

**KLASIFIKASI TUMOR OTAK BERDASARKAN CITRA MRI
MENGUNAKAN *GRAY LEVEL CO-OCCURRENCE*
MATRIX DAN *MULTILAYER PERCEPTRON***

SKRIPSI

Oleh :
GAVRILA PANDITA
NIM. 220605110004



**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS SAINS DAN TEKNOLOGI
UNIVERSITAS ISLAM NEGERI MAULANA MALIK IBRAHIM
MALANG
2025**

**KLASIFIKASI TUMOR OTAK BERDASARKAN CITRA MRI
MENGUNAKAN *GRAY LEVEL CO-OCCURRENCE*
MATRIX DAN *MULTILAYER PERCEPTRON***

SKRIPSI

Diajukan kepada:
Universitas Islam Negeri Maulana Malik Ibrahim Malang
Untuk memenuhi Salah Satu Persyaratan dalam
Memperoleh Gelar Sarjana Komputer (S.Kom)

Oleh :
GAVRILA PANDITA
NIM. 220605110004

**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS SAINS DAN TEKNOLOGI
UNIVERSITAS ISLAM NEGERI MAULANA MALIK IBRAHIM
MALANG
2025**

HALAMAN PERSETUJUAN

KLASIFIKASI TUMOR OTAK BERDASARKAN CITRA MRI MENGUNAKAN *GRAY LEVEL CO-OCCURRENCE* *MATRIX* DAN *MULTILAYER PERCEPTRON*

SKRIPSI

Oleh :
GAVRILA PANDITA
NIM. 220605110004

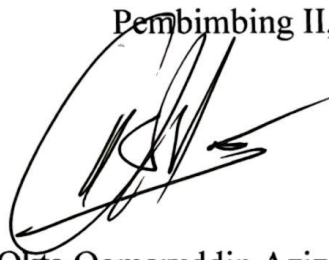
Telah Diperiksa dan Disetujui untuk Diuji:
Tanggal: 09 Desember 2025

Pembimbing I,



Tri Mukti Lestari, M.Kom
NIP. 199111082020122005

Pembimbing II,



Okta Qomaruddin Aziz, M.Kom
NIP. 199110192019031013

Mengetahui,
Ketua Program Studi Teknik Informatika
Fakultas Sains dan Teknologi
Universitas Islam Negeri Maulana Malik Ibrahim Malang



Supriyono, M. Kom
NIP. 19841010 201903 1 012

HALAMAN PENGESAHAN

KLASIFIKASI TUMOR OTAK BERDASARKAN CITRA MRI MENGUNAKAN *GRAY LEVEL CO-OCCURRENCE* *MATRIX* DAN *MULTILAYER PERCEPTRON*

SKRIPSI

Oleh :
GAVRILA PANDITA
NIM. 220605110004

Telah Dipertahankan di Depan Dewan Penguji Skripsi
dan Dinyatakan Diterima Sebagai Salah Satu Persyaratan
Untuk Memperoleh Gelar Sarjana Komputer (S.Kom)
Tanggal : 15 Desember 2025

Susunan Dewan Penguji

Ketua Penguji	: <u>Prof. Dr. Suhartono S.Si M.Kom</u> NIP. 196805192003121001
Anggota Penguji I	: <u>Ashri Shabrina Afrah, M.T</u> NIP. 199004302020122003
Anggota Penguji II	: <u>Tri Mukti Lestari, M.Kom</u> NIP. 199111082020122005
Anggota Penguji III	: <u>Okta Qomaruddin Aziz, M.Kom</u> NIP. 199110192019031013

()
()
()
()

Mengetahui dan Mengesahkan,
Ketua Program Studi Teknik Informatika
Fakultas Sains dan Teknologi
Universitas Islam Negeri Maulana Malik Ibrahim Malang




Suhartono, M. Kom
NIP. 19841010 201903 1 012

PERNYATAAN KEASLIAN TULISAN

Saya yang bertanda tangan di bawah ini:

Nama : Gavrila Pandita
NIM : 220605110004
Fakultas / Program Studi : Sains dan Teknologi / Teknik Informatika
Judul Skripsi : Klasifikasi Tumor Otak Berdasarkan Citra MRI
Menggunakan *Gray Level Co-occurrence Matrix*
dan *Multilayer Perceptron*

Menyatakan dengan sebenarnya bahwa Skripsi yang saya tulis ini benar-benar merupakan hasil karya saya sendiri, bukan merupakan pengambil alihan data, tulisan, atau pikiran orang lain yang saya akui sebagai hasil tulisan atau pikiran saya sendiri, kecuali dengan mencantumkan sumber cuplikan pada daftar pustaka.

Apabila dikemudian hari terbukti atau dapat dibuktikan skripsi ini merupakan hasil jiplakan, maka saya bersedia menerima sanksi atas perbuatan tersebut.

Malang, 31 Desember 2025
Yang membuat pernyataan,



Gavrila Pandita
NIM.220605110004

MOTTO

"Allah mengambil darimu sesuatu yang tidak pernah engkau sangka kehilangannya, maka Allah akan memberimu sesuatu yang tidak pernah engkau sangka akan memilikinya."

HALAMAN PERSEMBAHAN

Segala puji dan syukur penulis panjatkan ke hadirat Allah Subhanahu wa Ta'ala atas limpahan rahmat dan hidayah-Nya, sehingga skripsi ini dapat diselesaikan dengan baik.

Karya ini penulis dedikasikan kepada:

Mama tercinta, Atik Mudiharti, S.E.

Atas doa, dukungan, dan kasih sayang yang tidak pernah terputus, serta ketulusan Mama dalam menguatkan dan mendampingi penulis hingga tahap ini.

Ayah tercinta, Almarhum Aris Juana Yusuf, Lc., M.A.

Sosok panutan yang selalu dirindukan. Semoga Allah Subhanahu wa Ta'ala memberikan tempat terbaik di sisi-Nya.

Saudari kembar, Gavra Pandita

dan kakak pertama, Magda Achlaa Amany, S.Psi.

Terima kasih telah menjadi tempat berbagi cerita, mencari solusi, serta memberikan dukungan dan semangat di setiap kondisi.

KATA PENGANTAR

Assalamualaikum warahmatullahi wabarakatuh.

Dengan memanjatkan puji syukur ke hadirat Allah Subhanahu wa Ta'ala, penulis dapat menyelesaikan skripsi yang berjudul *“Klasifikasi Tumor Otak Berdasarkan Citra MRI Menggunakan Gray Level Co-Occurrence Matrix dan Multilayer Perceptron”*. Penulisan skripsi ini merupakan salah satu syarat akademik dalam menyelesaikan studi pada Program Studi Teknik Informatika, Fakultas Sains dan Teknologi. Shalawat dan salam semoga senantiasa tercurah kepada Nabi Muhammad Shallallahu ‘Alaihi Wasallam, beserta keluarga, sahabat, dan umatnya hingga akhir zaman.

Penulis menyadari bahwa tersusunnya skripsi ini tidak lepas dari peran serta berbagai pihak yang telah memberikan bantuan, arahan, dan dukungan. Oleh sebab itu, penulis menyampaikan ucapan terima kasih kepada:

1. Prof. Dr. Hj. Ilfi Nurdiana, S.Ag., M.Si., selaku Rektor Universitas Islam Negeri Maulana Malik Ibrahim Malang.
2. Dr. H. Agus Mulyono, S.Pd., M.Kes., selaku Dekan Fakultas Sains dan Teknologi Universitas Islam Negeri Maulana Malik Ibrahim Malang.
3. Supriyono, M.Kom., selaku Ketua Program Studi Teknik Informatika Universitas Islam Negeri Maulana Malik Ibrahim Malang.
4. Tri Mukti Lestari, M.Kom., selaku Dosen Pembimbing I yang telah memberikan bimbingan, arahan, dan masukan selama proses penyusunan skripsi ini.

5. Okta Qomaruddin Aziz, M.Kom., selaku Dosen Pembimbing II yang telah memberikan pendampingan, saran, serta koreksi yang sangat berarti dalam penyelesaian skripsi ini.
6. Prof. Dr. Suhartono, S.Si., M.Kom., selaku Ketua Penguji yang telah memberikan evaluasi dan masukan untuk penyempurnaan skripsi ini.
7. Ashri Shabrina Afrah, M.T., selaku Dosen Penguji yang telah memberikan saran dan kritik yang membangun dalam proses pengujian skripsi.
8. Dr. Agung Teguh Wibowo Almais, M.T., selaku Dosen Wali yang telah memberikan arahan dan bimbingan akademik selama penulis menempuh masa studi.
9. Nia Faricha, S.Si., selaku Admin Program Studi Teknik Informatika yang telah membantu dan memberikan informasi akademik selama masa studi.
10. Seluruh dosen, tenaga kependidikan, laboran, serta jajaran staf Program Studi Teknik Informatika yang telah memberikan dukungan dan bantuan selama penulis menempuh studi.
11. Mama Atik Mudiharti, S.E., yang selalu memberikan doa, dukungan, dan menjadi alasan utama bagi penulis untuk terus bersemangat menyelesaikan perkuliahan ini. Mama yang dengan penuh ketulusan senantiasa berusaha memenuhi segala kebutuhan anak-anaknya sebagai seorang *single parent*, sosok ibu yang kuat dan suportif, tidak pernah menuntut secara berlebihan, serta selalu merawat dan memastikan anak-anaknya berada dalam keadaan sehat dan tercukupi kebutuhannya.

12. Almarhum Ayah Aris Juana Yusuf, Lc., M.A., sosok yang selalu dirindukan oleh penulis. Ayah yang sangat cerdas, bijaksana, memiliki perjalanan pendidikan dan karier yang sangat keren. Ayah yang juga selalu menanamkan nilai-nilai keislaman melalui ajaran doa-doa dan kebiasaan islami yang baik, yang hingga kini terus menjadi pedoman dalam kehidupan penulis. Semoga Allah *Subhanahu wa Ta'ala* memberikan tempat terbaik di sisi-Nya.
13. Saudari kembar Gavra Pandita dan kakak pertama Magda Achlaa Amany, S.Psi., yang telah menjadi tempat berkeluh kesah, berdiskusi mencari solusi, pemberi hiburan di saat penat, serta selalu hadir membantu penulis dalam menghadapi berbagai kesulitan.
14. Yoza Setya Febrianti dan Sakila Aulia Maharani, teman dekat penulis dalam grup “Yupi” yang menemani penulis dalam berbagai momen selama masa perkuliahan, baik di kegiatan akademik maupun keseharian.
15. Teman-teman angkatan 2022 “Infinity”, khususnya Laudza Atsila Prasetyo dan Husnul Khatimah, yang senantiasa memberikan saran, masukan, serta solusi selama proses penulisan skripsi ini.
16. Ikmalunisa Annora, yang selalu menyempatkan diri memberikan semangat dan dukungan meskipun masing-masing memiliki kesibukan yang berbeda.
17. Teman-teman “Kamar 6”, yang senantiasa menemani dan menjadi tempat berbagi kebersamaan selama masa perkuliahan, terutama pada masa-masa penulis menjalani kehidupan di asrama.

18. Semua pihak yang telah berperan serta memberikan dukungan, baik secara langsung maupun tidak langsung, selama penulis menempuh perkuliahan sampai dengan selesainya penyusunan skripsi ini.
19. Terakhir, penulis mengucapkan terima kasih kepada diri sendiri karena terus berusaha, tetap konsisten, dan menyelesaikan skripsi ini meskipun tidak selalu berjalan mudah, dengan berbagai hal yang tidak dapat penulis ceritakan kepada siapa pun.

Dalam penyusunannya, skripsi ini masih memiliki keterbatasan yang perlu disempurnakan. Penulis sangat mengharapkan adanya kritik dan saran yang membangun sebagai bahan evaluasi di masa mendatang. Semoga karya ini dapat memberikan manfaat dan menjadi sumber pengetahuan bagi pembaca.

Malang, 25 Desember 2025

Penulis

DAFTAR ISI

HALAMAN JUDUL	i
HALAMAN PENGAJUAN	ii
HALAMAN PERSETUJUAN	iii
HALAMAN PENGESAHAN	iv
PERNYATAAN KEASLIAN TULISAN	v
MOTTO	vi
HALAMAN PERSEMBAHAN	vii
KATA PENGANTAR.....	viii
DAFTAR ISI.....	xii
DAFTAR GAMBAR.....	xiv
DAFTAR TABEL	xv
ABSTRAK	xvi
ABSTRACT	xvii
البحث مستخلص.....	xviii
BAB I PENDAHULUAN	1
1.1 Latar Belakang	1
1.2 Pernyataan Masalah	4
1.3 Batasan Masalah.....	5
1.4 Tujuan Penelitian	5
1.5 Manfaat Penelitian	5
BAB II STUDI PUSTAKA	6
2.1 Tumor Otak	6
2.2 <i>Pre-processing Data</i>	12
2.2.1 <i>Resize</i>	13
2.2.2 <i>Normalisasi</i>	15
2.2.3 <i>Noise Removal</i>	18
2.3 <i>Gray Level Co-occurrence Matrix</i>	20
2.4 <i>Multilayer Perceptron</i>	25
2.5 <i>Confusion matrix</i>	30
BAB III METODE PENELITIAN	33
3.1 Pengumpulan Data	33
3.2 Desain Sistem.....	34
3.3 <i>Pre-processing Data</i>	34
3.3.1 <i>Resize</i>	35
3.3.2 <i>Normalisasi</i>	36
3.3.3 <i>Noise Removal</i>	37
3.3.4 <i>Hasil Pre-processing</i>	39
3.4 <i>Ekstraksi Fitur Gray Level Co-occurrence Matrix</i>	40
3.5 Implementasi Algoritma <i>Multilayer Perceptron</i>	45
3.6 Skenario Pengujian.....	58
3.7 Evaluasi	59

BAB IV UJI COBA DAN PEMBAHASAN	62
4.1 Hasil Pengujian	62
4.1.1 Model A	62
4.1.2 Model B	63
4.1.3 Model C	64
4.1.4 Model D	65
4.1.5 Model E.....	66
4.1.6 Model F.....	67
4.1.7 Model G	68
4.1.8 Model H	69
4.1.9 Model I.....	70
4.1.10 Model J	71
4.2 Pembahasan.....	72
4.2.1 Pengaruh Parameter	74
4.2.1.1 <i>Hidden layer</i>	74
4.2.1.2 Neuron.....	75
4.2.1.3 <i>Batch size</i>	77
4.3 Analisis <i>K-Fold Cross Validation</i> pada Model Terbaik	79
4.3.1 Model Terbaik 3 <i>Hidden layer</i>	79
4.3.2 Model Terbaik 2 <i>Hidden layer</i>	80
4.3.3 Perbandingan Kestabilan Model Terbaik	82
4.4 Implementasi GUI.....	83
BAB V KESIMPULAN DAN SARAN	87
5.1 Kesimpulan	87
5.2 Saran	88
DAFTAR PUSTAKA	

DAFTAR GAMBAR

Gambar 2. 1 Arah orientasi sudut pada GLCM	20
Gambar 2. 2 Matriks <i>co-occurrence</i> pada GLCM.....	21
Gambar 2. 3 Arsitektur sederhana MLP	25
Gambar 3. 1 Desain Sistem.....	34
Gambar 3. 2 <i>Flowchart Pre-processing</i> Citra.....	34
Gambar 3. 3 Citra MRI sebelum dan sesudah proses <i>resize</i>	36
Gambar 3. 4 Citra MRI sebelum dan sesudah proses normalisasi	37
Gambar 3. 5 Citra MRI sebelum dan sesudah proses <i>denoise</i>	38
Gambar 3. 6 Hasil <i>pre-processing</i> pada citra MRI kelas tidak tumor.....	39
Gambar 3. 7 Hasil <i>pre-processing</i> pada citra MRI kelas tumor	39
Gambar 3. 8 <i>Flowchart</i> GLCM	40
Gambar 3. 9 Contoh matriks GLCM ternormalisasi.....	41
Gambar 3. 10 Arsitektur MLP dengan 3 <i>Hidden Layer</i>	45
Gambar 3. 11 Arsitektur MLP dengan 2 <i>Hidden Layer</i>	46
Gambar 4. 1 <i>Confusion matrix</i> model A.....	62
Gambar 4. 2 <i>Confusion matrix</i> model B	63
Gambar 4. 3 <i>Confusion matrix</i> model C	64
Gambar 4. 4 <i>Confusion matrix</i> model D.....	65
Gambar 4. 5 <i>Confusion matrix</i> model E	66
Gambar 4. 6 <i>Confusion matrix</i> model F.....	67
Gambar 4. 7 <i>Confusion matrix</i> model G.....	68
Gambar 4. 8 <i>Confusion matrix</i> model H.....	69
Gambar 4. 9 <i>Confusion matrix</i> model I.....	70
Gambar 4. 10 <i>Confusion matrix</i> model J	71
Gambar 4. 11 Grafik perbandingan matriks evaluasi seluruh model.....	73
Gambar 4. 12 <i>Boxplot</i> Hasil <i>K-Fold Cross Validation</i> pada Skenario Terbaik	79
Gambar 4. 13 GUI Sistem Klasifikasi Tumor Otak Berbasis Citra MRI.....	84
Gambar 4. 14 Tampilan Hasil Klasifikasi Citra MRI pada GUI.....	84

DAFTAR TABEL

Tabel 2. 1 Penelitian Terdahulu	9
Tabel 3. 1 Contoh Data Citra MRI Tumor Otak	33
Tabel 3. 2 Skenario Pengujian	59
Tabel 3. 3 Contoh <i>confusion matrix</i>	59
Tabel 4. 1 Hasil <i>confusion matrix</i> model A	62
Tabel 4. 2 Hasil <i>confusion matrix</i> model B	63
Tabel 4. 3 Hasil <i>confusion matrix</i> model C	64
Tabel 4. 4 Hasil <i>confusion matrix</i> model D	65
Tabel 4. 5 Hasil <i>confusion matrix</i> model E.....	66
Tabel 4. 6 Hasil <i>confusion matrix</i> model F.....	67
Tabel 4. 7 Hasil <i>confusion matrix</i> model G	68
Tabel 4. 8 Hasil <i>confusion matrix</i> model H	69
Tabel 4. 9 Hasil <i>confusion matrix</i> model I.....	70
Tabel 4. 10 Hasil <i>confusion matrix</i> model J.....	71
Tabel 4. 11 Hasil perfoma seluruh model	72
Tabel 4. 12 Rata-rata Akurasi Berdasarkan Jumlah <i>Hidden Layer</i>	74
Tabel 4. 13 Perbandingan Jumlah Neuron pada 3 <i>Hidden Layer</i>	75
Tabel 4. 14 Perbandingan Jumlah Neuron pada 2 <i>Hidden Layer</i>	76
Tabel 4. 15 Perbandingan <i>Batch size</i> pada 3 <i>Hidden Layer</i>	77
Tabel 4. 16 Perbandingan <i>Batch size</i> pada 2 <i>Hidden Layer</i>	78
Tabel 4. 17 Hasil <i>K-Fold Cross Validation</i> pada Skenario Terbaik	80

ABSTRAK

Pandita, Gavrila. 2025. **Klasifikasi Tumor Otak Berdasarkan Citra Mri Menggunakan *Gray Level Co-occurrence Matrix* dan *Multilayer Perceptron***. Skripsi. PROGRAM STUDI Teknik Informatika Fakultas Sains dan Teknologi Universitas Islam Negeri Maulana Malik Ibrahim Malang. Pembimbing: (I) Tri Mukti Lestari, M.Kom (II) Okta Qomaruddin Aziz, M.Kom.

Kata kunci: Klasifikasi tumor otak, Citra MRI, *Gray Level Co-Occurrence Matrix*, *Multilayer Perceptron*

Tumor otak merupakan pertumbuhan sel abnormal yang dapat menimbulkan gangguan neurologis serius dan banyak ditemukan pada usia produktif, sehingga diperlukan metode deteksi yang cepat dan akurat untuk mendukung penanganan dini. Penelitian ini bertujuan mengembangkan model klasifikasi untuk membedakan citra MRI bertumor dan tidak bertumor menggunakan *Gray Level Co-occurrence Matrix* (GLCM) dan *Multilayer Perceptron* (MLP). Tahapan penelitian diawali dengan preprocessing citra berupa penyeragaman ukuran, normalisasi intensitas, dan reduksi noise, kemudian dilanjutkan dengan ekstraksi fitur tekstur GLCM berupa nilai contrast, correlation, energy, dan homogeneity pada empat orientasi sudut. Fitur tersebut digunakan sebagai masukan pada beberapa skenario arsitektur MLP untuk memperoleh model terbaik. Hasil pengujian menunjukkan bahwa arsitektur dengan tiga *hidden layer* dan *batch size* 32 menghasilkan akurasi tertinggi sebesar 93,67% serta metrik presisi, *recall*, dan *F1-score* yang stabil. Temuan ini menunjukkan bahwa kombinasi GLCM dan MLP mampu mempelajari pola tekstur citra MRI secara efektif sehingga dapat menjadi pendekatan yang ringan dan representatif untuk mendukung deteksi awal tumor otak.

ABSTRACT

Pandita, Gavril. 2025. *Brain Tumor Classification Based on MRI Images Using Gray Level Co-occurrence Matrix and Multilayer Perceptron*. Thesis. Department of Informatics Engineering, Faculty of Science and Technology, Maulana Malik Ibrahim State Islamic University of Malang. Advisors: (I) Tri Mukti Lestari, M.Kom (II) Okta Qomaruddin Aziz, M.Kom.

Keywords: Brain tumor classification, MRI imaging, Gray Level Co-Occurrence Matrix, *Multilayer Perceptron*.

Brain tumors are abnormal cell growths that can cause serious neurological disorders and are commonly found in individuals of productive age, making fast and accurate detection essential for early medical intervention. This study aims to develop a classification model to distinguish between MRI images with and without tumors using the *Gray Level Co-occurrence Matrix* (GLCM) and a *Multilayer Perceptron* (MLP). The research process begins with image preprocessing, including size standardization, intensity normalization, and noise reduction, followed by the extraction of GLCM texture features—contrast, correlation, energy, and homogeneity—across four directional orientations. These features are used as inputs for several MLP architecture scenarios to obtain the most effective model. The results show that an architecture with three *hidden layers* and a *batch size* of 32 achieved the highest accuracy of 93.67%, along with stable precision, *recall*, and *F1-score* metrics. These findings demonstrate that the combination of GLCM and MLP can effectively learn texture patterns in MRI images, making it a lightweight and representative approach to support early detection of brain tumors.

البحث مستخلص

فانديتا، غفريلا. 2025. تصنيف أورام الدماغ بناءً على صور الرنين المغناطيسي باستخدام مصفوفة التواجد المشترك لمستويات الرمادي والبيرسيبترون متعدد الطبقات. أطروحة. قسم هندسة المعلوماتية، كلية العلوم والتكنولوجيا، جامعة مولانا مالك إبراهيم الإسلامية M.Kom.، أوكنا قمر الدين عزيز (II) M.Kom، تري موكتي ليستاري (I): الحكومية، مالانج. المشرفون

الكلمات المفتاحية : الكلمات المفتاحية: تصنيف أورام الدماغ، صور التصوير بالرنين المغناطيسي، مصفوفة التواجد المشترك لمستوى الرمادي، متعدد الطبقات بيرسيبترون

،تأورام الدماغ هي نمو غير طبيعي للخلايا يمكن أن يسبب اضطرابات عصبية خطيرة، وتوجد عادةً لدى الأشخاص في سن الإنتاج مما يتطلب طرقًا سريعة ودقيقة للكشف عنها لدعم العلاج المبكر. تهدف هذه الدراسة إلى تطوير نموذج تصنيف للتمييز بين صور والبيرسيبترون متعدد الطبقات (GLCM) الرنين المغناطيسي للأورام وغير الأورام باستخدام مصفوفة التواجد المشترك لمستوى الرمادي بدأت الدراسة بمعالجة الصور مسبقًا، والتي تضمنت توحيد الحجم وتوحيد الكثافة وتقليل الضوضاء، تلاها استخراج (MLP) في شكل قيم التباين والتربط والطاقة والتجانس في أربعة اتجاهات زاوية. تم استخدام هذه الخصائص GLCM خصائص نسيج للحصول على أفضل نموذج. أظهرت نتائج الاختبار أن البنية ذات الطبقات الخفية MLP كمدخلات في عدة سيناريوهات لهندسة تشير هذه النتائج إلى F1 الثلاث وحجم الدفعة 32 حققت أعلى دقة بنسبة 93.67٪ ومقاييس دقيقة ومستقرة للاستدعاء ودرجة قادر على تعلم أنماط نسيج صور التصوير بالرنين المغناطيسي بشكل فعال، مما يجعله نهجًا خفيًا MLP و GLCM أن الجمع بين وتمثيلًا لدعم الكشف المبكر عن أورام الدماغ.

BAB I

PENDAHULUAN

1.1 Latar Belakang

Pertumbuhan sel abnormal pada jaringan otak dikenal sebagai tumor otak dan dapat bersifat jinak atau ganas. Kondisi ini berpotensi mengganggu fungsi otak serta menyebabkan berbagai gejala neurologis, termasuk sakit kepala, gangguan penglihatan, dan kejang-kejang (Fikriah et al., 2024). Penelitian oleh Sembiring et al., (2025) melaporkan 101 pasien tumor otak dengan rata-rata usia 46,2 tahun, sedangkan Sari et al., (2025) menemukan rata-rata usia pasien sekitar 51 tahun dengan rentang terbanyak 41–61 tahun. Mayoritas kasus terjadi pada perempuan, dan tumor metastasis lebih banyak ditemukan pada pasien usia produktif.

Pengembangan metode komputasi di bidang kesehatan mencerminkan kemajuan teknologi sekaligus pentingnya ilmu dalam menjaga kualitas hidup manusia. Islam pun menempatkan ilmu pada derajat yang mulia, sebagaimana firman Allah SWT dalam QS. *Al-Mujādalah* ayat 11:

يَا أَيُّهَا الَّذِينَ آمَنُوا إِذَا قِيلَ لَكُمْ تَفَسَّحُوا فِي الْمَجَالِسِ فَافْسَحُوا يَفْسَحِ اللَّهُ لَكُمْ وَإِذَا قِيلَ انشُرُوا فَانْشُرُوا اللَّهُ الَّذِينَ آمَنُوا مِنْكُمْ وَالَّذِينَ أُوتُوا الْعِلْمَ دَرَجَاتٍ وَاللَّهُ بِمَا تَعْمَلُونَ خَبِيرٌ

“Wahai orang-orang yang beriman! Apabila dikatakan kepadamu, "Berilah kelapangan di dalam majelis-majelis," maka lapangkanlah, niscaya Allah akan memberi kelapangan untukmu. Dan apabila dikatakan, "Berdirilah kamu," maka berdirilah, niscaya Allah akan mengangkat (derajat) orang-orang yang beriman di antaramu dan orang-orang yang diberi ilmu beberapa derajat. Dan Allah Mahateliti atas apa yang kamu kerjakan.” (QS.Al-Mujādalah/58:11)

Tafsir Al-Misbah Volume 2 (Shihab, 2009), menegaskan bahwa Allah meninggikan derajat orang-orang yang beriman dan berilmu melalui sikap rendah hati, adab, dan kepatuhan dalam proses menuntut ilmu. Ilmu dipahami sebagai amanah yang harus digunakan secara bertanggung jawab dan bermanfaat, sehingga dalam konteks penelitian medis dan ilmiah, pengembangan pengetahuan yang ditujukan untuk kemaslahatan dan kesehatan manusia bernilai mulia ketika dilandasi iman dan etika yang benar.

Magnetic Resonance Imaging (MRI) memanfaatkan medan magnet dan gelombang radio untuk menghasilkan citra beresolusi tinggi dari jaringan otak (Kirana et al., 2023). Namun, perbedaan tekstur antara jaringan normal dan tumor sering sulit dikenali secara manual, terutama pada ukuran kecil atau kontras rendah (Na & Yang, 2021). Karena diagnosis oleh radiolog bersifat subjektif dan bervariasi (Cornelissen et al., 2024), dibutuhkan penerapan *Machine learning* dan *Deep Learning* agar klasifikasi tumor menjadi lebih konsisten, objektif, dan akurat (Kumar et al., 2024; Sravanthi Peddinti et al., 2021).

Ekstraksi fitur merupakan tahap penting dalam pengolahan citra untuk mengambil informasi utama yang merepresentasikan ciri khas objek, seperti bentuk, warna, atau tekstur (Barburiceanu et al., 2021). Pertama, metode statistik seperti *Gray Level Co-occurrence Matrix* (GLCM), *Local Binary Pattern* (LBP), *Haralick*, dan *Gray Level Dependence Matrix* (GLDM) menganalisis hubungan antar piksel berdasarkan tingkat keabuan. LBP efektif mengenali pola halus meski pencahayaan berubah (Sowrirajan & Balasubramanian, 2022), sedangkan GLCM

menghasilkan fitur penting seperti *contrast*, *homogeneity*, *energy*, dan *correlation* yang mendukung klasifikasi tumor otak (Malkauthekar et al., 2025).

Kedua, metode transformasi seperti *Discrete Wavelet Transform* (DWT) dan *Gabor transform* yang digunakan untuk menangkap informasi frekuensi dan pola tekstur pada berbagai skala. DWT terbukti dapat meningkatkan akurasi dalam mendeteksi kelainan pada otak dengan tingkat efisiensi komputasi yang baik (Sowrirajan & Balasubramanian, 2022).

Ketiga, metode berbasis bentuk atau geometris seperti analisis kontur, *fractal dimension*, dan *moment invariants* digunakan untuk menggambarkan bentuk serta struktur tumor, meskipun memerlukan segmentasi yang lebih kompleks (Gül & Kaya, 2024).

Penelitian ini memilih *Gray Level Co-occurrence Matrix* (GLCM) karena memiliki akurasi tinggi dalam pengenalan pola citra yang mencapai 88% pada klasifikasi glioma dan hingga 99% saat dikombinasikan dengan LBP (Dheepak et al., 2024). Selain akurat, GLCM efisien secara komputasi dan mudah diterapkan berkat parameter sederhana yang dapat disesuaikan, seperti jarak antar piksel, sudut, dan tingkat keabuan (Malkauthekar et al., 2025; Tahosin et al., 2023).

Setelah fitur citra diperoleh, tahap klasifikasi bertujuan mengelompokkan citra berdasarkan karakteristik yang telah diekstraksi. Pendekatan yang umum digunakan pada tahap ini adalah *Artificial Neural Network* (ANN), yang dirancang untuk mempelajari hubungan *non-linear* antar fitur dengan meniru mekanisme kerja jaringan saraf biologis (Vijithananda et al., 2022). Beragam arsitektur ANN telah dikembangkan, termasuk *Multilayer Perceptron* (MLP), *Convolutional*

Neural Network (CNN), *Radial Basis Function Network* (RBFN), dan *Vision Transformer* (ViT) (Pagadala et al., n.d.; Wang et al., 2023). MLP memodelkan relasi *non-linear* antar fitur (Al Haris et al., 2025), CNN mengekstraksi informasi spasial secara bertingkat, RBFN mengolah data *non-linear* (Mathi et al., 2025), sementara ViT menangkap dependensi global pada citra medis yang kompleks (Wang et al., 2023).

Dalam penelitian ini, algoritma *Multilayer Perceptron* (MLP) dipilih sebagai klasifikator karena fleksibel dalam memetakan hubungan kompleks antar fitur dan mampu menangani data berdimensi tinggi seperti fitur tekstur GLCM (Bhattacharya et al., 2022). MLP memiliki kapasitas pembelajaran kuat melalui lapisan tersembunyi yang menyesuaikan bobot secara *non-linear* (Fan et al., 2025), sehingga dinilai sesuai untuk menghasilkan akurasi tinggi dengan efisiensi komputasi.

Kinerja sistem dievaluasi menggunakan empat metrik evaluasi, yakni akurasi, presisi, *recall*, dan *F1-score*, yang digunakan untuk menilai ketepatan serta keseimbangan hasil deteksi tumor dan non-tumor (Yousaf et al., 2020).

Melalui penerapan GLCM dalam ekstraksi fitur tekstur dan MLP sebagai metode klasifikasi, penelitian ini diharapkan menghasilkan sistem yang mampu meningkatkan performa identifikasi tumor otak pada citra MRI secara otomatis, cepat, dan objektif.

1.2 Pernyataan Masalah

Bagaimana hasil penerapan metode *Gray Level Co-occurrence Matrix* (GLCM) sebagai teknik ekstraksi fitur tekstur dan algoritma *Multilayer Perceptron*

(MLP) sebagai klasifikator dalam mengklasifikasikan keberadaan tumor otak berdasarkan citra MRI?

1.3 Batasan Masalah

Agar penelitian lebih terarah, ditetapkan batasan sebagai berikut:

1. Dataset yang digunakan adalah *Br35H: Brain Tumor Detection 2020* dari Kaggle, berisi citra MRI otak dengan dua kelas, Tumor dan Tidak Tumor.
2. Penelitian ini menggunakan metode GLCM untuk ekstraksi fitur tekstur dan MLP untuk klasifikasi, dengan evaluasi menggunakan *confusion matrix*.

1.4 Tujuan Penelitian

Menerapkan metode *Gray Level Co-occurrence Matrix* sebagai teknik ekstraksi fitur tekstur dan algoritma *Multilayer Perceptron* sebagai klasifikator dalam mengklasifikasikan tumor otak berdasarkan citra MRI serta menganalisis performa model dalam klasifikasi kategori tumor dan tidak tumor.

1.5 Manfaat Penelitian

Penelitian ini diharapkan mampu memberikan sejumlah manfaat bagi:

1. Akademisi, yaitu dapat menjadi tambahan referensi ilmiah terkait penerapan metode ekstraksi fitur GLCM dan algoritma MLP dalam pengolahan citra medis.
2. Praktisi medis, yaitu dapat memberikan gambaran mengenai potensi pemanfaatan teknologi berbasis *machine learning* untuk membantu proses identifikasi tumor otak.

BAB II

STUDI PUSTAKA

2.1 Tumor Otak

Tumor otak merupakan kondisi pertumbuhan sel abnormal pada jaringan otak yang mengganggu fungsi sistem saraf. Hal ini umumnya terjadi akibat kegagalan mekanisme *apoptosis*, yaitu kemampuan alami sel untuk menghentikan diri saat sudah tidak dibutuhkan. Akibatnya, sel-sel yang rusak terus berkembang dan membentuk massa tumor yang berpotensi memengaruhi kesehatan fisik maupun mental penderita (Essianda et al., 2023).

Berdasarkan sifat keganasannya, tumor otak terbagi menjadi dua kelompok besar, yaitu tumor ganas (*malignant*) dan tumor jinak (*benign*). Tumor ganas bersifat agresif, cepat berkembang, serta berpotensi menyebar ke jaringan lain sehingga sering disebut kanker. Sebaliknya, tumor jinak cenderung tumbuh lambat dan jarang menyebar, namun tetap dapat menimbulkan efek karena menekan jaringan otak di sekitarnya. Contoh tumor ganas meliputi *glioblastoma*, *astrocytoma*, dan *medulloblastoma*, sedangkan *meningioma*, *adenoma pituitari*, dan *craniopharyngioma* termasuk jenis jinak (Essianda et al., 2023)

Sebagian besar faktor penyebab tumor otak bersifat *sporadic* atau secara acak, jarang, tidak berpola, dan biasanya bukan karena faktor keturunan. Namun ada pula yang terkait dengan kelainan genetik atau sindrom familial tertentu, seperti *neurofibromatosis*, *tuberous sclerosis*, dan *von Hippel-Lindau* (Ostrom et al., 2021). Faktor lingkungan seperti paparan radiasi juga dapat menjadi pemicu

walaupun jarang ditemukan. Kondisi ini kemudian menimbulkan manifestasi gejala yang beragam, mulai dari sakit kepala kronis, mual, muntah, hingga kejang.

Dari gejala yang mungkin timbul, pada dasarnya hanya memberikan indikasi awal dan belum cukup untuk memastikan diagnosis secara tepat. Oleh karena itu, diperlukan metode pencitraan medis yang mampu menggambarkan struktur otak secara rinci, terutama jaringan lunak yang kurang terlihat pada *CT scan*. *Magnetic Resonance Imaging* (MRI) menjadi modalitas utama karena mampu menampilkan kontras jaringan dengan detail tinggi serta membedakan jaringan normal dan abnormal secara akurat (Suta et al., 2019). Keunggulan ini menjadikan MRI banyak digunakan dalam identifikasi dan klasifikasi tumor otak. Seiring perkembangan teknologi, MRI tidak hanya dimanfaatkan untuk diagnosis visual, tetapi juga dikombinasikan dengan berbagai metode pengolahan citra dan teknik komputasi modern untuk mendukung analisis medis yang lebih cepat dan akurat (Husnal et al., 2023).

Berdasarkan hasil penelitian Amaliah Faradibah et al. (2023) yang membandingkan beberapa algoritma klasifikasi, terlihat bahwa setiap metode memiliki keunggulan tersendiri dalam aspek evaluasi model. Algoritma SVM memberikan akurasi tertinggi, sementara ANN unggul dalam presisi dan *F1-score*. Hal ini menunjukkan bahwa kinerja model sangat dipengaruhi oleh karakteristik data dan pendekatan ekstraksi fitur yang digunakan. Temuan tersebut memperkuat pentingnya pemilihan metode yang sesuai dengan jenis fitur yang diolah, sebagaimana penelitian ini menerapkan GLCM dan MLP untuk memperoleh hasil yang lebih representatif pada citra tekstur MRI.

Dalam penelitiannya, Wahyu Ardiantito S et al. (2023) menunjukkan bahwa pendekatan berbasis analisis tekstur mampu menghasilkan performa klasifikasi yang tinggi pada citra MRI otak. Hasil penelitian tersebut menegaskan pentingnya pemilihan metode ekstraksi fitur yang tepat untuk membedakan pola jaringan otak secara akurat. Meskipun pendekatan yang digunakan telah memberikan hasil yang baik, potensi peningkatan akurasi masih terbuka dengan penerapan metode lain yang dapat menggambarkan keterhubungan antar piksel, seperti GLCM yang digunakan dalam penelitian ini.

Penelitian oleh Citra R et al. (2024) menunjukkan bahwa penerapan *deep learning* berbasis *Convolutional Neural Network* (CNN) mampu memberikan hasil klasifikasi tumor otak yang sangat baik pada citra MRI. Keunggulan model ini terletak pada kemampuannya mengekstraksi fitur secara otomatis dari citra tanpa memerlukan tahap ekstraksi manual. Namun, pendekatan tersebut juga menuntut jumlah data pelatihan yang besar dan sumber daya komputasi tinggi. Oleh karena itu, penelitian ini menggunakan metode GLCM dan MLP sebagai alternatif yang lebih ringan namun tetap efektif dalam mengenali pola tekstur citra otak.

Sejumlah penelitian terdahulu telah dilakukan untuk mengidentifikasi dan mengklasifikasikan tumor otak menggunakan berbagai pendekatan *image processing* dan algoritma kecerdasan buatan. Penelitian-penelitian tersebut umumnya memanfaatkan citra MRI sebagai data utama karena mampu menampilkan struktur jaringan otak dengan kontras tinggi. Variasi metode yang digunakan meliputi teknik ekstraksi fitur seperti *Gray Level Co-occurrence Matrix* (GLCM) serta algoritma klasifikasi seperti *Multilayer Perceptron* (MLP).

Ringkasan beberapa penelitian relevan yang menjadi acuan dalam studi ini disajikan pada Tabel 2.1 sebagai bahan rujukan.

Tabel 2. 1 Penelitian Terdahulu

No	Peneliti	Judul	Objek	Metode	Hasil	Perbedaan Peneliti
1.	(Amaliah Faradibah et al., 2023)	<i>Comparison Analysis of Random Forest Classifier, Support Vector Machine, and Artificial Neural Network Performance in Multiclass Brain Tumor Classification</i>	Citra MRI otak (glioma, meningioma, pituitary)	Segmentasi <i>Canny</i> , ekstraksi fitur <i>Hu Moments</i> , klasifikasi dengan RF, SVM, ANN	SVM akurasi tertinggi 71%; ANN unggul di presisi & <i>F1-score</i>	Metode ekstraksi GLCM dan algoritma klasifikasi MLP.
2.	(Wahyu Ardiantito S et al., 2023)	Klasifikasi Tumor Otak Menggunakan <i>Local Binary Pattern</i> dan SVM <i>Classifier</i>	Dataset MRI otak (3 kelas)	Ekstraksi fitur LBP, klasifikasi dengan SVM (kernel RBF)	Akurasi 88%, presisi 86%, <i>recall</i> 87%; kombinasi LBP+SVM efektif untuk diagnosis dini	Metode ekstraksi GLCM dan algoritma klasifikasi MLP.
3.	(Citra R et al., 2024)	Klasifikasi Tumor Otak Berbasis <i>Magnetic Resonance Imaging</i> Menggunakan Algoritma <i>Convolutional Neural Network</i>	Dataset MRI otak (4 kelas: glioma, meningioma, pituitary, no tumor)	CNN (<i>AlexNet</i>) dengan dataset 7.022 citra MRI (4 kelas)	Akurasi 95%, presisi 86%, <i>recall</i> 90%, <i>F1-score</i> 88%; CNN unggul dalam klasifikasi otomatis	Model klasifikasi MLP berbasis fitur GLCM.
4.	(Al Haris et al., 2025)	<i>Multilayer Perceptron for Brain Tumor Picture Classification Based on GLCM Feature Extraction</i>	Dataset MRI otak (7023 citra, 4 kelas)	Pra-pemrosesan (<i>resize, scaling, normalization</i>), ekstraksi fitur GLCM, klasifikasi dengan MLP	Akurasi total 82,99%; per kelas: glioma 80,09%, meningioma 63,92%, pituitary 83,60%, tidak tumor 83,45%;	Dataset dan konfigurasi MLP.

					AUC tertinggi <i>pituitary</i> 0,88	
5.	(Fikriah et al., 2024b)	Klasifikasi Hasil MRI Tumor Otak dengan Ekstraksi Fitur GLCM	Citra MRI otak (dataset Kaggle: glioma, meningioma, pituitary, dan no-tumor)	Ekstraksi fitur GLCM (0°, 45°, 90°, 135°), klasifikasi dengan <i>Naïve Bayes</i> , C4.5, <i>Neural Network</i>	<i>Naïve Bayes</i> mencapai akurasi tertinggi 96,8%	Penerapan algoritma klasifikasi MLP.
6.	(Aggarwal, 2022)	<i>Learning Texture Features from GLCM for Classification of Brain Tumor MRI Images using Random Forest Classifier</i>	Citra MRI otak (245 citra: 154 dengan tumor, 91 tanpa tumor)	Ekstraksi fitur GLCM, klasifikasi dengan <i>Random Forest</i>	GLCM berhasil mengekstraksi ciri tekstur signifikan, menghasilkan akurasi tinggi pada klasifikasi tumor otak	Metode klasifikasi MLP.
7.	(Srivaishnavi et al., 2025)	<i>Feature Extraction from Segmented Brain MRI Images using Traditional Methods and Classification Analysis of Brain Tumor using ML Classifiers</i>	Citra MRI otak hasil segmentasi menggunakan model MACB	Segmentasi MACB, ekstraksi fitur (GLCM, <i>Wavelet</i> , LBP, <i>Shape</i> , <i>Intensity</i>), klasifikasi dengan ML	<i>Decision Tree</i> mencapai akurasi tertinggi 96,05%	Metode ekstraksi fitur GLCM dan konfigurasi MLP.
8.	(Alsalihi et al., 2022)	<i>GLCM and CNN Deep Learning Model for Improved MRI Breast Tumors Detection</i>	Citra MRI payudara (89 pasien dari <i>The Cancer Imaging Archive</i>)	Ekstraksi fitur GLCM + CNN, seleksi fitur ANOVA, klasifikasi CNN	Model gabungan GLCM+CNN mencapai akurasi 98,4%, lebih tinggi dari metode Tunggal	Objek penelitian citra MRI otak dan algoritma klasifikasi MLP.

9.	(Kohsasih et al., 2022)	<i>A deep learning model to detect the brain tumor based on magnetic resonance images</i>	Citra MRI otak dengan berbagai jenis tumor	CNN berbasis deep learning	Akurasi tinggi dalam mendeteksi dan mengklasifikasi tumor otak dari citra MRI	Model klasifikasi MLP berbasis GLCM.
10.	(Saeedi et al., 2023)	<i>MRI-based brain tumor detection using convolutional deep learning methods and chosen machine learning techniques</i>	Dataset MRI T1-weighted sebanyak 3264 citra (glioma, meningioma, pituitari, otak sehat)	CNN 2D, Auto-encoder, ML (MLP, KNN, SVM, dll.)	CNN 2D akurasi 96,47% (AUROC ~0,99); MLP akurasi terendah 28%	Fokus penelitian pada optimalisasi MLP.
11.	(Istianah & Sumarti, 2020)	<i>Classification of Pneumonia in Thoracic X-Ray Images Based on Texture Characteristics Using the MLP (Multi-Layer Perceptron) Method</i>	Citra X-ray dada pasien pneumonia (fungal, bacterial, dan lipoid pneumonia)	Pre-processing (cropping, resizing, grayscale), ekstraksi fitur Histogram & GLCM (energy, contrast, correlation, homogeneity), klasifikasi dengan MLP (WEKA)	Akurasi, sensitivitas, dan spesifisitas mencapai 100%; fitur paling berpengaruh yaitu correlation, energy, dan homogeneity	Berbeda pada metode ekstraksi fitur (Histogram + GLCM) dan objek penelitian (citra X-ray pneumonia)
12.	Penelitian ini	Klasifikasi Tumor Otak Berdasarkan Citra MRI Menggunakan Gray Level Co-occurrence Matrix dan Multilayer Perceptron	Dataset MRI otak Br35H (2 kelas: Tumor dan Tidak Tumor)	Pre-processing citra, ekstraksi fitur tekstur GLCM (contrast, correlation, energy, homogeneity), klasifikasi menggunakan MLP dengan variasi jumlah hidden layer, neuron, dan batch size	Model terbaik (3 hidden layer 256–128–64, batch size 32) mencapai akurasi 93,67%	Fokus pada analisis pengaruh kedalaman jaringan dan jumlah neuron MLP terhadap kinerja klasifikasi berbasis fitur GLCM pada skenario biner Tumor dan Tidak Tumor

2.2 Pre-processing Data

Pre-processing data merupakan tahap awal yang sangat penting dalam proses *machine learning* karena menentukan kualitas informasi yang akan dipelajari oleh model. Data mentah yang diperoleh dari berbagai sumber sering kali mengandung permasalahan seperti *noise*, data hilang (*missing values*), duplikasi, dan inkonsistensi format. Tanpa dilakukan proses pra-pemrosesan yang tepat, kualitas model dapat menurun secara signifikan karena data yang buruk akan menghasilkan hasil klasifikasi yang keliru (Maharana et al., 2022). Oleh karena itu, *pre-processing* bertujuan untuk menyiapkan data agar berada dalam bentuk yang bersih, seragam, dan dapat diolah secara efisien oleh algoritma pembelajaran mesin.

Menurut Maharana et al. (2022), proses *pre-processing* meliputi serangkaian langkah seperti *data cleaning*, *data integration*, *data transformation*, dan *data reduction* yang berfungsi untuk memperbaiki kualitas data dan mengurangi kompleksitas pemrosesan. Pada tahap *data cleaning*, data yang tidak relevan atau duplikat dihapus, sedangkan kesalahan struktural diperbaiki untuk menjaga konsistensi antar-atribut. Setelah itu, *data integration* dilakukan untuk menggabungkan beberapa sumber data menjadi satu format yang koheren, diikuti dengan *data transformation* untuk mengubah skala atau distribusi data agar sesuai dengan kebutuhan model, misalnya melalui proses *normalization* atau *standardization*.

Selain itu, Dagal et al. (2025) menekankan bahwa *pre-processing* tidak hanya berperan dalam pembersihan data, tetapi juga menjadi dasar bagi analisis visual dan peningkatan performa klasifikasi. Melalui tahap ini, data disusun ulang

agar lebih mudah divisualisasikan dan diinterpretasikan oleh manusia maupun mesin. Dalam penelitian tersebut, fungsi seperti *center function* digunakan untuk melakukan penyesuaian nilai rata-rata data menjadi nol, yang membantu meningkatkan *separability* antar kelas serta mempercepat konvergensi model. Dengan melakukan *pre-processing* yang tepat, proses pembelajaran model menjadi lebih stabil dan hasil klasifikasi lebih akurat.

Tahapan ini juga berfungsi untuk mengurangi efek *overfitting* dengan menyeimbangkan distribusi data dan menyingkirkan *outlier*. Dagal et al. (2025) menjelaskan bahwa *pre-processing* yang sistematis dapat memperbaiki representasi data, meminimalkan ketidakseimbangan kelas, serta meningkatkan efektivitas *sampling* dalam tahap pelatihan. Oleh karena itu, *pre-processing* bukan hanya langkah teknis untuk menyiapkan *dataset*, tetapi juga merupakan strategi penting untuk memastikan model mampu mengenali pola yang benar dan menghasilkan prediksi yang dapat diandalkan.

Dengan demikian, *pre-processing data* menjadi fondasi utama sebelum dilakukan proses lanjut seperti *feature extraction* atau klasifikasi. Proses ini menjamin bahwa data yang digunakan bersih, representatif, serta berada dalam format yang sesuai dengan kebutuhan algoritma, sehingga model pembelajaran mesin dapat mencapai performa optimal (Dagal et al., 2025; Maharana et al., 2022).

2.2.1 Resize

Tahap *resize* merupakan proses penyesuaian ukuran citra agar memiliki dimensi yang seragam sebelum dilakukan pelatihan model *deep learning*. Penyeragaman ukuran ini penting untuk menyesuaikan input terhadap arsitektur

jaringan saraf yang biasanya memerlukan ukuran piksel tetap. Dengan melakukan *resize*, data citra menjadi lebih efisien diproses dan membantu mempercepat proses pelatihan tanpa mengubah karakteristik utama objek di dalam gambar.

Secara matematis, proses *resize* dengan interpolasi bilinear ditunjukkan pada Persamaan 2.1. Nilai intensitas piksel baru $I'(x', y')$ diperoleh dari kombinasi berbobot empat piksel tetangga terdekat pada citra asal. Bobot w_{ij} merepresentasikan jarak relatif antara posisi piksel baru terhadap posisi piksel asal sebagaimana pada Persamaan 2.2.

Persamaan 2.1 menentukan nilai piksel hasil *resize*:

$$I'(x', y') = \sum_{i=0}^1 \sum_{j=0}^1 w_{ij} I(x + i, y + j) \quad (2.1)$$

Keterangan:

$I'(x', y')$: nilai intensitas piksel hasil *resize* pada koordinat baru (x', y') .
 $I(x + i, y + j)$: nilai intensitas piksel pada koordinat $(x + i, y + j)$ di citra asal.

Bobot dihitung menggunakan Persamaan (2.2):

$$w_{ij} = (1 - |x' - (x + i)|) (1 - |y' - (y + j)|) \quad (2.2)$$

Keterangan:

w_{ij} : bobot kontribusi tiap piksel tetangga terhadap piksel hasil *resize*.
 $x' - (x + i), y' - (y + j)$: jarak relatif piksel hasil *resize* terhadap posisi piksel asal.

Persamaan ini memastikan piksel baru merupakan interpolasi halus dari empat piksel di sekitarnya.

Menurut Hammad et al. (2023), proses *resize* pada citra kulit dilakukan untuk memastikan setiap gambar memiliki resolusi konsisten sehingga dapat diterapkan langsung ke dalam *Convolutional Neural Network (CNN)*. Ukuran citra yang tidak seragam dapat menyebabkan inkonsistensi pada lapisan konvolusi, yang berpotensi menurunkan akurasi model. Peneliti menggunakan metode *bilinear*

interpolation untuk mempertahankan ketajaman detail pada area lesi kulit, karena distorsi ukuran dapat menyebabkan hilangnya fitur penting yang menjadi pembeda antara eksim dan psoriasis.

Sementara itu, Oybek Kizi et al. (2025) menegaskan bahwa proses *resize* juga merupakan tahap penting dalam *pre-processing* citra darah mikroskopik sebelum dilakukan klasifikasi leukemia. Ukuran citra standar seperti 224×224 piksel, 227×227 piksel, dan 300×300 piksel digunakan agar kompatibel dengan berbagai arsitektur *deep learning* seperti *AlexNet*, *ResNet*, dan *Inception*. Selain itu, proses ini turut mengurangi kompleksitas komputasi tanpa mengorbankan informasi penting dari bentuk dan struktur sel darah. Beberapa studi yang mereka tinjau menunjukkan bahwa *resize* yang proporsional dapat membantu menyeimbangkan antara akurasi deteksi dan kecepatan inferensi model.

Selain menjaga keseragaman dimensi, tahap *resize* juga berfungsi sebagai bagian dari proses *data normalization pipeline* untuk menstabilkan ukuran batch dan meminimalkan *overfitting* ketika data latih memiliki variasi resolusi tinggi. Menurut kedua penelitian tersebut, langkah ini terbukti efektif dalam meningkatkan efisiensi komputasi dan memperkuat konsistensi performa model pada berbagai dataset citra medis (Hammad et al., 2023; Oybek Kizi et al., 2025).

2.2.2 Normalisasi

Proses normalisasi merupakan tahap penting dalam *pre-processing data* citra karena berfungsi untuk menyesuaikan skala intensitas piksel agar berada pada rentang nilai tertentu, biasanya antara 0 hingga 1 atau -1 hingga 1. Tujuan utamanya adalah memastikan bahwa setiap piksel memiliki kontribusi yang proporsional

terhadap proses pembelajaran model, sehingga mempercepat konvergensi dan menghindari dominasi nilai piksel tertentu pada tahap pelatihan jaringan saraf.

Untuk menyeragamkan skala intensitas antar citra, digunakan metode *min-max normalization* sebagaimana Persamaan 2.3. Persamaan ini mengubah setiap nilai piksel $I(x, y)$ menjadi nilai baru $I_{norm}(x, y)$ pada rentang 0–1.

$$I_{norm}(x, y) = \frac{I(x, y) - I_{min}}{I_{max} - I_{min}} \quad (2.3)$$

Keterangan:

$I_{norm}(x, y)$: nilai intensitas piksel setelah normalisasi.
 $I(x, y)$: nilai intensitas asli piksel pada citra sebelum normalisasi.
 I_{min} : nilai intensitas minimum dalam citra.
 I_{max} : nilai intensitas maksimum dalam citra.
 255 : nilai maksimum intensitas untuk citra 8-bit (skala 0–255).

Untuk citra 8-bit dengan rentang 0–255, bentuk sederhananya menjadi Persamaan 2.4:

$$I_{norm}(x, y) = \frac{I(x, y)}{255} \quad (2.4)$$

Keterangan:

$I_{norm}(x, y)$: nilai intensitas piksel setelah normalisasi.
 $I(x, y)$: nilai intensitas asli piksel pada citra sebelum normalisasi.
 255 : nilai maksimum intensitas untuk citra 8-bit (skala 0–255).

yang berarti setiap piksel dibagi oleh nilai maksimum 255 agar seluruh nilai berada dalam interval $[0, 1]$.

Menurut Albahli (2025), normalisasi dilakukan untuk menstandarkan representasi warna dan pencahayaan antar dataset yang berbeda, seperti ISIC 2020, HAM10000, dan PH2, guna mengurangi perbedaan distribusi domain pada citra kulit. Hal ini sangat penting karena variasi intensitas dan warna yang tidak terkendali dapat menurunkan kemampuan model dalam mengenali pola lesi secara konsisten.

Selain itu, normalisasi juga membantu menjaga kestabilan numerik dalam proses propagasi maju dan mundur pada jaringan *deep learning*. Dalam penelitian Albahli (2025), setiap citra hasil *pre-processing* diubah ke format 512×512 piksel dan nilai intensitas pikselnya dinormalisasi ke dalam rentang $[0, 1]$. Langkah ini dilakukan untuk memastikan bahwa perhitungan gradien selama pelatihan berjalan lebih efisien dan tidak menyebabkan *vanishing gradient problem*. Lebih lanjut, Albahli menambahkan bahwa proses *histogram matching* dilakukan setelah normalisasi guna menyeimbangkan distribusi warna antar dataset, sehingga model tidak bias terhadap sumber data tertentu.

Sementara itu, dalam kajian sistematis oleh Alayed (2024) mengenai pengenalan Bahasa Isyarat Arab, proses normalisasi juga menjadi tahap penting sebelum fitur diekstraksi dan diklasifikasikan. Banyak penelitian terdahulu yang menormalkan intensitas piksel citra tangan untuk mengatasi perbedaan pencahayaan dan kontras antar gambar. Teknik seperti *min-max normalization* dan *z-score normalization* digunakan untuk menjaga keseragaman data input, sehingga model dapat mempelajari representasi visual dengan lebih stabil. Alayed (2024) menekankan bahwa normalisasi menjadi fondasi bagi keberhasilan tahap selanjutnya, seperti ekstraksi fitur tekstur dan pelatihan model *deep learning*, karena dapat meningkatkan konvergensi serta akurasi pengenalan gerakan tangan.

Dengan demikian, normalisasi berperan penting dalam standarisasi nilai piksel dan peningkatan performa model, karena mampu mengurangi variasi pencahayaan dan warna serta meningkatkan generalisasi model terhadap berbagai kondisi citra (Alayed, 2024; Albahli, 2025).

2.2.3 Noise Removal

Proses *noise removal* merupakan langkah penting dalam *pre-processing* citra medis untuk meningkatkan kualitas visual dan akurasi analisis pada tahap selanjutnya. *Noise* atau derau sering muncul akibat gangguan selama proses perekaman, transmisi, maupun konversi sinyal digital yang menyebabkan distorsi intensitas pada citra. Dalam domain medis, jenis *noise* yang umum dijumpai meliputi *Gaussian noise*, *Rician noise*, dan *impulsive noise* seperti *salt and pepper noise*. Keberadaan gangguan tersebut dapat menutupi detail anatomi penting, mengurangi kontras, serta menurunkan kemampuan model *deep learning* dalam mengenali pola visual secara tepat (Alanazi et al., 2023). Oleh karena itu, penghilangan *noise* menjadi tahap fundamental sebelum dilakukan ekstraksi fitur atau klasifikasi.

Proses *noise removal* dilakukan menggunakan *Gaussian filter* yang merepresentasikan distribusi normal dua dimensi sebagaimana Persamaan 2.5. Fungsi Gaussian $G(x, y)$ memberikan bobot berdasarkan jarak piksel terhadap pusat kernel, diatur oleh parameter σ sebagai simpangan baku.

$$G(x, y) = \frac{1}{2\pi\sigma^2} \exp\left(-\frac{x^2+y^2}{2\sigma^2}\right) \quad (2.5)$$

Keterangan:

$G(x, y)$: fungsi kernel Gaussian dua dimensi.

σ : simpangan baku yang menentukan tingkat perataan filter.

$\exp$: fungsi eksponensial.

π : konstanta pi ($\approx 3,1416$)

Nilai intensitas baru setiap piksel $I'(x, y)$ dihitung melalui konvolusi kernel

Gaussian dengan citra sebagaimana Persamaan 2.6:

$$I'(x, y) = \sum_{i=-k}^k \sum_{j=-k}^k G(i, j) I(x + i, y + j) \quad (2.6)$$

Keterangan:

$I(x + i, y + j)$: nilai intensitas piksel pada posisi $(x + i, y + j)$ sebelum filter diterapkan.

$I'(x, y)$: nilai intensitas piksel setelah dilakukan *filtering*.

k : setengah ukuran kernel (misalnya $k = 1$ untuk kernel 3×3).

Persamaan ini menjelaskan bahwa setiap nilai piksel baru merupakan hasil penjumlahan berbobot dari piksel-piksel di sekitarnya berdasarkan distribusi Gaussian.

Penelitian oleh Alanazi et al. (2023) mengembangkan metode *Nested Filtering followed by Morphological Operation* (NFMO) untuk menangani *high-density impulsive noise* pada citra medis seperti MRI dan X-ray. Metode ini menggabungkan *Modified Laplacian Vector Median Filter* (MLVMF) untuk mendeteksi piksel rusak dengan dua lapisan *nested filtering* dan dilatasi morfologi guna memperkirakan nilai piksel yang hilang. Pendekatan tersebut berhasil mencapai *Peak Signal-to-Noise Ratio* (PSNR) sebesar 29,99 dB pada tingkat *noise* 90%, menunjukkan efektivitasnya dalam memulihkan citra tanpa kehilangan detail penting.

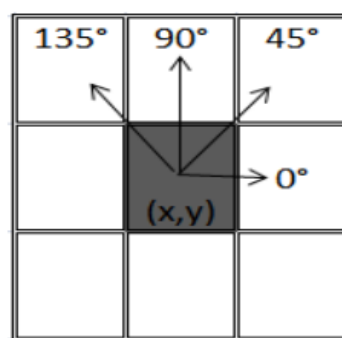
Metode *filtering* konvensional seperti median, *Gaussian*, dan *bilateral filter* memang sering digunakan untuk tujuan serupa, namun efektivitasnya menurun saat menghadapi citra dengan tingkat *noise* yang sangat tinggi. Keunggulan NFMO terletak pada kemampuannya memperkaya informasi lokal melalui investigasi berlapis dan operasi morfologi yang menjaga struktur citra. Upaya *noise removal* semacam ini terbukti berdampak signifikan pada performa sistem deteksi berbasis *deep learning*. Hammad et al. (2023) misalnya, menerapkan proses *filtering* pada

tahap awal sebelum melatih model CNN “Derma Care” untuk memastikan citra kulit bebas *noise* sehingga model dapat belajar dari pola yang bersih dan terfokus.

Secara keseluruhan, *noise removal* tidak hanya memperbaiki tampilan citra, tetapi juga menjaga kestabilan nilai piksel agar algoritma analisis citra bekerja lebih optimal. Penerapan metode adaptif seperti *NFMO* mampu mengurangi gangguan visual sekaligus mempertahankan informasi diagnostik penting pada citra medis. (Alanazi et al., 2023; Hammad et al., 2023).

2.3 Gray Level Co-occurrence Matrix

Metode *Gray Level Co-occurrence Matrix* (GLCM) digunakan untuk mengekstraksi fitur tekstur dengan menganalisis keterkaitan spasial antar piksel. Ullu et al. (2022) menjelaskan bahwa GLCM menghitung seberapa sering pasangan intensitas keabuan tertentu muncul secara bersamaan pada jarak dan orientasi yang telah ditentukan, seperti 0° , 45° , 90° , dan 135° . Representasi arah sudut dapat dilihat pada Gambar 2.1.



Gambar 2. 1 Arah orientasi sudut pada GLCM

Tahap awal dimulai dengan mengubah citra ke dalam format *grayscale*. Selanjutnya, dihitung frekuensi kemunculan pasangan piksel sehingga terbentuk *matrix co-occurrence*. Berdasarkan matriks ini, dilakukan perhitungan nilai abu-

abu antar pasangan piksel, kemudian diekstraksi fitur tekstur yang dihasilkan (Siregar, 2023). Ilustrasi pembentukan matriks *co-occurrence* Gambar 2.2

1	1	0	0
0	0	1	0
0	0	0	2
0	0	1	0

Gambar 2. 2 Matriks *co-occurrence* pada GLCM

Normalisasi dilakukan setelah matriks *co-occurrence* diperoleh, dengan tujuan mengonversi setiap nilai matriks menjadi probabilitas kemunculan pasangan piksel. Proses ini dilakukan melalui pembagian setiap elemen matriks dengan jumlah total pasangan piksel sesuai Persamaan 2.7.

$$P(i,j) = \frac{C(i,j)}{\sum_i \sum_j C(i,j)} \quad (2.7)$$

$P(i,j)$: probabilitas kemunculan pasangan piksel.
 $C(i,j)$: frekuensi pasangan piksel pada matriks *co-occurrence*.

Fitur tekstur yang diperkenalkan oleh (Haralick et al., 1973) diekstraksi dari matriks *Gray Level Co-occurrence (GLCM)* untuk merepresentasikan karakteristik tekstur citra secara kuantitatif. Fitur utama yang umum digunakan meliputi *contrast*, *correlation*, *energy*, dan *homogeneity* (Kabir et al., 2024). Adapun persamaan untuk menghitung nilai fitur tekstur berdasarkan matriks GLCM dapat dituliskan sebagai berikut.

1. *Contrast*

Contrast mengukur perbedaan intensitas antar piksel pada citra, di mana nilai *contrast* yang tinggi menunjukkan perbedaan intensitas yang tajam antar piksel, sebagaimana pada Persamaan 2.8:

$$\text{Contrast} = \sum_{i=0}^{N-1} \sum_{j=0}^{N-1} (i - j)^2 P(i, j) \quad (2.8)$$

Keterangan:

i, j : indeks level keabuan untuk baris (i) dan kolom (j) dalam matriks.
 N : jumlah level keabuan.
 $P(i, j)$: probabilitas kemunculan pasangan piksel.

2. *Correlation*

Correlation mengukur hubungan linear antar piksel, di mana nilai yang mendekati 1 menunjukkan tingkat keterkaitan yang kuat, sebagaimana pada Persamaan 2.9:

$$\text{Correlation} = \frac{\sum_{i=0}^{N-1} \sum_{j=0}^{N-1} (i - \mu_i)(j - \mu_j) P(i, j)}{\sigma_i \sigma_j} \quad (2.9)$$

Keterangan:

i, j : indeks level keabuan untuk baris (i) dan kolom (j) dalam matriks.
 N : jumlah level keabuan.
 $P(i, j)$: probabilitas kemunculan pasangan piksel.
 μ_i, μ_j : rata-rata nilai keabuan pada arah baris (i) dan kolom (j).
 σ_i, σ_j : simpangan baku untuk baris (i) dan kolom (j).

3. *Energy*

Energy mengukur tingkat keteraturan tekstur, di mana nilai yang tinggi menunjukkan pola intensitas yang teratur dan seragam, sebagaimana pada Persamaan 2.10:

$$\text{Energy} = \sum_{i=0}^{N-1} \sum_{j=0}^{N-1} (P(i, j))^2 \quad (2.10)$$

Keterangan:

i, j : indeks level keabuan untuk baris (i) dan kolom (j) dalam matriks.
 N : jumlah level keabuan.
 $P(i, j)$: probabilitas kemunculan pasangan piksel.

4. *Homogeneity*

Homogeneity mengukur tingkat keseragaman tekstur, di mana nilai yang tinggi menunjukkan bahwa intensitas piksel dalam citra semakin seragam, sebagaimana pada Persamaan 2.11:

$$\text{Homogeneity} = \sum_{i=0}^{N-1} \sum_{j=0}^{N-1} \frac{P(i,j)}{1+(i-j)^2} \quad (2.11)$$

Keterangan:

i, j : indeks level keabuan pada baris dan kolom dalam matriks GLCM
 N : jumlah level keabuan citra
 $P(i, j)$: probabilitas kemunculan pasangan piksel dengan nilai keabuan i dan j
 $(i - j)^2$: selisih kuadrat antara dua level keabuan yang merepresentasikan jarak antar Piksel

Dalam bidang medis, GLCM berperan penting sebagai metode ekstraksi fitur tekstur untuk membedakan jaringan patologis dari jaringan normal pada citra diagnostik. Misalnya, pada penelitian (Pattanaik et al., 2022) penggunaan fitur GLCM seperti *contrast*, *energy*, *entropy*, dan *homogeneity* dari MRI berhasil memisahkan area tumor otak dari jaringan sehat secara efektif. Demikian pula, studi yang dilakukan oleh (Vijayalakshmi et al., 2025) menerapkan GLCM (*contrast*, *correlation*, *entropy*, *energy*, *homogeneity*) pada mamogram untuk mendeteksi kanker payudara, yang membuktikan bahwa fitur-tekstur GLCM tetap relevan dalam mendukung diagnosis dini dan klasifikasi lesi jinak vs ganas. Kombinasi fitur GLCM dengan metode pembelajaran mesin atau *radiomics* pada penelitian (Karami et al., 2025) memperkuat kemampuan sistem CAD (*Computer-Aided Diagnosis*) dalam memperbaiki akurasi dan sensitivitas diagnosis berbasis citra medis.

Berdasarkan penelitian-penelitian yang dirangkum pada Tabel 2.1, metode *Gray Level Co-occurrence Matrix* (GLCM) telah terbukti efektif dalam mengekstraksi fitur tekstur citra medis. Fikriah et al. (2024) menunjukkan bahwa

variasi arah sudut dalam pembentukan matriks GLCM mampu menangkap pola tekstur yang kompleks pada citra MRI tumor otak. Hasil penelitian tersebut menegaskan bahwa GLCM dapat meningkatkan akurasi klasifikasi dengan memanfaatkan hubungan spasial antar piksel, sehingga relevan digunakan sebagai metode ekstraksi fitur pada penelitian ini.

Penelitian oleh Aggarwal, (2022) memperkuat efektivitas GLCM dalam membedakan karakteristik tekstur antara citra otak normal dan citra dengan tumor. Hasil analisis menunjukkan bahwa fitur-fitur yang dihasilkan dari GLCM mampu menangkap pola intensitas dan variasi tekstur secara signifikan, sehingga mendukung proses klasifikasi dengan tingkat akurasi yang tinggi. Temuan ini menegaskan bahwa GLCM memiliki kemampuan representatif yang baik dalam mengidentifikasi perbedaan struktural pada citra MRI otak.

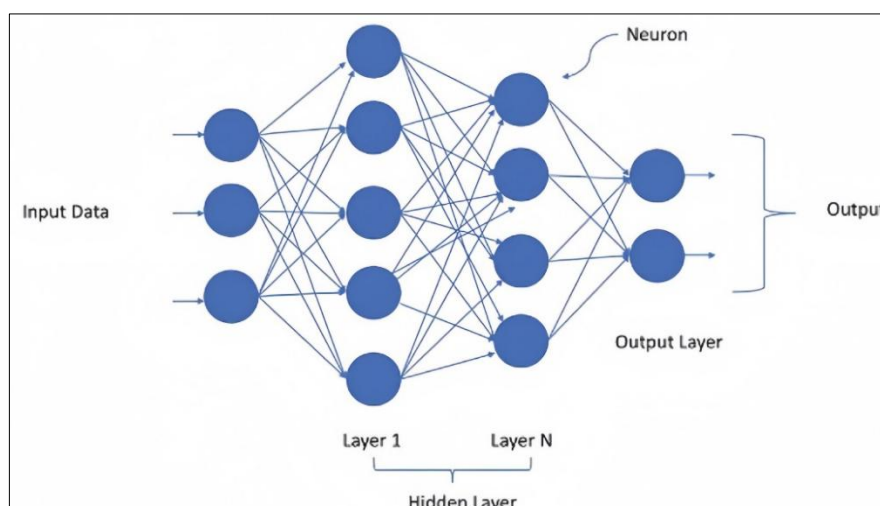
Srivaishnavi et al. (2025) memperluas penerapan GLCM dengan mengombinasikannya bersama metode ekstraksi fitur lain seperti *wavelet* dan *Local Binary Pattern* (LBP). Hasil penelitian menunjukkan bahwa meskipun digunakan bersama berbagai pendekatan tradisional, GLCM tetap memberikan kontribusi signifikan terhadap peningkatan performa klasifikasi. Temuan ini menegaskan bahwa GLCM masih relevan digunakan dalam analisis citra medis modern karena kemampuannya mengekspresikan informasi tekstur yang tidak mudah ditangkap oleh metode lain.

Penelitian oleh Alsalihi et al. (2022) menunjukkan bahwa penggabungan metode klasik seperti GLCM dengan pendekatan *deep learning* mampu meningkatkan performa deteksi citra medis secara signifikan. Integrasi GLCM

dengan *Convolutional Neural Network* (CNN) menghasilkan model yang lebih akurat dibandingkan penggunaan salah satu metode secara tunggal. Temuan ini mengindikasikan bahwa GLCM memiliki fleksibilitas tinggi untuk dikombinasikan dengan berbagai algoritma modern, sekaligus memperkuat perannya sebagai komponen penting dalam tahap ekstraksi fitur pada sistem klasifikasi berbasis citra.

2.4 Multilayer Perceptron

Dalam *Artificial Neural Network* (ANN), *Multilayer Perceptron* (MLP) dikenal sebagai arsitektur dasar yang beroperasi secara *feedforward* dan efektif untuk memodelkan masalah *non-linear* serta *non-deterministic* (Santoso et al., 2020). Arsitektur MLP umumnya terdiri dari *input layer*, *hidden layer*, dan *output layer*. Setiap neuron pada lapisan yang berbeda terhubung secara *fully connected*, sedangkan neuron dalam lapisan yang sama tidak memiliki koneksi langsung (Sildir et al., 2020).



Gambar 2. 3 Arsitektur sederhana MLP

Pada arsitektur MLP, data diproses secara berurutan mulai dari *input layer*, diteruskan melalui beberapa *hidden layer*, hingga mencapai *output layer*. Proses

pembobotan dan penerapan fungsi aktivasi pada setiap neuron memungkinkan jaringan mempelajari pola *non-linear* guna menghasilkan keluaran yang relevan untuk tugas klasifikasi atau prediksi.

1. *Input Layer*, berfungsi sebagai penghubung dengan lingkungan eksternal dan menerima data masukan yang direpresentasikan sebagai vektor fitur :

$$\mathbf{x} = [x_1, x_2, \dots, x_n] \quad (2.12)$$

Keterangan :

$\mathbf{x}$: vektor input.

2. *Hidden layer*, berisi neuron dengan bobot serta fungsi aktivasi non-linear, sehingga jaringan dapat memodelkan pola yang kompleks. Setiap neuron melakukan penjumlahan berbobot terhadap input dan bias seperti pada Persamaan 2.12, kemudian hasilnya diproses melalui fungsi aktivasi sebagaimana ditunjukkan pada Persamaan 2.13 dan 2.14.

$$z_j = \sum_{i=1}^n w_{ij} x_i + b_j \quad (2.13)$$

$$a_j = f(z_j) \quad (2.14)$$

Keterangan:

z_j : hasil penjumlahan berbobot pada neuron ke- j (*weighted sum*)
 w_{ij} : bobot koneksi dari neuron input ke- i menuju neuron ke- j
 x_i : nilai input ke- i
 b_j : nilai bias pada neuron ke- j
 a_j : keluaran (output) dari neuron ke- j setelah fungsi aktivasi
 $f(z_j)$: fungsi aktivasi non-linier (misalnya *sigmoid* atau *ReLU*)
 n : jumlah neuron pada lapisan sebelumnya

Fungsi aktivasi sigmoid mengubah nilai input menjadi rentang antara 0 dan 1, sehingga cocok untuk pemodelan probabilitas. Persamaannya ditunjukkan pada Persamaan 2.15:

$$f(z) = \frac{1}{1+e^{-z}} \quad (2.15)$$

Keterangan:

$f(z)$: nilai keluaran setelah fungsi aktivasi sigmoid diterapkan
 z : hasil penjumlahan berbobot (*weighted sum*) dari neuron
 e : bilangan eksponensial ($\approx 2,718$), basis dari logaritma natural

Sementara itu, fungsi aktivasi ReLU (*Rectified Linear Unit*) hanya meneruskan nilai positif dan mengubah nilai negatif menjadi nol. Fungsi ini mampu mempercepat proses pelatihan jaringan karena mengatasi masalah *vanishing gradient*, sebagaimana ditunjukkan pada Persamaan 2.16:

$$f(z) = \max(0, z) \quad (2.16)$$

Keterangan :

$f(z)$: nilai keluaran setelah fungsi aktivasi ReLU diterapkan
 z : hasil penjumlahan berbobot (*weighted sum*) dari neuron

3. *Output Layer*, menyalurkan pola hasil pemrosesan dari *hidden layer* berupa hasil klasifikasi atau prediksi. Proses ini diawali dengan perhitungan nilai masukan bersih z_k menggunakan Persamaan 2.17, yang kemudian diteruskan ke fungsi aktivasi untuk menghasilkan keluaran akhir y_k sebagaimana pada Persamaan 2.18.

$$z_k = \sum_{j=1}^m w_{jk} h_j + b_k \quad (2.17)$$

$$y_k = f(z_k) \quad (2.18)$$

Keterangan :

z_k : hasil penjumlahan berbobot pada neuron keluaran ke- k (*weighted sum*)
 w_{jk} : bobot koneksi dari neuron tersembunyi ke- j menuju neuron keluaran ke- k
 h_j : keluaran dari neuron tersembunyi ke- j (*hidden layer output*)
 b_k : bias pada neuron keluaran ke- k
 y_k : nilai keluaran akhir dari neuron ke- k setelah fungsi aktivasi diterapkan
 m : jumlah neuron pada lapisan tersembunyi
 $f(z_k)$: fungsi aktivasi pada lapisan keluaran (misalnya *sigmoid* atau *softmax*)

Dalam jaringan saraf tiruan jenis *Multilayer Perceptron* (MLP), prinsip kerja dimulai dari propagasi maju (*forward propagation*), di mana data masukan dihitung melalui bobot dan bias, kemudian ditransformasikan oleh neuron-neuron pada lapisan tersembunyi. Proses ini diperkuat oleh penggunaan fungsi aktivasi non-linear, seperti sigmoid atau ReLU, yang memungkinkan jaringan memodelkan pola kompleks dan non-linier. Setelah *output* diperoleh, dilakukan perbandingan dengan target menggunakan fungsi kerugian untuk menghitung error. Selanjutnya, propagasi mundur (*backpropagation*) digunakan untuk menyebarkan *error* tersebut kembali ke lapisan-lapisan sebelumnya, sehingga bobot dan bias diperbarui melalui algoritma optimasi berbasis gradien. Mekanisme ini memungkinkan MLP belajar secara bertahap dan meningkatkan akurasi klasifikasi maupun prediksi (Waldo, 2022).

Multilayer Perceptron (MLP) memiliki sejumlah kelebihan dalam tugas klasifikasi, di antaranya kemampuan untuk menangani data yang bersifat non-linear melalui penggunaan fungsi aktivasi non-linear pada lapisan tersembunyinya. Selain itu, fleksibilitas arsitektur MLP memungkinkan penyesuaian jumlah lapisan maupun neuron sesuai dengan kompleksitas permasalahan, sehingga jaringan ini dapat menggeneralisasi pola data yang beragam. Namun, MLP juga memiliki keterbatasan, seperti kecenderungan mengalami *overfitting* ketika jumlah parameter terlalu besar dibandingkan dengan data latih yang tersedia, serta kebutuhan akan data dalam jumlah besar agar proses pembelajaran menghasilkan model yang stabil dan akurat. Hal ini menjadikan MLP memerlukan perhatian

khusus pada aspek regulasi, jumlah data, dan teknik optimisasi agar dapat mencapai performa klasifikasi yang optimal (Asif Khan et al., 2021).

Berdasarkan penelitian-penelitian yang telah dirangkum pada Tabel 2.1, algoritma *Multilayer Perceptron* (MLP) menjadi salah satu pendekatan yang banyak digunakan dalam klasifikasi citra medis, khususnya citra MRI otak. MLP termasuk ke dalam keluarga *Artificial Neural Network* (ANN) yang mampu mempelajari pola non-linear melalui lapisan tersembunyi (*hidden layer*). Sejumlah penelitian terdahulu telah membuktikan efektivitas MLP dan arsitektur jaringan saraf lainnya, seperti CNN, dalam mendeteksi serta mengklasifikasikan tumor otak.

Penelitian oleh Kohsasih et al. (2022) menunjukkan efektivitas pendekatan *deep learning* dalam mendeteksi tumor otak berbasis citra MRI. Model yang dikembangkan menggunakan arsitektur CNN untuk mengenali pola visual dan perbedaan struktur jaringan otak secara otomatis. Hasil penelitian ini memperlihatkan bahwa jaringan saraf tiruan memiliki kemampuan yang kuat dalam mengekstraksi ciri dan melakukan klasifikasi citra medis. Temuan tersebut menjadi dasar penting bagi pengembangan model berbasis *Artificial Neural Network* (ANN), termasuk MLP, yang mampu menangani pola non-linear pada data citra.

(Saeedi et al., 2023) melakukan perbandingan beberapa algoritma klasifikasi berbasis citra MRI, termasuk MLP, untuk mendeteksi jenis tumor otak. Dari hasil analisis, metode berbasis CNN menunjukkan performa terbaik, sedangkan MLP masih menghasilkan akurasi yang relatif rendah. Kondisi ini menunjukkan bahwa efektivitas MLP sangat dipengaruhi oleh desain arsitektur dan

parameter pelatihan yang digunakan, sehingga diperlukan optimasi lebih lanjut agar kinerjanya dapat meningkat.

Al Haris et al. (2025) membuktikan bahwa integrasi metode GLCM dan algoritma MLP dapat menghasilkan klasifikasi citra MRI tumor otak dengan tingkat ketepatan yang baik. Penelitian tersebut menunjukkan bahwa kemampuan MLP dalam mengenali pola tekstur sangat dipengaruhi oleh karakteristik fitur yang dihasilkan dari GLCM. Temuan ini menegaskan pentingnya sinergi antara proses ekstraksi fitur dan model klasifikasi dalam analisis citra medis.

Istianah & Sumarti, (2020) mengemukakan bahwa pendekatan berbasis analisis tekstur menggunakan GLCM yang dipadukan dengan algoritma MLP efektif dalam mengidentifikasi jenis pneumonia dari citra X-ray. Melalui pemanfaatan jaringan saraf tiruan, penelitian ini membuktikan bahwa pola intensitas dan keteraturan piksel pada citra mampu menjadi penentu penting dalam proses klasifikasi. Hasilnya menegaskan bahwa penerapan pembelajaran mesin berbasis fitur tekstur dapat meningkatkan keakuratan sistem diagnosis otomatis pada bidang radiologi.

2.5 Confusion matrix

Evaluasi kinerja model klasifikasi dilakukan dengan memanfaatkan *confusion matrix*, yang membandingkan keluaran prediksi model dengan label aktual atau *ground truth* (Müller et al., 2022). Matriks ini menampilkan distribusi prediksi benar dan salah dalam format tabel, sehingga interpretasi performa model menjadi lebih sistematis.

Dalam skenario klasifikasi biner, *confusion matrix* memiliki struktur tabel 2×2 yang memuat empat elemen, yakni True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN) (Rainio et al., 2024).

1. *True positive* (TP) adalah kondisi ketika model berhasil mengenali data positif, seperti citra MRI bertumor, sebagai tumor.
2. *True negative* (TN) adalah kondisi ketika data negatif, misalnya citra MRI pasien sehat, diprediksi dengan benar sebagai tidak tumor.
3. *False positive* (FP) merupakan kesalahan prediksi pada data negatif yang diklasifikasikan sebagai positif.
4. *False negative* (FN) merupakan kesalahan prediksi pada data positif yang diklasifikasikan sebagai negatif, sehingga tumor tidak terdeteksi.

Dari keempat komponen *confusion matrix*, dapat dihitung beberapa metrik evaluasi performa model, yaitu *accuracy*, *precision*, *recall*, dan *F1-score* (Müller et al., 2022).

1. *Accuracy* (Akurasi)

Akurasi digunakan untuk menilai proporsi prediksi yang tepat terhadap seluruh data uji (Rainio et al., 2024), sebagaimana pada persamaan 2.19:

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (2.19)$$

Akurasi menunjukkan tingkat kebenaran keseluruhan prediksi model, namun kurang representatif pada distribusi data yang tidak seimbang.

2. *Precision* (Presisi)

Presisi menunjukkan ketepatan model dalam mengidentifikasi data positif dari seluruh prediksi (Jeong et al., 2025), sebagaimana pada persamaan 2.20:

$$Precision = \frac{TP}{TP+FP} \quad (2.20)$$

Nilai presisi yang tinggi menunjukkan rendahnya kesalahan *false positive*.

3. *Recall* (Sensitivitas / *True positive Rate*)

Recall mengukur kemampuan model dalam mendeteksi data positif secara benar (Müller et al., 2022), sebagaimana pada persamaan 2.21:

$$Recall = \frac{TP}{TP+FN} \quad (2.21)$$

Nilai *recall* yang tinggi menunjukkan minimnya kesalahan *false negative*.

4. *F1-score*

F1-score digunakan untuk menilai keseimbangan antara presisi dan *recall* melalui rata-rata harmonik keduanya (Santamato & Marengo, 2025), sebagaimana pada persamaan 2.22:

$$F1-Score = \frac{2 \times Precision \times Recall}{Precision + Recall} \quad (2.22)$$

Nilai *F1-score* mendekati 1 menunjukkan keseimbangan yang baik antara presisi dan *recall*.

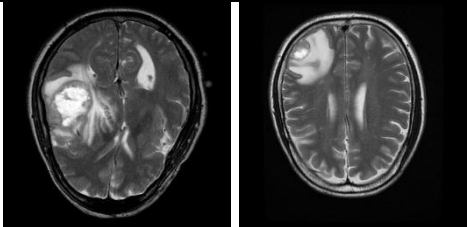
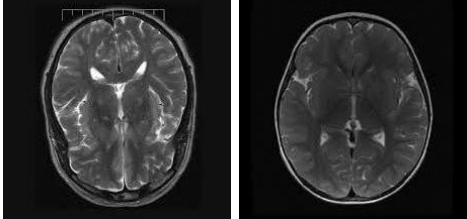
BAB III

DESAIN DAN IMPLEMENTASI

3.1 Pengumpulan Data

Data yang digunakan ini merupakan data sekunder dari dataset *Br35H :: Brain Tumor Detection 2020* (Hamada, 2020) yang diperoleh melalui *platform Kaggle*. Dataset ini diakses dan diunduh pada 1 September 2025 dan terdiri atas 3.000 citra MRI otak berformat JPG yang terbagi secara seimbang menjadi dua kategori, yaitu 1.500 citra bertumor (folder *yes*) dan 1.500 citra tidak tumor (folder *no*). Seluruh citra telah terlabel dengan jelas sehingga memudahkan proses ekstraksi fitur dan klasifikasi pada tahap penelitian.

Tabel 3. 1 Contoh Data Citra MRI Tumor Otak

No	Label	Citra MRI	
1.	Tumor		
2.	Tidak tumor		

Citra MRI diklasifikasikan ke dalam dua kategori, yakni tumor dan tidak tumor, yang selanjutnya dimanfaatkan sebagai data masukan dalam proses *pre-*

processing, ekstraksi fitur *Gray Level Co-occurrence Matrix* (GLCM), dan pelatihan model klasifikasi berbasis *Multilayer Perceptron* (MLP).

3.2 Desain Sistem

Desain sistem merupakan tahapan perancangan yang bertujuan untuk menggambarkan bagaimana suatu sistem bekerja dalam memproses data masukan hingga menghasilkan keluaran yang sesuai dengan kebutuhan pengguna.



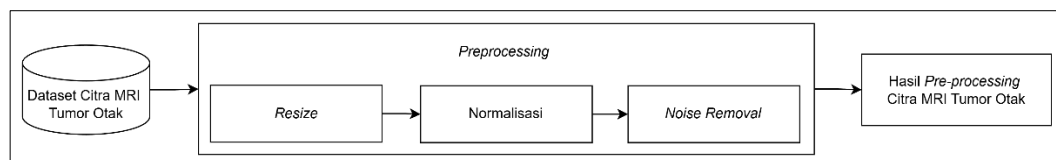
Gambar 3. 1 Desain Sistem

Gambar 3.1 desain sistem dimana alur kerja dimulai dari proses masukan citra MRI, tahap *pre-processing* dan ekstraksi fitur, hingga proses klasifikasi menggunakan model yang telah dilatih untuk menghasilkan keluaran berupa status tumor pada citra.

3.3 Pre-processing Data

Pengolahan awal data (*pre-processing*) bertujuan menyiapkan data mentah agar memiliki kualitas dan konsistensi yang memadai untuk analisis lanjutan.

Diagram alur proses tersebut ditampilkan pada Gambar 3.2.



Gambar 3. 2 Flowchart Pre-processing Citra

Alur *pre-processing* tersebut mencakup beberapa proses utama, seperti penyesuaian ukuran citra, normalisasi nilai piksel, serta pengurangan noise,

sehingga citra hasil *pre-processing* memiliki kualitas yang lebih baik dan siap digunakan pada tahap ekstraksi fitur.

3.3.1 *Resize*

Pada tahap ini, seluruh citra diseragamkan ukurannya agar dapat diproses oleh sistem dengan dimensi yang sama. Citra MRI yang digunakan dalam penelitian memiliki ukuran awal 256×197 piksel. Karena citra tersebut tidak berbentuk persegi, maka terlebih dahulu dilakukan *padding* hitam di sisi kiri dan kanan hingga menjadi 256×255 piksel. Setelah itu, citra diubah ukurannya menjadi 224×224 piksel menggunakan metode interpolasi bilinear sebagaimana dijelaskan pada Persamaan 2.1.

Metode ini bekerja dengan memperkirakan nilai intensitas piksel baru berdasarkan kombinasi empat piksel terdekat, dengan bobot yang ditentukan oleh jarak relatif antar koordinat piksel sebagaimana dirumuskan pada Persamaan 2.2.

Sebagai contoh, piksel hasil *resize* pada koordinat $(x', y') = (112, 112)$ diperoleh dari area sekitar $(x, y) = (127, 5, 128)$ pada citra sebelum *resize*. Nilai empat piksel tetangga di sekitar koordinat tersebut adalah:

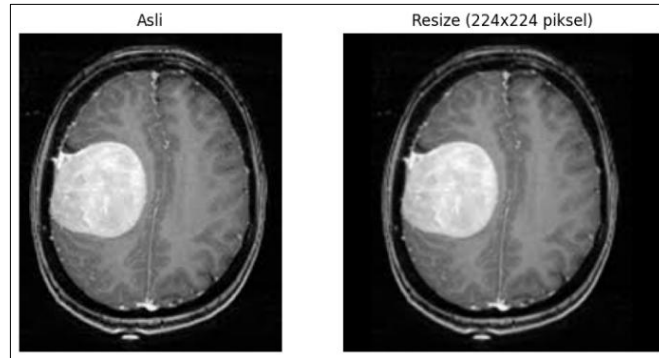
$$I(127, 127) = 253, I(127, 128) = 251, I(128, 127) = 255, I(128, 128) = 251$$

Dengan jarak horizontal 0,5 dan vertikal 0, maka nilai piksel baru dihitung menggunakan interpolasi bilinear sebagai berikut:

$$\begin{aligned} I'(112, 112) &= (1 - 0,5)(1 - 0) \cdot 253 + 0,5(1 - 0) \cdot 251 \\ &= 0,5(253) + 0,5(251) \\ &= 126,5 + 125,5 = 252 \end{aligned}$$

Dari hasil tersebut diperoleh bahwa nilai piksel baru pada posisi (112,112) adalah 252, yang merupakan hasil interpolasi halus dari keempat piksel tetangga.

Dengan demikian, proses *resize* ini mampu mempertahankan detail citra otak tanpa menghasilkan tepi yang pecah atau kehilangan informasi penting.



Gambar 3. 3 Citra MRI sebelum dan sesudah proses *resize*

Gambar 3.3 menggambarkan hasil penyeragaman ukuran citra MRI melalui proses *resize*. Ukuran citra diubah menjadi 224×224 piksel dengan metode interpolasi bilinear, sehingga detail visual citra tetap terjaga.

3.3.2 Normalisasi

Penyesuaian nilai intensitas piksel dilakukan melalui tahap normalisasi setelah ukuran citra disamakan. Normalisasi bertujuan menempatkan nilai piksel pada rentang 0–1 agar proses pembelajaran berjalan lebih seimbang. Metode *min–max normalization* digunakan pada tahap ini sesuai dengan Persamaan 2.3.

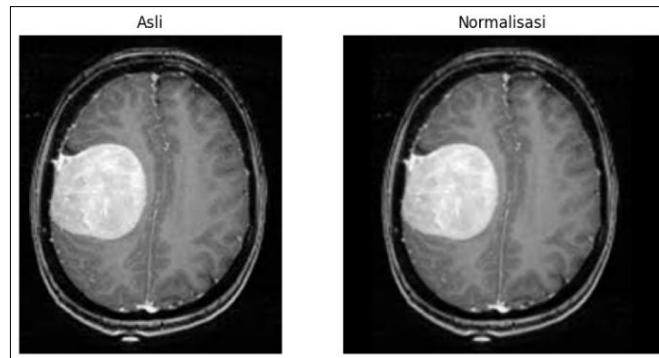
Sebagai ilustrasi, pada citra hasil *resize* yang memiliki nilai intensitas maksimum $I_{max} = 255$, nilai intensitas setiap piksel dinormalisasi dengan rumus:

$$I_{norm}(x, y) = \frac{I(x, y)}{255}$$

Jika nilai intensitas piksel hasil *resize* pada posisi (112,112) adalah 252, maka hasil normalisasinya dihitung sebagai berikut:

$$I_{norm}(112,112) = \frac{252}{255} = 0,988$$

Nilai 0,988 menunjukkan bahwa intensitas piksel tersebut berada pada 98,8% dari nilai maksimum. Melalui normalisasi ini, seluruh citra memiliki rentang nilai yang sama, yaitu antara 0 dan 1, sehingga proses perbandingan dan pembelajaran antar citra menjadi lebih stabil dan adil.



Gambar 3. 4 Citra MRI sebelum dan sesudah proses normalisasi

Gambar 3.4 menunjukkan perbandingan citra sebelum dan sesudah proses normalisasi. Setelah normalisasi, nilai intensitas piksel berada pada rentang yang lebih seragam, sehingga perbedaan kontras antar area pada citra menjadi lebih stabil dan memudahkan proses ekstraksi fitur selanjutnya.

3.3.3 *Noise Removal*

Tahap terakhir adalah *noise removal*, yang dilakukan untuk menghilangkan gangguan visual atau bintik acak (*noise*) yang dapat mengganggu proses ekstraksi fitur. Pada penelitian ini digunakan filter Gaussian dengan ukuran kernel 3×3 sebagaimana dijelaskan pada Persamaan 2.5 dan Persamaan 2.6. Filter Gaussian dipilih karena mampu mereduksi *noise* tanpa mengaburkan tepi objek pada citra.

Kernel Gaussian yang digunakan memiliki nilai bobot sebagai berikut:

$$\frac{1}{16} \begin{bmatrix} 1 & 2 & 1 \\ 2 & 4 & 2 \\ 1 & 2 & 1 \end{bmatrix}$$

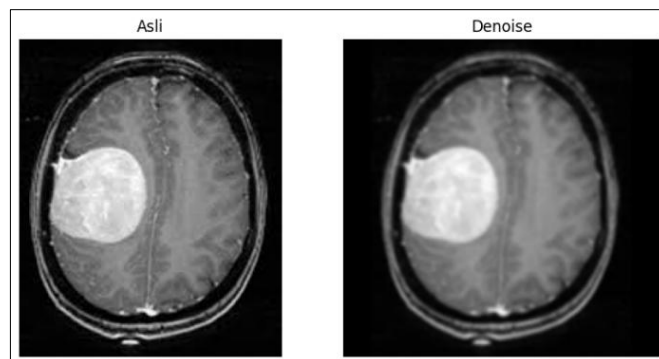
Proses *filtering* dilakukan dengan mengalikan kernel tersebut terhadap setiap piksel dan tetangganya pada citra hasil normalisasi. Sebagai contoh, pada salah satu area citra tidak tumor, nilai intensitas ternormalisasi di sekitar piksel pusat adalah:

$$\begin{bmatrix} 0,6353 & 0,6549 & 0,6588 \\ 0,6118 & 0,6667 & 0,6824 \\ 0,6314 & 0,6549 & 0,6314 \end{bmatrix}$$

Setiap nilai kemudian dikalikan dengan bobot kernel Gaussian dan dijumlahkan seluruhnya untuk mendapatkan nilai hasil *filtering*:

$$\begin{aligned} I'(x, y) &= \frac{1}{16} [(1)(0,6353) + (2)(0,6549) + (1)(0,6588) \\ &\quad + (2)(0,6118) + (4)(0,6667) + (2)(0,6824) \\ &\quad + (1)(0,6314) + (2)(0,6549) + (1)(0,6314)] \\ &= 0,6519 \end{aligned}$$

Hasil perhitungan menunjukkan nilai piksel pusat setelah *Gaussian filtering* sebesar 0,6519, yang lebih halus dari sebelumnya, menandakan berkurangnya noise tanpa menghilangkan struktur anatomi otak yang penting.

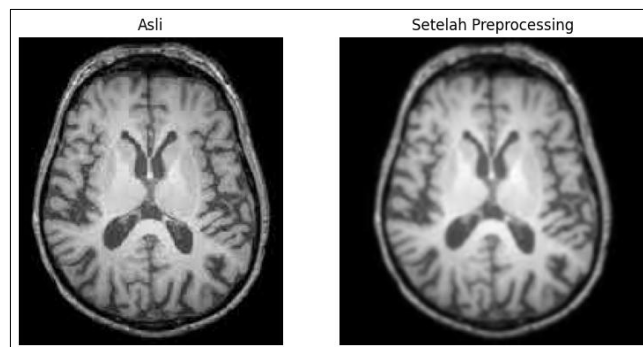


Gambar 3. 5 Citra MRI sebelum dan sesudah proses *denoise*

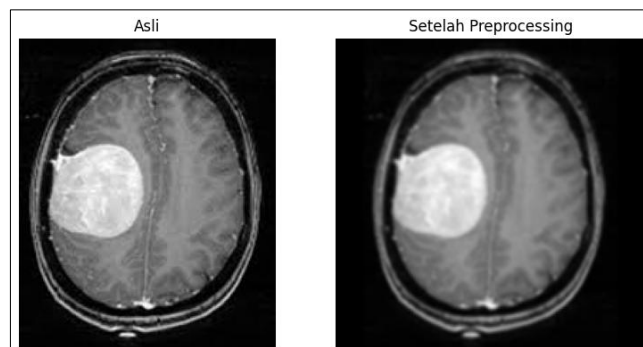
Gambar 3.5 menunjukkan citra sebelum dan sesudah proses *denoise* menggunakan *Gaussian blur*, yang memperlihatkan peredaman derau halus tanpa menghilangkan struktur penting pada citra otak sebelum tahap ekstraksi fitur.

3.3.4 Hasil *Pre-processing*

Setelah melalui ketiga tahap *pre-processing*, citra MRI memiliki ukuran yang seragam, rentang intensitas yang telah dinormalisasi, serta tampilan yang lebih bersih dari *noise*. Gambar 3.6 dan Gambar 3.7 perbandingan citra sebelum dan sesudah dilakukan *pre-processing*, yang menunjukkan bahwa hasil akhir memiliki tingkat kontras dan kehalusan yang lebih baik untuk proses ekstraksi fitur pada tahap berikutnya.



Gambar 3. 6 Hasil *pre-processing* pada citra MRI kelas tidak tumor

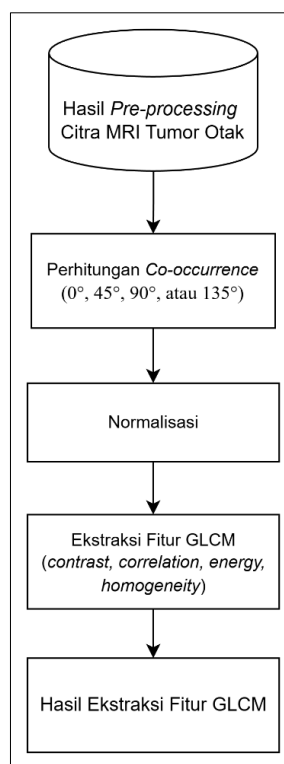


Gambar 3. 7 Hasil *pre-processing* pada citra MRI kelas tumor

Hasil *pre-processing* tersebut menghasilkan citra dengan karakteristik visual yang lebih konsisten pada kedua kelas, sehingga perbedaan tekstur dapat direpresentasikan dengan lebih jelas. Kondisi ini mendukung proses ekstraksi fitur *Gray Level Co-occurrence Matrix* (GLCM), karena distribusi nilai intensitas dan pola tekstur pada citra telah berada dalam kondisi yang lebih stabil dan seragam.

3.4 Ekstraksi Fitur *Gray Level Co-occurrence Matrix*

Sebagai metode ekstraksi tekstur, *Gray Level Co-occurrence Matrix* (GLCM) memanfaatkan informasi keterkaitan antar piksel berdasarkan tingkat keabuan, jarak, dan orientasi tertentu. Dari proses tersebut diperoleh fitur *contrast*, *correlation*, *energy*, dan *homogeneity* yang digunakan untuk merepresentasikan karakteristik tekstur citra.



Gambar 3. 8 *Flowchart GLCM*

Proses ekstraksi fitur GLCM pada Gambar 3.8 diawali dengan penggunaan citra MRI tumor otak yang telah diproses sebelumnya. Hubungan antar piksel dianalisis melalui perhitungan nilai *co-occurrence* pada sudut 0°, 45°, 90°, dan 135°. Selanjutnya, nilai hasil perhitungan dinormalisasi untuk menyeragamkan skala data, sebelum diekstraksi menjadi fitur *contrast*, *correlation*, *energy*, dan

homogeneity yang digunakan sebagai representasi tekstur citra dalam proses klasifikasi.

Gambar 3.9 menggambarkan penerapan metode GLCM untuk ekstraksi fitur tekstur pada sub-matriks 4×4 yang berasal dari data 28×28 , dengan konfigurasi jarak $d = 1$ dan arah horizontal (0°).

$$I_{4 \times 4} = \begin{bmatrix} 1 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 1 \\ 1 & 0 & 0 & 1 \end{bmatrix}$$

Gambar 3. 9 Contoh matriks GLCM ternormalisasi

Tahap pertama dilakukan dengan membangun matriks ko-okurensi C berdasarkan pasangan piksel bertetangga ke kanan (arah 0° dengan jarak $d = 1$). Dari hasil penghitungan diperoleh frekuensi pasangan $(0,0) = 3$, $(0,1) = 5$, $(1,0) = 4$, dan $(1,1) = 0$. Sehingga matriks ko-okurensi dapat ditulis sebagai berikut.

$$C = \begin{bmatrix} 3 & 5 \\ 4 & 0 \end{bmatrix}, \quad \sum C = 12$$

Proses normalisasi dilakukan dengan mengonversi matriks menjadi matriks probabilitas P melalui pembagian setiap elemen terhadap jumlah total pasangan piksel, sebagaimana dinyatakan pada Persamaan 2.7.

$$P(i, j) = \frac{C(i, j)}{\sum_i \sum_j C(i, j)}$$

$$P = \frac{C}{12} = \begin{bmatrix} 0.25 & 0.4167 \\ 0.3333 & 0 \end{bmatrix}$$

Setelah nilai dari hasil normalisasi didapat, matriks probabilitas tersebut digunakan sebagai dasar untuk menghitung berbagai fitur tekstur seperti *contrast*, *energy*, *homogeneity*, dan *correlation* pada tahap selanjutnya.

1. *Contrast*

Nilai *contrast* dihitung berdasarkan matriks probabilitas yang telah diperoleh sebelumnya sesuai dengan Persamaan 2.8 dan diperoleh hasil :

$$\begin{aligned}
 \text{Contrast} &= \sum_{i=0}^{N-1} \sum_{j=0}^{N-1} (i - j)^2 P(i, j) \\
 \text{Contrast} &= (0 - 0)^2 \cdot P(0,0) + (0 - 1)^2 \cdot P(0,1) + (1 - 0)^2 \cdot P(1,0) + \\
 &\quad (1 - 1)^2 \cdot P(1,1) \\
 &= 0 \cdot \frac{3}{12} + 1 \cdot \frac{5}{12} + 1 \cdot \frac{4}{12} + 0 \cdot \frac{0}{12} \\
 &= 0 + \frac{5}{12} + \frac{4}{12} + 0 \\
 &= \frac{9}{12} = \frac{3}{4} \approx 0.75
 \end{aligned}$$

Nilai *contrast* sebesar 0.75 menunjukkan bahwa ada perbedaan intensitas piksel yang cukup jelas pada arah yang dihitung. Ini menunjukkan bahwa, bagian citra ini memiliki pola yang kontras dan tidak terlalu seragam, sehingga teksturnya tampak lebih bervariasi.

2. *Correlation*

Untuk menghitung nilai *correlation*, diperlukan terlebih dahulu nilai rata-rata keabuan (μ) dan simpangan baku (σ) dari baris dan kolom matriks probabilitas. Rata-rata keabuan dihitung menggunakan Persamaan 3.1 dan Persamaan 3.2, sedangkan simpangan baku dihitung menggunakan Persamaan 3.3 dan Persamaan 3.4.

$$\mu_i = \sum_i i \sum_j P(i, j) \quad (3.1)$$

$$\mu_i = (0)(0.25 + 0.4167) + (1)(0.3333 + 0) = 0.3333$$

$$\mu_j = \sum_j j \sum_i P(i, j) \quad (3.2)$$

$$\mu_j = (0)(0.25 + 0.3333) + (1)(0.4167 + 0) = 0.4167$$

$$\sigma_i = \sqrt{\sum_i (i - \mu_i)^2 \sum_j P(i, j)} \quad (3.3)$$

$$\sigma_i = \sqrt{(0 - 0.3333)^2(0.25 + 0.4167) + (1 - 0.3333)^2(0.3333 + 0)} = 0.4714$$

$$\sigma_j = \sqrt{\sum_j (j - \mu_j)^2 \sum_i P(i, j)} \quad (3.4)$$

$$\sigma_j = \sqrt{(0 - 0.4167)^2(0.25 + 0.3333) + (1 - 0.4167)^2(0.4167 + 0)} = 0.4916$$

Kemudian, nilai *correlation* dihitung dengan Persamaan 2.9:

$$\begin{aligned} \text{Correlation} &= \frac{\sum_{i=0}^{N-1} \sum_{j=0}^{N-1} (i - \mu_i)(j - \mu_j)P(i, j)}{\sigma_i \sigma_j} \\ &= \frac{(0-0.3333)(0-0.4167)(0.25)+(0-0.3333)(1-0.4167)(0.4167)+(1-0.3333)(0-0.4167)(0.3333)+(1-0.3333)(1-0.4167)(0)}{0.4714 \times 0.4916} \\ &= \frac{(-0.3333)(-0.4167)(0.25)+(-0.3333)(0.5833)(0.4167)+(0.6667)(-0.4167)(0.3333)+0}{0.2317} \\ &= \frac{0.0347-0.0809-0.0926+0}{0.2317} \\ &= \frac{-0.1388}{0.2317} = -0.60 \end{aligned}$$

Nilai *correlation* sebesar -0.60 menunjukkan adanya hubungan negatif antara pasangan piksel, artinya ketika intensitas piksel meningkat, tetangganya cenderung menurun. Dengan kata lain, pola tekstur pada arah ini menunjukkan variasi kontras yang cukup jelas dan tidak seragam.

3. *Energy*

Berdasarkan Persamaan 2.10, nilai *energy* dihitung menggunakan matriks probabilitas yang telah diperoleh sebelumnya. Adapun hasil perhitungannya adalah sebagai berikut :

$$\begin{aligned}
\text{Energy} &= \sum_{i=0}^{N-1} \sum_{j=0}^{N-1} (P(i,j))^2 \\
\text{Energy} &= (0.25)^2 + (0.4167)^2 + (0.3333)^2 + (0)^2 \\
&= 0.0625 + 0.1736 + 0.1111 + 0 \\
&= 0.3472
\end{aligned}$$

Nilai *energy* sebesar 0.3472 menunjukkan bahwa tekstur citra memiliki tingkat keteraturan yang sedang. Semakin tinggi nilai *energy*, semakin seragam atau halus pola tekstur citra, sedangkan nilai yang tidak terlalu tinggi seperti ini mengindikasikan adanya sedikit variasi intensitas antar piksel.

4. *Homogeneity*

Berdasarkan Persamaan 2.11, nilai *homogeneity* dihitung menggunakan matriks probabilitas yang telah diperoleh sebelumnya. Adapun hasil perhitungannya adalah sebagai berikut :

$$\begin{aligned}
\text{Homogeneity} &= \sum_{i=0}^{N-1} \sum_{j=0}^{N-1} \frac{P(i,j)}{1+(i-j)^2} \\
\text{Homogeneity} &= \frac{0.25}{1+(0-0)^2} + \frac{0.4167}{1+(0-1)^2} + \frac{0.3333}{1+(1-0)^2} + \frac{0}{1+(1-1)^2} \\
&= \frac{0.25}{1} + \frac{0.4167}{2} + \frac{0.3333}{2} + \frac{0}{1} \\
&= 0.25 + 0.2083 + 0.1667 + 0 \\
&= 0.625
\end{aligned}$$

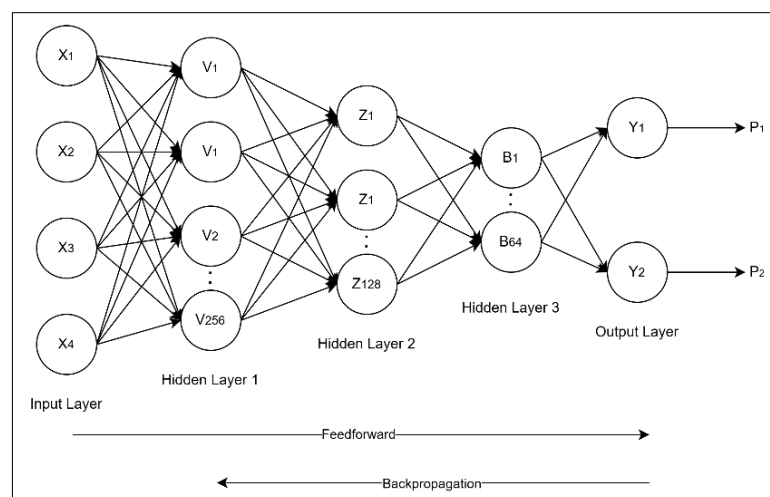
Nilai *homogeneity* sebesar 0.625 menunjukkan bahwa tekstur citra memiliki tingkat keseragaman yang cukup baik. Semakin tinggi nilai *homogeneity*, semakin mirip atau seragam intensitas antar piksel yang berdekatan,

sedangkan nilai sedang seperti ini menandakan adanya sedikit variasi pada pola tekstur citra.

Empat parameter utama hasil ekstraksi GLCM, yakni *contrast*, *correlation*, *energy*, dan *homogeneity*, diperoleh dari proses ekstraksi fitur tekstur. Hasil perhitungan menunjukkan nilai *contrast* 0,75, *correlation* -0,60, *energy* 0,3472, serta *homogeneity* 0,625. Parameter-parameter ini digunakan sebagai representasi tekstur citra dan menjadi masukan pada proses klasifikasi berikutnya.

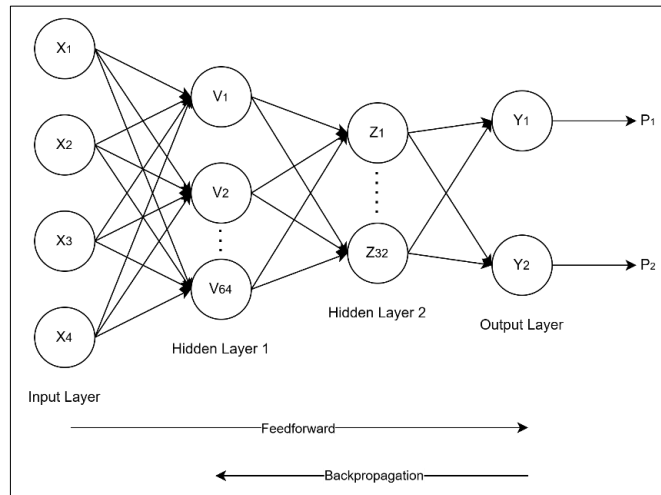
3.5 Implementasi Algoritma *Multilayer Perceptron*

Perancangan arsitektur *Multilayer Perceptron* (MLP) pada penelitian ini disusun dalam beberapa konfigurasi, seperti jumlah *hidden layer* untuk melihat pengaruh kedalaman jaringan terhadap proses klasifikasi citra MRI otak.



Gambar 3. 10 Arsitektur MLP dengan 3 *Hidden layer*

Gambar 3.10 Arsitektur *Multilayer Perceptron* (MLP) dengan tiga *hidden layer* yang digunakan untuk menganalisis pengaruh kedalaman jaringan terhadap hasil klasifikasi.



Gambar 3. 11 Arsitektur MLP dengan 2 *Hidden layer*

Gambar 3.11 Arsitektur *Multilayer Perceptron* (MLP) dengan dua *hidden layer* yang digunakan untuk membandingkan kinerja jaringan dengan jumlah *hidden layer* yang lebih sedikit. Penjelasan mengenai setiap *layer* pada arsitektur MLP diuraikan sebagai berikut:

1. *Input Layer*

Nilai *contrast*, *correlation*, *energy*, dan *homogeneity* yang diperoleh dari ekstraksi GLCM digunakan sebagai vektor masukan sebagaimana dirumuskan pada Persamaan 2.12.

$$x = [x_1, x_2, x_3, x_4]$$

$$x = [0.75, -0.60, 0.3472, 0.625]$$

Vektor ini merepresentasikan karakteristik utama citra hasil pengolahan GLCM. Vektor ini menjadi titik awal pemrosesan pada jaringan saraf tiruan, di mana setiap komponennya berperan sebagai informasi numerik yang menggambarkan pola tekstur citra. Selanjutnya, seluruh elemen vektor input akan diteruskan ke setiap neuron pada lapisan tersembunyi untuk dihitung

melalui kombinasi bobot dan bias guna membentuk representasi fitur yang lebih kompleks.

2. *Hidden layer*

Tahap berikutnya setelah pembentukan vektor masukan adalah pengolahan data pada *hidden layer*. Lapisan ini bertugas mengubah data input menjadi representasi yang lebih bermakna. Penelitian ini menerapkan dua *hidden layer* yang bekerja secara berurutan untuk meningkatkan kedalaman proses pembelajaran pola.

a. *Hidden layer* Pertama

Setiap neuron tersembunyi menerima seluruh komponen *input*, kemudian mengalikan masing-masing komponen dengan bobot yang menghubungkannya serta menambahkan nilai *bias* b_j , yang berfungsi untuk menggeser nilai aktivasi agar jaringan tidak selalu bernilai nol. Proses ini dilakukan sesuai dengan Persamaan 2.13.

Perhitungan neuron ke-1 :

$$z_j = \sum_{i=1}^n w_{ij} x_i + b_j$$

$$(x_1 \times w_{1j}), \quad (x_2 \times w_{2j}), \quad (x_3 \times w_{3j}), \quad (x_4 \times w_{4j})$$

$$z_1 = (0.75 \times w_{11}) + (-0.60 \times w_{21}) + (0.3472 \times w_{31}) + (0.625 \times w_{41}) + b_1$$

Setelah hasil perkalian diperoleh, keempat nilai tersebut dijumlahkan untuk mendapatkan jumlah berbobot total (*weighted sum*) sebelum penambahan bias:

$$\begin{aligned} z_1 &= (0.75 \times 0.43210) + (-0.60 \times 0.65789) + (0.3472 \times 0.32145) \\ &\quad + (0.625 \times 0.58941) + 0.18452 \end{aligned}$$

Selanjutnya, hasil penjumlahan ini ditambahkan dengan nilai bias b_1 membentuk nilai net input z_1 :

$$z_1 = s + b_1 = 0.40932969 + 0.18452 = 0.59384969$$

Tahap berikutnya adalah menerapkan fungsi aktivasi *ReLU* untuk menentukan apakah neuron pada lapisan tersembunyi akan aktif atau tidak. Fungsi *ReLU* akan meneruskan nilai positif apa adanya dan mengubah nilai negatif menjadi nol, sebagaimana dijelaskan pada Persamaan 2.14 dan diformulasikan dalam Persamaan 2.16. Karena nilai z_1 bernilai positif, maka neuron tersebut tetap aktif dan nilai keluarannya sama dengan nilai inputnya.

$$\begin{aligned} a_1 &= f(z_1) \\ &= \max(0, z_1) = 0.59384969 \end{aligned}$$

Perhitungan neuron ke-2 :

$$z_2 = s + b_2 = 0.3663535 + 0.16540 = 0.5317535$$

$$a_2 = f(z_2) = \max(0, z_2) = 0.5317535$$

Perhitungan neuron ke-3 :

$$z_3 = s + b_3 = 0.3016965 + 0.17860 = 0.4802965$$

$$a_3 = f(z_3) = \max(0, z_3) = 0.4802965$$

Perhitungan neuron ke-4 :

$$z_4 = s + b_4 = 0.4460609 + 0.19320 = 0.6392609$$

$$a_4 = f(z_4) = \max(0, z_4) = 0.6392609$$

Setelah keempat neuron menghitung nilai aktivasi masing-masing, diperoleh keluaran :

$$A^{(1)} = [a_1, a_2, a_3, a_4]$$

$$A^{(1)} = [0.59385, 0.53175, 0.48030, 0.63926]$$

Nilai-nilai ini merepresentasikan aktivasi dari masing-masing neuron yang akan menjadi masukan untuk *layer* berikutnya.

b. *Hidden layer* Kedua

Pada *Hidden layer* kedua, masukan berasal dari keluaran lapisan pertama yang telah diaktivasi dengan ReLU. Setiap neuron menghitung kombinasi linier antara nilai masukan dan bobot, ditambah dengan bias untuk membentuk representasi fitur yang lebih kompleks. Hasilnya kemudian diproses kembali menggunakan fungsi aktivasi ReLU guna menentukan neuron yang aktif dan siap diteruskan ke lapisan berikutnya.

$$A^{(1)} = [0.59385, 0.53175, 0.48030, 0.63926]$$

Perhitungan neuron ke-1 :

Masing-masing nilai masukan dikalikan dengan bobot yang terhubung ke neuron pertama, kemudian ditambahkan dengan nilai bias neuron pertama

$$\begin{aligned} z_1^{(2)} &= (0.59385 \times 0.15) + (0.53175 \times 0.05) + (0.48030 \times (-0.10)) \\ &\quad + (0.63926 \times 0.40) + 0.01 \\ &= (0.0890775) + (0.0265875) + (-0.04803) + (0.255704) + 0.01 \\ &= 0.333339 \end{aligned}$$

Nilai $z_1^{(2)}$ diproses menggunakan fungsi aktivasi ReLU :

$$a_1^{(2)} = f(z_1^{(2)}) = \max(0, z_1^{(2)}) = 0.333339$$

Perhitungan neuron ke-2 :

$$\begin{aligned}
 z_2^{(2)} &= (0.59385 \times (-0.25)) + (0.53175 \times 0.20) + (0.48030 \times 0.30) + \\
 &\quad (0.63926 \times (-0.05)) - 0.03 \\
 &= (-0.1484625) + (0.10635) + (0.14409) + (-0.031963) - 0.03 \\
 &= 0.0400145
 \end{aligned}$$

Nilai $z_2^{(2)}$ diproses menggunakan fungsi aktivasi ReLU :

$$a_2^{(2)} = f(z_2^{(2)}) = \max(0, z_2^{(2)}) = 0.0400145$$

Setelah kedua neuron menghitung nilai aktivasi masing-masing, diperoleh keluaran:

$$\begin{aligned}
 A^{(2)} &= [a_1^{(2)}, a_2^{(2)}] \\
 A^{(2)} &= [0.333339, 0.0400145]
 \end{aligned}$$

3. Output Layer

Lapisan keluaran merupakan tahap akhir dalam jaringan *Multilayer Perceptron* yang berfungsi menghasilkan nilai prediksi. Pada tahap ini, neuron keluaran menerima sinyal dari seluruh neuron pada *Hidden layer* terakhir dan mengolahnya untuk menghasilkan output akhir sesuai kelas yang diprediksi.

Langkah pertama pada proses perhitungan lapisan keluaran adalah menghitung nilai pra-aktivasi (z_k) untuk setiap neuron keluaran. Proses perhitungan diawali dengan mengombinasikan bobot koneksi dan nilai aktivasi setiap neuron pada *hidden layer* terakhir, kemudian menambahkan nilai bias sebagaimana dirumuskan dalam Persamaan 2.17.

Nilai pra-aktivasi yang diperoleh selanjutnya diproses melalui fungsi aktivasi pada Persamaan 2.18 untuk menghasilkan keluaran akhir jaringan.

$$z_k = \sum_{j=1}^m w_{jk} h_j + b_k$$

$$y_k = f(z_k)$$

Dengan dua masukan dari *Hidden layer* sebelumnya, yaitu

$$A^{(2)} = [0.333339, 0.0400145]$$

maka persamaannya dapat dituliskan sebagai berikut dengan nilai w_{jk} dan

b_k diambil dari hasil pelatihan model.

$$z_k = w_{1k} (0.333339) + w_{2k} (0.0400145) + b_k$$

Selanjutnya diterapkan fungsi aktivasi *sigmoid* untuk mengubah nilai pra-aktivasi (z) menjadi nilai keluaran (y) dengan rentang antara 0 hingga 1, sebagaimana pada Persamaan 2.16. Fungsi ini digunakan untuk mengonversi hasil perhitungan sebelumnya menjadi probabilitas yang merepresentasikan tingkat keyakinan model terhadap kelas tertentu.

$$\begin{aligned} z &= w_{1o}(h_1) + w_{2o}(h_2) + b_o \\ &= w_{1o}(0.333339) + w_{2o}(0.0400145) + b_o \\ &= (0.333339 \times 0.721) + (0.0400145 \times 0.530) + 0.312 \\ &= 0.57412 \end{aligned}$$

$$y = \frac{1}{1 + e^{-z}} = \frac{1}{1 + e^{-0.57412}} = 0.6391$$

Dengan demikian, nilai keluaran akhir pada *output layer* diperoleh sebagai berikut:

$$z = 0.57412 \Rightarrow y = 0.6391$$

Probabilitas hasil prediksi jaringan ditunjukkan oleh nilai keluaran pada lapisan output. Nilai y yang dihasilkan berada pada interval 0–1 dan mencerminkan tingkat keyakinan model terhadap kelas positif. Untuk menentukan hasil klasifikasi, digunakan nilai ambang batas 0,5. Nilai keluaran yang melampaui batas tersebut diinterpretasikan sebagai kelas positif (tumor), sedangkan nilai yang lebih kecil dikategorikan sebagai kelas negatif (tidak tumor). Berdasarkan hasil yang diperoleh, nilai keluaran berada di atas 0,5 sehingga citra MRI diprediksi termasuk kategori tumor.

4. Menghitung *Loss* Dengan *Binary Cross-Entropy*

Fungsi ini digunakan untuk mengukur sejauh mana hasil prediksi jaringan y berbeda dari nilai target sebenarnya t . Nilai selisih antara keduanya menunjukkan tingkat kesalahan model dalam melakukan prediksi. Semakin besar perbedaan antara output prediksi dan target, maka semakin besar pula nilai *loss* yang dihasilkan, yang menandakan bahwa model masih belum mampu memprediksi dengan akurat. Secara umum, rumusnya dituliskan pada Persamaan 3.7:

$$L = -[t \ln(y) + (1 - t) \ln(1 - y)] \quad (3.7)$$

Dari perhitungan sebelumnya:

$$y = 0.6391$$

Nilai ini disubstitusikan ke dalam persamaan 3.7 dengan target $t = 1$ yang menunjukkan bahwa citra tergolong kelas *tumor*. Perhitungannya adalah sebagai berikut :

$$\begin{aligned}
L &= -[1 \cdot \ln(0.6391) + (1 - 1) \ln(1 - 0.6391)] \\
&= -\ln(0.6391) \\
&= -(-0.4473) = 0.4473
\end{aligned}$$

Nilai *loss* sebesar 0,4473 diperoleh dari hasil perhitungan, yang menandakan masih terdapat selisih antara hasil prediksi model dan nilai target. Nilai ini digunakan sebagai indikator kesalahan model dan berperan penting dalam proses *backpropagation* untuk memperbarui bobot dan bias jaringan. Melalui proses tersebut, model dapat melakukan penyesuaian secara bertahap sehingga menghasilkan prediksi yang lebih akurat pada tahap pelatihan selanjutnya. Semakin rendah nilai *loss* yang dicapai, semakin baik kinerja jaringan dalam mengenali pola data.

5. *Backpropagation*

Setelah nilai *loss* (L) diperoleh, jaringan saraf melakukan penyesuaian bobot agar kesalahan prediksi berkurang pada iterasi berikutnya. Proses ini disebut *backpropagation*, yaitu perhitungan turunan (*gradien*) dari fungsi *loss* terhadap bobot untuk menentukan arah dan besarnya perubahan yang diperlukan agar nilai *loss* menurun.

Neuron keluaran pada tahap ini menggunakan fungsi aktivasi sigmoid untuk mengubah nilai pra-aktivasi (z) menjadi nilai keluaran (y) dengan rentang 0–1. Turunan fungsi aktivasi ini kemudian digunakan sebagai bagian dari mekanisme pembaruan bobot, yang dinyatakan pada Persamaan 3.8:

$$f'(z) = y(1 - y) \quad (3.8)$$

Selanjutnya dilakukan perhitungan turunan error terhadap keluaran jaringan untuk menentukan arah perubahan bobot. Pada kasus klasifikasi biner dengan fungsi aktivasi *sigmoid*, nilai error pada neuron keluaran dihitung menggunakan Persamaan 3.9:

$$\delta_o = (t - y) \cdot f'(z) \quad (3.9)$$

Selanjutnya dilakukan substitusi nilai hasil perhitungan sebelumnya ke dalam persamaan untuk memperoleh nilai error pada lapisan keluaran. Berdasarkan nilai yang diketahui, target $t = 1$ dan keluaran jaringan $y = 0.6391$. Langkah pertama adalah menghitung turunan fungsi aktivasi sigmoid:

$$f'(z) = 0.6391(1 - 0.6391) = 0.2305$$

Selanjutnya, nilai error pada neuron keluaran dihitung menggunakan persamaan:

$$\delta_o = (1 - 0.6391) \times 0.2305 = 0.0833$$

Diperoleh nilai $\delta_o = 0.0833$, yang menunjukkan besarnya *error* yang dikirim kembali ke *hidden layer* untuk memperbarui bobot pada setiap neuron.

6. Optimasi

Fungsi proses optimasi, yaitu memperbarui bobot dan bias menggunakan nilai gradien error δ yang diperoleh dari hasil *backpropagation*. Pembaruan dilakukan dengan menambahkan perubahan bobot Δw yang dihitung berdasarkan laju pembelajaran η , nilai error, dan aktivasi neuron pada lapisan sebelumnya. Bobot pada iterasi berikutnya

diperbarui berdasarkan hasil perhitungan gradien error menggunakan Persamaan 3.10:

$$w_{jo}^{(t+1)} = w_{jo}^{(t)} + \Delta w_{jo} \quad (3.10)$$

Sedangkan pembaruan nilai *bias* dilakukan menggunakan Persamaan 3.11:

$$b_o^{(t+1)} = b_o^{(t)} + \eta \cdot \delta_o \quad (3.11)$$

Dari persamaan tersebut, maka dilakukan perhitungan dengan nilai yang sudah diperoleh, yaitu :

$$\eta = 0.1, \quad \delta_o = 0.0833, \quad h_1 = 0.333339, \quad h_2 = 0.0400145$$

Bobot lama :

$$w_{1o}^{(t)} = 0.721, \quad w_{2o}^{(t)} = 0.530, \quad b_o^{(t)} = 0.312$$

Proses dimulai dengan melakukan perhitungan perubahan bobot melalui Persamaan 3.12:

$$\Delta w_{jo} = \eta \cdot \delta_o \cdot h_j \quad (3.12)$$

Dari hasil substitusi diperoleh:

$$\Delta w_{1o} = \eta \cdot \delta_o \cdot h_1 = 0.1 \times 0.0833 \times 0.333339 = 0.00278$$

$$\Delta w_{2o} = \eta \cdot \delta_o \cdot h_2 = 0.1 \times 0.0833 \times 0.0400145 = 0.00033$$

Selanjutnya, perubahan bias dihitung menggunakan Persamaan 3.13:

$$\Delta b_o = \eta \cdot \delta_o \quad (3.13)$$

sehingga:

$$\Delta b_o = \eta \cdot \delta_o = 0.1 \times 0.0833 = 0.00833$$

Setelah nilai perubahan diperoleh, bobot dan bias baru dihitung dengan menambahkan nilai perubahan tersebut pada bobot dan bias sebelumnya.

Hasil pembaruannya adalah:

$$w_{10}^{(t+1)} = 0.721 + 0.00278 = 0.72378$$

$$w_{20}^{(t+1)} = 0.530 + 0.00033 = 0.53033$$

$$b_o^{(t+1)} = 0.312 + 0.00833 = 0.32033$$

Hasil perhitungan menunjukkan bahwa bobot dan bias diperbarui secara bertahap selama proses pelatihan. Proses pembaruan ini berlangsung pada setiap *epoch* hingga jaringan mencapai kestabilan, yang ditandai dengan menurunnya nilai *loss* dan meningkatnya akurasi dalam mengklasifikasikan citra MRI.

7. *Epoch Iteration*

Pada tahap *epoch iteration*, jaringan saraf menjalankan proses pelatihan secara berulang melalui sejumlah *epoch* yang telah ditetapkan. Satu *epoch* merepresentasikan satu siklus penuh mulai dari propagasi maju, perhitungan error, propagasi balik, hingga pembaruan bobot dan bias. Dengan demikian, setiap *epoch* memberikan kesempatan bagi jaringan untuk menyesuaikan bobotnya secara bertahap terhadap pola data yang tersedia pada set pelatihan.

Meskipun nilai *epoch* tidak ditentukan melalui perhitungan khusus, jumlah iterasi pembaruan bobot dalam satu *epoch* ditentukan oleh ukuran *batch* yang digunakan. Hubungan antara jumlah data pelatihan dan *batch size* dirumuskan oleh Ian Goodfellow *et al.* (2016) pada Persamaan 3.14.

$$I = \left\lceil \frac{N}{B} \right\rceil \quad (3.14)$$

I = iterations per epoch
 N = jumlah sampel data latih
 B = batch size

Dengan jumlah data latih sebanyak 2.399 dan *batch size* sebesar 32, maka jumlah iterasi pembaruan bobot dalam satu *epoch* adalah:

$$I = \left\lceil \frac{2399}{32} \right\rceil = 75$$

Artinya, pada setiap *epoch* proses pembaruan bobot berlangsung sekitar 75 kali. Jika pelatihan dijalankan selama 1000 *epoch*, maka total pembaruan bobot yang dilakukan oleh jaringan dihitung menggunakan Persamaan 3.15:

$$U = I \times E \quad (3.15)$$

U = total update bobot
 I = *iterations per-epoch*
 E = jumlah *epoch*

$$U = 75 \times 1000 = 75,000$$

Jumlah iterasi yang besar ini memungkinkan jaringan memperbarui bobot secara bertahap hingga pola data dapat dipelajari dengan lebih baik. Dalam implementasinya, jumlah *epoch* awal ditetapkan sebanyak 1000, namun proses pelatihan tidak selalu mencapai seluruh *epoch* tersebut karena mekanisme *early stopping* digunakan untuk menghentikan pelatihan lebih awal ketika nilai loss tidak menunjukkan perbaikan yang berarti. Dengan demikian, *epoch iteration* berfungsi sebagai proses pengulangan terstruktur yang memastikan pembaruan bobot berlangsung secara konsisten hingga jaringan mencapai kondisi konvergen dan mampu mengenali pola citra MRI dengan lebih stabil.

8. *Early Stopping*

Early stopping merupakan teknik regulasi yang diterapkan untuk mengurangi risiko *overfitting* selama proses pelatihan jaringan saraf. Mekanisme ini bekerja dengan memantau nilai *loss* pada data pelatihan dan data validasi di setiap *epoch*. Apabila nilai *loss* tidak mengalami penurunan yang signifikan atau justru meningkat dalam beberapa *epoch* berturut-turut, maka proses pelatihan dihentikan sebelum jumlah *epoch* maksimum tercapai.

Pendekatan ini memungkinkan model berhenti belajar pada kondisi optimal, sehingga terhindar dari pembelajaran *noise* atau pola yang tidak relevan. Dengan demikian, *early stopping* berperan dalam menjaga kemampuan generalisasi model serta mencegah pembaruan bobot yang tidak memberikan peningkatan kinerja klasifikasi. Pada penelitian ini, penerapan *early stopping* menyebabkan proses pelatihan berhenti sebelum 1000 *epoch* ketika nilai *loss* telah stabil, sehingga pelatihan menjadi lebih efisien dan terkendali.

3.6 Skenario Pengujian

Pengujian dilakukan dengan memvariasikan arsitektur jaringan saraf pada MLP, mencakup jumlah *hidden layer*, jumlah neuron di setiap lapisan tersembunyi, serta ukuran *batch size* yang digunakan saat pelatihan. Kombinasi parameter tersebut dievaluasi untuk melihat pengaruh kedalaman jaringan, kapasitas neuron, dan strategi pemrosesan data per-*batch* terhadap akurasi dan kestabilan model dalam proses klasifikasi.

Tabel 3. 2 Skenario Pengujian

<i>Split Data</i>	<i>Model</i>	<i>Hidden layer</i>	<i>Neuron Hidden layer</i>	<i>Batch size</i>	<i>Max Iteration</i>
80:20	Model A	3	64, 32, 16	32	1000
	Model B		128, 64, 32		
	Model C		256, 128, 64		
	Model D		64, 32, 16	64	
	Model E		128, 64, 32		
	Model F		256, 128, 64		
	Model G	2	64, 32	32	
	Model H		32, 16		
	Model I		64, 32	64	
	Model J		32, 16		

Tabel 3.2 rincian skenario pengujian yang diterapkan dalam penelitian ini. Setiap skenario merepresentasikan konfigurasi model yang berbeda untuk keperluan evaluasi.

3.7 Evaluasi

Sebagai penerapan dari metrik evaluasi yang telah dijelaskan pada Bab 2.5, bagian ini menampilkan contoh perhitungan menggunakan data *confusion matrix* sederhana. Perhitungan ini digunakan untuk memperoleh nilai akurasi, presisi, *recall*, dan *F1-score* berdasarkan prediksi model terhadap data pengujian.

Tabel 3. 3 Contoh *confusion matrix*

		Aktual	
		Tumor	Tidak tumor
Prediksi	Tumor	TP = 40	FP = 5
	Tidak tumor	FN = 10	TN = 45

Dari tabel 3.3, diperoleh hasil :

1. *True positive (TP)* sebanyak 40 data, yaitu citra tumor diprediksi benar sebagai tumor oleh model.
2. *True negative (TN)* sebanyak 45 data, yaitu citra tidak tumor diprediksi benar sebagai tidak tumor oleh model.
3. *False positive (FP)* sebanyak 5 data, yaitu citra tidak tumor diprediksi sebagai tumor.
4. *False negative (FN)* sebanyak 10 data, yaitu citra tumor diprediksi sebagai tidak tumor.

Selanjutnya dilakukan perhitungan metrik evaluasi sebagai berikut :

1. *Accuracy*

Nilai akurasi diperoleh dengan membandingkan jumlah prediksi yang tepat terhadap keseluruhan data uji, dengan perhitungan mengacu pada Persamaan 2.19.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} = \frac{40 + 45}{40 + 45 + 5 + 10} = \frac{85}{100} = 0.85$$

Akurasi sebesar 85% menunjukkan bahwa sebagian besar data uji berhasil diklasifikasikan dengan benar oleh model.

2. *Precision*

Presisi digunakan untuk mengukur proporsi prediksi positif yang benar dari seluruh data yang diprediksi sebagai positif oleh model, dengan perhitungan mengacu pada Persamaan 2.20.

$$Precision = \frac{TP}{TP + FP} = \frac{40}{40 + 5} = \frac{40}{45} = 0.8889$$

Nilai presisi sebesar 88,9% mengindikasikan bahwa mayoritas prediksi positif yang dihasilkan model sesuai dengan kondisi sebenarnya.

3. *Recall*

Recall digunakan untuk menilai kemampuan model dalam mengenali data positif secara benar, dengan perhitungan mengacu pada Persamaan 2.21

$$Recall = \frac{TP}{TP + FN} = \frac{40}{40 + 10} = \frac{40}{50} = 0.80$$

Nilai *recall* sebesar 80% menunjukkan bahwa sebagian besar data positif dapat terdeteksi oleh model, meskipun masih terdapat sejumlah data positif yang belum berhasil dikenali.

4. *F1-score*

F1-score digunakan untuk menggambarkan keseimbangan antara nilai *presisi* dan *recall*, dengan perhitungan mengacu pada Persamaan 2.22.

$$F1-Score = \frac{2 \times Precision \times Recall}{Precision + Recall} = \frac{2 \times 0.8889 \times 0.80}{0.8889 + 0.80} = 0.842$$

Nilai *F1-score* sebesar 84,2% menunjukkan bahwa model memiliki keseimbangan yang cukup baik antara ketepatan prediksi dan kemampuan dalam mendeteksi kelas positif.

Hasil evaluasi menunjukkan bahwa model memperoleh akurasi 85%, presisi 88,9%, *recall* 80%, dan *F1-score* 84,2%. Capaian tersebut mengindikasikan kinerja yang cukup baik dengan keseimbangan antara akurasi prediksi dan kemampuan model dalam mengenali kelas positif.

BAB IV

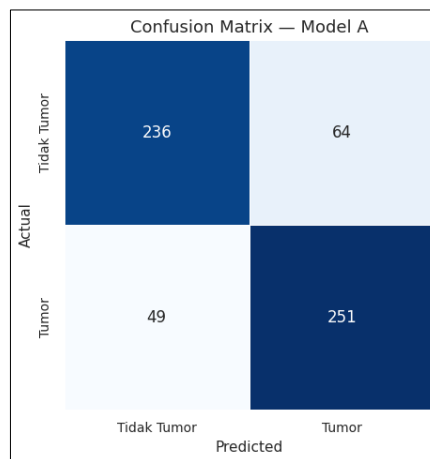
HASIL DAN PEMBAHASAN

4.1 Hasil Pengujian

Berdasarkan Tabel 3.2, sepuluh konfigurasi model diuji dengan variasi *hidden layer*, jumlah neuron, dan *batch size*, menggunakan pembagian data 80:20, serta dievaluasi melalui *confusion matrix* untuk menentukan konfigurasi dengan akurasi tertinggi.

4.1.1 Model A

Model A diuji menggunakan pembagian data 80:20 yang menghasilkan 2.399 data latih dan 600 data uji. Arsitektur yang digunakan terdiri atas tiga *hidden layer* berukuran (64, 32, 16) dengan *batch size* 32, dan hasil pengujiannya ditampilkan pada Gambar 4.1 serta dirangkum pada Tabel 4.1.



Gambar 4. 1 *Confusion matrix* model A

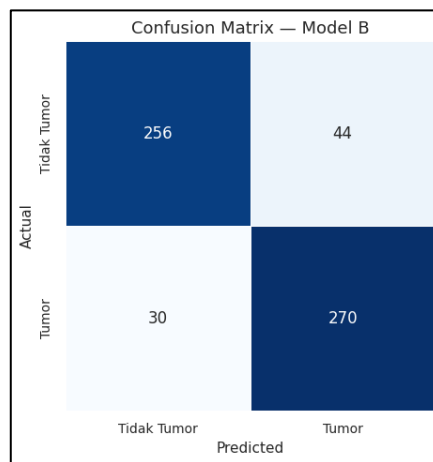
Tabel 4. 1 Hasil *confusion matrix* model A

Akurasi (%)	Presisi (%)	Recall (%)	F1-score (%)
81.17	79.68	83.67	81.63

Berdasarkan Gambar 4.1, Model A mampu mengklasifikasikan 236 citra tidak tumor sebagai *true negative* dan 251 citra tumor sebagai *true positive*, dengan kesalahan berupa 64 *false positive* dan 49 *false negative*. Secara keseluruhan, model mencapai akurasi 81,17%, presisi 79,68%, *recall* 83,67%, dan *F1-score* 81,63%, yang menunjukkan keseimbangan performa yang cukup baik.

4.1.2 Model B

Model B diuji menggunakan pembagian data 80:20 yang menghasilkan 2.399 data latih dan 600 data uji. Arsitektur yang digunakan terdiri atas tiga *hidden layer* berukuran (128, 64, 32) dengan *batch size* 32, dan hasil pengujiannya ditampilkan pada Gambar 4.2 serta dirangkum pada Tabel 4.2.



Gambar 4. 2 Confusion matrix model B

Tabel 4. 2 Hasil *confusion matrix* model B

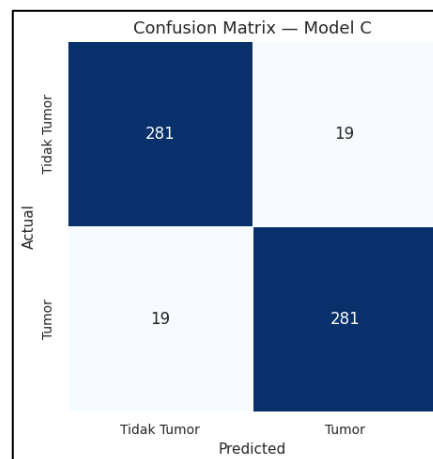
Akurasi (%)	Presisi (%)	Recall (%)	F1-score (%)
87.67	85.99	90	87.95

Berdasarkan Gambar 4.2, Model B mampu mengklasifikasikan 256 citra tidak tumor sebagai *true negative* dan 270 citra tumor sebagai *true positive*, dengan kesalahan berupa 44 *false positive* dan 30 *false negative*. Secara keseluruhan, model

mencapai akurasi 87,67%, dengan presisi 85,99% dan *recall* 90% yang menunjukkan ketepatan serta kemampuan model dalam mengenali citra tumor, serta *F1-score* sebesar 87,95% yang menandakan keseimbangan antara presisi dan *recall*.

4.1.3 Model C

Model C diuji menggunakan pembagian data 80:20 yang menghasilkan 2.399 data latih dan 600 data uji. Arsitektur yang digunakan terdiri atas tiga *hidden layer* berukuran (256, 128, 64) dengan *batch size* 32, dan hasil pengujiannya ditampilkan pada Gambar 4.3 serta dirangkum pada Tabel 4.3.



Gambar 4. 3 Confusion matrix model C

Tabel 4. 3 Hasil *confusion matrix* model C

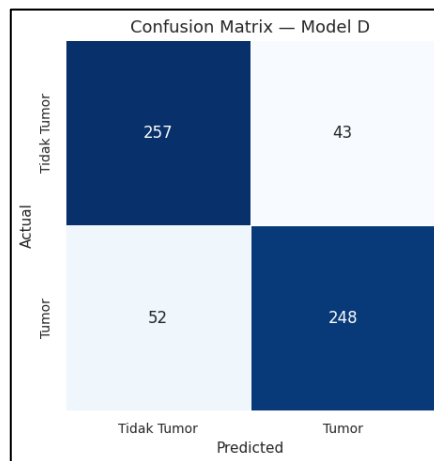
Akurasi (%)	Presisi (%)	Recall (%)	F1-score (%)
93.67	93.67	93.67	93.67

Berdasarkan Gambar 4.3, Model C mampu mengklasifikasikan 281 citra tidak tumor sebagai *true negative* dan 281 citra tumor sebagai *true positive*, dengan kesalahan berupa 19 *false positive* dan 19 *false negative*. Secara keseluruhan, model

mencapai akurasi, presisi, dan *recall* sebesar 93,67%, serta *F1-score* 93,67%, yang menunjukkan keseimbangan performa yang sangat baik antara presisi dan *recall*.

4.1.4 Model D

Model D diuji menggunakan pembagian data 80:20 yang menghasilkan 2.399 data latih dan 600 data uji. Arsitektur yang digunakan terdiri atas tiga *hidden layer* berukuran (64, 32, 16) dengan *batch size* 64, dan hasil pengujiannya ditampilkan pada Gambar 4.4 serta dirangkum pada Tabel 4.4.



Gambar 4. 4 *Confusion matrix* model D

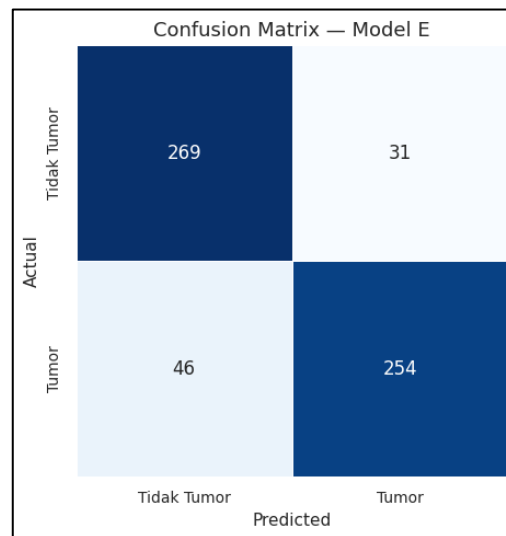
Tabel 4. 4 Hasil *confusion matrix* model D

Akurasi (%)	Presisi (%)	Recall (%)	F1-score (%)
84.17	85.22	82.67	83.93

Berdasarkan Gambar 4.4, Model D mampu mengklasifikasikan 257 citra tidak tumor sebagai *true negative* dan 248 citra tumor sebagai *true positive*, dengan kesalahan berupa 43 *false positive* dan 52 *false negative*. Secara keseluruhan, model mencapai akurasi 84,17%, dengan presisi 85,22% dan *recall* 82,67% yang menunjukkan ketepatan prediksi yang cukup baik, serta *F1-score* 83,93% yang mencerminkan keseimbangan antara presisi dan *recall*.

4.1.5 Model E

Model E diuji menggunakan pembagian data 80:20 yang menghasilkan 2.399 data latih dan 600 data uji. Arsitektur yang digunakan terdiri atas tiga *hidden layer* berukuran (128, 64, 32) dengan *batch size* 64, dan hasil pengujiannya ditampilkan pada Gambar 4.5 serta dirangkum pada Tabel 4.5.



Gambar 4. 5 *Confusion matrix* model E

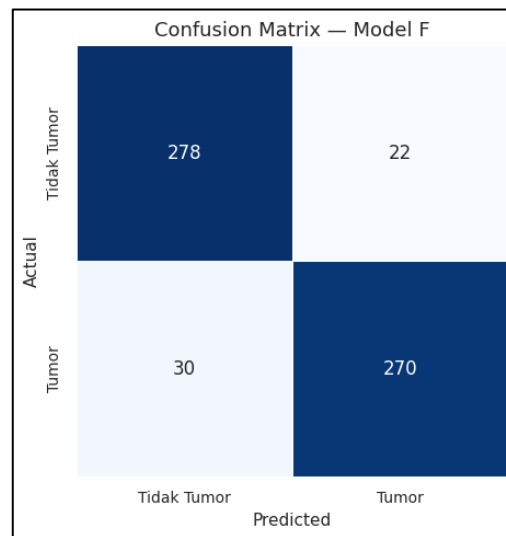
Tabel 4. 5 Hasil *confusion matrix* model E

Akurasi (%)	Presisi (%)	Recall (%)	F1-score (%)
87.17	89.12	84.67	86.84

Berdasarkan Gambar 4.5, Model E mampu mengklasifikasikan 269 citra tidak tumor sebagai *true negative* dan 254 citra tumor sebagai *true positive*, dengan kesalahan berupa 31 *false positive* dan 46 *false negative*. Secara keseluruhan, model mencapai akurasi 87,17%, dengan presisi 89,12% dan *recall* 84,67% yang menunjukkan ketepatan prediksi yang baik, serta *F1-score* 86,84% yang mencerminkan keseimbangan antara presisi dan *recall*.

4.1.6 Model F

Model F diuji menggunakan pembagian data 80:20 yang menghasilkan 2.399 data latih dan 600 data uji. Arsitektur yang digunakan terdiri atas tiga *hidden layer* berukuran (256, 128, 64) dengan *batch size* 64, dan hasil pengujiannya ditampilkan pada Gambar 4.6 serta dirangkum pada Tabel 4.6.



Gambar 4. 6 *Confusion matrix* model F

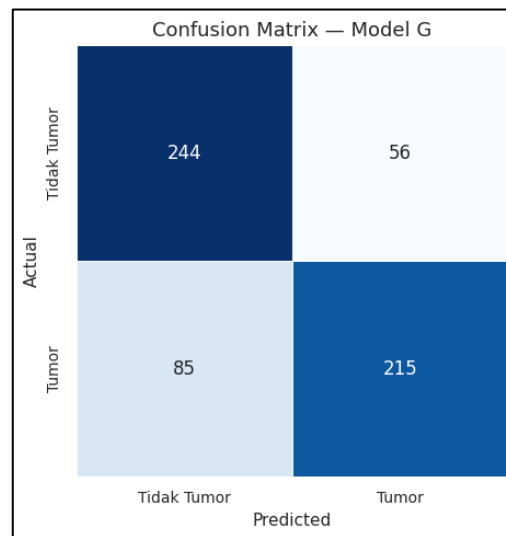
Tabel 4. 6 Hasil *confusion matrix* model F

Akurasi (%)	Presisi (%)	Recall (%)	F1-score (%)
91.33	92.47	90	91.22

Berdasarkan Gambar 4.6, Model F mampu mengklasifikasikan 278 citra tidak tumor sebagai *true negative* dan 270 citra tumor sebagai *true positive*, dengan kesalahan berupa 22 *false positive* dan 30 *false negative*. Secara keseluruhan, model mencapai akurasi 91,33%, dengan presisi 92,47% dan *recall* 90% yang menunjukkan ketepatan prediksi yang tinggi, serta *F1-score* 91,22% yang mencerminkan keseimbangan antara presisi dan *recall*.

4.1.7 Model G

Model G diuji menggunakan pembagian data 80:20 yang menghasilkan 2.399 data latih dan 600 data uji. Arsitektur yang digunakan terdiri atas dua *hidden layer* berukuran (64, 32) dengan *batch size* 32, dan hasil pengujiannya ditampilkan pada Gambar 4.7 serta dirangkum pada Tabel 4.7.



Gambar 4. 7 *Confusion matrix* model G

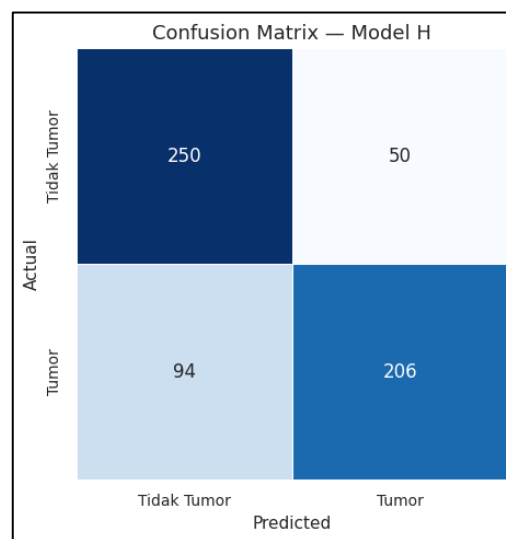
Tabel 4. 7 Hasil *confusion matrix* model G

Akurasi (%)	Presisi (%)	Recall (%)	F1-score (%)
76.5	79.34	71.67	75.31

Berdasarkan Gambar 4.7, Model G mampu mengklasifikasikan 244 citra tidak tumor sebagai *true negative* dan 215 citra tumor sebagai *true positive*, dengan kesalahan berupa 56 *false positive* dan 85 *false negative*. Secara keseluruhan, model mencapai akurasi 76,50%, dengan presisi 79,34% dan *recall* 71,67% yang menunjukkan ketepatan prediksi yang cukup, serta *F1-score* 75,31% yang mencerminkan keseimbangan antara presisi dan *recall*.

4.1.8 Model H

Model H diuji menggunakan pembagian data 80:20 yang menghasilkan 2.399 data latih dan 600 data uji. Arsitektur yang digunakan terdiri atas dua *hidden layer* berukuran (32, 16) dengan *batch size* 32, dan hasil pengujiannya ditampilkan pada Gambar 4.8 serta dirangkum pada Tabel 4.8.



Gambar 4. 8 *Confusion matrix* model H

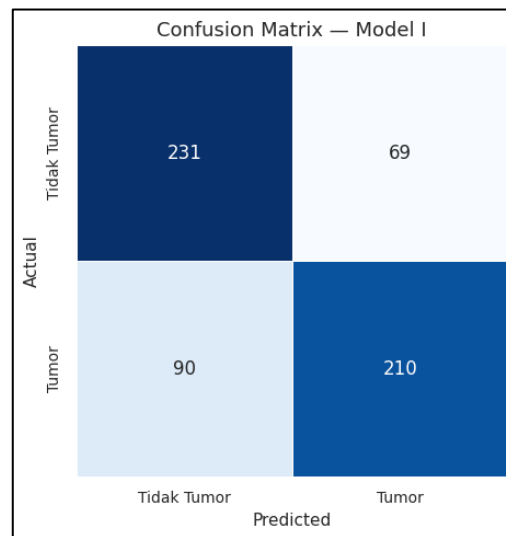
Tabel 4. 8 Hasil *confusion matrix* model H

Akurasi (%)	Presisi (%)	Recall (%)	F1-score (%)
76	80.47	68.67	74.1

Berdasarkan Gambar 4.8, Model H mampu mengklasifikasikan 250 citra tidak tumor sebagai *true negative* dan 206 citra tumor sebagai *true positive*, dengan kesalahan berupa 50 *false positive* dan 94 *false negative*. Secara keseluruhan, model mencapai akurasi 76%, dengan presisi 80,47% dan *recall* 68,67% yang menunjukkan ketepatan prediksi yang cukup, serta *F1-score* 74,10% yang mencerminkan keseimbangan antara presisi dan *recall*.

4.1.9 Model I

Model I diuji menggunakan pembagian data 80:20 yang menghasilkan 2.399 data latih dan 600 data uji. Arsitektur yang digunakan terdiri atas dua *hidden layer* berukuran (64, 32) dengan *batch size* 64, dan hasil pengujiannya ditampilkan pada Gambar 4.9 serta dirangkum pada Tabel 4.9.



Gambar 4. 9 *Confusion matrix* model I

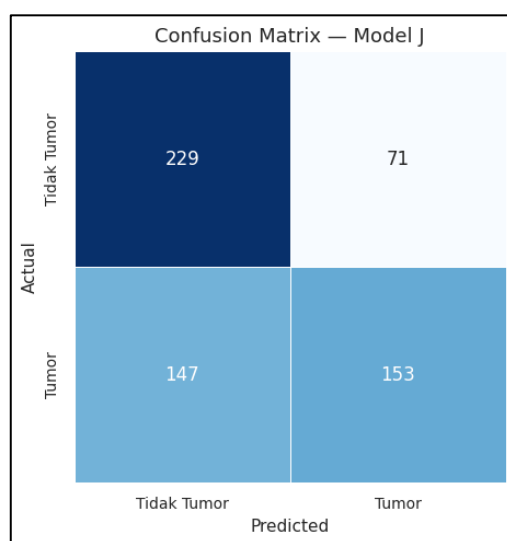
Tabel 4. 9 Hasil *confusion matrix* model I

Akurasi (%)	Presisi (%)	Recall (%)	F1-score (%)
73.5	75.27	70	72.54

Berdasarkan Gambar 4.9, Model I mampu mengklasifikasikan 231 citra tidak tumor sebagai *true negative* dan 210 citra tumor sebagai *true positive*, dengan kesalahan berupa 69 *false positive* dan 90 *false negative*. Secara keseluruhan, model mencapai akurasi 73,50%, dengan presisi 75,27% dan *recall* 70% yang menunjukkan ketepatan prediksi yang cukup, serta *F1-score* 72,54% yang mencerminkan keseimbangan antara presisi dan *recall*.

4.1.10 Model J

Model J diuji menggunakan pembagian data 80:20 yang menghasilkan 2.399 data latih dan 600 data uji. Arsitektur yang digunakan terdiri atas dua *hidden layer* berukuran (32, 16) dengan *batch size* 64, dan hasil pengujiannya ditampilkan pada Gambar 4.10 serta dirangkum pada Tabel 4.10.



Gambar 4. 10 *Confusion matrix* model J

Tabel 4. 10 Hasil *confusion matrix* model J

Akurasi (%)	Presisi (%)	Recall (%)	F1-score (%)
63.67	68.3	51	58.4

Berdasarkan Gambar 4.10, Model J mampu mengklasifikasikan 229 citra tidak tumor sebagai *true negative* dan 153 citra tumor sebagai *true positive*, dengan kesalahan berupa 71 *false positive* dan 147 *false negative*. Secara keseluruhan, model mencapai akurasi 63,67%, dengan presisi 68,30% dan *recall* 51% yang menunjukkan ketepatan prediksi yang terbatas, serta *F1-score* 58,40% yang mencerminkan keseimbangan antara presisi dan *recall*.

4.2 Pembahasan

Hasil pengujian sepuluh model MLP dengan variasi arsitektur jaringan, dimana Tabel 4.11 merangkum nilai akurasi, presisi, *recall*, dan *F1-score* dari seluruh skenario untuk menunjukkan performa masing-masing model.

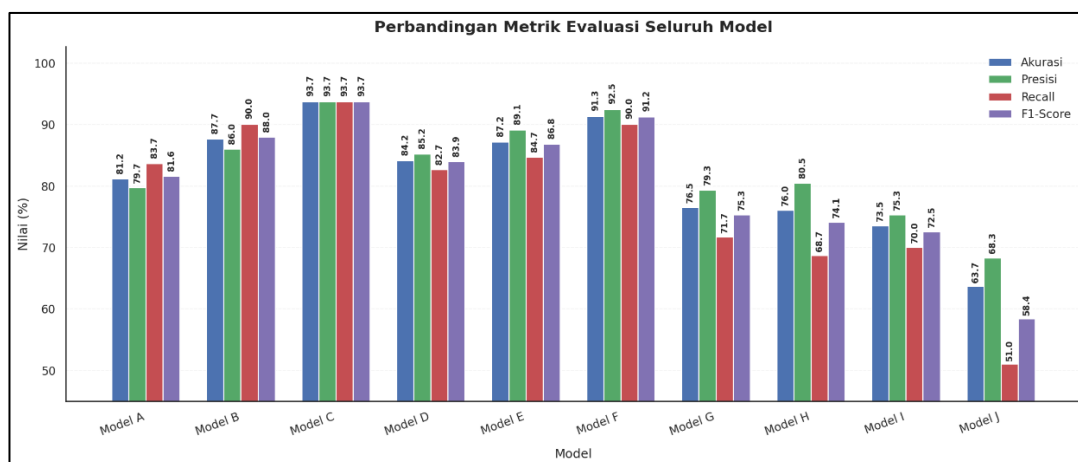
Tabel 4. 11 Hasil perfoma seluruh model

No	Model	Hidden layer	Batch	Max Iter	Akurasi	Presisi	Recall	F1-score	
1.	Model A	(64, 32, 16)	32	1000	81.17	79.68	83.67	81.63	
2.	Model B	(128, 64, 32)			87.67	85.99	90	87.95	
3.	Model C	(256, 128, 64)			93.67	93.67	93.67	93.67	
4.	Model D	(64, 32, 16)	64		84.17	85.22	82.67	83.93	
5.	Model E	(128, 64, 32)			87.17	89.12	84.67	86.84	
6.	Model F	(256, 128, 64)			91.33	92.47	90	91.22	
7.	Model G	(64, 32)	32		76.5	79.34	71.67	75.31	
8.	Model H	(32, 16)			76	80.47	68.67	74.1	
9.	Model I	(64, 32)	64		73.5	75.27	70	72.54	
10.	Model J	(32, 16)			63.67	68.3	51	58.4	

Tabel 4.11 menunjukkan bahwa model dengan arsitektur besar, seperti Model C dan Model F dengan tiga *hidden layer* berukuran (256, 128, 64), memiliki performa paling unggul pada seluruh metrik evaluasi. Model berukuran menengah, seperti Model B dan Model E, berada pada tingkat performa sedang dengan akurasi sekitar 87%, sedangkan model berarsitektur kecil, yaitu Model G hingga Model J, menunjukkan performa terendah. Secara umum, tabel tersebut memperlihatkan bahwa peningkatan kapasitas jaringan berkorelasi positif dengan peningkatan akurasi, presisi, *recall*, dan *F1-score*.

Untuk membandingkan kinerja antar model, digunakan empat metrik evaluasi utama, yaitu akurasi, presisi, *recall*, dan *F1-score*. Akurasi menunjukkan tingkat kebenaran prediksi secara keseluruhan, presisi mengukur ketepatan prediksi

pada kelas positif, sedangkan *recall* menilai kemampuan model dalam mengidentifikasi sampel positif yang sesungguhnya. Sementara itu, *F1-score* berfungsi sebagai indikator keseimbangan antara presisi dan *recall* yang dihitung melalui rata-rata harmonik.



Gambar 4. 11 Grafik perbandingan matriks evaluasi seluruh model

Gambar 4.11 memperlihatkan hasil perbandingan kinerja model A hingga J menggunakan metrik akurasi, presisi, *recall*, dan *F1-score*. Dari hasil tersebut dapat diketahui bahwa peningkatan kapasitas arsitektur berpengaruh positif terhadap performa model, di mana Model C mencapai capaian tertinggi pada seluruh metrik dengan nilai sekitar 93–94%, sehingga menunjukkan keunggulan konfigurasi tiga *hidden layer* dengan jumlah neuron yang besar.

Pada Model D hingga Model F, performa relatif stabil meskipun terjadi fluktuasi kecil, dengan Model F menunjukkan keseimbangan terbaik antar metrik. Sebaliknya, Model G hingga Model J mengalami penurunan performa yang signifikan, terutama pada *recall*, dengan Model J mencatat nilai terendah sekitar 51%, disertai penurunan *F1-score* dan akurasi, yang menegaskan keterbatasan arsitektur berkapasitas kecil dalam merepresentasikan kompleksitas fitur.

4.2.1 Pengaruh Parameter

Subbab ini menjelaskan pengaruh parameter yang digunakan dalam penelitian terhadap kinerja model.

4.2.1.1 *Hidden layer*

Perbandingan akurasi untuk arsitektur MLP dengan 3 *hidden layer* dan 2 *hidden layer* Tabel 4.12.

Tabel 4. 12 Rata-rata Akurasi Berdasarkan Jumlah *Hidden layer*

No	3 <i>Hidden layer</i>		2 <i>Hidden layer</i>	
	Model	Akurasi (%)	Model	Akurasi (%)
1.	Model A	81.17	Model G	76.50
2.	Model B	87.67	Model H	76.00
3.	Model C	93.67	Model I	73.50
4.	Model D	84.17	Model J	63.67
5.	Model E	87.17		
6.	Model F	91.33		
	Rata-rata	87.20	Rata-rata	72.42

Hasil pengujian menunjukkan bahwa model dengan 3 *hidden layer* secara konsisten memiliki performa lebih baik dibandingkan model dengan 2 *hidden layer*. Rata-rata akurasi kelompok 3 *hidden layer* mencapai 87,20%, sedangkan kelompok 2 *hidden layer* hanya sebesar 72,42%.

Model C dan Model F menjadi yang paling menonjol pada kelompok 3 *hidden layer* dengan akurasi masing-masing 93,67% dan 91,33%, jauh melampaui model terbaik pada kelompok 2 *hidden layer*, yaitu Model G dengan akurasi 76,50%. Sebaliknya, Model J dengan konfigurasi paling sederhana mencatat akurasi terendah sebesar 63,67%.

Performa yang lebih tinggi pada model dengan 3 *hidden layer* menunjukkan bahwa arsitektur jaringan yang lebih dalam memiliki kemampuan representasi yang

lebih besar dalam mempelajari hubungan non-linier yang kompleks dari data. Jaringan yang lebih dalam memungkinkan proses *hierarchical feature learning*, di mana setiap *layer* mengekstraksi representasi fitur pada tingkat abstraksi yang berbeda sehingga meningkatkan akurasi klasifikasi secara keseluruhan (Mienye & Swart, 2024).

4.2.1.2 Neuron

Subbab ini membahas pengaruh jumlah neuron terhadap kinerja model klasifikasi.

(1) 3 Hidden layer

Perbandingan akurasi model dengan arsitektur 3 *hidden layer* yang dibedakan berdasarkan jumlah neuron Tabel 4.13.

Tabel 4. 13 Perbandingan Jumlah Neuron pada 3 *Hidden layer*

No	Neuron 64–32–16		Neuron 128–64–32		Neuron 256–128–64	
	Model	Akurasi (%)	Model	Akurasi (%)	Model	Akurasi (%)
1.	Model A	81.17	Model B	87.67	Model C	93.67
2.	Model D	84.17	Model E	87.17	Model F	91.33
	Rata-rata	82.67	Rata-rata	87.42	Rata-rata	92.50

Dari Tabel 4.13 terlihat bahwa semakin besar konfigurasi neuron, semakin tinggi akurasi yang diperoleh. Variasi 256–128–64 memberikan rata-rata akurasi tertinggi sebesar 92,50%, diikuti konfigurasi 128–64–32 dengan rata-rata 87,42%. Sementara itu, konfigurasi 64–32–16 menghasilkan rata-rata akurasi paling rendah sebesar 82,67%.

Pola peningkatan akurasi seiring bertambahnya jumlah neuron pada 3 *hidden layer* menunjukkan bahwa kapasitas representasi yang lebih besar memungkinkan jaringan menangkap karakteristik fitur yang lebih kompleks. Konfigurasi dengan neuron lebih banyak mampu mempelajari hubungan non-linier secara lebih efektif, sehingga menghasilkan akurasi prediksi yang lebih tinggi dibandingkan konfigurasi dengan neuron lebih sedikit (Alagoz et al., 2025).

(2) 2 *Hidden layer*

Perbandingan akurasi model dengan arsitektur 2 *hidden layer* yang dibedakan berdasarkan jumlah neuron Tabel 4.14.

Tabel 4. 14 Perbandingan Jumlah Neuron pada 2 *Hidden layer*

No	Neuron 64–32		Neuron 32–16	
	Model	Akurasi (%)	Model	Akurasi (%)
1.	Model G	76.50	Model H	76.00
2.	Model I	73.50	Model J	63.67
	Rata-rata	75.00	Rata-rata	69.83

Tabel 4.14 menunjukkan konfigurasi 64–32 mencatat rata-rata akurasi tertinggi sebesar 75,00%, sedangkan konfigurasi 32–16 menghasilkan akurasi lebih rendah yaitu dengan rata-rata 69,83%. Hal ini menunjukkan bahwa pada arsitektur 2 *hidden layer*, jumlah neuron yang lebih besar memberikan kapasitas representasi yang lebih memadai sehingga jaringan mampu memodelkan hubungan non-linier dalam data secara lebih efektif. Studi menunjukkan bahwa peningkatan jumlah neuron pada jaringan *multilayer perceptron* memperluas kemampuan model dalam menangkap variasi pola data, yang berdampak langsung pada peningkatan akurasi klasifikasi (Yoshida et al., 2021).

4.2.1.3 *Batch size*

Subbab ini membahas pengaruh *batch size* terhadap kinerja model klasifikasi.

(1) *3 Hidden layer*

Perbandingan akurasi model dengan arsitektur *3 hidden layer* berdasarkan penggunaan *batch size* Tabel 4.15.

Tabel 4. 15 Perbandingan *Batch size* pada *3 Hidden layer*

No	3 Hidden layer			
	Batch size 32		Batch size 64	
	Model	Akurasi (%)	Model	Akurasi (%)
1.	Model A	81.17	Model D	84.17
2.	Model B	87.67	Model E	87.17
3.	Model C	93.67	Model F	91.33
	Rata-rata	87.50	Rata-rata	87.56

Berdasarkan Tabel 4.15 terlihat bahwa perbedaan *batch size* pada arsitektur *3 hidden layer* tidak menghasilkan variasi akurasi yang signifikan. *Batch size 32* memiliki rata-rata akurasi 87,50%, sedangkan *batch size 64* sedikit lebih tinggi yaitu dengan rata-rata 87,56%.

Perbedaan akurasi yang sangat tipis antara *batch size 32* dan *64* pada arsitektur *3 hidden layer* menunjukkan bahwa kapasitas representasi model yang besar membuat pengaruh ukuran batch menjadi kurang signifikan. Pada jaringan yang lebih dalam, faktor arsitektur seperti jumlah *layer* dan neuron lebih dominan dalam menentukan performa, karena mekanisme pembelajaran telah mampu menangkap fitur data secara efektif (Hwang et al., 2024).

(2) *2 Hidden layer*

Perbandingan akurasi model dengan arsitektur *2 hidden layer* berdasarkan variasi *batch size* Tabel 4.16.

Tabel 4. 16 Perbandingan *Batch size* pada *2 Hidden layer*

No	2 Hidden layer			
	Batch size 32		Batch size 64	
	Model	Akurasi (%)	Model	Akurasi (%)
1.	Model G	76.50	Model I	73.50
2.	Model H	76.00	Model J	63.67
	Rata-rata	76.25	Rata-rata	68.59

Berdasarkan Tabel 4.16, penggunaan *batch size* 32 memberikan akurasi yang lebih tinggi pada seluruh model *2 hidden layer*, dengan rata-rata 76,25%. Sebaliknya, *batch size* 64 menunjukkan penurunan performa yang cukup jelas, dengan rata-rata 68,59%. Pola ini mengindikasikan bahwa untuk arsitektur *2 hidden layer* yang kapasitasnya lebih kecil, *batch size* yang lebih besar cenderung mengurangi stabilitas pembelajaran sehingga penangkapan pola citra MRI menjadi kurang optimal. Dengan demikian, *batch size* yang lebih kecil lebih sesuai untuk arsitektur *2 hidden layer* karena mampu mempertahankan kualitas pembaruan bobot selama pelatihan (Hwang et al., 2024).

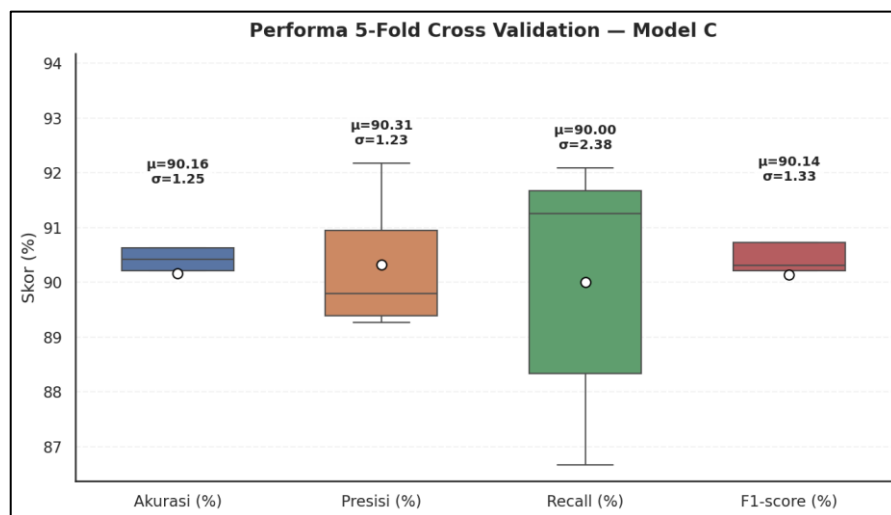
Hasil evaluasi menunjukkan bahwa konfigurasi *Multilayer Perceptron* (MLP) dengan tiga *hidden layer* berukuran 256–128–64 neuron dan *batch size* 32 memberikan performa paling optimal. Keunggulan model ini tercermin dari capaian nilai akurasi, presisi, *recall*, dan *F1-score* tertinggi serta kestabilan kinerja pada data uji, yang menekankan peran penting pemilihan arsitektur dan parameter pelatihan dalam proses klasifikasi citra.

4.3 Analisis *K-Fold Cross Validation* pada Model Terbaik

K-Fold Cross Validation diterapkan pada model terbaik untuk menguji kestabilan performa pada pembagian data yang berbeda. Data dibagi ke dalam beberapa *fold*, kemudian model dilatih dan diuji secara bergantian hingga setiap bagian pernah menjadi data uji. Pendekatan ini memastikan evaluasi tidak bergantung pada satu skenario pembagian data, sehingga mampu menunjukkan konsistensi dan keandalan performa model secara menyeluruh.

4.3.1 Model Terbaik 3 *Hidden layer*

Pengujian *K-Fold Cross Validation* dilakukan pada Model C guna menilai kestabilan kinerjanya terhadap variasi pembagian data. Visualisasi hasil pengujian ditampilkan melalui boxplot pada Gambar 4.12, sedangkan ringkasan nilai evaluasi disajikan pada Tabel 4.17.



Gambar 4. 12 *Boxplot* Hasil 5-Fold Cross Validation pada Model C

Gambar 4.12 *boxplot* hasil 5-fold cross validation, bahwa nilai akurasi, presisi, dan F1-score memiliki sebaran yang relatif sempit, menandakan bahwa

performa model cenderung stabil pada setiap fold. Sebaran pada metrik recall terlihat sedikit lebih lebar dibandingkan metrik lainnya, namun perbedaannya masih berada dalam batas yang wajar dan tidak menunjukkan penurunan performa yang signifikan.

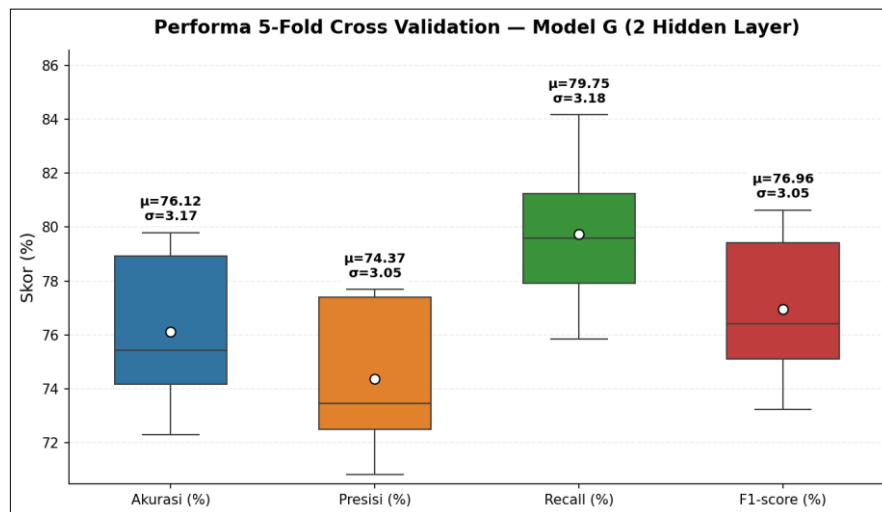
Tabel 4. 17 Hasil 5-Fold Cross Validation pada Model C

<i>Fold</i>	Akurasi (%)	Presisi (%)	Recall (%)	F1-score (%)
0	90.42	92.17	88.33	90.21
1	90.63	89.80	91.67	90.72
2	90.21	89.39	91.25	90.31
3	91.46	90.95	92.08	91.51
4	88.10	89.27	86.67	87.95
Rata-rata	90.16	90.31	90.00	90.14
Standar Deviasi	1.25	1.23	2.38	1.33

Tabel 4.17 memperlihatkan bahwa Model C mampu mencapai rata-rata akurasi 90,16%, presisi 90,31%, *recall* 90,00%, dan *F1-score* 90,14%. Rendahnya nilai standar deviasi yang berada pada kisaran 1,23% hingga 2,38% mengindikasikan variasi performa antar *fold* yang minimal. Kondisi tersebut menunjukkan bahwa kinerja Model C relatif konsisten dan tidak bergantung pada satu skenario pengujian tertentu.

4.3.2 Model Terbaik 2 *Hidden layer*

Pengujian 5-Fold Cross Validation diterapkan pada model terbaik dengan dua *hidden layer* guna mengevaluasi kestabilan performanya. Visualisasi hasil pengujian disajikan melalui boxplot pada Gambar 4.13, sementara ringkasan nilai evaluasi ditampilkan pada Tabel 4.18.



Gambar 4. 13 *Boxplot* Hasil 5-Fold Cross Validation pada Model G

Gambar 4.13 Berdasarkan hasil 5-fold cross validation, Model G menunjukkan performa yang cukup beragam pada setiap fold, yang tercermin dari sebaran nilai pada boxplot untuk keempat metrik evaluasi. Variasi ini mengindikasikan bahwa model dengan dua *hidden layer* masih memiliki keterbatasan dalam mempertahankan performa yang konsisten ketika dihadapkan pada pembagian data latih dan data uji yang berbeda.

Tabel 4. 18 Hasil 5-Fold Cross Validation pada Model G

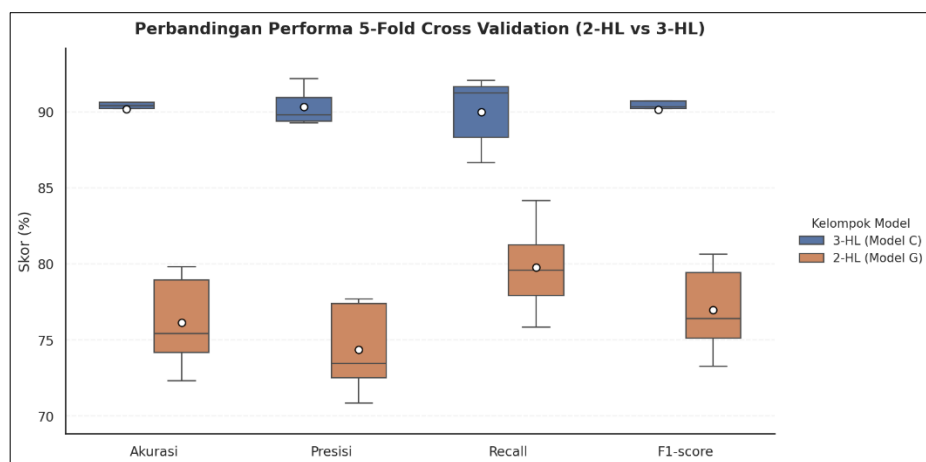
<i>Fold</i>	Akurasi (%)	Presisi (%)	Recall (%)	F1-score (%)
0	72.29	70.82	75.83	73.24
1	79.79	77.39	84.17	80.64
2	75.42	73.46	79.58	76.40
3	74.17	72.48	77.92	75.10
4	78.91	77.69	81.25	79.43
Rata-rata	76.12	74.37	79.75	76.96
Standar Deviasi	3.17	3.05	3.18	3.05

Berdasarkan Tabel 4.18, Model G memperoleh rata-rata akurasi 76,12%, presisi 74,37%, *recall* 79,75%, dan *F1-score* 76,96%. Nilai standar deviasi pada seluruh metrik yang berada di kisaran $\pm 3\%$ menunjukkan adanya variasi performa

antar *fold*, sehingga hasil prediksi model relatif sensitif terhadap perbedaan pembagian data. Kondisi ini mengindikasikan bahwa arsitektur dua *hidden layer* masih memiliki keterbatasan dalam menangkap pola tekstur citra MRI secara konsisten, sehingga kestabilan performa model belum optimal dan performa yang dihasilkan cenderung bervariasi pada setiap skenario pengujian.

4.3.3 Perbandingan Kestabilan Model Terbaik

Untuk memperjelas perbedaan kestabilan performa antara model dengan dua *hidden layer* dan tiga *hidden layer*, dilakukan visualisasi menggunakan *boxplot* berdasarkan hasil 5-Fold Cross Validation. Visualisasi ini digunakan untuk melihat sebaran nilai, median, serta variasi performa pada setiap metrik evaluasi secara lebih intuitif.



Gambar 4. 14 Perbandingan Performa Model 3-HL dan 2-HL Berdasarkan Boxplot Hasil 5-Fold Cross Validation

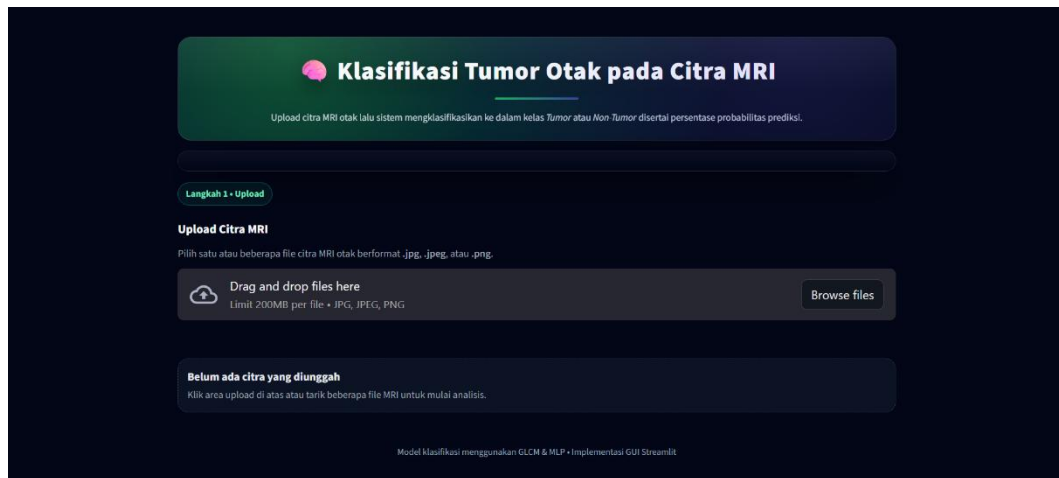
Berdasarkan Gambar 4.14, melalui penerapan *K-Fold Cross Validation*, Model C dengan arsitektur tiga *hidden layer* terbukti memiliki tingkat kestabilan performa yang lebih baik dibandingkan Model G dengan dua *hidden layer*. Indikasi tersebut terlihat dari rata-rata nilai akurasi, presisi, *recall*, dan F1-score yang lebih

tinggi serta variasi nilai yang lebih rendah pada setiap metrik. Visualisasi *boxplot* pada Gambar 4.14 menunjukkan sebaran nilai Model C yang lebih sempit pada seluruh metrik, sehingga menegaskan bahwa kinerja model tidak bergantung pada satu pembagian data tertentu, sejalan dengan prinsip evaluasi *k-fold cross-validation* (Teodorescu & Obreja Braşoveanu, 2025).

Selain itu, kestabilan performa Model C mengindikasikan bahwa arsitektur jaringan yang lebih dalam mampu mempelajari pola tekstur citra MRI hasil ekstraksi GLCM secara lebih baik dan konsisten pada setiap *fold*. Sebaliknya, Model G dengan kapasitas jaringan yang lebih kecil cenderung mengalami fluktuasi performa antar *fold*, yang pada Gambar 4.14 ditunjukkan oleh rentang dan sebaran nilai yang lebih lebar, serta standar deviasi yang lebih besar. Faktor pengaturan pelatihan seperti *batch size* juga berpengaruh terhadap kestabilan pembelajaran dan kemampuan generalisasi model, sebagaimana dilaporkan oleh Hwang et al., (2024).

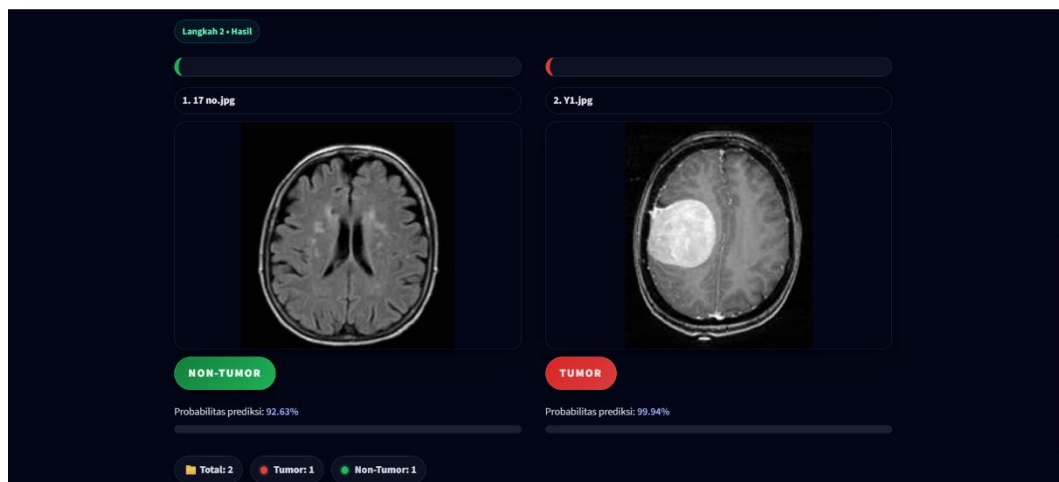
4.4 Implementasi GUI

Model klasifikasi tumor otak selanjutnya diaplikasikan dalam bentuk aplikasi berbasis web melalui perancangan *Graphical User Interface* (GUI) sebagai tahap akhir penelitian. GUI berfungsi sebagai media interaksi pengguna dalam menguji citra MRI dan menampilkan hasil prediksi sistem. Aplikasi ini dikembangkan menggunakan *Streamlit*, yaitu *framework Python* yang memungkinkan integrasi model *machine learning* dengan antarmuka web secara sederhana dan efisien tanpa pengembangan *front-end* yang kompleks. Tampilan antarmuka utama sistem Gambar 4.15.



Gambar 4. 15 GUI Sistem Klasifikasi Tumor Otak Berbasis Citra MRI

Pengguna mengunggah citra MRI otak melalui sistem, yang selanjutnya diproses melalui tahap *pre-processing* dan ekstraksi fitur tekstur menggunakan *Gray Level Co-Occurrence Matrix* (GLCM). Fitur yang dihasilkan kemudian menjadi masukan bagi model *Multilayer Perceptron* (MLP) untuk menentukan label klasifikasi citra.



Gambar 4. 16 Tampilan Hasil Klasifikasi Citra MRI pada GUI

Gambar 4.16 menampilkan halaman hasil klasifikasi, di mana setiap citra ditampilkan bersama label prediksi Tumor atau Tidak Tumor serta nilai probabilitas

dalam bentuk persentase. Perbedaan kelas ditampilkan secara visual untuk memudahkan interpretasi hasil. Selain itu, sistem juga menyediakan ringkasan berupa jumlah total citra yang diuji serta jumlah citra pada masing-masing kelas. Berdasarkan hasil implementasi, GUI yang dikembangkan mampu menjalankan fungsi sistem dengan baik dan konsisten dengan *pipeline* pengujian model, sehingga model klasifikasi tumor otak dapat diimplementasikan dalam bentuk aplikasi yang mudah digunakan.

Proses pengembangan sistem klasifikasi citra MRI dalam penelitian ini menunjukkan bahwa performa terbaik dicapai melalui tahapan pengujian yang dilakukan secara sistematis dan terukur. Variasi *hidden layer*, jumlah neuron, dan *batch size* menghasilkan kinerja yang berbeda, sehingga diperlukan konsistensi dalam penentuan konfigurasi untuk memperoleh model yang optimal. Sebagaimana disebutkan dalam QS. Al-Qamar ayat 49, Allah Swt. berfirman bahwa :

إِنَّا كُلَّ شَيْءٍ خَلَقْنَاهُ بِقَدَرٍ

“*Sesungguhnya Kami menciptakan segala sesuatu sesuai dengan ukuran*”
(QS.Al-Qamar/54:49)

Tafsir Al-Misbah Volume 14 (Shihab, 2009), menegaskan bahwa seluruh ciptaan Allah berlangsung berdasarkan *qadar*, yaitu ukuran, ketentuan, dan sistem yang telah ditetapkan secara presisi dan proporsional. Konsep ini menunjukkan bahwa tidak ada proses penciptaan yang bersifat acak, melainkan tersusun dalam keteraturan yang saling berkaitan. Dalam konteks penelitian ini, prinsip tersebut tercermin pada pentingnya keseimbangan antarparameter, di mana setiap

komponen memiliki fungsi dan peran yang terukur serta saling mendukung dalam membentuk kinerja model secara keseluruhan.

Selain itu, upaya mengembangkan sistem pendeteksi tumor otak juga merupakan bagian dari ikhtiar manusia dalam menghadapi penyakit. Hal ini ditegaskan dalam sabda Nabi ﷺ :

مَا أَنْزَلَ اللَّهُ دَاءً إِلَّا أَنْزَلَ لَهُ شِفَاءً

“Tidaklah Allah menurunkan suatu penyakit, melainkan Allah juga menurunkan obatnya.” (HR. Bukhari, no. 5678)

Hadis HR. Al-Bukhari No. 5678 yang disyarahkan oleh Ibnu Hajar al-‘Asqalani dalam *Fath al-Bari bi Sharh Sahih al-Bukhari* menegaskan bahwa setiap penyakit diciptakan bersamaan dengan potensi obatnya, sehingga manusia diperintahkan untuk melakukan ikhtiar dalam mencari pengobatan tanpa menafikan keyakinan bahwa kesembuhan berasal dari Allah SWT. Dalam konteks penelitian ini, pengembangan model klasifikasi citra MRI yang lebih akurat merupakan bentuk ikhtiar ilmiah untuk menemukan dan memanfaatkan sarana deteksi dini tumor otak. Dengan demikian, seluruh proses penelitian tidak hanya bernilai akademik, tetapi juga merepresentasikan upaya pemanfaatan ilmu pengetahuan untuk kemaslahatan manusia sesuai dengan tuntunan syariat Islam (Ibn Baz).

BAB V

KESIMPULAN DAN SARAN

5.1 Kesimpulan

Berdasarkan hasil penelitian dan pengujian yang telah dilakukan, diperoleh beberapa kesimpulan sebagai berikut:

1. Tahapan *pre-processing* mencakup *resize*, *normalization*, dan *noise removal*, serta proses ekstraksi fitur tekstur dengan *Gray Level Co-Occurrence Matrix* (GLCM) terbukti meningkatkan kualitas data. Proses ini menghasilkan fitur-fitur tekstur yang lebih stabil dan informatif sehingga membantu model dalam membedakan citra tumor dan tidak tumor.
2. Kedalaman arsitektur jaringan dan jumlah neuron yang lebih besar memberikan kapasitas pembelajaran yang lebih baik. Dengan arsitektur tersebut, jaringan mampu mempelajari pola tekstur secara lebih menyeluruh sehingga kinerja klasifikasi mengalami peningkatan.
3. Model C dengan tiga *hidden layer* berukuran 256, 128, dan 64 neuron serta *batch size* 32 adalah konfigurasi paling optimal dengan akurasi 93,67%. Susunan tersebut memberikan kapasitas representasi yang memadai untuk menangkap variasi tekstur citra secara menyeluruh.

5.2 Saran

Berdasarkan evaluasi terhadap sepuluh skenario model MLP, diperoleh sejumlah rekomendasi sebagai dasar pengembangan penelitian klasifikasi tumor otak berbasis citra MRI selanjutnya.

1. Eksplorasi konfigurasi model yang lebih luas.

Penelitian selanjutnya dapat mengembangkan arsitektur *Multilayer Perceptron* (MLP) dengan memvariasikan *hidden layer*, jumlah neuron, *activation function*, serta teknik *regularization*, disertai optimasi *hyperparameter* menggunakan *GridSearchCV*, *RandomizedSearchCV*, atau *Bayesian Optimization* guna memperoleh konfigurasi yang lebih optimal.

2. Penambahan ragam fitur tekstur.

Pengembangan selanjutnya dapat mengeksplorasi fitur *Haralick* secara lengkap, metode *Local Binary Pattern (LBP)*, atau kombinasi beberapa metode *feature extraction* guna menghasilkan representasi citra yang lebih kaya dan informatif.

3. Eksperimen pada variasi pembagian data.

Uji coba dengan menerapkan perbandingan data latih dan data uji yang berbeda, seperti 70:30, 75:25, atau 90:10, dapat dilakukan untuk menilai pengaruh ketersediaan data latih terhadap stabilitas dan konsistensi performa model.

DAFTAR PUSTAKA

- Aggarwal, A. K. (2022). Learning Texture Features from GLCM for Classification of Brain Tumor MRI Images using Random Forest Classifier. *WSEAS TRANSACTIONS ON SIGNAL PROCESSING*, 18, 60–63. <https://doi.org/10.37394/232014.2022.18.8>
- Al Haris, M., Aulia, S., & Wasono, R. (2025). Multi-layer Perceptron for Brain Tumor Picture Classification Based on GLCM Feature Extraction. In Y. Ben Olina, F. Nur Damayanti, & Y. Tursinawati (Eds.), *3rd Lawang Sewu International Symposium on Medical and Health Sciences (3rd-LEWIS-MHS)* (Vol. 85, pp. 268–283). Atlantis Press International BV. https://doi.org/10.2991/978-94-6463-760-1_24
- Alagoz, B. B., Keles, C., Ates, A., Özdemir, E., & Alkhulaifi, N. (2025). Optimal deep neural network architecture design with improved generalization for data-driven cooling load estimation problem. *Neural Computing and Applications*, 37(19), 13597–13616. <https://doi.org/10.1007/s00521-025-11212-7>
- Alanazi, T. M., Berriri, K., Albekairi, M., Ben Atitallah, A., Sahbani, A., & Kaaniche, K. (2023). New Real-Time High-Density Impulsive Noise Removal Method Applied to Medical Images. *Diagnostics*, 13(10), 1709. <https://doi.org/10.3390/diagnostics13101709>
- Alayed, A. (2024). Machine Learning and Deep Learning Approaches for Arabic Sign Language Recognition: A Decade Systematic Literature Review. *Sensors*, 24(23), 7798. <https://doi.org/10.3390/s24237798>
- Albahli, S. (2025). A Robust YOLOv8-Based Framework for Real-Time Melanoma Detection and Segmentation with Multi-Dataset Training. *Diagnostics*, 15(6), 691. <https://doi.org/10.3390/diagnostics15060691>
- Alsalihi, A., K. Aljobouri, H., & ALTameemi, E. A. K. (2022). GLCM and CNN Deep Learning Model for Improved MRI Breast Tumors Detection. *International Journal of Online and Biomedical Engineering (iJOE)*, 18(12), 123–137. <https://doi.org/10.3991/ijoe.v18i12.31897>
- Amaliah Faradibah, Dewi Widyawati, A Ulfah Tenripada Syahar, & Sitti Rahmah Jabir. (2023). Comparison Analysis of Random Forest Classifier, Support Vector Machine, and Artificial Neural Network Performance in Multiclass Brain Tumor Classification. *Indonesian Journal of Data and Science*, 4(2), 54–63. <https://doi.org/10.56705/ijodas.v4i2.73>

- Asif Khan, M. K., Ali, F., Irfan, M., Muhammad, F., Althobiani, F., Ali, A., Khan, S., Rahman, S., Perun, G., & Glowacz, A. (2021). Mitigation of Phase Noise and Nonlinearities for High Capacity Radio-over-Fiber Links. *Electronics*, 10(3), 345. <https://doi.org/10.3390/electronics10030345>
- Barburiceanu, S., Terebes, R., & Meza, S. (2021). 3D Texture Feature Extraction and Classification Using GLCM and LBP-Based Descriptors. *Applied Sciences*, 11(5), 2332. <https://doi.org/10.3390/app11052332>
- Bhattacharya, S., Bennet, L., Davidson, J. O., & Unsworth, C. P. (2022). Multi-layer perceptron classification & quantification of neuronal survival in hypoxic-ischemic brain image slices using a novel gradient direction, grey level co-occurrence matrix image training. *PLOS ONE*, 17(12), e0278874. <https://doi.org/10.1371/journal.pone.0278874>
- Citra R, F., Indriyani, F., & Rahadjeng, I. R. (2024). Klasifikasi Tumor Otak Berbasis Magnetic Resonance Imaging Menggunakan Algoritma Convolutional Neural Network. *Digital Transformation Technology*, 3(2), 918–924. <https://doi.org/10.47709/digitech.v3i2.3469>
- Cornelissen, S., Schouten, S. M., Langenhuizen, P. P. J. H., Lie, S. T., Kunst, H. P. M., De With, P. H. N., & Verheul, J. B. (2024). Defining tumor growth in vestibular schwannomas: A volumetric inter-observer variability study in contrast-enhanced T1-weighted MRI. *Neuroradiology*, 66(11), 2033–2042. <https://doi.org/10.1007/s00234-024-03416-w>
- Dagal, I., Harrison, A., Ibrahim, A.-W., & Mbasso, W. F. (2025). Comprehensive evaluation of data preprocessing and visualization techniques for enhanced classification and sampling. *Cluster Computing*, 28(7), 476. <https://doi.org/10.1007/s10586-025-05512-9>
- Dheepak, G., J., A. C., & Vaishali, D. (2024). Brain tumor classification: A novel approach integrating GLCM, LBP and composite features. *Frontiers in Oncology*, 13. <https://doi.org/10.3389/fonc.2023.1248452>
- Essianda, V., Indrasari, A. D., Widyastuti, P., Syahla, T., & Rohadi, R. (2023). Brain Tumor: Molecular Biology, Pathophysiology, and Clinical Symptoms. *Jurnal Biologi Tropis*, 23(4), 260–269. <https://doi.org/10.29303/jbt.v23i4.5585>
- Fan, X., Zhang, Y., Peng, Y., Li, Q., Wei, X., Wang, J., & Zou, F. (2025). LO-MLPRNN: A Classification Algorithm for Multispectral Remote Sensing Images by Fusing Selective Convolution. *Sensors*, 25(8), 2472. <https://doi.org/10.3390/s25082472>
- Fikriah, F. K., Ariyanto, A. D. P., & Setyawan, A. F. (2024a). KLASIFIKASI HASIL MRI TUMOR OTAK DENGAN EKTRAKSI FITUR GRAY

- LEVEL CO-OCCURANCE MATRIX (GLCM). *Rabit : Jurnal Teknologi dan Sistem Informasi Univrab*, 9(2), 343–350. <https://doi.org/10.36341/rabit.v9i2.4793>
- Fikriah, F. K., Ariyanto, A. D. P., & Setyawan, A. F. (2024b). KLASIFIKASI HASIL MRI TUMOR OTAK DENGAN EKTRAKSI FITUR GRAY LEVEL CO-OCCURANCE MATRIX (GLCM). *Rabit : Jurnal Teknologi dan Sistem Informasi Univrab*, 9(2), 343–350. <https://doi.org/10.36341/rabit.v9i2.4793>
- Gül, M., & Kaya, Y. (2024). Comparing of brain tumor diagnosis with developed local binary patterns methods. *Neural Computing and Applications*, 36(13), 7545–7558. <https://doi.org/10.1007/s00521-024-09476-6>
- Hammad, M., Pławiak, P., ElAffendi, M., El-Latif, A. A. A., & Latif, A. A. A. (2023). Enhanced Deep Learning Approach for Accurate Eczema and Psoriasis Skin Detection. *Sensors*, 23(16), 7295. <https://doi.org/10.3390/s23167295>
- Haralick, R. M., Shanmugam, K., & Dinstein, I. (1973). Textural Features for Image Classification. *IEEE Transactions on Systems, Man, and Cybernetics, SMC-3*(6), 610–621. <https://doi.org/10.1109/TSMC.1973.4309314>
- Husnal, F., Wijaya, F. A., Fauzi, F., Paryudi, I., Veritawati, I., & Nursari, S. R. C. (2023). *ANALISIS GAMBAR MRI OTAK UNTUK MENDETEKSI TUMOR OTAK MENGGUNAKAN ALGORITMA CNN. 4.*
- Hwang, J.-S., Lee, S.-S., Gil, J.-W., & Lee, C.-K. (2024). Determination of Optimal *Batch size* of Deep Learning Models with Time Series Data. *Sustainability*, 16(14), 5936. <https://doi.org/10.3390/su16145936>
- Ibn Baz, A. A. (-). *FATHUL BAARI. 10.*
- Istianah, L., & Sumarti, H. (2020). Classification of Pneumonia in Thoracic X-Ray images based on texture characteristics using the MLP (Multi-Layer Perceptron) method. *Journal of Natural Sciences and Mathematics Research*, 6(2), 78–84. <https://doi.org/10.21580/jnsmr.2020.6.2.11228>
- Jeong, J., Bang, S., Jung, Y., & Jo, J. (2025). A Study on the Performance Comparison of Brain MRI Image-Based Abnormality Classification Models. *Life*, 15(10), 1614. <https://doi.org/10.3390/life15101614>
- Kabir, M., Unal, F., Akinci, T. C., Martinez-Morales, A. A., & Ekici, S. (2024). Revealing GLCM Metric Variations across a Plant Disease Dataset: A Comprehensive Examination and Future Prospects for Enhanced Deep Learning Applications. *Electronics*, 13(12), 2299. <https://doi.org/10.3390/electronics13122299>

- Karami, A., Rezaei, S. M., Motamedfar, A., Shahin, B., Nazari, P., & Cheki, M. (2025). Developing a predictive model for breast cancer detection using radiomics-based mammography and machine learning. *Egyptian Journal of Radiology and Nuclear Medicine*, 56(1), 165. <https://doi.org/10.1186/s43055-025-01588-w>
- Khan, M. A., & Auvee, R. B. Z. (2024). *Comparative Analysis of Resource-Efficient CNN Architectures for Brain Tumor Classification* (No. arXiv:2411.15596). arXiv. <https://doi.org/10.48550/arXiv.2411.15596>
- Kirana, K. C., Nidhom, A. M., Fadhlullah, A. F., Siregar, G. C. P., & Begananda, H. B. (2023). *Klasifikasi Penyakit Tumor Otak Menggunakan K-Nearest Neighbour Berbasis Grey Level Coocurrence Matrix*. 33(1).
- Kohsasih, K. L., Agung Rizky, M. D., Rosnelly, R., & Widjaja, W. W. (2022). A deep learning model to detect the brain tumor based on magnetic resonance images. *JURNAL INFOTEL*, 14(3), 203–208. <https://doi.org/10.20895/infotel.v14i3.793>
- Kumar, K., Jyoti, K., & Kumar, K. (2024). Machine learning for brain tumor classification: Evaluating feature extraction and algorithm efficiency. *Discover Artificial Intelligence*, 4(1), 112. <https://doi.org/10.1007/s44163-024-00214-4>
- Maharana, K., Mondal, S., & Nemade, B. (2022). A review: Data pre-processing and data augmentation techniques. *Global Transitions Proceedings*, 3(1), 91–99. <https://doi.org/10.1016/j.gltp.2022.04.020>
- Malkauthekar, M., Gulve, A., & Deshmukh, R. (2025). A novel Probabilistic Bi-Level Teaching–Learning-Based Optimization (P-BTLBO) algorithm for hybrid feature extraction and multi-class brain tumor classification using ResNet-50 and GLCM. *Journal of Engineering and Applied Science*, 72(1). <https://doi.org/10.1186/s44147-025-00671-3>
- Mathi, S., Deb, D., Chaudhuri, A. K., Joseph, L., & Biswas, D. S. (2025). A Review Of Convolutional Neural Networks For Medical Image Analysis: Trends And Future Directions. *Journal of Neonatal Surgery*, 14(7).
- Mienye, I. D., & Swart, T. G. (2024). A Comprehensive Review of Deep Learning: Architectures, Recent Advances, and Applications. *Information*, 15(12), 755. <https://doi.org/10.3390/info15120755>
- Müller, D., Soto-Rey, I., & Kramer, F. (2022). Towards a guideline for evaluation metrics in medical image segmentation. *BMC Research Notes*, 15(1), 210. <https://doi.org/10.1186/s13104-022-06096-y>

- Na, L., & Yang, Z. (2021). *International Journal of Cognitive Informatics and Natural Intelligence*. 15(4). <https://doi.org/10.4018/ijcini>
- Ostrom, Q. T., Francis, S. S., & Barnholtz-Sloan, J. S. (2021). Epidemiology of Brain and Other CNS Tumors. *Current Neurology and Neuroscience Reports*, 21(12), 68. <https://doi.org/10.1007/s11910-021-01152-9>
- Oybek Kizi, R. F., Theodore Armand, T. P., & Kim, H.-C. (2025). A Review of Deep Learning Techniques for Leukemia Cancer Classification Based on Blood Smear Images. *Applied Biosciences*, 4(1), 9. <https://doi.org/10.3390/applbiosci4010009>
- Pagadala, P. K., Omkari, D. Y., Lakshmi, P. S., & Sutresno, S. A. (n.d.). AUTOMATED BRAIN TUMOR DETECTION WITH GLCM- BASED FEATURE EXTRACTION AND PCA FOR DIMENSION REDUCTION AND CLASSIFICATION USING MACHINE LEARNING. . . *Vol.*, 12.
- Pattanaik, B., Anitha, K., Rathore, S., Biswas, P., Sethy, P., & Behera, S. (2022). Brain tumor magnetic resonance images classification based machine learning paradigms. *Współczesna Onkologia*, 26(4), 268–274. <https://doi.org/10.5114/wo.2023.124612>
- Rainio, O., Teuho, J., & Klén, R. (2024). Evaluation metrics and statistical tests for machine learning. *Scientific Reports*, 14(1), 6086. <https://doi.org/10.1038/s41598-024-56706-x>
- Saeedi, S., Rezayi, S., Keshavarz, H., & R. Niakan Kalhori, S. (2023). MRI-based brain tumor detection using convolutional deep learning methods and chosen machine learning techniques. *BMC Medical Informatics and Decision Making*, 23(1), 16. <https://doi.org/10.1186/s12911-023-02114-6>
- Santamato, V., & Marengo, A. (2025). Multilabel Classification of Radiology Image Concepts Using Deep Learning. *Applied Sciences*, 15(9), 5140. <https://doi.org/10.3390/app15095140>
- Santoso, M. H., Larasati, D. A., & Muhathir. (2020). Wayang Image Classification Using MLP Method and GLCM Feature Extraction. *Journal of Computer Science, Information Technology and Telecommunication Engineering*. <https://doi.org/10.30596/jcositte.v1i2.5131>
- Sari, R. N., Aninditha, T., Rachman, A., Ranakusuma, T. A. S., Andriani, R., Yesnuar Safri, A., Sofyan, H. R., & Wiratman, W. (2025). Characteristics of Brain Tumor Metastases at dr. Cipto Mangunkusumo National Referral Hospital. *eJournal Kedokteran Indonesia*, 12(3), 279. <https://doi.org/10.23886/ejki.12.873.279>

- Sembiring, S., Sofyan, H. R., Madjid, I. S., Alvonsius Silalahi, R. A. N., Albertin, E., & Aninditha, T. (2025). Headache Characteristic in Brain Tumor Patients. *Indonesian Journal of Cancer*, 19(1), 34–39. <https://doi.org/10.33371/ijoc.v19i1.1197>
- Shihab, M. Q. (2009a). *Tafsir al-mishbah: Pesan, kesan dan keserasian al-Qur'an* (Edisi baru, Vols. 1–2). Lentera Hati.
- Shihab, M. Q. (2009b). *Tafsir al-mishbah: Pesan, kesan dan keserasian al-Qur'an* (Edisi baru, Vols. 1–14). Lentera Hati.
- Sildir, H., Aydin, E., & Kavzoglu, T. (2020). Design of Feedforward Neural Networks in the Classification of Hyperspectral Imagery Using Superstructural Optimization. *Remote Sensing*, 12(6), 956. <https://doi.org/10.3390/rs12060956>
- Siregar, N. L. (2023). *EKSTRAKSI FITUR CITRA BERDASARKAN TEKSTUR DENGAN GLCM (GRAY LEVEL CO-OCCURRENCE)*. 3(1).
- Sowrirajan, S. R., & Balasubramanian, S. (2022). Brain Tumor Classification Using Machine Learning and Deep Learning Algorithms. *International Journal of Electrical and Electronics Research*, 10(4), 999–1004. <https://doi.org/10.37391/ijeer.100441>
- Sravanthi Peddinti, A., Maloji, S., & Manepalli, K. (2021). Evolution in diagnosis and detection of brain tumor – review. *Journal of Physics: Conference Series*, 2115(1), 012039. <https://doi.org/10.1088/1742-6596/2115/1/012039>
- Srivaishnavi, K. R., Perumal, T. P., & Anishiya, P. (2025). Feature Extraction from Segmented Brain MRI Images using Traditional Methods and Classification Analysis of Brain Tumor using Machine Learning Classifiers. *Indian Journal Of Science And Technology*, 18(17), 1309–1320. <https://doi.org/10.17485/IJST/v18i17.1902>
- Suta, I. B. L. M., Hartati, R. S., & Divayana, Y. (2019). Diagnosa Tumor Otak Berdasarkan Citra MRI (Magnetic Resonance Imaging). *Majalah Ilmiah Teknologi Elektro*, 18(2). <https://doi.org/10.24843/MITE.2019.v18i02.P01>
- Tahosin, M. S., Sheakh, M. A., Islam, T., Lima, R. J., & Begum, M. (2023). Optimizing brain tumor classification through feature selection and hyperparameter tuning in machine learning models. *Informatics in Medicine Unlocked*, 43, 101414. <https://doi.org/10.1016/j.imu.2023.101414>
- Teodorescu, V., & Obreja Braşoveanu, L. (2025). Assessing the Validity of k-Fold Cross-Validation for Model Selection: Evidence from Bankruptcy Prediction Using Random Forest and XGBoost. *Computation*, 13(5), 127. <https://doi.org/10.3390/computation13050127>

- Ullu, H. H., Baso, B., Risald, R., Manek, P. G., & Chrisinta, D. (2022). Ekstraksi Fitur Berbasis Tekstur Pada Citra Tenun Timor Menggunakan Metode Gray Level Co-occurrence Matrix (GLCM). *Journal of Information and Technology*, 2(2), 70–74. <https://doi.org/10.32938/jitu.v2i2.3245>
- Vijayalakshmi, S., Pandey, B. K., Pandey, D., & Lelisho, M. E. (2025). Innovative deep learning classifiers for breast cancer detection through hybrid feature extraction techniques. *Scientific Reports*, 15(1), 22212. <https://doi.org/10.1038/s41598-025-06669-4>
- Vijithananda, S. M., Jayatilake, M. L., Hewavithana, B., Gonçalves, T., Rato, L. M., Weerakoon, B. S., Kalupahana, T. D., Silva, A. D., & Dissanayake, K. D. (2022). Feature extraction from MRI ADC images for brain tumor classification using machine learning techniques. *BioMedical Engineering OnLine*, 21(1), 52. <https://doi.org/10.1186/s12938-022-01022-6>
- Wahyu Ardiantito S, Stacyana Jesika Surianto, Suci Ramadhani, & Willy Pramudia Ananta. (2023). Klasifikasi Tumor Otak Menggunakan Local Binary Pattern dan SVM Classifier. *Student Research Journal*, 1(6), 182–190. <https://doi.org/10.55606/srjyappi.v1i6.823>
- Waldo, J. (2022). *A comparative study of back propagation and its alternatives on multilayer perceptrons* (No. arXiv:2206.06098). arXiv. <https://doi.org/10.48550/arXiv.2206.06098>
- Wang, N., Liang, C., Zhang, X., Sui, C., Gao, Y., Guo, L., & Wen, H. (2023). Brain structure–function coupling associated with cognitive impairment in cerebral small vessel disease. *Frontiers in Neuroscience*, 17. <https://doi.org/10.3389/fnins.2023.1163274>
- Yoshida, H., Hasegawa, Y., Matsushima, M., Sugiyama, T., Kawabe, T., & Shikida, M. (2021). Miniaturization of Respiratory Measurement System in Artificial Ventilator for Small Animal Experiments to Reduce Dead Space and Its Application to Lung Elasticity Evaluation. *Sensors*, 21(15), 5123. <https://doi.org/10.3390/s21155123>
- Yousaf, S., Anwar, S. M., RaviPrakash, H., & Bagci, U. (2020). *Brain Tumor Survival Prediction using Radiomics Features* (No. arXiv:2009.02903). arXiv. <https://doi.org/10.48550/arXiv.2009.02903>