

BAB 2

LANDASAN TEORI

2.1 Penelitian Terdahulu

Penelitian mengenai penerapan *deep learning* untuk diagnosis pada citra medis telah berkembang secara signifikan. Studi-studi terdahulu yang ditinjau mencakup perbedaan pada beberapa aspek, seperti modalitas citra yang digunakan, jenis algoritma, dan arsitektur model. Tinjauan terhadap studi-studi tersebut menjadi landasan untuk memetakan posisi dan kontribusi penelitian ini di antara penelitian yang sudah ada. Perbandingan menyeluruh antar penelitian terkait disajikan pada Tabel 2.1.

Tabel 2.1. Matriks Perbandingan Penelitian Terkait dan Penelitian Ini

No	Aspek / Fitur Utama	[1] Aykaç et al. (2025)	[2] Kamencay et al. (2025)	[5] Han et al. (2025)	[6] Kim et al. (2024)	[19] Portase & Ardelean (2026)	Penelitian Ini (2026)
1	Modalitas citra TOF-MRA	✓	X	X	X	X	✓
2	Menggunakan model YOLO26	X	X	X	X	✓	✓
3	Menggunakan arsitektur <i>single-stage detector</i>	X	✓	✓	X	✓	✓
4	Deteksi objek dengan <i>bounding box</i>	X	✓	✓	X	✓	✓
5	Tanpa modifikasi arsitektur jaringan secara manual	✓	✓	X	✓	✓	✓
6	Menyertakan sampel negatif (pasien sehat) dalam pelatihan	✓	X	X	✓	X	✓
7	Target deteksi aneurisma intrakranial	✓	✓	X	✓	X	✓

Keterangan: Simbol ✓ = ada, X = tidak ada.

Aykaç et al. [1] mengembangkan sistem deteksi dan segmentasi aneurisma intrakranial menggunakan model Mask R-CNN pada citra TOF-MRA. Model yang dikembangkan mencapai nilai *Dice Similarity Coefficient* sebesar 0,8832, presisi 0,9404, dan *recall* 0,8677 pada *dataset* yang terdiri dari 447 subjek. Namun, arsitektur Mask R-CNN bekerja dalam dua tahap terpisah, yaitu tahap pertama menghasilkan kandidat area yang perlu diperiksa, lalu tahap kedua melakukan analisis lebih mendalam pada setiap kandidat tersebut. Proses dua tahap ini membuat waktu pemrosesan per gambar menjadi lebih lama dan membutuhkan

perangkat keras dengan kapasitas komputasi yang besar.

Kamencay et al. [2] mengevaluasi empat varian model YOLO, yaitu YOLOv8n, YOLOv8m, YOLOv11n, dan YOLOv11m, untuk mendeteksi aneurisma intrakranial pada citra CTA. Model-model berbasis *single-stage detector* ini mampu memproses gambar secara cepat dengan jumlah parameter yang relatif kecil. Dari keempat varian, YOLOv8m menghasilkan keseimbangan terbaik antara presisi dan sensitivitas dengan nilai *F1-score* tertinggi sebesar 0,6232, sementara YOLOv11n mencatat nilai *recall* tertinggi sebesar 0,5814. Namun, *dataset* yang digunakan hanya berasal dari satu rumah sakit dan seluruh gambar merupakan pasien dengan aneurisma, tanpa satu pun gambar pasien sehat sebagai pembanding, sehingga kemampuan model dalam membedakan pembuluh darah normal dari lesi aneurisma tidak diuji dalam studi ini.

Han et al. [6] memodifikasi arsitektur YOLOv11 untuk meningkatkan akurasi deteksi tumor otak pada citra MRI. Studi ini berfokus pada tumor otak, bukan aneurisma intrakranial, sehingga konteks klinis dan karakteristik objek yang dideteksi berbeda dari penelitian ini. Modifikasi dilakukan dengan menyisipkan tiga modul atensi *custom* secara manual ke dalam lapisan jaringan. *Spatial attention* bekerja dengan menentukan area mana dalam gambar yang paling perlu diperhatikan oleh model. *Shuffle3D attention* mengubah kombinasi urutan fitur secara dinamis, sehingga model mempelajari hubungan antar-karakteristik klinis gambar secara lebih luas. *Dual-channel attention* memproses informasi dari dua jalur penyaringan gambar yang berbeda secara bersamaan agar dapat mengumpulkan detail fitur medis yang lebih kaya. Hasilnya, model yang dimodifikasi mencapai nilai *mAP₅₀* sebesar 96,8% pada *dataset* Kaggle publik. Akan tetapi, modifikasi struktur jaringan secara manual seperti ini menuntut penyesuaian kode program yang cukup rumit dan dapat menimbulkan ketidakstabilan selama proses pelatihan model, terutama ketika diterapkan pada *dataset* medis yang volumenya terbatas.

Kim et al. [5] mengembangkan metode penyaringan tambahan berbasis logika struktur tubuh untuk menekan angka *false positive* pada model deteksi aneurisma dari citra CTA. Metode ini bekerja dengan cara menghapus tebakan model AI yang muncul di lokasi yang tidak masuk akal secara medis, seperti tebakan yang berada di area pembuluh darah balik (vena) atau yang terletak di luar wilayah otak. Berdasarkan uji coba, pendekatan ini berhasil membuang 70,6% prediksi salah pada model AI CPM-Net tanpa melewatkan satu pun penyakit asli yang seharusnya terdeteksi. Meski efektif, pendekatan tersebut memerlukan modul

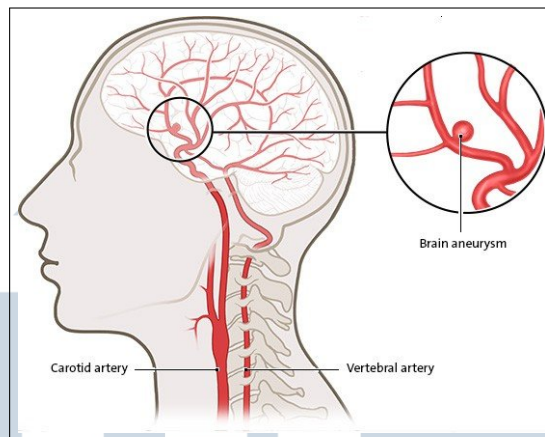
komputasi tambahan yang terpisah di luar jaringan kecerdasan buatan utama. Hal ini meningkatkan kerumitan alur kerja sistem serta memperpanjang total waktu pemrosesan yang dibutuhkan hingga hasil akhir diagnosis keluar.

Portase dan Ardelean [8] mengevaluasi sembilan model deteksi objek, termasuk YOLO26, untuk mendeteksi polip usus pada citra kolonoskopi menggunakan tiga dataset publik. Perbandingan dilakukan dalam kondisi pelatihan yang sama tanpa modifikasi arsitektur tambahan pada model mana pun. Hasil ini menunjukkan YOLO26 sudah diuji pada tugas deteksi objek berbasis *bounding box* di modalitas citra medis lain, yaitu kolonoskopi. Penerapan YOLO26 pada citra TOF-MRA untuk mendeteksi aneurisma intrakranial belum ditemukan pada penelitian mana pun yang ditinjau di atas.

Berdasarkan tinjauan pustaka di atas, sebagian besar penelitian terdahulu masih mengandalkan model *two-stage* yang lambat, modifikasi arsitektur *custom* manual yang rumit, atau modul *post-processing* yang memperpanjang rantai komputasi, sedangkan penerapan YOLO26 pada citra medis masih sangat terbatas. Guna mengisi celah riset tersebut, penelitian ini menerapkan model Ultralytics YOLO26 varian *Small* (`yolo26s.pt`) versi standar tanpa modifikasi lapisan eksternal pada citra TOF-MRA. Melalui pemanfaatan fitur bawaan model generasi terbaru yang sudah dioptimalkan secara arsitektural untuk menghasilkan alur prediksi langsung tanpa proses penyaringan tradisional, penelitian ini diharapkan mampu mencapai performa akurasi deteksi kelainan vaskular yang optimal dengan efisiensi waktu pemrosesan yang tinggi.

2.2 Aneurisma Intrakranial

Aneurisma intrakranial adalah kondisi di mana dinding pembuluh darah arteri di dalam otak mengalami pelebaran atau penonjolan secara abnormal akibat melemahnya struktur dinding pembuluh tersebut [9]. Kondisi ini dapat dipicu oleh berbagai faktor, mulai dari riwayat genetik keluarga, tekanan darah tinggi, hingga kebiasaan merokok [10]. Secara statistik, diperkirakan 3 hingga 5 persen dari total populasi umum memiliki aneurisma intrakranial [11]. Sebagian besar penderita tidak merasakan gejala apa pun, sehingga kondisi ini kerap ditemukan secara tidak sengaja saat pemeriksaan untuk penyakit lain [9]. Ilustrasi bentuk aneurisma intrakranial pada pembuluh darah otak dapat dilihat pada Gambar 2.1.



Gambar 2.1. Ilustrasi bentuk aneurisma intrakranial pada pembuluh darah otak.

Sumber: [12]

Walaupun sering kali tidak bergejala, aneurisma intrakranial menyimpan potensi bahaya apabila dibiarkan hingga pecah. Bahaya utama aneurisma intrakranial adalah pecahnya pembuluh darah yang menonjol tersebut [9]. Pecahnya aneurisma memicu pendarahan akut di area sekitar otak, yang merupakan kondisi darurat dengan tingkat kematian yang mencapai 50 persen [11]. Pasien yang berhasil selamat pun sering kali harus menanggung dampak kelumpuhan atau kerusakan saraf permanen [1]. Oleh karena itu, langkah pemeriksaan dan deteksi dini menjadi langkah yang sangat menentukan kualitas dan keberlangsungan hidup pasien.

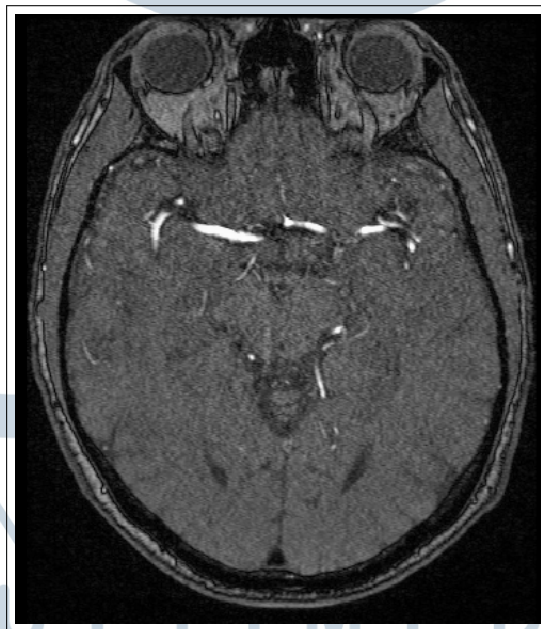
Salah satu parameter utama yang digunakan dokter untuk menilai risiko dan menentukan penanganan adalah ukuran diameter aneurisma [7]. Berdasarkan ukuran diameter maksimalnya, aneurisma secara umum dapat diklasifikasikan ke dalam empat kategori utama [7]. Kategori pertama adalah aneurisma berukuran sangat kecil (*tiny*) yang memiliki diameter di bawah 3 mm. Kategori kedua merupakan aneurisma berukuran kecil (*small*) dengan rentang diameter mulai dari 3 mm hingga di bawah 5 mm. Selanjutnya, terdapat kategori aneurisma sedang (*medium*) yang ukuran diameternya berkisar antara 5 mm hingga di bawah 7 mm. Kategori yang terakhir adalah aneurisma berukuran besar (*large*) yang memiliki diameter mencapai 7 mm atau lebih besar [7].

Pengelompokan ukuran ini membantu para dokter untuk mengenali karakteristik pembengkakan, di mana aneurisma yang berukuran lebih besar umumnya memiliki risiko pecah yang lebih tinggi [7]. Namun, tidak menutup kemungkinan bahwa aneurisma yang berukuran sangat kecil juga memiliki risiko pecah, sehingga keberadaannya tidak boleh diabaikan dalam pemeriksaan [7].

Sayangnya, aneurisma dengan ukuran yang sangat kecil ini merupakan kelainan yang paling rentan terlewat saat dokter memeriksa irisan gambar secara manual [1]. Oleh karena itu, penggunaan bantuan teknologi kecerdasan buatan menjadi penting dalam proses deteksi semua variasi ukuran aneurisma [7].

2.3 Time-of-Flight Magnetic Resonance Angiography (TOF-MRA)

Time-of-Flight Magnetic Resonance Angiography (TOF-MRA) adalah teknik pencitraan medis yang secara khusus dirancang untuk memvisualisasikan aliran pembuluh darah di dalam otak [13]. Metode ini bekerja dengan memanfaatkan pergerakan aliran darah itu sendiri, di mana darah yang bergerak menghasilkan sinyal lebih terang pada citra dibandingkan jaringan otak yang diam di sekitarnya [14]. Karena mekanisme ini sepenuhnya mengandalkan pergerakan aliran darah, TOF-MRA dapat menghasilkan gambaran struktur pembuluh darah yang jelas tanpa membutuhkan penyuntikan cairan agen kontras ke dalam tubuh pasien [13]. Contoh hasil pencitraan TOF-MRA dari dataset yang digunakan dalam penelitian ini dapat dilihat pada Gambar 2.2.



Gambar 2.2. Contoh hasil pencitraan TOF-MRA dari dataset yang digunakan dalam penelitian ini.
Sumber: [13]

Dibandingkan dengan CTA dan DSA yang juga umum digunakan untuk mendiagnosis aneurisma intrakranial, TOF-MRA tidak melibatkan radiasi ion

karena mengandalkan medan magnet untuk menilai aliran darah, sedangkan *Computed Tomography Angiography* (CTA) dan *Digital Subtraction Angiography* (DSA) menghasilkan citra menggunakan radiasi *X-ray* [2]. DSA berperan sebagai *gold standard* untuk konfirmasi obliterasi aneurisma, namun prosedurnya mengharuskan akses arteri secara invasif sehingga jauh lebih agresif dibanding CTA maupun MRA. CTA menawarkan alternatif yang lebih tidak invasif dibanding DSA, tetapi tetap memaparkan radiasi *X-ray* kepada pasien [2]. Karakteristik tanpa radiasi ini menjadikan TOF-MRA pilihan yang lebih aman untuk pemeriksaan berulang atau skrining rutin dibandingkan CT maupun CTA [2].

Dalam praktik klinis saat ini, TOF-MRA sudah digunakan secara luas sebagai metode standar untuk mendeteksi kelainan pembuluh darah seperti aneurisma intrakranial [13]. Metode ini dipilih karena mampu menghasilkan citra tiga dimensi dengan resolusi spasial yang memadai sekaligus terjamin aman, karena tidak memaparkan radiasi kepada pasien [14]. Resolusi yang cukup tinggi ini penting bagi dokter untuk mengamati pembuluh darah berukuran kecil dan memantau bentuk fisik aneurisma secara akurat [14]. Dengan pencitraan yang detail tersebut, dokter dapat melakukan diagnosis sedini mungkin untuk mencegah pecahnya aneurisma yang berujung pada pendarahan fatal di otak [13].

Meski demikian, evaluasi TOF-MRA secara manual tetap menghadirkan hambatan tersendiri bagi tenaga medis. Dokter radiologi harus menelusuri irisan-irisan gambar tiga dimensi dari berbagai sudut pandang untuk menemukan lokasi aneurisma [13]. Proses pencarian visual yang mendetail ini memakan waktu cukup lama, dan risiko terlewatnya aneurisma tetap ada, terutama untuk aneurisma berukuran sangat kecil [13]. Keterbatasan inilah yang kemudian mendorong para peneliti untuk mengembangkan sistem deteksi otomatis berbasis kecerdasan buatan sebagai pendamping dokter dalam menganalisis citra TOF-MRA [13].

2.4 Preprocessing Data

Citra medis seperti TOF-MRA pada dasarnya merupakan data tiga dimensi yang berukuran besar. Memproses data sebesar itu secara langsung akan memberatkan komputer, sehingga data 3D tersebut biasanya diubah terlebih dahulu menjadi kumpulan gambar dua dimensi sebelum dianalisis oleh model kecerdasan buatan [15]. Cara paling umum adalah memotong volume otak secara horizontal dari atas ke bawah, menghasilkan tumpukan irisan gambar mendatar yang disebut irisan aksial [1]. Proses konversi ini membantu mempercepat proses komputasi,

namun memiliki kelemahan berupa setiap irisan dianalisis secara terpisah sehingga informasi tentang posisi dan hubungan antar-irisan tidak digunakan [15].

Untuk mengatasi kelemahan tersebut, para peneliti mengembangkan pendekatan 2.5D. Alih-alih menganalisis satu irisan gambar saja, pendekatan ini bekerja dengan cara menerima beberapa irisan yang berurutan sekaligus, misalnya irisan di atas, irisan target, dan irisan di bawahnya [15]. Dengan cara ini, model dapat menangkap gambaran dari area sekitar struktur pembuluh darah tanpa harus memproses seluruh volume 3D [16], sehingga kebutuhan daya komputasi untuk pendekatan ini tetap jauh lebih ringan dibanding pendekatan 3D penuh [16].

Di luar cara penyusunan gambar, komposisi data yang digunakan dalam pelatihan juga berpengaruh pada kemampuan model. Salah satu strategi yang umum diterapkan adalah menyertakan gambar dari pasien sehat yang tidak memiliki aneurisma, yang disebut *negative sampling* [1]. Langkah ini penting dilakukan untuk mengajarkan model agar mampu mengenali seperti apa pembuluh darah yang normal. Dengan memberikan data gambar pasien sehat, model tidak akan mudah tertipu oleh anatomi pembuluh darah biasa [1].

2.5 Deteksi Objek

Computer vision adalah bidang yang menerapkan ilmu komputer dan matematika untuk menganalisis dan menginterpretasi citra secara otomatis. Dalam konteks pencitraan medis, bidang ini memadukan keahlian tersebut dengan keahlian medis untuk mengembangkan solusi analisis citra, termasuk untuk mendeteksi kelainan seperti aneurisma dan tumor otak [17]. Di antara berbagai tugas *computer vision*, deteksi objek merupakan salah satu tugas paling mendasar dan menantang [18]. Deteksi objek bertujuan mengidentifikasi kategori sekaligus lokasi setiap objek dalam gambar melalui kotak pembatas (*bounding box*) [4]. Pada konteks pengenalan aneurisma intrakranial, pendekatan ini memungkinkan dokter melokalisasi area kelainan secara cepat melalui kotak pembatas yang dihasilkan model [4].

Deteksi objek berbeda dari klasifikasi gambar dan segmentasi dalam cara menyampaikan informasi lokasi. Klasifikasi gambar hanya menentukan status suatu gambar secara keseluruhan, misalnya ada atau tidaknya aneurisma, tanpa menunjukkan lokasinya [4]. Segmentasi menandai lokasi objek hingga tingkat piksel sehingga batas tepi lesi dapat digambarkan secara presisi, namun ketika beberapa aneurisma saling berdekatan atau menempel, segmentasi mengenali

seluruh area tersebut sebagai satu kesatuan sehingga tidak dapat membedakan lesi secara individual [4]. Deteksi objek menandai lokasi setiap objek dalam kotak pembatas tersendiri, sehingga lesi yang berdekatan tetap dapat dibedakan sebagai kandidat yang terpisah [4].

2.6 YOLO26

You Only Look Once (YOLO) adalah algoritma deteksi objek yang memproses seluruh gambar sekaligus dalam satu langkah. Cara kerja satu langkah ini membuat YOLO jauh lebih cepat dan cocok untuk analisis citra medis secara *real-time*. Sejak diperkenalkan pertama kali pada 2015, setiap generasi baru YOLO membawa perubahan arsitektur yang memperbaiki keseimbangan antara kecepatan dan akurasi deteksi [18]. YOLO26 adalah generasi terbaru dari algoritma ini yang dirancang agar dapat berjalan efisien pada perangkat dengan daya komputasi terbatas [19].

Perubahan paling signifikan pada YOLO26 adalah penghilangan *Non-Maximum Suppression* (NMS). Pada generasi sebelumnya, NMS bekerja sebagai tahap *post-processing* yang memilih satu kotak prediksi terbaik dari banyak kandidat yang tumpang tindih, dan beban komputasinya naik seiring bertambahnya kepadatan objek dalam gambar [20]. YOLO26 menghapus tahap tersebut dengan mengadopsi arsitektur *end-to-end* yang langsung menghasilkan satu prediksi per objek tanpa penyaringan tambahan [8].

YOLO26 juga menghadirkan dua metode pelatihan baru untuk meningkatkan kemampuan mendeteksi objek berukuran kecil, yaitu *Progressive Loss Balancing* (ProgLoss) dan *Small-Target-Aware Label Assignment* (STAL) [8]. ProgLoss mengatur prioritas pembelajaran model secara bertahap selama proses pelatihan. Pada tahap awal pelatihan, model difokuskan untuk mengenali jenis objek secara umum. Menjelang akhir pelatihan, fokus beralih ke ketepatan posisi kotak pembatas [20]. Sedangkan STAL bekerja dengan cara memberi perhatian lebih pada objek-objek kecil selama pelatihan, karena model cenderung lebih mudah mendeteksi objek besar dan mengabaikan yang kecil [20].

Dengan tersedianya ProgLoss dan STAL yang sudah terintegrasi secara arsitektural pada YOLO26, penelitian ini menggunakan model standar tanpa modifikasi struktur jaringan secara manual. Modifikasi arsitektur secara manual, seperti penyisipan modul atensi *custom*, menuntut penyesuaian kode program yang rumit dan berpotensi menimbulkan ketidakstabilan selama pelatihan, terutama

pada *dataset* medis yang volumenya terbatas [6]. Penambahan modul komputasi tambahan di luar jaringan utama juga meningkatkan kompleksitas alur kerja sistem dan memperpanjang waktu pemrosesan yang dibutuhkan [5].

Portase dan Ardelean [8] mengevaluasi arsitektur dasar YOLO dalam bentuk standarnya tanpa modifikasi khusus domain, dengan alasan pendekatan ini menghasilkan metrik dasar yang dapat direplikasi peneliti lain dan membantu mengidentifikasi arsitektur dasar yang paling menjanjikan sebelum dilakukan adaptasi *domain-specific* lebih lanjut. Kedua alasan tersebut turut mendasari penelitian ini. Dengan menggunakan model YOLO26 standar tanpa modifikasi, hasil evaluasi yang diperoleh mencerminkan kemampuan fitur bawaan arsitektur tersebut dalam menangani objek berukuran kecil dan data tidak seimbang, sebagai dasar sebelum adaptasi *domain-specific* dipertimbangkan pada penelitian selanjutnya.

2.7 Metrik Evaluasi

Dalam mengukur seberapa baik kinerja sebuah sistem kecerdasan buatan, diperlukan berbagai tolak ukur khusus yang disebut metrik evaluasi. *Precision* mengukur seberapa sering tebakan sistem terbukti benar, sehingga mencerminkan seberapa jarang sistem membunyikan alarm palsu [8]. *Precision* diformulasikan pada Rumus 2.1 sebagai berikut:

$$Precision = \frac{TP}{TP + FP} \quad (2.1)$$

di mana TP (*True Positive*) adalah jumlah deteksi yang benar, dan FP (*False Positive*) adalah jumlah deteksi yang salah. *Recall* atau sensitivitas mengukur seberapa banyak aneurisma yang berhasil ditemukan dari seluruh aneurisma yang ada, sehingga mencerminkan seberapa jarang sistem melewati penyakit [8]. Formulasi *Recall* ditunjukkan pada Rumus 2.2:

$$Recall = \frac{TP}{TP + FN} \quad (2.2)$$

di mana FN (*False Negative*) adalah jumlah aneurisma yang gagal terdeteksi. *F1-Score* menggabungkan *Precision* dan *Recall* menjadi satu angka agar kedua sisi tersebut dapat dinilai sekaligus [8]. Rumus 2.3 menunjukkan cara penghitungan

F1-Score:

$$F1-Score = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (2.3)$$

Sebelum mAP dapat dihitung, diperlukan pemahaman mengenai *Intersection over Union* (IoU), yaitu metrik yang mengevaluasi akurasi lokalisasi dengan menghitung rasio antara luas irisan (*intersection area*) dan luas gabungan (*union area*) dari kotak prediksi dan kotak anotasi *ground truth* [8]. Semakin besar nilai IoU, semakin tepat posisi kotak prediksi menempel pada lokasi aneurisma yang sesungguhnya. Nilai IoU ini kemudian dijadikan ambang batas untuk menentukan benar atau salahnya suatu deteksi dalam perhitungan *Average Precision* (AP), yaitu luas area di bawah kurva *precision-recall* [8]. *Mean Average Precision* (mAP) merupakan rata-rata nilai AP. Metrik mAP@50 dihitung dari deteksi yang memiliki IoU minimal 0,5 terhadap anotasi *ground truth*, sedangkan mAP@50-95 dihitung dengan merata-ratakan AP pada sepuluh ambang IoU berbeda, mulai dari 0,5 hingga 0,95 dengan interval 0,05. Pada ambang IoU yang tinggi, deteksi yang kotaknya hanya sedikit tumpang tindih dengan lokasi aneurisma yang sesungguhnya tidak lagi terhitung sebagai deteksi benar, sehingga nilai mAP@50-95 menjadi lebih rendah dibanding mAP@50 [8].

Meskipun metrik tersebut umum digunakan dalam deteksi objek, penggunaannya saja belum cukup untuk bidang medis. Dokter lebih memperhatikan berapa banyak tebakan salah yang muncul per pasien, bukan sekadar angka akurasi global. Oleh karena itu, penelitian deteksi aneurisma umumnya menggunakan kurva *Free-Response Receiver Operating Characteristic* (FROC), yang secara bersamaan mengamati sensitivitas deteksi dan jumlah tebakan salah per kasus pasien [13]. Sistem yang mampu menemukan banyak aneurisma tetapi juga sering salah tebak justru akan menyulitkan dokter, karena setiap tanda yang muncul tetap harus diperiksa ulang secara manual.

Evaluasi juga perlu dilakukan pada tingkat lesi, yaitu menilai apakah sistem berhasil menemukan dan menandai lokasi setiap aneurisma satu per satu, bukan sekadar menyimpulkan bahwa seorang pasien memiliki aneurisma atau tidak [10]. Karena gambar MRA diperiksa lapis demi lapis dalam bentuk dua dimensi, kotak deteksi dari berbagai irisan gambar yang berdekatan perlu digabungkan menjadi satu temuan utuh dalam ruang tiga dimensi [1]. Tanpa penggabungan ini, satu aneurisma yang muncul di beberapa irisan berturut-turut bisa dihitung sebagai banyak aneurisma berbeda, padahal semuanya menunjukkan objek yang sama.

Evaluasi juga perlu dikelompokkan berdasarkan ukuran fisik aneurisma. Dalam pedoman klinis, ukuran aneurisma intrakranial dibagi ke dalam empat kategori, yaitu sangat kecil (di bawah 3 mm), kecil (3 hingga kurang dari 5 mm), sedang (5 hingga kurang dari 7 mm), dan besar (7 mm atau lebih) [7]. Pemisahan ini penting karena sensitivitas sistem biasanya turun tajam pada aneurisma yang sangat kecil, sementara aneurisma sekecil itu tetap berpotensi pecah dan membahayakan pasien [7]. Dengan melihat hasil metrik pada tiap kategori ukuran secara terpisah, pengembang dan tenaga medis dapat mengetahui batas kemampuan sistem sebelum diterapkan dalam penanganan pasien [7]. Pada akhirnya, nilai evaluasi yang tinggi di atas kertas belum cukup untuk memastikan sistem tersebut siap digunakan oleh dokter. Karena masih perlu diuji coba secara langsung pada banyak pasien dari berbagai rumah sakit untuk membuktikan bahwa kemampuannya benar-benar dapat diandalkan di dunia nyata [21].

