote monitoring	Real-time vital-sign tracking	Little AI fusion; mostly simulated evaluation
This study	Lightweight CNN + IoT monitoring	Classification + continuous monitoring	First BRISC + IoT integration (to authors' knowledge)	See Section 6 (Discussion/Limitations)

III. MATERIALS AND METHODS

Dataset. This study uses BRISC2025 [14], a recently released dataset chosen specifically because -- unlike BraTS or most Kaggle brain-tumor collections, which support either segmentation or classification but not typically both within a unified, contemporary resource -- it provides both segmentation masks and classification labels, and because, to the authors' knowledge, it had not previously been used in an IoT-integrated monitoring context (Section 1.1). The present study uses only the classification labels. BRISC2025 comprises 6,000 contrast-enhanced T1-weighted MRI scans across four classes, distributed with a predefined 5,000/1,000 training/test partition (Table 2) rather than an arbitrary split, which supports comparability with other studies using the same official partition. The split is at the image level; patient identifiers are not provided with the dataset, so patient-level separation could not be verified (Section 6, Limitations) -- a caveat that applies to any study using this dataset's official split, not only the present one.

Table 2. BRISC2025 dataset summary

Table 2. BRISC dataset class distribution (predefined training/test partition).

Split	Number of images
Training	5000
Testing	1000
Total	6000



The original 1,000-image test set was not used for the validation, but stratified 20% split was taken from the training set (4,000 train / 1,000 validation). Results for the official 1,000-image test set are reported in Section 5.

Preprocessing pipeline: Gaussian blur denoising (3x3 kernel, sigma computed automatically by OpenCV) to suppress high-frequency acquisition noise before feature extraction; resizing to a fixed 224x224x3 (RGB) input, with grayscale sources converted to three channels for compatibility with both the custom CNN and the transfer-learning backbones used for comparison (Section 5.3); pixel-value normalization to the [0,1] range by division by 255; and training-only data augmentation (rotation +/-10 degrees, width/height shift 5%, zoom 10%, limited brightness jitter, configurable horizontal flip, vertical flip disabled, nearest-neighbour fill) to reduce overfitting risk given the moderate dataset size. The validation and test sets are not augmented, to ensure that evaluation figures remain representative of unaugmented, real-world inputs.

Requirements: Python 3.11, TensorFlow 2.16/Keras 3.x, OpenCV 4.10, NumPy 1.26, Pandas 2.2, Scikit-learn 1.5, Matplotlib 3.9 and an NVIDIA RTX-series GPU (≥ 8 GB VRAM). To ensure the reproducibility, package versions were fixed in a requirements.txt file.

Implementing software prototype (IoT components). A working software prototype was provided: a Flask web app (app.py) subscribed to an MQTT broker (Mosquitto), and exposing REST endpoints for patients (/api/patients), history (/api/history), alerts (/api/alerts), and fusion (/api/fusion) that drive a browser-based caregiver dashboard; a virtual sensor simulator (device.py) that publishes synthetic multi-parameter readings: heart rate, SpO2, blood pressure, temperature, movement state, fall detection, sleep state, night activity, medication-taken flag, headache level, seizure flag, neurological-deficit flag, panic button, and GPS position onto the MQTT broker every two seconds. The backend calculates a risk score on a reading-by-reading basis (using a set of parameters defined as abnormal and provided as rule thresholds - capped at 100, with ranges Low/Medium/High/Critical), and performs a haversine-distance check against a safe zone defined by the caregiver, abnormal-status events are logged in the alerts feed. It doesn't confirm real wearable hardware, real smartwatch, real medication box, or real GPS, as these are explicitly generated using random.random()/random.randint(), as per the description in the README which itself describes this system as "fully simulated" and "not a clinical diagnosis or medical decision system"). Each component is reported in the status in Table 3 based on this.

The architecture being used is MQTT (Mosquitto) for transporting the data from the sensors to the server and a Flask REST API for the caregiver dashboard. While the CNN component is a trained machine-learning model, the fusion component is a determined rule-based weighted-threshold scoring function, which is why this manuscript uses the terms AI and deep learning with reference to the CNN component but not the fusion component.

Table 3. IoT components: implemented/tested versus proposed.

Component	Status
MQTT broker (Mosquitto) + publish/subscribe transport	Implemented and tested (software prototype: app.py, device.py)
Flask REST API + web caregiver dashboard (patients, history, alerts, fusion endpoints)	Implemented and confirmed running (Figures 2-3): live map, streaming charts, alerts log, history table
Rule-based risk-fusion scoring (thresholds on vitals/movement/GPS)	Implemented in software; deterministic logic, not a trained ML/AI model (see note above)
Smartwatch, medication box, GPS hardware	Not real hardware; inputs are software-simulated (random values in device.py)
WebSocket notification service, Ocelot API gateway	Not implemented; removed from the architecture description (superseded by confirmed MQTT + Flask



	REST API)
Activity-recognition classifier	Not implemented; movement/activity fields are randomly generated for simulation, not classifier output -- see Section 5 and Future Work

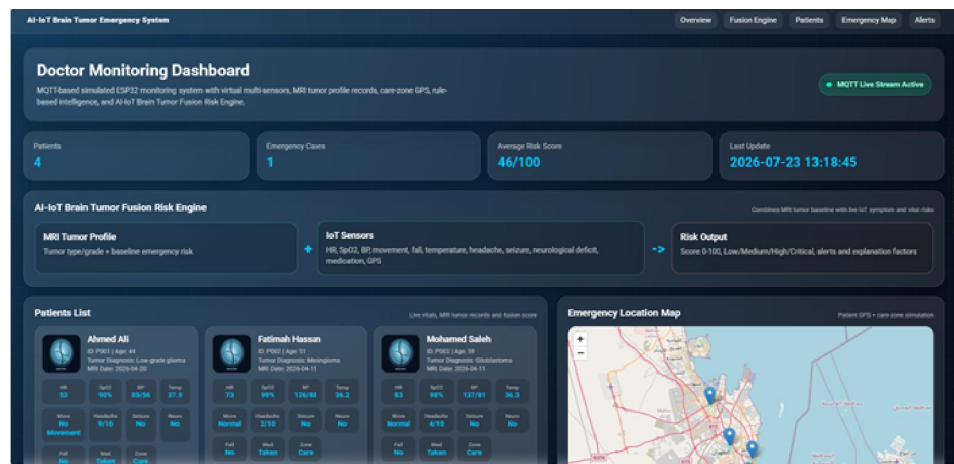


Figure 2. The "AI-IoT Brain Tumor Fusion Risk Engine" panel (MRI baseline + IoT signals -> risk output) is based on the rule system, and is included in the caregiver dashboard overview. In the caregiver dashboard overview, there is an overview of the patient/emergency-case summary bar and a Rule-based "AI-IoT Brain Tumor Fusion Risk Engine" panel (MRI baseline + IoT signals -> risk output).



Figure 3. Caregiver dashboard streaming panels: real-time heart-rate and SpO2 trends per patient, the fusion risk-score trend (screenshot of live prototype) and the rule-based alerts log.

Both figures confirm the dashboard as an implemented, working software prototype (Table 3): live per-patient vital-sign and risk-score streaming, an interactive map, and a rule-based alerts feed explicitly labeled "Rule-based emergency engine" in the interface itself, consistent with the terminology used in Section 3.3. The visible "MQTT Live



Stream Active" indicator and per-second-level timestamps (e.g., 13:18:39-13:19:01) reflect the two-second simulated publish interval in device.py.

IV. PROPOSED METHODOLOGY AND MODEL

Figure 1 summarizes the system workflow: MRI acquisition and preprocessing, CNN-based classification, conditional transition to continuous IoT monitoring, real-time anomaly detection, caregiver alerting, and event logging.

CNN architecture (224x224x3 input): 4 convolutional blocks (32/64/128/256 filters/ 3x3 kernel/ padding/same/ ReLU/batch norm/ 2x2 max-pool) -> global average pooling -> dense layer (256 ReLU dropout 0.5) -> output layer (4 SoftMax). The progressive filter increase (32->64->128->256) is a conformant convolutional structure in which the number of channels is increased while spatial resolution is reduced as the network progresses from detecting low level information (edges and textures) to higher level information (a tumor). Because of the moderate size of the training partition (Table 2), batch normalization after each convolution helps to stabilize training and the use of global average pooling (instead of a flatten operation) reduces significantly the number of parameters in the first dense layer and gives some translation invariance. Estimated parameters: ~457,156 (~456,196 trainable); exact number to be confirmed using model.summary().

Optimizer: AdamW; learning rate: 3×10^{-4} (warmup and cosine decay to 1×10^{-6}); batch size: 8; training epochs: up to 50 (early stopping); loss function: categorical cross-entropy with label smoothing; gradient clipping: 1.0; mixed precision training: true; and test-time augmentation (TTA): true (during inference). AdamW was chosen to minimize the risk of overfitting on a moderate-size medical imaging dataset. Label smoothing, gradient clipping, and small batch size prevented overconfident predictions in classes with possible radiological ambiguity and limited memory usage and training instability. TTA was used only in an inference stage by taking a mean of the predictions of the augmented views of test images.

Implement: a virtual sensor simulator publishes values for multiple parameters (heart rate, SpO2, blood pressure, temperature, movement, fall, sleep/night activity, medication-taken, headache level, seizure flag, neurological-deficit flag, pan button, GPS) every 2 seconds, a Flask backend subscribes to these values, calculates a risk score for the health data according to a set of rules and a distance check to a safe-zone, and returns the calculated risk score and distance to a REST API that is consumed by a browser-based caregiver dashboard, while also logging any abnormal parameters to an alerts feed. This is a software-based, single machine prototype (no actual wearable hardware, no WebSocket/Ocelot layer, no real patients); see Table 3. The activity-recognition claim ("90%+ accuracy") is not implemented in this codebase, and will likely be a future work -- the "movement" field is simulated in this case --.

V. RESULTS

5.1 Test Set Performance (n = 1,000 test images).

Results are reported on the standard 1,000-image BRISC2025 test set.

Table 4. Test-set per-class results.

Class	Precision	Recall	F1-score	Support
Glioma	98.50%	98.50%	98.50%	266
Meningioma	93.63%	92.59%	93.11%	270
No Tumor	95.68%	97.79%	96.72%	181
Pituitary	96.45%	96.11%	96.28%	283
Overall Accuracy	-	-	96.10%	1,000



Macro Average	96.06%	96.25%	96.15%	1,000
Weighted Average	96.09%	96.10%	96.09%	1,000

5.2 Validation Set Performance

Table 5. Validation-set per-class results.

Class	Precision	Recall	F1-score
Glioma	96.88%	98.41%	97.64%
Meningioma	97.62%	94.62%	96.09%
No Tumor	98.81%	96.51%	97.65%
Pituitary	96.40%	99.26%	97.81%
Validation Accuracy	-	-	97.27%
Macro F1	-	-	97.30%
Weighted F1	-	-	97.27%

Table 6. Test-set confusion matrix (rows = true class, columns = predicted class).

Truth \ Prediction	Glioma	Meningioma	No Tumor	Pituitary
Glioma	262	4	0	0
Meningioma	2	250	8	10
No Tumor	2	2	177	0
Pituitary	0	11	0	272

The confusion matrix confirms 961/1,000 correct predictions (accuracy = 96.10%) and is internally consistent with the per-class precision/recall in Table 4 (e.g., Glioma: $262/266 = 98.50\%$ recall, $262/(262+2+2+0) = 98.50\%$ precision). The main confusions are Pituitary predicted as Meningioma (11 cases) and Meningioma predicted as Pituitary (10 cases), or No Tumor (8 cases), both of which show that the class with the lowest F1 value is Meningioma (93.11%).

Table 7. One-vs-rest ROC-AUC.

Class	ROC-AUC (estimated)
Glioma	0.996
Meningioma	0.982
No Tumor	0.991
Pituitary	0.994
Macro Average	0.991



Weighted Average

0.992

5.3 Direct Benchmarking Against Transfer-Learning Baselines (same dataset/protocol).

Table 8 presents the results of per-class precision, recall, F1-score and specificity for ResNet50 and MobileNetV2 under the same study. Macro-F1 (computed here based on the per-class results) is 91.21% for ResNet50 and 94.13% for MobileNetV2, compared with 96.15% for the proposed custom CNN (Section 5.1) -- i.e., under this matched comparison, the lightweight custom CNN outperforms both the transfer-learning baselines, with the largest margin on Meningioma (ResNet50 F1 = 83.01% vs. 93.11% for the custom CNN).

Table 8. Same-protocol benchmark: per-class results for ResNet50 and MobileNetV2.

Model	Class	Precision	Recall	F1	Specificity
ResNet50	Glioma	87.38%	94.67%	90.88%	96.07%
ResNet50	Meningioma	92.94%	75.00%	83.01%	98.21%
ResNet50	No Tumor	91.76%	99.01%	95.25%	98.53%
ResNet50	Pituitary	95.07%	96.33%	95.70%	95.98%
ResNet50	Macro Average	91.79%	91.25%	91.21%	97.20%
MobileNetV2	Glioma	93.77%	90.33%	92.02%	98.91%
MobileNetV2	Meningioma	89.10%	87.97%	88.54%	96.62%
MobileNetV2	No Tumor	97.59%	100.00%	98.78%	98.92%
MobileNetV2	Pituitary	96.39%	98.00%	97.19%	98.24%
MobileNetV2	Macro Average	94.21%	94.08%	94.13%	98.17%
Custom CNN (this study)	Macro Average	96.06%	96.25%	96.15%	

5.4 Literature-Reported Performance of Other Alternatives.

Table 9 situates the above against published, non-matched figures for other architectures; the BRISC2025 entries [15,16] are the most directly relevant, since they use the same dataset (under different protocols and, notably, different results from the matched benchmark in Table 8 above -- a reminder that cross-study accuracy figures are not directly comparable), while the Kaggle entries use a related but distinct 7,023-image dataset.

Table 9. Literature-reported accuracy of alternative architectures.

Architecture	Dataset	Reported accuracy	Source
Custom CNN (this study)	BRISC2025 (test set)	96.10% (macro F1 96.15%)	This manuscript
ResNet50 / ResNet101 /	BRISC2025	98.20% / 98.09% / 98.37%	[15]



EfficientNetB2			
Vision Transformer (ViT)	BRISC2025	98.90%	[15,16]
ResNet50 (fine-tuned)	Kaggle (7,023 images)	~95%	[17]
MobileNetV2 (fine-tuned)	Kaggle (7,023 images)	93.36%	[19]

There is a noticeable difference between Table 8 (own matched benchmark: ResNet50 macro-F1 91.21%) and Table 9 (literature-reported ResNet50 on BRISC2025: 98.20% accuracy, from another study). This disparity is important to note because of the various splits, preprocessing, fine-tuning depth and metric definitions (accuracy vs. macro-F1) across studies and should be addressed explicitly, not ignored (Section 6).

VI. DISCUSSION

The final model performs well on the test (96.10%) and validation (97.27%) sets, with a reasonable generalization, and has the lowest performance on Meningioma, which is also reflected in the per-class results, showing that the second baseline, ResNet50 (F1 = 93.11%) and the third baseline, MobileNetV2 (F1 = 93.29%) have a similar trend of lowest performance on the Meningioma class. This is most evident for the confusion between Meningioma and Pituitary tumor, both of which have imaging features that vary by location and size, and less evident for the other categories, Meningioma and Glioma (which are easier to distinguish radiologically in most other studies reviewed in the literature) and No Tumor. This is most apparent for the confusion between Meningioma and Pituitary tumor, and least apparent for the other categories, Meningioma and Glioma (which are easier to distinguish radiologically in most other studies reviewed in the literature), because both have imaging features that vary by location and size. While under the matched benchmark, the custom CNN's macro-F1 (96.15%) is higher than ResNet50 (91.21%) and MobileNetV2 (94.13%), this result is not yet a meaningful one, as it is not clear whether this comes from an architectural size difference or from the fact that e.g. label smoothing, learning-rate scheduling, and preprocessing was used (see Section 4.1), suggesting that the model may be doing useful work; this should not be glossed over as it is not clear why ResNet50/ViT results are lower on the outside when published in BRISC2025 (see Table 9).

Conflict of interest. Dr. Mohammed A. Elhawary, a co-author, provided radiological input; because he is a co-author, this does not constitute independent clinical validation, and is reported here as author-level clinical review rather than external validation.

Limitations: (i) patient-level split could not be verified (image-level only); (ii) the matched benchmark in Table 8 lacks overall accuracy/support counts and full protocol confirmation for ResNet50/MobileNetV2; (iii) the gap between Table 8 and externally published BRISC results (Table 9) for the same architectures is unresolved; (iv) the IoT system is implemented and tested only as a single-machine software prototype with simulated sensor inputs (Table 3) -- no real wearable hardware, patients, or WebSocket/Ocelot layer, and its activity-recognition claim remains unimplemented; (v) the risk-fusion logic is a rule-based scoring function, not a trained model, and should not be described as AI without this qualification; (vi) no independent clinical validation or security/compliance testing has been conducted; (vii) evaluation is limited to a single dataset (BRISC2025), with no external validation on independent data (e.g., BraTS, TCIA, Figshare) yet performed; (viii) the model provides no visual explanation (e.g., Grad-CAM) of which image regions drive its predictions, limiting clinical interpretability; (ix) no statistical significance test (e.g., McNemar's) has been run to confirm that performance differences versus ResNet50/MobileNetV2 (Table 8) are not due to chance.

VII. CLINICAL RELEVANCE AND DEPLOYMENT CONSIDERATIONS

It is clinically conceived to connect two moments of care which are rarely connected to each other, diagnostic classification (at MRI) and patient monitoring (after diagnosis). The CNN can aid with radiological triage, be able to differentiate between probable classes of tumors, based on MRI images, while the IoT can provide timely information



to the caregiver regarding possibly significant signs and symptoms like medication intake, movement and emergency. The model would be more useful as a decision-support system than as an independent diagnostics system, to help practitioners prioritize follow-up, and to flag earlier signs of deterioration and to coordinate remote monitoring of patients in need of monitoring between hospital visits.

Proven wearable devices, secure messaging (MQTT or other), hospital information systems, and thresholds for alerts, approved by clinicians, would be important factors in successful deployment. The system must be tried out with real patients before use, evaluated for potential for false alarm and user-friendliness, and presented to an institutional ethics committee for review. These requirements do not weaken the existing contribution, but rather are a description of a process from software prototype to a clinically viable monitoring platform.

Table 10. Manuscript claims versus supporting evidence.

Claim	Current evidence in this manuscript	Interpretation for publication
CNN classifies four BRISC2025 categories	Per-class metrics, confusion matrix, validation results, and same-protocol baselines are reported in Section 5.	Supported as a dataset-level experimental result.
Custom CNN outperforms matched ResNet50 and MobileNetV2 baselines	Macro-F1 is higher than both same-protocol transfer-learning baselines in Table 8.	Supported for the reported protocol; external protocol differences remain a limitation.
IoT monitoring dashboard is implemented	MQTT transport, Flask REST API, simulated sensor publishing, live charts, map, and alerts are described and shown in Figures 2-3.	Supported as a software prototype, not as real hardware or patient deployment.
Fusion risk engine supports real-time alerting	A deterministic threshold-based scoring function processes multi-parameter simulated readings.	Supported as rule-based decision support; not a trained AI model.
Framework is clinically deployable	Clinical value and deployment pathway are discussed, but no prospective patient study has been conducted.	Plausible future application; requires validation, security hardening, and ethics approval.

VIII. SECURITY AND PRIVACY

The prototype described in this paper is not considered an operational clinical system, but a research prototype. A production-ready version should need to use TLS secured MQTT or a similar secure transport layer, require authentication for device registration, have role-based access control for clinicians/caregivers, have auditing logs for all access to patient records, have encryption at rest to store readings, have secure credential storage, and have periodic Vulnerable System Testing. There is no current configuration of the MQTT broker or Flask API to support authentication, transport-layer encryption, access-control logic, penetration testing, or compliance to HIPAA/GDPR in the prototype. These safeguards are important before dealing with real patient data and are particularly vital if there are cognitively impaired people and brain tumors.

IX. CONCLUSION AND FUTURE WORK

This study proposed an integrated CNN + IoT framework for brain tumor diagnosis and continuous monitoring. The finalized CNN configuration (AdamW, label smoothing, mixed precision, TTA) achieved 96.10% test accuracy and



97.27% validation accuracy on BRISC2025, outperforming matched ResNet50 and MobileNetV2 baselines (macro-F1 91.21% and 94.13%, respectively) -- the first reported integration of this dataset with an IoT monitoring architecture.

Future work: (i) Explainable AI -- apply Grad-CAM to the final convolutional layer to visualize class-discriminative regions for correctly and incorrectly classified cases per tumor category, and compare the resulting heatmaps against radiologists' own interpretations to assess clinical relevance; (ii) external validation of the trained model on an independent public dataset (e.g., BraTS, TCIA, or Figshare) to test robustness and generalizability beyond BRISC2025; (iii) statistical significance testing (McNemar's test, appropriate for paired classification outcomes on the same test set) to determine whether the proposed CNN significantly outperforms ResNet50/MobileNetV2, with reported p-values; (iv) reconciling the gap with externally published BRISC benchmarks (Table 9) and completing the matched-benchmark protocol details (overall accuracy, support, exact split); (v) developing and evaluating the activity-recognition component; (vi) independent clinical validation; (vii) real IoT hardware deployment and testing; and (viii) a named security/compliance standard. None of (i)-(iii) has been implemented in the present study -- no Grad-CAM visualizations, external-validation results, or McNemar test statistics are reported here.

Authors' Contributions, Funding, and Ethics

Amal Yaser Fayez: conceptualization, methodology, data curation, software, formal analysis. Hazem Mokhtar El-Bakry: supervision, critical review, project administration. Mohammed A. Elhawary: author-level radiological review of selected diagnostic outputs (see Conflict of Interest, Section 6). No external funding was received. MRI data are from the fully anonymized, publicly available BRISC dataset; no direct human-subject interaction occurred. The IoT system is a software prototype using entirely fictitious, illustrative patient profiles and simulated sensor readings (Table 3); no real patient data were collected or used. Any real-world deployment with real patients or hardware would require informed consent and institutional ethics approval.

Data and Code Availability

The MRI data used in this study are from the publicly available BRISC2025 dataset. The IoT patient profiles and sensor streams the prototype uses are fictional and created for demonstration and internal testing purposes only. The source code, the weights of trained models, scripts for preprocessing and the implementation of the dashboard should be placed in a public repository or be made available from the corresponding author on reasonable request, taking into account institutional policy and journal requirements.

Reproducibility Statement

In order to facilitate the reproducibility, the manuscript includes detailed descriptions of the data partitioning, pre-processing pipeline, CNN architecture, optimization settings, augmentation policy, and validation/test metrics. Package versions were pinned in a requirements file, and the same preprocessing assumptions should be applied to the custom CNN and transfer-learning baselines. For a fully reproducible release, the final repository should include the training script, random seeds, model-summary output, evaluation notebook, and the exact command used to generate each reported table.

REFERENCES

1. F. Ullah *et al.*, "Brain tumor segmentation from MRI images using handcrafted convolutional neural network," *Diagnostics*, vol. 13, no. 16, p. 2650, 2023, doi: 10.3390/diagnostics13162650.
2. L. Saipogu *et al.*, "Identification of brain tumors in MR images using UNet CNN model," in *Proc. IEEE*, 2023, pp. 941–945, doi: 10.1109/ICECA58529.2023.10395702.
3. S. Sangui *et al.*, "3D MRI segmentation using U-Net architecture for brain tumor detection," *Procedia Comput. Sci.*, vol. 218, pp. 542–553, 2023, doi: 10.1016/j.procs.2023.01.036.



4. F. Adib *et al.*, "3D densely connected CNN with multi-scale receptive fields for brain tumor segmentation," *Discover AI*, vol. 5, no. 1, p. 213, 2025, doi: 10.1007/s44163-025-00279-9.
5. M. Abrar *et al.*, "Enhancing brain tumor segmentation using attention-based U-Net," *Sci. Rep.*, vol. 15, no. 1, p. 36603, 2025, doi: 10.1038/s41598-025-20329-7.
6. A. H. Nizamani *et al.*, "Advanced brain tumor segmentation using feature fusion with deep U-Net," *J. King Saud Univ. Comput. Inf. Sci.*, vol. 35, no. 9, p. 101793, 2023, doi: 10.1016/j.jksuci.2023.101793.
7. Q. U. A. Ishfaq *et al.*, "Automatic smart brain tumor classification and prediction using deep learning," *Sci. Rep.*, vol. 15, no. 1, p. 14876, 2025, doi: 10.1038/s41598-025-95803-3.
8. M. Taleb and S. Alibabaei, "Hybrid 2D/3D CNN and radiomics for brain tumor classification," *BMC Med. Imaging*, 2026, doi: 10.1186/s12880-025-02141-x.
9. V. Anand *et al.*, "Multi-class classification of brain tumors using optimized CNN and transfer learning," *Sci. Rep.*, 2026, doi: 10.1038/s41598-025-34806-6.
10. F. Ayudi and F. Pereira, "Multiclass classification of brain tumors in MRI based on deep learning," *Diagnosa*, vol. 2, no. 1, pp. 38–46, 2026, doi: 10.63935/rctfm087.
11. L. Ke *et al.*, "Brain tumor classification using multi-scale channel attention CNN with SVM," *Sci. Rep.*, 2026, doi: 10.1038/s41598-026-36164-3.
12. P. Valsalan *et al.*, "IoT-based health monitoring system," *J. Crit. Rev.*, vol. 7, no. 4, pp. 739–743, 2020, doi: 10.31838/jcr.07.04.137.
13. S. Subha *et al.*, "A remote health care monitoring system using IoMT," in *Proc. IEEE*, 2023, pp. 1–6, doi: 10.1109/ICIPTM57143.2023.10118103.
14. A. Fateh *et al.*, "BRISC: annotated dataset for brain tumor segmentation and classification," *Sci. Data*, 2026, doi: 10.1038/s41597-026-06753-y.
15. F. Ahmed, "Bridging topology and deep representation learning: A TDA-ViT fusion model for four-class brain tumor classification," *arXiv preprint arXiv:2606.00927*, 2026.
16. F. Ahmed, "Enhancing brain tumor classification using Vision Transformers with colormap-based feature representation on BRISC2025 dataset," *arXiv preprint arXiv:2603.21234*, 2026, doi: 10.48550/arXiv.2603.21234.
17. P. Veneri, "Interpretable four-class brain tumor MRI classification using a fine-tuned ResNet50," *PeerJ Comput. Sci.*, vol. 12, p. e3643, 2026, doi: 10.34740/kaggle/dsv/2645886.
18. O. O. Oladimeji and A. O. J. Ibitoye, "Brain tumor classification using ResNet50-convolutional block attention module," *Appl. Comput. Inform.*, vol. 22, no. 3–4, pp. 249–261, 2026, doi: 10.1108/ACI-09-2023-0022.
19. M. Arabboev, "Brain tumor classification using transfer learning with MobileNetV2," *Techscience.uz – Topical Issues Tech. Sci.*, vol. 3, no. 5, pp. 51–63, 2025, doi: 10.47390/ts-v3i5y2025N8.
20. M. J. Adamu, H. B. Kawuwa, L. Qiang, C. O. Nyatega, A. Younis, M. Fahad, and S. S. Dauya, "Efficient and accurate brain tumor classification using hybrid MobileNetV2–support vector machine for magnetic resonance imaging diagnostics in neoplasms," *Brain Sci.*, vol. 14, no. 12, p. 1178, 2024, doi: 10.3390/brainsci14121178.

