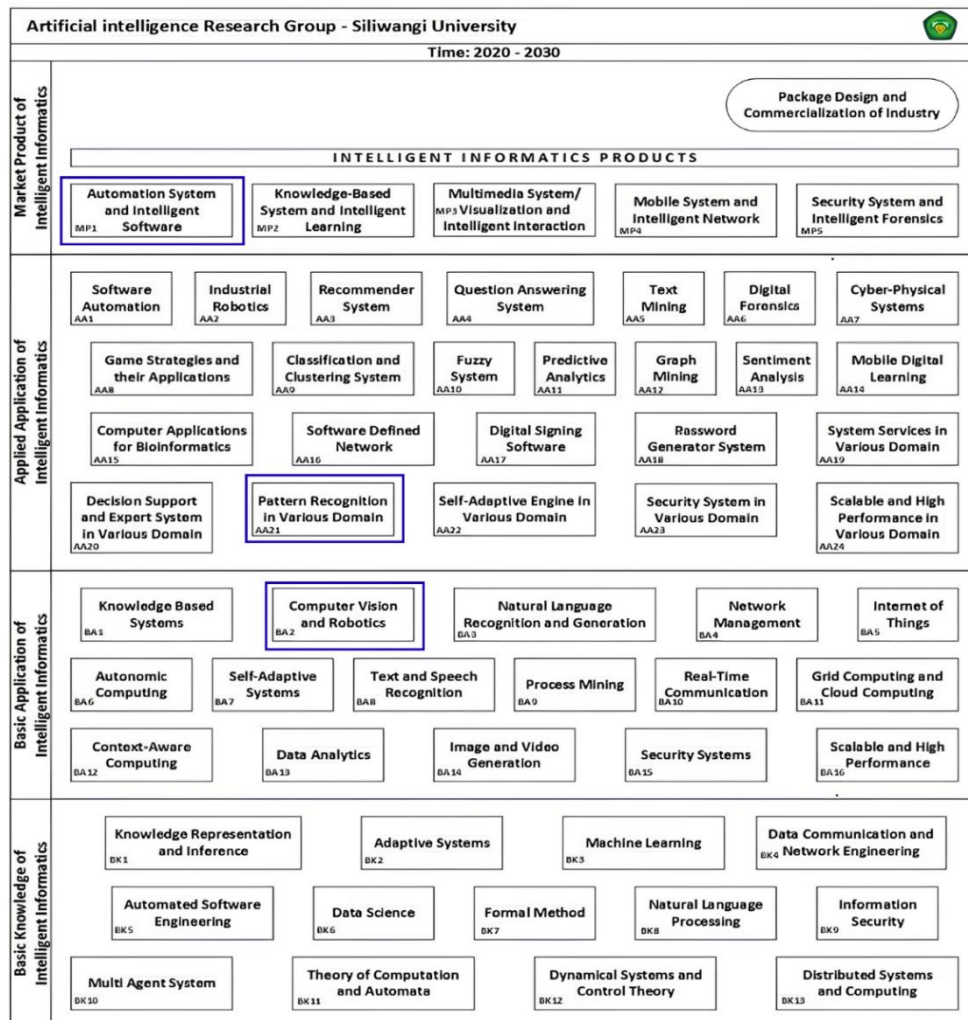


BAB III

METODOLOGI PENELITIAN

3.1 Roadmap Penelitian

Roadmap atau peta jalan penelitian yang digunakan pada penelitian ini menggunakan *roadmap* penelitian kelompok *Artificial Intelligence Research Group* Universitas Siliwangi (AIS) 2020-2030 yang terdapat pada Gambar 3.1.

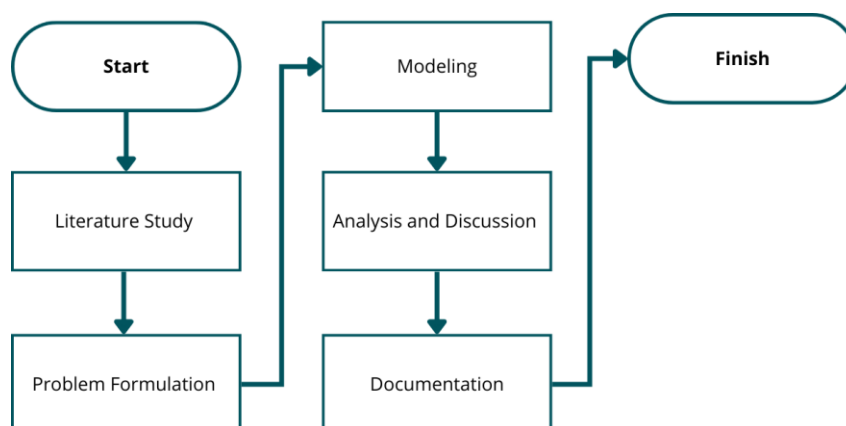


Gambar 3.1. Roadmap Artificial Intelligence Research Group Universitas Siliwangi 2020-2030 (AIS, 2020)

Berdasarkan *Roadmap AIS* yang terdapat pada Gambar 3.1, penelitian ini berfokus pada implementasi *Segment Anything Model* (SAM) untuk mempercepat pembuatan anotasi segmentasi pada gambar yang bertujuan untuk memudahkan model YOLOv11 dalam mengenali pola segmentasi citra untuk deteksi kecacatan pada biji kopi. Dan juga implementasi *quantization* untuk mengurangi beban kerja model yang dapat diterapkan pada perangkat *low-end*. Maka dari itu, *roadmap* penelitian berkaitan dengan *Computer Vision* (BA2) dan juga *Pattern Recognition in Various Domains* (AA21). Penelitian ini juga memiliki potensi untuk menghasilkan sebuah model yang dapat dikembangkan lebih lanjut ke dalam ranah *Intelligent Software Systems* (MP1) dan *Internet of Things* (BA5), sebagai implementasi model ke dalam *software* otomatis deteksi cacat biji kopi, baik *software* dengan perangkat terbatas maupun dengan memanfaatkan IoT.

3.2 Tahapan Penelitian

Tahapan penelitian yang digunakan mengadopsi prosedur dari penelitian (Masyora dkk., 2024; Utami dkk., 2024) yang telah dimodifikasi sesuai dengan kebutuhan penelitian. Rincian tahapan penelitian terdapat pada Gambar 3.2.



Gambar 3.2 Tahapan Penelitian

3.2.2 Studi Literatur (*Literature Study*)

Tahapan pertama yang dilakukan adalah studi literatur, tahapan ini bertujuan untuk mencari permasalahan dan referensi yang berkaitan pada deteksi obeejek kecacatan biji kopi (Masyora dkk., 2024). Indikator pada tahapan ini adalah memahami informasi penelitian terkait deteksi kecacatan biji kopi, dengan luaran yang dapat disampaikan pada Bab 1 Latar Belakang Penelitian dan Bab 2 Tinjauan Pustaka.

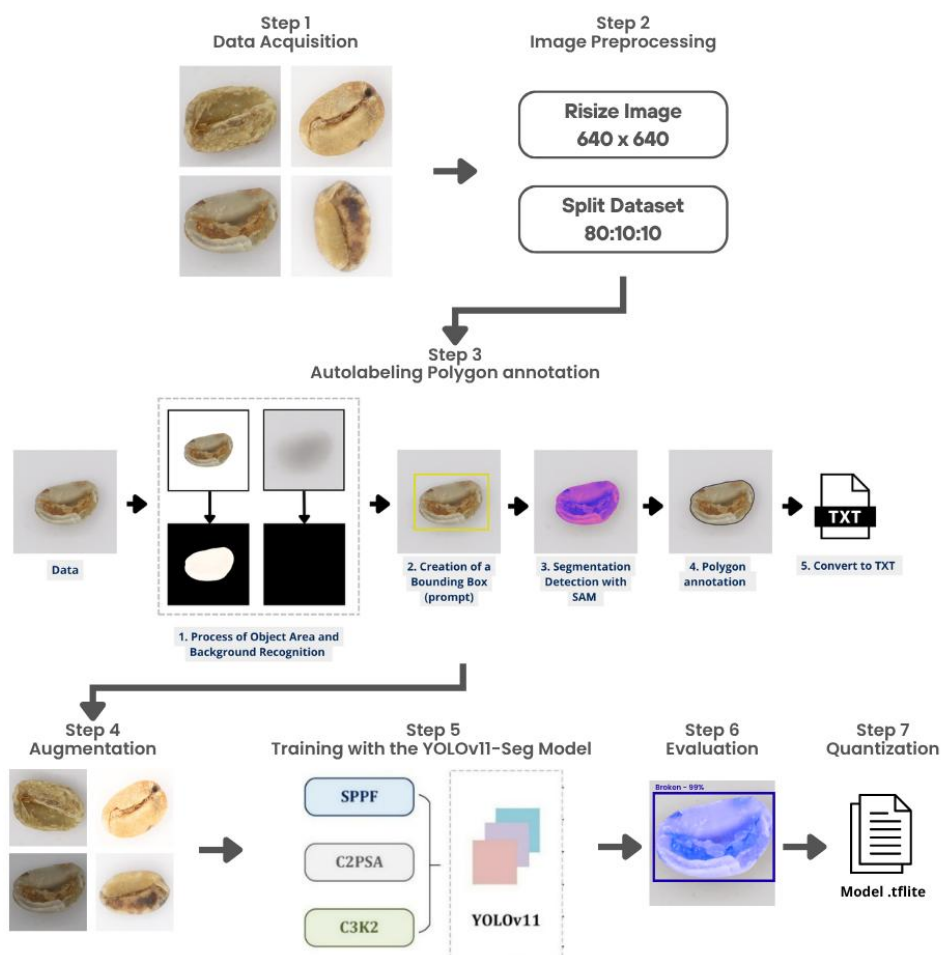
3.2.3 Perumusan Masalah (*Problem Formulation*)

Tahapan kedua yang dilakukan adalah perumusan masalah yang bertujuan untuk mengidentifikasi kendala yang dihadapi (Masyora dkk., 2024; Utami dkk., 2024). Setelah melakukan studi literatur, terdapat beberapa tantangan utama, yaitu sulitnya melakukan anotasi data segmentasi secara manual yang presisi serta bagaimana menjalankan model pada perangkat *low-end*. Pada tahap ini, penulis merumuskan masalah mengenai cara mengoptimalkan deteksi kecacatan biji kopi menggunakan pendekatan segmentasi yang akurat namun efisien.

Cara mengatasi kompleksitas pelabelan, penulis merencanakan penggunaan *Segment Anything Model* (SAM) sebagai instrumen utama dalam tahapan pembuatan anotasi. Hal ini dilakukan untuk memastikan bahwa *ground truth* yang dihasilkan memiliki tingkat presisi tinggi guna mendukung keberhasilan proses pelatihan model. Fokus perumusan juga diarahkan pada teknik kuantisasi untuk menekan beban komputasi model. Indikator pada tahapan ini adalah terpetakannya masalah teknis dalam anotasi dan optimasi model, dengan luaran berupa daftar pertanyaan penelitian dan batasan masalah yang dicantumkan pada Bab 1.

3.2.4 Pemodelan (*Modeling*)

Tahapan ketiga adalah pemodelan atau rancangan model, tahapan ini bertujuan untuk membuat, melatih, dan menguji model *Segment Anything Model* (SAM), YOLOv11-Seg, dan *Quantization* untuk deteksi kecacatan biji kopi yang dapat diterapkan pada perangkat *low-end* (Afinda, 2024). Gambar 3.3 menunjukkan alur pemodelan yang telah diadaptasi dan dimodifikasi dari penelitian (Badgujar dkk., 2024). Modifikasi ini dilakukan untuk memastikan setiap tahapan sesuai dengan ruang lingkup dan kebutuhan.



Gambar 3.3. Diagram Alur Penelitian Pembuatan Model Deteksi Kopi Cacat.

3.2.4.1. Akuisisi Data (*Data Acquisition*)

Tahapan pertama sebelum pembuatan model adalah akuisisi data. Akuisisi data merupakan teknik mengumpulkan dan menyimpan data untuk dikelola agar menghasilkan data yang siap digunakan pada tahap selanjutnya (Dirhamsyah dkk., 2023; Wahyudi dkk., 2022). Data yang digunakan dalam penelitian ini merupakan data sekunder yang berasal dari penelitian terdahulu (Arwatchananukul dkk., 2024), yaitu "[Coffee Green Bean with 17 Defects \(original\)](#)". Data ini mencakup citra biji kopi hijau dengan 17 kategori cacat dan total citra 979, tanpa label *Bounding Box* dan label anotasi poligon.

Selain data kecacatan, penelitian ini juga menambahkan satu kelas data tambahan berupa citra biji kopi tanpa cacat (*non-defective coffee beans*) yang diambil dari data publik yang memiliki karakteristik yang sama dengan data utama dan dapat diakses pada tautan berikut: [coffee beans dataset - v1 rotated Dataset](#). Indikator pada tahapan ini adalah terkumpulnya citra kecacatan biji kopi yang representatif terhadap 17 jenis cacat dan 1 kelas tidak cacat (*Good*).

Penggunaan kelas *good* untuk model deteksi dan klasifikasi biji kopi baik dan tidak baik (*defect*) sudah digunakan pada penelitian sebelumnya, yang menggunakan dua kelas data, yakni kelas baik (*good*) dan tidak baik (*defect*) (Jameel John M. Reales dkk., 2024; Juarto & Yulianto, 2026). Serta terdapat penelitian yang menggunakan kelas *good* dengan istilah *premium* yang disandingkan dengan kelas *defect* lainnya (Febriana dkk., 2022). Tujuan dari adanya kelas *good* dan juga *defect* ini adalah agar dalam penerapan model, model

tidak hanya mengenal kopi cacat saja, melainkan dapat mengenal kopi yang baik juga (Jameel John M. Reales dkk., 2024).

3.2.4.2. Praproses

Pada tahap praproses terdapat 2 tahapan yang dilakukan, yaitu mengubah ukuran gambar (*resize image*) dan pembagian data (*splitting data*). Tahapan pertama adalah *resize* citra atau pengubah ukuran gambar. Tahapan ini bertujuan untuk menstandarkan ukuran gambar dengan menyesuaikan dengan model yang akan digunakan, yakni berukuran 640x640 (Gope dkk., 2024; Ultralytics, 2025). Tahap kedua adalah pembagian dataset (*data splitting*), yaitu tahap pembagian data menjadi subset *training*, *validation*, dan *testing*. Pembagian dilakukan dengan 2 skema uji 80:10:10 (Gope dkk., 2024; Zhong dkk., 2025). Dikarenakan pada tahap augmentasi dan anotasi citra yang dihasilkan masih disimpan dalam folder per kelas. Maka pembagian data dilakukan dengan penerapan *Grouped Splitting* (pembagian berbasis grup), yaitu pembagian data bersifat per kelas atau per folder. Tujuannya agar tidak terjadi kesalahan pembagian data.

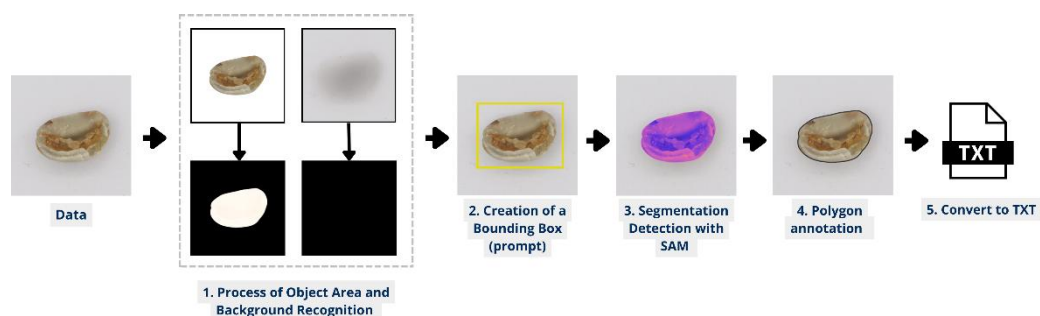
3.2.4.3. Otomatisasi Label Anotasi Poligon (*Autolabeling Anotation Polygon*)

Hasil data yang telah dibagi ke dalam beberapa subset pelatihan akan dianalisis dengan pembuatan anotasi poligon secara otomatis dengan bantuan model Segment Anything (SAM). Pembuatan otomatisasi anotasi poligon diperlukan karena model YOLO varian segmentasi memerlukan anotasi poligon untuk menunjukkan area segmentasi pada objek yang mengikuti konturnya dan akan digunakan dalam pelatihan model (Hurtik dkk., 2022; Jocher dkk., 2023). Pembuatan anotasi poligon dilakukan pada tahap ini, karena data utama yang

bersumber dari penelitian (Arwatchananukul dkk., 2024) tidak memiliki label anotasi baik poligon maupun *Bounding Box* dan hanya terdapat label kelas. Langkah-langkah dalam pembuatan *autolabeling* anotasi poligon untuk area segmentasi biji kopi terdapat pada Algoritma 3.1.

Algoritma 3.1. Anotasi Otomatis Segmentasi Menggunakan SAM	
Input	: Dataset gambar D, Model SAM <i>Pre-trained</i> , Parameter pra-pemrosesan dan padding
Proses	: Untuk setiap citra I dalam D lakukan langkah-langkah berikut:
	1. Pra-pemrosesan gambar : Memisahkan antara latar belakang dengan objek menggunakan informasi warna, agar tepi objek dapat terlihat dengan jelas.
	2. <i>Bounding Box</i> : Menentukan lokasi objek utama untuk membentuk <i>Bounding Box</i> sebagai prompt segmentasi.
	3. Segmentasi Objek (SAM) : Melakukan segmentasi semantik objek menggunakan model SAM berdasarkan <i>Bounding Box</i> yang dihasilkan.
	4. Konversi ke format YOLO : Mengubah area segmentasi anotasi poligon dan menyimpan dalam format YOLO.
Output	: Dataset anotasi notasi dalam Format YOLO

Ide perancangan Algoritma 3.1 mengacu pada konsep dasar anotasi, yaitu proses pemisahan objek dari latar belakang agar model dapat memahami informasi visual secara terstruktur (Minaee dkk., 2021; Yu dkk., 2023). Untuk visualisasi alur proses pembuatan anotasi otomatis terdapat pada Gambar 3.4.



Gambar 3.4. Alur Anotasi Otomatis Anotasi menggunakan *Hole Filling* dan SAM.

Tahap awal pembuatan anotasi poligon yang terdapat pada Algoritma 3.1 dan Gambar 3.4 dimulai dengan pra-pemrosesan citra dan masking untuk mempertegas kontur serta perbedaan antara objek dan latar belakang. Pada penelitian ini, proses preprocessing memanfaatkan pemisahan berbasis warna dengan menggunakan ruang warna HSV, khususnya kanal saturasi. Penggunaan kanal saturasi bertujuan untuk menonjolkan area objek yang memiliki tingkat kejenuhan warna lebih tinggi dibandingkan dengan latar belakang, serta mengurangi pengaruh variasi pencahayaan yang umum terjadi pada citra alami (Jintasuttisak dkk., 2025).

Selain penerapan HSV, diterapkan juga metode *hole filling* pada tahap preprocessing untuk menutup celah atau lubang di dalam area objek yang dapat muncul akibat cacat objek, bayangan, atau variasi tekstur. Metode *hole filling* umum digunakan dalam pemrosesan citra untuk memperbaiki bentuk objek hasil segmentasi awal dengan cara mengisi bagian internal yang seharusnya termasuk ke dalam objek, sehingga menghasilkan representasi objek yang lebih utuh dan solid (Sarsembayeva dkk., 2025; Zheng dkk., 2023). Tahapan ini bertujuan untuk memastikan *Bounding Box* yang dihasilkan mencakup seluruh area objek secara konsisten.

Selanjutnya, pembuatan *Bounding Box* dilakukan dengan skema penambahan padding adaptif. Adaptif pada penelitian ini merujuk pada definisi menurut KBBI, yaitu dapat menyesuaikan diri sesuai dengan keadaan (KBBI, 2025a). Tujuan pembuatan padding adaptif ini adalah agar kotak pembatas atau *Bounding Box* dapat memberikan jarak terhadap objek dan dapat menyesuaikan ukuran berdasarkan dimensi objek yang telah terdeteksi. Maka dari itu, langkah pertama

untuk membuat *Bounding Box* yang berjarak adalah pembuatan padding adaptif. Persamaan pencarian nilai padding terdapat pada persamaan 3.1.

$$p = \max(\max(\min(\lfloor (\max(w, h) \times p_{ratio}) \rfloor, p_{max}), g_{min})) \quad (3.1)$$

Setelah nilai padding diperoleh, langkah berikutnya adalah mentransformasikan koordinat objek awal menjadi koordinat *Bounding Box* akhir. Proses ini menggunakan pendekatan *Region of Interest* (ROI) dengan memastikan bahwa koordinat baru tidak melampaui batas dimensi citra asli (W dan H). Perhitungan koordinat akhir (x_1, y_1) dan (x_2, y_2) didefinisikan dalam Persamaan 3.2 hingga 3.5 (Nihranz, 2026).

$$x_1 = \max(0, x - p) \quad (3.2)$$

$$y_1 = \max(0, y - p) \quad (3.3)$$

$$x_2 = \min(W, x + w + p) \quad (3.4)$$

$$y_2 = \min(H, y + h + p) \quad (3.5)$$

Penggunaan fungsi $\max(0, \dots)$ dan $\min(W/H, \dots)$ berfungsi sebagai mekanisme *clipping* yang bertujuan untuk memastikan setiap *Bounding Box* yang dihasilkan tetap berada di dalam ruang koordinat citra (valid), sehingga mencegah terjadinya kesalahan.

Bounding Box yang diperoleh selanjutnya digunakan sebagai *prompt* awal bagi *Segment Anything Model* (SAM) untuk menentukan lokasi objek. Pendekatan ini didasarkan pada konsep *prompt-based segmentation*, di mana informasi awal seperti kotak pembatas digunakan untuk mengarahkan model agar fokus pada area objek tertentu dan mengurangi ambiguitas latar belakang (Kirillov dkk., 2023).

Tahapan berikutnya adalah segmentasi semantik menggunakan SAM, yang memanfaatkan *prompt* berupa *Bounding Box* untuk menghasilkan mask objek secara otomatis. SAM dirancang untuk melakukan segmentasi yang bersifat *general-purpose* dengan kemampuan memisahkan objek dan latar belakang secara presisi tanpa pelatihan ulang pada domain spesifik dan menghasilkan area segmentasi atau garis tepi objek (Kirillov dkk., 2023; Moghimi dkk., 2024).

Setelah mendapatkan area segmentasi pada objek yang dideteksi, area tersebut kemudian dikonversi ke dalam format poligon menggunakan algoritma *Douglas-peucker* melalui fungsi `approxPolyDP`. Pembuatan format poligon ini dikendalikan oleh parameter toleransi yang dihitung berdasarkan Persamaan 3.6 (Sapkota dkk., 2025).

$$\varepsilon = \alpha P_k \quad (3.6)$$

Dalam persamaan ini, P_k mewakili perimeter (area) dalam satuan piksel dari kontur asli instance I_k , sementara α (alpha) adalah parameter toleransi tak berdimensi yang menentukan jumlah simpul poligon sesuai dengan kompleksitas bentuk objek. Merujuk pada penelitian sebelumnya yang menggunakan nilai alpha 0,001 untuk objek apel (Sapkota dkk., 2025), penelitian ini akan melakukan eksperimen dengan variasi nilai alpha yang berbeda. Ini didasarkan pada perbedaan karakteristik morfologis dari setiap kelas objek yang digunakan. Nilai alpha yang diuji dalam penelitian ini mencakup 0.001, 0.0008, 0.0006, 0.0004, dan 0.0002 untuk mengevaluasi korelasi antara jumlah titik vertex poligon dan kinerja akurasi deteksi model.

Setelah proses penyederhanaan simpul selesai, semua koordinat poligon yang dihasilkan masih dalam satuan piksel asli. Koordinat awal ini akan dilakukan normalisasi untuk mematuhi standar input model YOLOv11. Normalisasi koordinat dilakukan dengan membagi nilai x dan y dengan dimensi lebar dan tinggi gambar sesuai dengan Persamaan 3.7 (Sapkota et al., 2025).

$$\tilde{x}_i = \frac{x_i}{W}, \tilde{y}_i = \frac{y_i}{H} \quad (3.7)$$

Hasil normalisasi menggunakan persamaan 3.2 menghasilkan nilai *floating-point* dalam rentang 0 hingga 1. Tahap akhir dari proses ini adalah konversi data ke dalam file teks (.txt), di mana setiap instance objek direpresentasikan dalam satu baris yang dimulai dengan indeks kelas (*class_id*) diikuti oleh serangkaian koordinat poligon yang dinormalisasi $(\tilde{x}_1, \tilde{y}_1, \dots, \tilde{x}_n, \tilde{y}_n)$.

3.2.4.4. Augmentasi (*Augmentation*)

Tahap augmentasi data bertujuan untuk meningkatkan variasi data latih serta mengatasi permasalahan ketidakseimbangan jumlah data antarkelas (*class imbalance*) yang mengakibatkan model bias terhadap kelas mayoritas (Erlin dkk., 2022; Iksatya dkk., 2026). Proses augmentasi dilakukan pada dua tahap: pertama, setelah tahap anotasi yang bertujuan untuk menyamakan jumlah data setiap kelasnya, dan kedua, pada tahap pelatihan. Parameter yang digunakan pada teknik augmentasi terdapat pada tabel 3.1.

Tabel 3.1. Parameter Augmentasi

No	Parameter	Nilai/Teknik	Metode	Tahap
<i>Gemoetric</i>				
1	Rotasi	45°, 135°, 270°	Fixed	1

No	Parameter	Nilai/Teknik	Metode	Tahap
2	<i>Degrees</i>	20	Random	2
3	<i>Shear</i>	5.0	Random	2
4	<i>Perspective</i>	0.0002	Random	2
5	<i>Flipud</i>	0.5	Random	2
6	<i>Fliplr</i>	0.5	Random	2
<i>Color Space</i>				
7	<i>Hue</i>	Rentang 0.015	Random	2
8	<i>Saturation</i>	Rentang 0.5	Random	2
9	<i>Brightness</i>	Rentang 0.5	Random	2
<i>Image Combination</i>				
10	<i>Mosaic</i>	1.0	Random	2
11	<i>Mixup</i>	0.1	Random	2
12	<i>Copy-Paste</i>	0.5	Random	2

Berdasarkan parameter yang ditunjukkan pada Tabel 3.1, teknik augmentasi yang diimplementasikan mencakup transformasi geometris, penyesuaian ruang warna (color space), serta teknik image combination. Teknik augmentasi geometris meliputi rotasi, *flipping*, *degrees*, shear, dan perspektif yang diterapkan untuk meningkatkan invariansi model terhadap perubahan orientasi, skala, dan posisi objek (Arwatchananukul et al., 2024; Zhong et al., 2025).

Sementara itu, augmentasi color space melalui penyesuaian *hue*, *saturation* dan *value* dilakukan untuk mensimulasikan variasi kondisi pencahayaan. Hal ini bertujuan agar model lebih robust terhadap fluktuasi intensitas cahaya di lingkungan nyata (Chazar dkk., 2025). Pada aspek segmentasi, dilakukan beberapa teknik augmentasi tambahan, yakni kombinasi gambar seperti *mosaic*, *mixup*, dan

copy-paste untuk memperkaya konteks spasial dan membantu model dalam mengenali objek pada data baru (Puspita, Naufal, dan Al Zami, 2025).

3.2.4.5. Pelatihan, Tes, dan Evaluasi Model

Pelatihan model YOLOv11-seg dilakukan dengan optimasi parameter untuk mendapatkan keseimbangan antara akurasi dan kecepatan inferensi. Setiap parameter yang akan digunakan untuk pelatihan model, terdapat pada Tabel 3.2.

Tabel 3.2. Parameter *Training Model*

Parameter	Nilai/Teknik
<i>Epochs</i>	300
<i>Batch Size</i>	16
<i>Optimizer</i>	AdamW
<i>Initial Learning Rate Awal (Lr0)</i>	0.0005
<i>Initial Learning Rate (Lrf)</i>	0.001
<i>Momentum</i>	0.937
<i>Weight Decay</i>	0.0005
<i>Patience</i>	50
<i>Warmup Epochs</i>	10
<i>Automatic Mixed Precision (AMP)</i>	<i>True</i>

Pemilihan jumlah epoch sebanyak 300 bertujuan untuk dapat lebih mempelajari fitur data di saat pelatihan terus menunjukkan performa yang baik (Yin dkk., 2026). Batch size sebesar 16 digunakan karena merupakan konfigurasi standar pada YOLO dan mampu menjaga keseimbangan antara stabilitas pelatihan dan keterbatasan memori GPU (Cao dkk., 2025; Gope dkk., 2024; Thai dkk., 2024; Ultralytics, 2025). Optimizer AdamW digunakan karena terbukti efektif dalam menangani parameter dengan fitur *weight decay* = 0.0005 yang mampu mengurangi risiko overfitting (He dkk., 2024; Sun dkk., 2023).

Selanjutnya digunakan juga parameter tambahan seperti *initial learning rate awal*, *initial learning rate*, *momentum*, dan *weight decay*. *Initial learning rate awal* (Lr0) sebesar 0.0005 dan Lrf = 0.001 dipilih karena terbukti dapat meningkatkan kinerja pada kelas dengan frekuensi menengah dan dapat meningkatkan akurasi mAP50 (Zhao dkk., 2025). Selain itu, diterapkan fase warmup untuk menstabilkan gradien di awal pelatihan guna mencegah terjadinya lonjakan bobot yang drastis (*divergence*) (Sun dkk., 2023).

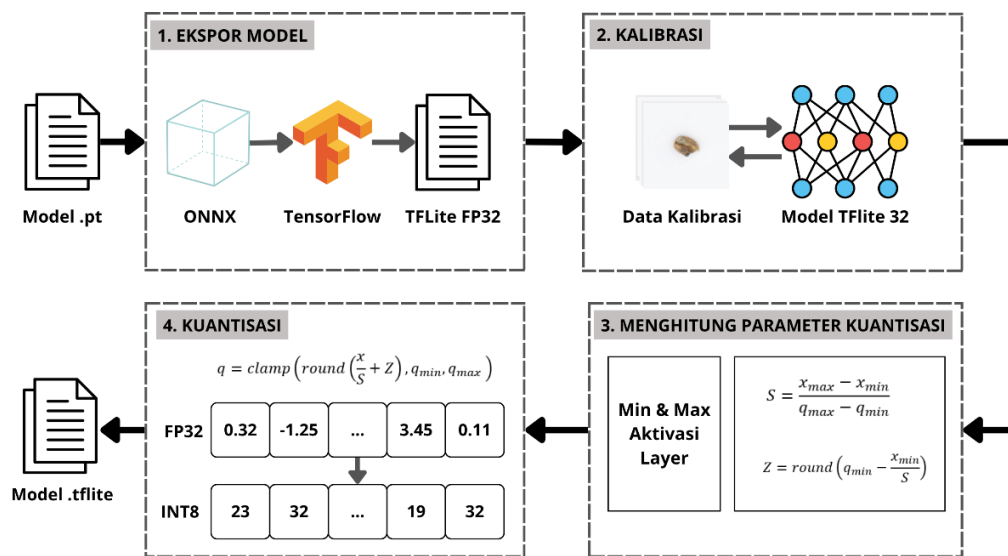
Fitur *Automatic Mixed Precision* (AMP) diaktifkan untuk mendukung efisiensi pada perangkat dengan cara mengurangi beban memori (Alahdal et al., 2024). Serta, mekanisme *early stopping* dengan parameter *patience* sebesar 50 epoch diterapkan untuk menghentikan proses pelatihan secara otomatis jika tidak terjadi peningkatan performa pada dataset validasi, sehingga mencegah terjadinya overfitting serta menghemat waktu komputasi (Sun dkk., 2023).

Saat proses pelatihan dan juga tes dilakukan secara bersamaan dengan evaluasi model. Metrik evaluasi model mencakup *Confusion Matrix*, *Intersection over Union* (IoU), *Precision*, *Recall*, *F1-score*, serta mAP (*mean Average Precision*) untuk mengukur seberapa akurat model dalam melakukan segmentasi pada objek biji kopi yang cacat.

3.2.4.6. Optimasi Model

Tahap optimasi model dilakukan untuk meningkatkan efisiensi komputasi agar model dapat dijalankan pada perangkat dengan keterbatasan sumber daya. Pada penelitian ini digunakan metode *Post-Training Quantization* (PTQ) dengan pendekatan INT8 dan FP16. Gambar 3.5, merupakan alur lengkap proses kuantisasi

yang terdapat pada Gambar 3.5, khususnya posisi tahap konversi FP 32 ke INT 8, merujuk pada penelitian (Rami dkk., 2025) yang telah disesuaikan dengan alur proses penelitian yang telah dilakukan.



Gambar 3.5. Tahapan Teknik Kuantisasi Model

Pada Gambar 3.5, proses diawali dengan mengonversi model dari PyTorch ke dalam model yang akan digunakan, contohnya TensorFlow Lite dengan FP32. Model FP32 hasil konversi menjadi model yang akan dikuantisasi ke dalam INT8. Proses kuantisasi melibatkan tahap kalibrasi yang memanfaatkan representative dataset untuk menjalankan model dalam format FP32 guna memperoleh distribusi nilai aktivasi pada setiap layer. Hasil dari proses ini berupa nilai minimum dan maksimum yang kemudian digunakan untuk menghitung parameter kuantisasi, yaitu *scale* dan *zero-point*. Parameter tersebut selanjutnya digunakan dalam proses konversi bobot dan aktivasi dari representasi *floating-point* 32-bit menjadi bilangan integer 8-bit.

3.2.5 Analisis dan Pembahasan (*Analysis and Discussion*)

Setelah seluruh rangkaian pengujian selesai, tahap selanjutnya adalah melakukan analisis serta pembahasan terhadap data yang telah diperoleh. Tujuan dari tahapan ini adalah untuk memastikan apakah model yang dibuat sesuai dengan standar kualitas yang diharapkan atau jika terdapat aspek lain yang memerlukan perbaikan lebih lanjut (Masyora dkk., 2024). Pada tahap ini, hasil dari setiap skenario pengujian, khususnya pembahasan mengenai perbandingan antara model dasar YOLOv11-seg (FP32) dengan model yang telah dioptimasi melalui kuantisasi FP16 dan INT8.

3.2.6 Dokumentasi (*Documentation*)

Tahap terakhir adalah dokumentasi atau penyusunan laporan penelitian mulai dari tujuan penelitian hingga hasil dan kesimpulan yang dilengkapi dengan saran yang bermanfaat untuk penelitian ke depannya (Suganda dkk., 2025). Luaran dari penyusunan laporan penelitian ini adalah skripsi dan publikasi ilmiah dengan target jurnal internasional atau lokal minimal Sinta 3.