

MEDACT: FUNGSI AKTIVASI HIBRID BARU UNTUK MODEL RESNET PADA KLASIFIKASI CITRA MEDIS

Muhammad Syarif Baital¹, Jasman², Nur Inda³

¹STMIK Catur Sakti Kendari, ²STIKOM 22 Januari Kendari, ³ITBM Polewali Mandar
e-mail: ¹syarif.baital@gmail.com, ²jasman@gmail.com, ³nurinda@itbmpolman.ac.id

Penelitian ini menyoroti pentingnya pemilihan fungsi aktivasi pada model deep learning untuk klasifikasi tumor otak berbasis citra MRI. Sebanyak sebelas fungsi aktivasi, termasuk ReLU, Leaky ReLU, ELU, Swish, Mish, PReLU, GELU, SELU, HardSwish, MedAct Fixed, dan MedAct Learnable, dievaluasi pada arsitektur ResNet. Hasil menunjukkan hampir semua fungsi mencapai akurasi pengujian $\geq 99\%$ dengan perbedaan relatif kecil. ReLU, Swish, Mish, dan MedAct Learnable menempati posisi terbaik dengan hanya dua kesalahan klasifikasi, sedangkan PReLU dan MedAct Fixed menunjukkan kelemahan dengan tujuh kesalahan. Temuan penting adalah bahwa fungsi aktivasi baru (MedAct Learnable), mampu menyamai performa fungsi aktivasi modern terbaik dan menunjukkan peningkatan signifikan dibandingkan MedAct Fixed. Hal ini menegaskan bahwa sifat adaptif parameter a memberikan kontribusi positif terhadap generalisasi model.

Kata Kunci—tumor otak, MRI, ResNet, fungsi aktivasi, deep learning, MedAct.

I. PENDAHULUAN

Tumor otak merupakan salah satu penyakit yang memiliki dampak besar terhadap kualitas hidup pasien, dengan tingkat morbiditas dan mortalitas yang cukup tinggi. Keberhasilan terapi sangat dipengaruhi oleh kemampuan deteksi dini dan klasifikasi jenis tumor secara akurat. Oleh karena itu, metode diagnostik yang lebih presisi menjadi kebutuhan mendesak dalam praktik klinis. Pencitraan resonansi magnetik (MRI) saat ini merupakan modalitas utama yang banyak digunakan dalam proses diagnosis tumor otak [1], [2] karena mampu menampilkan struktur jaringan otak dengan resolusi tinggi dan tanpa radiasi ionisasi.

Meskipun MRI memberikan informasi radiologis yang detail, interpretasi citra secara manual oleh radiolog sering kali menghadapi tantangan. Faktor subjektivitas, kelelahan, dan keterbatasan pengalaman dapat menyebabkan inkonsistensi dalam penentuan jenis tumor. Kesalahan diagnosis berpotensi mengakibatkan pemilihan terapi yang tidak tepat, sehingga diperlukan metode berbasis komputasi yang mampu mendukung proses

klasifikasi secara objektif, konsisten, dan dapat direplikasi [3]. Dalam konteks ini, pendekatan berbasis kecerdasan buatan, khususnya *deep learning*, menjadi solusi yang menjanjikan.

Deep learning telah terbukti efektif dalam berbagai aplikasi pengolahan citra medis, termasuk deteksi dan klasifikasi tumor otak. *Convolutional Neural Networks* (CNN) merupakan salah satu algoritma paling populer yang memiliki kemampuan mengekstraksi fitur kompleks secara otomatis dari citra tanpa memerlukan rekayasa fitur manual. Di antara berbagai arsitektur CNN, ResNet menjadi salah satu yang paling banyak digunakan karena keberhasilannya dalam mengatasi masalah degradasi performa pada jaringan dalam melalui mekanisme residual learning [4], [5], [6].

Meskipun ResNet mampu memberikan hasil yang andal, performa model *deep learning* secara keseluruhan tidak hanya ditentukan oleh arsitektur jaringan, tetapi juga sangat dipengaruhi oleh pemilihan fungsi aktivasi. Fungsi aktivasi berperan penting dalam memperkenalkan non-linearitas pada jaringan, sehingga model dapat mempelajari representasi kompleks dari data. ReLU [1], [7], [8] telah lama menjadi fungsi aktivasi standar karena kesederhanaan dan efisiensinya, namun memiliki keterbatasan seperti masalah *dying ReLU* yang dapat menghambat pembelajaran *neuron* tertentu.

Untuk mengatasi keterbatasan ReLU, berbagai fungsi aktivasi baru telah diperkenalkan. Leaky ReLU [8], [9] dan PReLU [4], [5], [6] dikembangkan untuk mengurangi masalah *neuron* mati. Selanjutnya, fungsi seperti ELU [10], [11], Swish [12], [13], Mish [12], [13], GELU [14], SELU [1], [7], dan HardSwish [15] muncul sebagai alternatif yang menawarkan berbagai keunggulan, mulai dari stabilitas gradien, kehalusan fungsi, hingga peningkatan performa generalisasi. Variasi ini menunjukkan bahwa pemilihan fungsi aktivasi bukan sekadar aspek teknis, melainkan komponen strategis dalam merancang model *deep learning* yang optimal.

Berdasarkan hal tersebut, penelitian ini difokuskan pada evaluasi perbandingan fungsi aktivasi klasik dan modern pada arsitektur ResNet untuk klasifikasi tumor otak berbasis MRI. Selain itu, penelitian ini juga memperkenalkan fungsi aktivasi baru bernama MedAct,

yang diuji dalam dua varian, yaitu *Fixed* dan *Learnable*. Melalui eksperimen ini, diharapkan dapat diketahui apakah MedAct mampu meningkatkan generalisasi model serta mengurangi tingkat kesalahan klasifikasi dibandingkan fungsi aktivasi konvensional. Dengan demikian, penelitian ini tidak hanya bertujuan mengevaluasi performa teknis, tetapi juga memberikan kontribusi praktis dalam pengembangan sistem diagnosis berbantuan komputer di bidang medis.

II. METODE PENELITIAN

A. Dataset

Dataset yang digunakan dalam penelitian ini terdiri dari 6.012 citra MRI tumor otak, yang telah banyak digunakan dalam studi klasifikasi berbasis *deep learning* [3], [16], [17]. Citra-citra tersebut mewakili tiga kelas utama tumor otak, yaitu glioma, meningioma, dan pituitary tumor. Ketiga jenis tumor ini dipilih karena memiliki karakteristik radiologis yang berbeda namun sering kali menunjukkan tumpang tindih visual pada MRI, sehingga menjadi tantangan menarik dalam pengembangan model klasifikasi otomatis.

Untuk memastikan proses pelatihan dan evaluasi berjalan secara adil, dataset dibagi ke dalam tiga subset, yaitu *training* (80%), *validasi* (10%), dan *pengujian* (10%). Pembagian ini dilakukan secara stratifikasi agar distribusi tiap kelas tetap proporsional pada setiap subset. Data *training* digunakan untuk melatih model ResNet dengan berbagai fungsi aktivasi, data *validasi* berfungsi memantau performa model selama proses *training*, sedangkan data *pengujian* digunakan sebagai tolok ukur akhir untuk menilai kemampuan generalisasi model terhadap data yang belum pernah dilihat sebelumnya.

B. Arsitektur Model

Model dasar yang digunakan dalam penelitian ini adalah ResNet-18, salah satu varian dari *Residual Network* yang diperkenalkan oleh Microsoft Research [1]. ResNet dikenal luas karena kemampuannya mengatasi permasalahan *vanishing gradient* dan degradasi performa pada jaringan yang sangat dalam melalui penerapan *residual learning*. Pada arsitektur ini, setiap blok *residual* memiliki koneksi pintas (*skip connection*) yang memungkinkan gradien mengalir lebih lancar selama proses pelatihan, sehingga model dapat mencapai konvergensi lebih cepat dan stabil. Dengan kedalaman 18 lapisan, ResNet-18 dipilih karena memiliki kompleksitas yang seimbang antara kapasitas representasi dan efisiensi komputasi, menjadikannya cocok untuk eksperimen klasifikasi citra medis.

Dalam penelitian ini, ResNet-18 dijadikan kerangka dasar, dan modifikasi utama dilakukan pada bagian fungsi aktivasi yang terintegrasi di dalam setiap lapisan non-linear. Sebelas varian fungsi aktivasi diuji secara bergantian untuk menggantikan fungsi aktivasi standar, sehingga dapat dibandingkan pengaruhnya terhadap kinerja model. Pendekatan ini memungkinkan evaluasi yang adil karena semua faktor lain, termasuk jumlah

lapisan, parameter, serta konfigurasi pelatihan, tetap dijaga konstan. Dengan demikian, perbedaan performa yang muncul dapat lebih jelas diatribusikan pada perbedaan fungsi aktivasi yang digunakan.

C. Fungsi Aktivasi yang Dievaluasi

1. Rectified Linear Unit (ReLU)

ReLU merupakan fungsi aktivasi paling populer dan menjadi standar dalam banyak arsitektur *deep learning*. Fungsi ini sederhana, didefinisikan sebagai $f(x) = \max(0, x)$, sehingga hanya mempertahankan nilai positif dan mematikan nilai negatif. Keunggulan utama ReLU adalah efisiensi komputasi dan kemampuannya mengurangi masalah *vanishing gradient*. Namun, ReLU memiliki kelemahan berupa fenomena *dying ReLU*, di mana *neuron* dapat berhenti belajar jika *output* selalu negatif [1].

2. Leaky ReLU

Leaky ReLU dikembangkan sebagai perbaikan dari ReLU untuk mengatasi masalah *neuron mati*. Fungsi ini memberikan gradien kecil pada nilai negatif dengan mendefinisikan $f(x) = x$ jika $x > 0$, dan $f(x) = \alpha x$ jika $x \leq 0$ dengan α bernilai kecil (misalnya 0.01). Dengan cara ini, *neuron* tetap dapat belajar meskipun menerima input negatif. Leaky ReLU sering digunakan pada kasus di mana stabilitas pembelajaran lebih penting daripada kesederhanaan fungsi [9].

3. Exponential Linear Unit (ELU)

ELU memperkenalkan pendekatan dengan membuat bagian negatif dari fungsi lebih halus menggunakan eksponensial. Definisinya adalah $f(x) = x$ jika $x > 0$, dan $f(x) = \alpha(\exp(x) - 1)$ jika $x \leq 0$. ELU mampu menghasilkan *output* yang lebih dekat dengan nol, sehingga mempercepat pembelajaran dan meningkatkan generalisasi. Fungsi ini terbukti lebih stabil dibanding ReLU maupun Leaky ReLU, meskipun lebih mahal secara komputasi [9].

4. Swish

Swish adalah fungsi aktivasi yang ditemukan melalui pencarian otomatis. Definisinya adalah $f(x) = x * \text{sigmoid}(x)$. Fungsi ini bersifat halus dan non-monotonik, yang memungkinkan representasi lebih kaya dibandingkan fungsi monoton seperti ReLU. Swish terbukti meningkatkan akurasi dalam berbagai model *deep learning* besar, terutama pada tugas klasifikasi citra. Kelebihannya terletak pada kemampuan menjaga gradien tetap stabil tanpa mengorbankan kompleksitas representasi [1].

5. Mish

Mish merupakan fungsi aktivasi non-monotonik yang didefinisikan sebagai $f(x) = x * \tanh(\text{softplus}(x))$. Fungsi ini mirip dengan Swish tetapi dengan kurva lebih halus pada bagian negatif, sehingga mendorong distribusi gradien yang lebih baik. Mish dilaporkan memberikan peningkatan performa signifikan pada berbagai *benchmark* visi komputer dibandingkan ReLU, Swish, maupun ELU. Sifat *self-regularized* pada Mish juga membuatnya lebih

tangguh terhadap *overfitting* [7].

6. Parametric ReLU (PReLU)

PReLU merupakan varian dari Leaky ReLU dengan parameter α yang dapat dipelajari selama *training*. Fungsi ini memungkinkan model untuk secara adaptif menentukan kemiringan bagian negatif sesuai kebutuhan data. Dengan fleksibilitas tersebut, PReLU berhasil meningkatkan performa pada dataset skala besar seperti ImageNet, bahkan melampaui ReLU standar. Namun, penambahan parameter juga dapat meningkatkan risiko *overfitting* jika *dataset* relatif kecil [1], [4].

7. Gaussian Error Linear Unit (GELU)

GELU mengombinasikan keunggulan ReLU dan sigmoid dengan fungsi probabilistik. Definisinya adalah $f(x) = x * \Phi(x)$, di mana $\Phi(x)$ adalah fungsi distribusi normal kumulatif. GELU bersifat halus dan memungkinkan *input* negatif tetap berkontribusi secara proporsional. Fungsi ini banyak digunakan pada model NLP modern seperti BERT dan Transformer, serta terbukti memberikan stabilitas dan akurasi tinggi dalam pelatihan jaringan [8].

8. Scaled Exponential Linear Unit (SELU)

SELU dirancang untuk mendukung *self-normalizing neural networks*. Fungsi ini mirip dengan ELU namun dengan skala tambahan yang membuat aktivasi secara otomatis menormalkan *output* menuju distribusi nol rata-rata dan varians satu. Dengan sifat ini, SELU dapat menjaga stabilitas distribusi aktivasi sepanjang lapisan jaringan, sehingga mempercepat konvergensi. Namun, penerapannya memerlukan kondisi khusus seperti penggunaan *AlphaDropout* [11].

9. HardSwish

HardSwish adalah varian dari Swish yang lebih efisien secara komputasi. Fungsi ini didefinisikan sebagai $f(x) = x * \text{ReLU6}(x + 3) / 6$, yang mendekati Swish tetapi menggunakan operasi sederhana. HardSwish diperkenalkan dalam MobileNetV3 untuk perangkat seluler, di mana efisiensi menjadi prioritas utama. Dengan sifatnya yang ringan, HardSwish menjadi pilihan tepat untuk aplikasi dengan keterbatasan sumber daya tanpa mengorbankan akurasi secara signifikan [14], [15], [16].

10. MedAct

Fungsi aktivasi MedAct yang diusulkan dalam penelitian ini didefinisikan sebagai kombinasi antara komponen *linier* dan *non-linier*, yaitu $f(x) = \text{ReLU}(x) + \alpha \cdot (\text{Swish}(x) \cdot \tanh(\text{Mish}(x)))$ dengan parameter pengatur α . Fungsi ini dirancang dalam dua mode berbeda, yaitu *Fixed α* , di mana nilai α ditentukan sejak awal (misalnya 0.5) dan tidak berubah selama proses pelatihan, serta *Learnable α* , di mana α diperlakukan sebagai parameter *trainable* yang diperbarui melalui *backpropagation* bersama bobot jaringan. Mode *Fixed α* memberikan kestabilan eksperimen karena kontribusi komponen halus tetap

konstan, namun kurang fleksibel dalam menyesuaikan dengan karakteristik layer atau data. Sebaliknya, mode *Learnable α* memungkinkan tiap layer menyesuaikan besaran kontribusi komponen *smooth* secara adaptif, sehingga lebih fleksibel meskipun berpotensi meningkatkan risiko *overfitting* pada *dataset* kecil. Dalam penelitian ini, kedua mode diuji secara paralel untuk memberikan gambaran menyeluruh mengenai stabilitas dan fleksibilitas pendekatan MedAct.

D. Prosedur Eksperimen

Proses pelatihan model dilakukan dengan konfigurasi yang seragam untuk seluruh fungsi aktivasi agar hasil perbandingan bersifat adil. Setiap model dilatih selama 10 *epoch* dengan ukuran *batch* tetap, sehingga setiap iterasi dapat memanfaatkan jumlah sampel yang sama. *Optimizer* yang digunakan adalah Adam dengan *learning rate* awal 0.001, yang dipilih karena kestabilannya dalam menangani gradien adaptif dan kemampuannya mencapai konvergensi lebih cepat pada data kompleks seperti citra MRI.

Tabel 1.
Konfigurasi model arsitektur ResNet

Hyper parameter	Baseline
Activation Function	ReLU, Leaky ReLU, ELU, Swish, Mish, PReLU, GELU, SELU, HardSwish, MedAct
Epoch	10
Weight	ImageNet
BatchSize	32
Optimizer	Adam
Augmentation	Yes (Resize, RandomHorizontalFlip, RandomRotation, ColorJitter)
Drop Out	Yes (0.5)
Data Split	80:20
Learning Rate	Yes (0.001)
Transfer Learning	Yes

Evaluasi model dilakukan secara berlapis dengan menggunakan data validasi selama proses *training*, serta data pengujian sebagai tolok ukur akhir untuk menilai kemampuan generalisasi. Berbagai metrik performa diukur, meliputi akurasi, *precision*, *recall*, *F1-score*, dan *area under the curve* (AUC). Selain itu, jumlah kesalahan klasifikasi dari masing-masing kelas juga dianalisis melalui *confusion matrix* untuk memberikan gambaran lebih rinci mengenai distribusi *error* antar kelas. Dengan pendekatan ini, evaluasi tidak hanya menyoroti kinerja model secara keseluruhan, tetapi juga sensitivitasnya terhadap setiap kategori tumor otak.

III. HASIL DAN PEMBAHASAN

A. Hasil Penelitian

1. Rectified Linear Unit (ReLU)

Training model yang dilakukan untuk menguji model ResNet-18 dengan fungsi aktivasi ReLU.

Tabel 2.
Classification Report training ResNet-18 dengan aktivasi ReLU

Precision	Recall	F1-Score	Support
-----------	--------	----------	---------

0.99	1	0.99	194	Glioma
1	0.99	1	212	Menin
1	1	1	201	Tumor
		1	607	Accuracy
1	1	1	607	Macro Avg
1	1	1	607	Weight Avg
		1	607	Macro AUC

Hasil eksperimen dengan fungsi aktivasi ReLU menunjukkan performa yang sangat tinggi baik pada tahap pelatihan, validasi, maupun pengujian. Selama 10 *epoch*, *training loss* turun drastis dari 0.19 menjadi 0.01, sementara *validation loss* tetap konsisten rendah di kisaran 0.02–0.04 tanpa indikasi *overfitting*. Akurasi pada data validasi mencapai lebih dari 98% sejak *epoch* kedua dan stabil di atas 99% hingga akhir pelatihan, dengan *best validation accuracy* $\approx 99.7\%$. Kurva konvergensi memperlihatkan bahwa baik akurasi maupun *loss* pada *training* dan validasi mengikuti pola yang selaras, menandakan generalisasi model yang sangat baik. Waktu pelatihan relatif efisien, sekitar 17 menit untuk 10 *epoch*, sehingga ReLU terbukti efektif sekaligus praktis untuk kasus klasifikasi tumor otak berbasis MRI.

Evaluasi pada data uji (607 sampel) memperkuat temuan tersebut dengan capaian akurasi keseluruhan 100% dan nilai rata-rata *precision*, *recall*, *F1-score*, serta AUC makro sama dengan 1.0. Hampir semua kelas terklasifikasi sempurna, kecuali terdapat dua kasus *false negative* (FN) pada kelas meningioma yang salah diprediksi sebagai glioma. Hal ini mengindikasikan bahwa citra meningioma memiliki beberapa kesamaan fitur radiologis dengan glioma, sehingga berpotensi menimbulkan ambiguitas dalam klasifikasi. Secara umum, ReLU dapat dijadikan *baseline* yang sangat kuat karena sudah mencapai performa mendekati sempurna dengan kemampuan generalisasi baik. Implikasinya, fungsi aktivasi lain seperti Mish, Swish, atau MedAct hanya dapat dianggap signifikan apabila mampu menunjukkan keunggulan tambahan, misalnya dalam hal stabilitas konvergensi, *robustness* terhadap *noise*, atau konsistensi pada *dataset* lintas institusi yang lebih menantang.

2. Leaky ReLU

Training model yang dilakukan untuk menguji model ResNet-18 dengan fungsi aktivasi Leaky ReLU.

Tabel 3.
Classification Report training ResNet-18 dengan aktivasi Leaky ReLU

Precision	Recall	F1-Score	Support	
1	0.99	0.99	194	Glioma
0.99	1	0.99	212	Menin
1	1	1	201	Tumor
		0.99	607	Accuracy
0.99	0.99	0.99	607	Macro Avg
0.99	0.99	0.99	607	Weight Avg
		1	607	Macro AUC

Eksperimen dengan Leaky ReLU menunjukkan performa yang hampir setara dengan ReLU, meskipun terdapat beberapa perbedaan penting dalam dinamika

pelatihan dan hasil akhir. Selama *training*, *loss* menurun stabil dari 0.2043 hingga 0.0123, dan akurasi meningkat dari 92,3% menjadi 99,6%. Namun, pada data validasi terlihat adanya fluktuasi yang lebih jelas dibanding ReLU, terutama pada *epoch* 2 hingga 6, dengan *validation loss* yang sempat bergelombang dan akurasi validasi yang menurun sebelum kembali stabil. Hal ini mengindikasikan bahwa meskipun model mampu mencapai *best validation accuracy* 99,7%, proses validasi dengan Leaky ReLU cenderung lebih *noisy*. Dari sisi efisiensi, waktu pelatihan yang dibutuhkan hanya sekitar 7,5 menit untuk 10 *epoch*, jauh lebih singkat dibandingkan ReLU yang memerlukan hampir 17 menit, sehingga fungsi ini menawarkan keuntungan dari segi kecepatan komputasi.

Evaluasi pada data uji (607 sampel) memperlihatkan akurasi keseluruhan sebesar 99% dengan nilai rata-rata *precision*, *recall*, dan *F1-score* mencapai 0,99 serta *macro AUC* 1.0. Kesalahan klasifikasi berjumlah empat kasus, dua glioma diprediksi sebagai meningioma, satu meningioma sebagai tumor, dan satu tumor sebagai meningioma. Dibandingkan dengan ReLU yang hanya menghasilkan dua kesalahan, Leaky ReLU menunjukkan kinerja yang sedikit lebih rendah, meskipun tetap berada pada level yang sangat tinggi. Perbedaan ini dapat dikaitkan dengan sifat Leaky ReLU yang memberikan gradien kecil pada input negatif, yang secara teoretis bermanfaat untuk mengatasi masalah *dying ReLU*, tetapi pada *dataset* medis yang besar dan relatif bersih seperti ini, keunggulan tersebut tidak begitu terlihat. Dengan demikian, meskipun Leaky ReLU dapat dijadikan *baseline* tambahan dalam studi komparatif, ReLU tetap lebih unggul dari sisi konsistensi generalisasi. Namun, catatan penting untuk penelitian lanjutan adalah efisiensi waktu yang dimiliki Leaky ReLU, yang dapat relevan untuk eksperimen berskala besar atau pada lingkungan komputasi dengan sumber daya terbatas.

3. Exponential Linear Unit (ELU)

Training model yang dilakukan untuk menguji model ResNet-18 dengan fungsi aktivasi ELU.

Tabel 4.
Classification Report training ResNet-18 dengan aktivasi ELU

Precision	Recall	F1-Score	Support	
1	1	1	194	Glioma
0.99	1	1	212	Menin
1	0.99	0.99	201	Tumor
		1	607	Accuracy
1	1	1	607	Macro Avg
1	1	1	607	Weight Avg
		1	607	Macro AUC

Eksperimen dengan ELU memperlihatkan kinerja yang sangat kompetitif, bahkan sedikit lebih unggul dibandingkan ReLU dan Leaky ReLU. Selama pelatihan, *training loss* menurun tajam dari 0.2063 menjadi 0.0173 dengan peningkatan akurasi dari 91,7% hingga 99,4%. Pada data validasi, model menunjukkan konsistensi yang baik dengan *validation loss* yang stabil di rentang 0.056

hingga 0.009, meskipun terdapat sedikit lonjakan pada *epoch* ketujuh yang segera terkoreksi. Akurasi validasi secara konsisten berada di atas 99% sejak *epoch* kedua, bahkan mencapai puncak 99,8%. Pola konvergensi ini menegaskan bahwa ELU mampu menjaga stabilitas proses pembelajaran, berbeda dengan Leaky ReLU yang menunjukkan fluktuasi lebih besar. Dari sisi efisiensi, durasi pelatihan yang hanya sekitar 7 menit 27 detik membuat ELU sebanding dengan Leaky ReLU dan jauh lebih efisien dibandingkan ReLU.

Evaluasi pada data uji menunjukkan hasil yang hampir sempurna, dengan akurasi keseluruhan mencapai 100% dan nilai rata-rata *precision*, *recall*, *F1-Score*, serta AUC masing-masing berada pada angka 1,0. Meski demikian, masih terdapat tiga kesalahan klasifikasi, yaitu satu meningioma yang diprediksi sebagai tumor serta dua tumor yang diprediksi sebagai meningioma. Sementara itu, kelas glioma berhasil diklasifikasikan tanpa kesalahan sama sekali. Hasil ini menempatkan ELU sebagai fungsi aktivasi dengan performa terbaik sejauh ini, karena mampu menyatukan akurasi tinggi, stabilitas validasi yang halus, dan efisiensi waktu pelatihan. Walaupun masih ada sedikit kelemahan pada pemisahan antara meningioma dan tumor, kinerja ELU secara keseluruhan menunjukkan potensi besar untuk diaplikasikan pada klasifikasi citra medis, terutama dalam konteks diagnosis tumor otak berbasis MRI.

4. Swish

Training model yang dilakukan untuk menguji model ResNet-18 dengan fungsi aktivasi Swish.

Tabel 5.
Classification Report training ResNet-18 dengan aktivasi Swish

Precision	Recall	F1-Score	Support	
0.99	1	1	194	Glioma
1	1	1	212	Menin
1	1	1	201	Tumor
		1	607	Accuracy
1	1	1	607	Macro Avg
1	1	1	607	Weight Avg
		1	607	Macro AUC

Eksperimen dengan Swish memperlihatkan bahwa fungsi aktivasi ini mampu mencapai kinerja yang sangat tinggi, meskipun tidak sepenuhnya melampaui ReLU maupun ELU. Selama pelatihan, *training loss* menurun stabil dari 0.1923 menjadi 0.0199, dengan akurasi yang meningkat hingga 99,5%. Namun, pada data validasi terlihat pola yang lebih fluktuatif dibandingkan fungsi aktivasi lain. *Validation loss* mengalami lonjakan pada *epoch* kedua serta pada rentang *epoch* ketujuh hingga kedelapan, yang diikuti dengan penurunan akurasi validasi hingga sekitar 96%. Walaupun demikian, secara keseluruhan akurasi validasi tetap tinggi, dengan puncak di angka 99,5%. Hal ini menunjukkan bahwa Swish menjaga performa kompetitif, tetapi stabilitasnya sedikit lebih rentan dibandingkan ELU yang cenderung lebih halus. Dari sisi efisiensi, waktu pelatihan sekitar 7 menit 26 detik

setara dengan ELU dan Leaky ReLU, sehingga tidak menimbulkan beban komputasi tambahan.

Evaluasi pada data uji memperlihatkan hasil yang sangat baik dengan akurasi keseluruhan mencapai 100%, serta nilai *precision*, *recall*, *F1-Score*, dan AUC sempurna pada rata-rata makro. Meski demikian, terdapat dua kesalahan klasifikasi, yaitu satu kasus glioma yang diprediksi sebagai meningioma dan satu kasus meningioma yang diprediksi sebagai glioma, sementara kelas tumor berhasil dikenali tanpa kesalahan. Distribusi *error* ini berbeda dari ReLU, meskipun jumlah total kesalahan sama, sehingga memperlihatkan bahwa tantangan utama terletak pada perbedaan karakteristik citra antara glioma dan meningioma yang memang kerap sulit dibedakan secara radiologis. Dengan demikian, Swish dapat dinilai kompetitif dan sejajar dengan ReLU, tetapi tidak memberikan peningkatan yang substansial pada dataset ini. Temuan ini menegaskan bahwa faktor *dataset* dan kemiripan fitur antar kelas mungkin lebih berpengaruh terhadap kesalahan klasifikasi daripada pemilihan fungsi aktivasi semata.

5. Mish

Training model yang dilakukan untuk menguji model ResNet-18 dengan fungsi aktivasi Mish.

Tabel 6.
Classification Report training ResNet-18 dengan aktivasi Mish

Precision	Recall	F1-Score	Support	
1	1	1	194	Glioma
0.99	0.99	0.99	212	Menin
0.99	0.99	0.99	201	Tumor
		0.99	607	Accuracy
0.99	0.99	0.99	607	Macro Avg
0.99	0.99	0.99	607	Weight Avg
		0.99	607	Macro AUC

Eksperimen dengan Mish menghasilkan performa yang sangat kompetitif, bahkan setara dengan fungsi aktivasi terbaik lain seperti ELU. Selama proses pelatihan, *training loss* menurun stabil dari 0.1969 menjadi 0.0122, sementara *training accuracy* meningkat secara konsisten hingga mencapai 99,6%. Pada data validasi, terlihat adanya fluktuasi di awal terutama pada *epoch* kedua dan kelima yang sempat menurunkan akurasi hingga 95,5%. Namun, setelah fase adaptasi awal tersebut, performa validasi meningkat tajam dan stabil, dengan *best validation accuracy* mencapai 99,83%, lebih tinggi dibanding ReLU, Leaky ReLU, dan Swish. Hal ini menunjukkan bahwa Mish membutuhkan waktu lebih lama untuk menyesuaikan diri pada fase awal, tetapi mampu memberikan hasil yang sangat kuat ketika pelatihan dilanjutkan hingga *epoch* akhir. Dari sisi efisiensi, waktu pelatihan sekitar 7 menit 25 detik, sebanding dengan ELU dan Swish, sehingga Mish tetap praktis digunakan pada eksperimen berskala besar.

Evaluasi pada data uji memperkuat temuan ini, dengan akurasi keseluruhan sebesar 100% dan nilai makro untuk *precision*, *recall*, *F1-Score*, serta AUC yang semuanya sempurna. Meski demikian, masih terdapat dua kesalahan

klasifikasi, yaitu satu kasus meningioma yang diprediksi sebagai tumor, serta satu kasus tumor yang diprediksi sebagai meningioma. Sementara itu, kelas glioma berhasil teridentifikasi dengan akurasi sempurna tanpa kesalahan. Jumlah kesalahan ini sama dengan ReLU, tetapi lebih sedikit dibandingkan ELU maupun Leaky ReLU. Perbedaan distribusi *error* menunjukkan bahwa tantangan utama tetap terletak pada *overlap* fitur antara meningioma dan tumor, bukan pada glioma. Dengan demikian, Mish dapat dipandang sebagai alternatif modern yang menawarkan kombinasi antara akurasi akhir yang sangat tinggi, stabilitas pada fase pelatihan jangka panjang, dan robustnes terhadap kesalahan klasifikasi, meskipun membutuhkan waktu adaptasi yang lebih lama pada *epoch* awal.

6. Parametric ReLU (PReLU)

Training model yang dilakukan untuk menguji model ResNet-18 dengan fungsi aktivasi PReLU.

Tabel 7.
Classification Report training ResNet-18 dengan aktivasi PReLU

Precision	Recall	F1-Score	Support	
0.98	1	0.99	194	Glioma
0.99	0.98	0.98	212	Menin
0.99	0.99	0.99	201	Tumor
		0.99	607	Accuracy
0.99	0.99	0.99	607	Macro Avg
0.99	0.99	0.99	607	Weight Avg
		0.98	607	Macro AUC

Eksperimen dengan PReLU memperlihatkan hasil yang kompetitif namun tidak lebih unggul dibandingkan fungsi aktivasi lain seperti ELU atau Mish. Selama pelatihan, *training loss* menurun secara konsisten hingga sekitar 0.016 dengan akurasi mencapai 99,4%, menunjukkan proses optimisasi yang berjalan baik. Akan tetapi, pada data validasi terlihat adanya fluktuasi yang lebih nyata, terutama pada *epoch* keempat hingga keenam, ketika *validation loss* meningkat hingga 0.123 dan akurasi sempat turun drastis menjadi sekitar 95%. Meskipun performa kemudian pulih dengan *best validation accuracy* sebesar 99,67%, setara dengan ReLU dan Leaky ReLU, pola ini mengindikasikan bahwa PReLU sedikit kurang stabil dalam menjaga konsistensi selama pelatihan. Dari sisi efisiensi, waktu training sekitar 7 menit 24 detik menempatkan PReLU pada level yang sama dengan ELU, Mish, dan Swish.

Evaluasi pada data uji menunjukkan akurasi keseluruhan sebesar 99% dengan nilai *precision*, *recall*, dan *F1-score* rata-rata 0,99, serta *macro AUC* mendekati sempurna pada angka 0,98. Namun, jumlah kesalahan klasifikasi relatif lebih tinggi, yaitu tujuh kasus, dibandingkan fungsi aktivasi lain yang hanya menghasilkan dua hingga tiga kesalahan. Sebagian besar *error* terjadi pada kelas meningioma yang salah diprediksi sebagai glioma atau tumor, serta beberapa kasus tumor yang salah dipetakan ke meningioma. Hal ini menandakan bahwa meskipun PReLU menawarkan fleksibilitas melalui

parameter slope negatif yang dapat dipelajari, pada konteks klasifikasi tumor otak berbasis MRI fungsi ini justru cenderung kurang konsisten. Dengan demikian, PReLU dapat dipandang sebagai baseline tambahan yang relevan, namun untuk aplikasi medis yang menuntut stabilitas dan akurasi tinggi, performanya masih berada di bawah Mish dan ELU.

7. Gaussian Error Linear Unit (GELU)

Training model yang dilakukan untuk menguji model ResNet-18 dengan fungsi aktivasi GELU.

Tabel 8.
Classification Report training ResNet-18 dengan aktivasi GELU

Precision	Recall	F1-Score	Support	
1	1	1	194	Glioma
0.99	1	0.99	212	Menin
			201	Tumor
		1	607	Accuracy
1	1	1	607	Macro Avg
1	1	1	607	Weight Avg
		0.98	607	Macro AUC

Eksperimen dengan GELU memperlihatkan karakteristik pelatihan yang stabil dan halus, sesuai dengan sifat matematisnya yang mengombinasikan *smoothness* dari fungsi sigmoid dengan linearitas ReLU. Selama proses *training*, *loss* menurun konsisten hingga sekitar 0.011 dengan *training accuracy* mencapai 99,7%, sementara *validation loss* tetap rendah di kisaran 0.016–0.038. *Validation accuracy* juga relatif stabil antara 98,8–99,5%, meski tidak setinggi yang dicapai oleh Mish atau ELU. Dari segi efisiensi, waktu pelatihan sekitar 7 menit 28 detik menjadikan GELU setara dengan fungsi aktivasi lain yang lebih ringan secara komputasi, seperti Swish dan ELU. Pola kurva akurasi yang mulus tanpa fluktuasi besar menunjukkan bahwa GELU dapat memberikan *gradient flow* yang lebih terkontrol, sehingga cocok untuk pelatihan model yang lebih dalam dengan jumlah *epoch* lebih besar.

Evaluasi pada data uji menghasilkan akurasi keseluruhan sempurna, dengan nilai *precision*, *recall*, dan *F1-score* yang konsisten tinggi (0,99–1,0), serta *macro AUC* mencapai 0,9998. Meskipun demikian, masih ditemukan tiga kesalahan klasifikasi, yakni satu kasus meningioma yang diprediksi sebagai tumor, serta dua kasus tumor yang diprediksi sebagai meningioma. Jumlah *error* ini setara dengan ELU, lebih sedikit dibanding PReLU (7 kesalahan), tetapi sedikit lebih banyak dibanding ReLU, Mish, maupun Swish (masing-masing 2 kesalahan). Dengan demikian, GELU dapat dipandang sebagai fungsi aktivasi yang stabil dan dapat diandalkan, namun pada *dataset* tumor otak ini keunggulannya belum cukup signifikan dibanding Mish atau ELU, terutama mengingat kompleksitas komputasi GELU yang lebih tinggi.

8. Scaled Exponential Linear Unit (SELU)

Training model yang dilakukan untuk menguji model ResNet-18 dengan fungsi aktivasi SELU.

Tabel 9.

Classification Report training ResNet-18 dengan aktivasi SELU

Precision	Recall	F1-Score	Support	
1	0.99	0.99	194	Glioma
0.99	1	0.99	212	Menin
1	1	1	201	Tumor
		0.99	607	Accuracy
0.99	0.99	0.99	607	Macro Avg
0.99	0.99	0.99	607	Weight Avg
		0.99	607	Macro AUC

Eksperimen dengan SELU menunjukkan performa yang sangat kompetitif, terutama pada metrik validasi. *Best validation accuracy* mencapai 99,83%, setara dengan ELU dan Mish yang sebelumnya menempati posisi teratas. Selama pelatihan, *training loss* menurun konsisten hingga 0,012 dengan akurasi mencapai 99,6%. Namun, pada data validasi terlihat adanya fluktuasi cukup jelas di awal, khususnya pada *epoch* kedua dan ketiga ketika *validation loss* sempat naik hingga 0,0927 dan akurasi turun ke 96,9%. Meskipun demikian, model mampu pulih dan menutup *training* dengan *validation loss* yang sangat rendah (0,0063) dan akurasi mendekati sempurna. Dari sisi efisiensi, waktu pelatihan sekitar 7 menit 27 detik, sebanding dengan Mish, ELU, dan GELU. Hal ini mengindikasikan bahwa SELU tetap mampu menjaga efisiensi komputasi meski memiliki sifat *self-normalizing* yang lebih kompleks.

Evaluasi pada data uji memperlihatkan akurasi keseluruhan sebesar 99% dengan *precision*, *recall*, dan *F1-score* rata-rata 0,99, serta *macro AUC* 0,99. Meski performa secara umum sangat baik, terdapat empat kesalahan klasifikasi, dua pada glioma yang diprediksi sebagai meningioma, satu meningioma yang diprediksi sebagai tumor, dan satu tumor yang diprediksi sebagai meningioma. Jumlah kesalahan ini sedikit lebih tinggi dibanding ReLU, Mish, atau Swish (2 *error*), serta GELU (3 *error*). Dengan demikian, meskipun SELU menawarkan keunggulan teoretis berupa stabilisasi jaringan melalui mekanisme *self-normalizing*, pada dataset ini masih terlihat adanya instabilitas kecil dalam validasi dan sedikit peningkatan jumlah kesalahan klasifikasi. Hasil ini menegaskan bahwa SELU merupakan kandidat kuat, tetapi belum memberikan keunggulan yang konsisten dibanding fungsi aktivasi terbaik lainnya.

9. HardSwish

Training model yang dilakukan untuk menguji model ResNet-18 dengan fungsi aktivasi HardSwish.

Tabel 10.

Classification Report training ResNet-18 dengan aktivasi HardSwish

Precision	Recall	F1-Score	Support	
0.99	1	1	194	Glioma
0.99	1	0.99	212	Menin
1	0.99	0.99	201	Tumor
		1	607	Accuracy
1	1	1	607	Macro Avg
1	1	1	607	Weight Avg

1 607 Macro AUC

Eksperimen dengan HardSwish memperlihatkan performa yang sangat menjanjikan, dengan *best validation accuracy* mencapai 99,83%, setara dengan Mish dan SELU sebagai nilai tertinggi yang diperoleh sejauh ini. Selama pelatihan, *training loss* menurun tajam hingga 0,0038, disertai dengan akurasi hampir sempurna sebesar 99,9%. *Validation loss* juga menunjukkan tren menurun yang stabil, berakhir pada 0,0052 dengan akurasi validasi konsisten tinggi. Meski demikian, terdapat fluktuasi yang cukup mencolok pada fase awal pelatihan, khususnya di *epoch* kedua dan ketiga, ketika *validation loss* melonjak ke 0,1358 dan akurasi validasi turun ke 95,4%. Setelah melewati fase ini, model mampu pulih dengan cepat dan mencapai konvergensi yang sangat baik pada *epoch-epoch* berikutnya. Waktu pelatihan yang hanya sekitar tujuh setengah menit menegaskan bahwa HardSwish tetap efisien, setara dengan Mish, ELU, dan GELU.

Hasil evaluasi pada data uji menunjukkan kinerja yang nyaris sempurna, dengan akurasi keseluruhan 100%, nilai *precision*, *recall*, dan *F1-score* berada pada rentang 0,99–1,00, serta *macro AUC* yang sempurna. Dari sisi kesalahan klasifikasi, total *error* hanya tiga kasus, satu meningioma yang diprediksi sebagai glioma, serta dua tumor yang diprediksi sebagai meningioma. Jumlah kesalahan ini setara dengan GELU, lebih sedikit dibanding PReLU dan SELU, meski sedikit lebih tinggi dibanding Mish, ReLU, dan Swish yang hanya mencatat dua *error*. Dengan demikian, HardSwish dapat dipandang sebagai fungsi aktivasi yang kuat dan efisien, meskipun membutuhkan mekanisme stabilisasi tambahan di awal pelatihan untuk mengurangi fluktuasi. Potensi penggunaannya pada jaringan yang lebih dalam atau pada skenario dengan jumlah *epoch* yang lebih banyak tampaknya sangat besar, mengingat kemampuan konvergensi akhirnya yang konsisten dan akurasi uji yang sempurna.

10. MedAct Fixed

Training model yang dilakukan untuk menguji model ResNet-18 dengan fungsi aktivasi hibrid baru MedAct Fixed.

Tabel 11.

Classification Report training ResNet-18 dengan aktivasi hibrid baru MedAct Fixed

Precision	Recall	F1-Score	Support	
0.98	1	0.99	194	Glioma
0.98	0.99	0.98	212	Menin
1	0.98	0.99	201	Tumor
		0.99	607	Accuracy
0.99	0.99	0.99	607	Macro Avg
0.99	0.99	0.99	607	Weight Avg
		0.98	607	Macro AUC

Eksperimen dengan MedAct Fixed (α tetap) memperlihatkan performa yang *solid* meskipun belum melampaui fungsi aktivasi konvensional seperti Mish atau ReLU. Selama proses pelatihan, *training loss* menurun

secara konsisten hingga sekitar 0,0051, dengan akurasi pelatihan mencapai 99,8%. *Validation loss* sempat menunjukkan fluktuasi yang cukup tajam, terutama pada *epoch* kelima dan kedelapan, yang menyebabkan akurasi validasi turun hingga sekitar 96,7%. Namun, setelah fase tersebut, model kembali stabil dan menutup pelatihan dengan *best validation accuracy* sebesar 99,67%, setara dengan ReLU dan PReLU. Waktu pelatihan yang efisien, yaitu sekitar tujuh setengah menit, menunjukkan bahwa penambahan komponen fungsi aktivasi baru ini tidak mengorbankan kecepatan komputasi.

Evaluasi pada data uji menghasilkan akurasi keseluruhan sebesar 99%, dengan nilai *precision*, *recall*, dan *F1-score* rata-rata berada pada angka 0,99, serta *macro AUC* 0,9998. Akan tetapi, jumlah kesalahan klasifikasi relatif tinggi, mencapai tujuh kasus, tiga meningioma yang salah diprediksi sebagai glioma dan empat tumor yang diprediksi sebagai meningioma. Jumlah ini setara dengan PReLU dan lebih banyak dibandingkan fungsi aktivasi lain seperti ReLU, Mish, maupun Swish. Dengan demikian, meskipun MedAct Fixed mampu mempertahankan performa tinggi secara umum, varian ini belum sepenuhnya menunjukkan keunggulan nyata dibandingkan aktivasi standar. Hal ini mengindikasikan bahwa penggunaan α tetap mungkin kurang fleksibel, sehingga kontribusinya terhadap generalisasi model masih terbatas.

11. MedAct Learnable

Training model yang dilakukan untuk menguji model ResNet-18 dengan fungsi aktivasi hibrid baru MedAct Learnable.

Tabel 12.
Classification Report training ResNet-18 dengan aktivasi hibrid baru MedAct Learnable

Precision	Recall	F1-Score	Support	
0.99	1	0.99	194	Glioma
1	0.99	1	212	Menin
1	1	1	201	Tumor
		1	607	Accuracy
1	1	1	607	Macro Avg
1	1	1	607	Weight Avg
		1	607	Macro AUC

Eksperimen dengan MedAct Learnable (α dapat dilatih) menunjukkan hasil yang sangat kompetitif sekaligus memperbaiki kelemahan varian *fixed*. Selama pelatihan, *training loss* menurun cepat sejak awal dan terus berkurang hingga mencapai 0,011 pada *epoch* ke-10, disertai akurasi pelatihan hampir sempurna (99,6%). Pada data validasi, meskipun sempat terjadi fluktuasi pada *epoch* keempat dan kelima, performa kembali stabil mulai *epoch* ketujuh hingga akhir, dengan *best validation accuracy* sebesar 99,67%. Pola ini mengindikasikan proses konvergensi yang baik, tanpa adanya indikasi *overfitting* berarti, karena kesenjangan antara *training* dan *validation* relatif kecil. Waktu pelatihan juga efisien, sekitar tujuh setengah menit, sebanding dengan fungsi aktivasi modern lain seperti Mish, ELU, dan HardSwish.

Evaluasi pada data uji memperlihatkan bahwa MedAct Learnable berhasil menurunkan jumlah kesalahan klasifikasi secara signifikan dibandingkan varian *fixed*. Dari total 607 sampel, hanya dua kesalahan yang tercatat, keduanya berupa kasus meningioma yang salah diprediksi sebagai glioma. Hal ini menghasilkan akurasi keseluruhan sekitar 99,67% dengan nilai *precision*, *recall*, dan *F1-score* rata-rata mendekati 1, serta *macro AUC* sempurna. Jika dibandingkan dengan MedAct Fixed yang mencatat tujuh kesalahan, versi *learnable* lebih adaptif dalam mengatur kontribusi komponen *smooth* pada tiap *layer*, sehingga memberikan generalisasi yang lebih baik. Dengan performa setara ReLU, Swish, dan Mish, serta stabilitas konvergensi yang baik, MedAct Learnable dapat dipandang sebagai pengembangan yang relevan untuk meningkatkan akurasi klasifikasi citra medis, meskipun perbedaan antar fungsi aktivasi pada dataset ini tetap berada pada margin yang sangat kecil.

Fungsi aktivasi hibrid baru MedAct dirancang untuk menggabungkan kekuatan *piecewise linearity* ReLU dengan keluwesan fungsi-fungsi non-linear yang lebih halus seperti Swish dan Mish. Berdasarkan hasil eksperimen, MedAct telah menunjukkan performa yang cukup menjanjikan untuk dipertimbangkan sebagai fungsi aktivasi alternatif dalam klasifikasi citra medis. Varian MedAct Learnable khususnya berhasil mencapai akurasi dan stabilitas yang sebanding dengan fungsi aktivasi mapan seperti ReLU, Mish, dan Swish, bahkan menurunkan jumlah kesalahan klasifikasi secara signifikan dibandingkan varian *fixed*. Hal ini menunjukkan bahwa konsep hibrid yang mengombinasikan komponen *linear* dan *smooth* adaptif mampu memberikan generalisasi yang kuat tanpa mengorbankan efisiensi pelatihan. Dengan demikian, MedAct dalam bentuk *learnable α* dapat dikatakan layak diuji lebih luas, terutama pada tugas klasifikasi medis lain yang memiliki kompleksitas visual serupa.

Namun, untuk dapat benar-benar dianggap matang, pengembangan lebih lanjut tetap diperlukan. Pertama, eksperimen ini masih terbatas pada satu *dataset* tumor otak; uji lintas *dataset*, multi-institusi, dan kondisi citra dengan kualitas beragam (misalnya citra dengan *noise* atau artefak klinis) penting dilakukan untuk memastikan *robustness* MedAct. Kedua, analisis sensitivitas terhadap nilai α serta perilaku MedAct pada arsitektur yang lebih dalam (misalnya ResNet-50/101 atau EfficientNet) dapat memperlihatkan sejauh mana fungsi ini tetap stabil. Ketiga, kajian teoretis mengenai gradien, kompleksitas komputasi, dan pengaruhnya terhadap *vanishing/exploding gradient* juga akan memperkuat justifikasi akademiknya. Dengan kata lain, MedAct sudah menunjukkan layak digunakan dalam konteks terbatas, tetapi masih memerlukan pengembangan pada aspek generalitas, *robustness*, dan analisis teoritis agar dapat diadopsi secara lebih luas di komunitas medis dan kecerdasan buatan.

IV. KESIMPULAN DAN SARAN

A. KESIMPULAN

Hasil eksperimen terhadap 11 fungsi aktivasi pada arsitektur ResNet-18 untuk klasifikasi tumor otak berbasis citra MRI menunjukkan bahwa hampir semua fungsi mampu mencapai performa sangat tinggi dengan akurasi uji mendekati atau setara dengan 100%. Aktivasi klasik seperti ReLU tetap menjadi *baseline* yang kuat, terbukti sederhana namun menghasilkan *error* minimal. Aktivasi modern seperti Mish, ELU, Swish, GELU, dan HardSwish memberikan performa yang sebanding, dengan Mish dan ELU menunjukkan kestabilan jangka panjang, sementara Swish dan HardSwish memperlihatkan fluktuasi awal namun konvergensi yang sangat baik di akhir pelatihan. Di sisi lain, PReLU dan SELU relatif kurang konsisten, dengan jumlah *error* sedikit lebih tinggi serta validasi yang lebih fluktuatif dibandingkan rekan-rekannya. Fungsi hibrid baru MedAct menambah perspektif menarik, varian *Fixed α* kurang optimal karena menghasilkan kesalahan lebih banyak, tetapi *Learnable α* mampu menurunkan *error* hingga hanya dua kasus, menjadikannya setara atau bahkan lebih adaptif dibandingkan fungsi aktivasi yang sudah mapan.

B. SARAN

Mengingat performa antar fungsi aktivasi pada dataset ini berbeda tipis, maka klaim perbedaan signifikan perlu diuji melalui replikasi eksperimen, *cross-validation* multi-institusi, dan evaluasi pada *dataset* yang lebih kompleks atau berisik. Selain itu, untuk MedAct khususnya, penelitian lanjutan sebaiknya diarahkan pada tiga aspek utama, yaitu:

1. Generalitas, dengan menguji pada arsitektur yang lebih dalam atau berbeda (misalnya EfficientNet, DenseNet, atau Vision Transformer);
2. Robustnes, dengan menguji ketahanan terhadap *noise*, *data imbalance*, dan variasi kualitas citra medis;
3. Analisis teoretis, termasuk studi perilaku gradien dan interpretabilitas kontribusi α pada tiap *layer*.

Dengan demikian, meskipun MedAct Learnable sudah menunjukkan potensi layak digunakan, penguatan melalui eksperimen lebih luas dan kajian matematis akan semakin memperkuat posisinya sebagai fungsi aktivasi alternatif dalam domain kecerdasan buatan untuk kesehatan.

Tumor Classification and Segmentation Using a Multiscale Convolutional Neural Network,” *Healthcare*, vol. 9, no. 2, p. 153, Feb. 2021, doi: 10.3390/healthcare9020153.

- [4] M. Aamir *et al.*, “Brain Tumor Detection and Classification Using an Optimized Convolutional Neural Network,” *Diagnostics*, vol. 14, no. 16, 2024, doi: 10.3390/diagnostics14161714.
- [5] K. Sharma, V. Jaiswal, T. K. Dey, S. S. Chauhan, A. N. Hati, and M. Jangid, “Brain Tumor Detection and Classification Using Attention Based Residual Learning,” in *2025 International Conference on Next Generation Communication & Information Processing (INCIP)*, 2025, pp. 892–897. doi: 10.1109/INCIP64058.2025.11019958.
- [6] S. N. Yousafzai, I. M. Nasir, S. Tehsin, and J. A. Khan, “MRA-Net: Multiscale Residual Attention Network for Multiclass Alzheimer Disease Classification,” in *2024 5th International Conference on Innovative Computing (ICIC)*, 2024, pp. 1–8. doi: 10.1109/ICIC63915.2024.11115990.
- [7] Z. Liu *et al.*, “Swin Transformer: Hierarchical Vision Transformer using Shifted Windows,” in *2021 IEEE/CVF International Conference on Computer Vision (ICCV)*, 2021, pp. 9992–10002. doi: 10.1109/ICCV48922.2021.00986.
- [8] N. Shazeer, “GLU Variants Improve Transformer.” 2020. [Online]. Available: <https://arxiv.org/abs/2002.05202>
- [9] M. Tan and Q. Le, “EfficientNet: Rethinking Model Scaling for Convolutional Neural Networks,” in *Proceedings of the 36th International Conference on Machine Learning*, K. Chaudhuri and R. Salakhutdinov, Eds., in Proceedings of Machine Learning Research, vol. 97. PMLR, Jun. 2019, pp. 6105–6114. [Online]. Available: <https://proceedings.mlr.press/v97/tan19a.html>
- [10] S. Panigrahi, D. R. D. Adhikary, and B. K. Pattanayak, “Brain tumor classification: a blend of ensemble learning and fine-tuned pre-trained models,” *Discov. Appl. Sci.*, vol. 7, no. 4, p. 274, Mar. 2025, doi: 10.1007/s42452-025-06695-x.
- [11] A. Khan, A. Sohail, U. Zahoor, and A. S. Qureshi, “A survey of the recent architectures of deep convolutional neural networks,” *Artif. Intell. Rev.*, vol. 53, no. 8, pp. 5455–5516, Dec. 2020, doi: 10.1007/s10462-020-09825-6.
- [12] “Swin Transformer Hierarchical Vision Transformer using Shifted Windows.”
- [13] “Brain Tumor Classification in MRI Insights from LIME and Grad-CAM Explainable AI Techniques.”
- [14] S. Elfving, E. Uchibe, and K. Doya, “Sigmoid-Weighted Linear Units for Neural Network Function Approximation in Reinforcement Learning.” 2017. [Online]. Available: <https://arxiv.org/abs/1702.03118>
- [15] T. Berghout, “The Neural Frontier of Future Medical Imaging: A Review of Deep Learning for Brain Tumor Detection,” *J. Imaging*, vol. 11, no. 1, 2025, doi: 10.3390/jimaging11010002.
- [16] M. A. Khan *et al.*, “Transfer learning for accurate brain tumor classification in MRI: a step forward in medical diagnostics,” *Discov. Oncol.*, vol. 16, no. 1, p. 1040, Jun. 2025, doi: 10.1007/s12672-025-02671-4.
- [17] A. Mondal and V. K. Shrivastava, “A novel Parametric Flatten-p Mish activation function based deep CNN model for brain tumor classification,” *Comput. Biol. Med.*, vol. 150, p. 106183, 2022, doi: <https://doi.org/10.1016/j.compbiomed.2022.106183>.

DAFTAR PUSTAKA

- [1] Y. Chel and L. Poh, “Brain Tumor Classification in MRI: Insights from LIME and Grad-CAM Explainable AI Techniques,” *IEEE Access*, vol. PP, pp. 1–1, Jan. 2025, doi: 10.1109/ACCESS.2025.3603272.
- [2] R. Alexander, A. Lakshana, T. Lavanya, G. S. Fathima, N. Kavya, and M. V. Janet A, “Deep Learning-based MR Image Segmentation for Brain Tumor Detection using EAU-Net,” in *2025 3rd International Conference on Inventive Computing and Informatics (ICICI)*, 2025, pp. 1120–1126. doi: 10.1109/ICICI65870.2025.11069594.
- [3] F. J. Díaz-Pernas, M. Martínez-Zarzuela, M. Antón-Rodríguez, and D. González-Ortega, “A Deep Learning Approach for Brain