

BAB III

METODE PENELITIAN

3.1 Objek Penelitian

Objek penelitian ini adalah siswa sekolah menggunakan APD yang tersedia disekolah. *Dataset* tersebut mencakup siswa yang menggunakan APD lengkap yaitu kacamata, dan sarung tangan serta siswa yang tidak menggunakan APD lengkap atau menggunakan salah satunya.

Penelitian ini bertujuan untuk mendeteksi siswa yang menggunakan APD lengkap dan tidak lengkap untuk mencegah serta mengurangi kecelakaan kerja saat praktikum, Berikut ini adalah contoh objek penelitiannya :

a. Kacamata *Welding*



Gambar 3. 1 Kacamata *Welding*

b. Sarung Tangan *Welding*

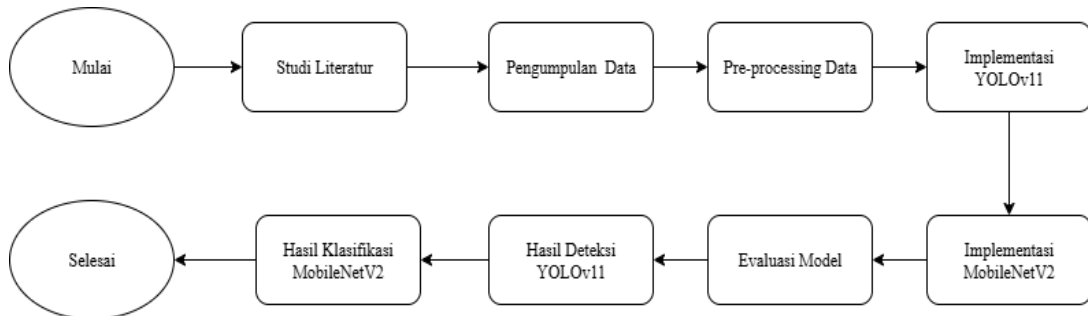


Gambar 3. 2 Sarung Tangan *Welding*

3.2 Prosedur Penelitian

Prosedur penelitian ditunjukkan pada Gambar 3.3 Penelitian ini disusun secara terstruktur dan sistematis untuk memastikan setiap langkah berjalan dengan optimal. Setiap tahapan dirancang secara berurutan agar saling melengkapi dalam

mencapai tujuan yang telah ditetapkan. Untuk memberikan gambaran yang lebih jelas dan rinci, alur penelitian disajikan dalam bentuk *flowchart* yang mengilustrasikan keseluruhan proses penelitian



Gambar 3. 3 *FlowChart* Prosedur Penelitian

3.3 Studi Literatur

Studi literatur dilakukan untuk meninjau implementasi algoritma *YOLO* dalam deteksi objek, khususnya dalam konteks keselamatan kerja dan APD. Penelitian terdahulu menunjukkan bahwa *YOLO* memiliki keunggulan dalam kecepatan deteksi dan akurasi tinggi. Selain itu, dikaji pula arsitektur *CNN* seperti *MobileNetV2* sebagai solusi ringan untuk klasifikasi citra, serta pengaplikasian teknik pre-processing dan augmentasi data dalam pelatihan model deep learning.

3.4 Pengumpulan Data

Pengumpulan data dilakukan di SMK Plus Laboratorium Indonesia dengan total 434 sampel citra siswa yang menggunakan dan tidak menggunakan APD saat praktik pengelasan. Pengambilan gambar dilakukan secara langsung di lapangan menggunakan dua perangkat smartphone, yaitu *Redmi Note 10* dan *iPhone 14*, untuk memperoleh variasi kualitas dan sudut pandang gambar. Data yang terkumpul kemudian dikelompokkan ke dalam empat kategori, yaitu 'Kacamata', 'Sarung_Tangan', 'Tanpa_Kacamata', dan 'Tanpa_SarungTangan'. Kategori objek tersebut ditunjukkan pada Gambar 3.4.



(A) APD Lengkap



(B) APD Tidak Lengkap

Gambar 3. 4 APD lengkap dan Tidak Lengkap

3.5 Pre-Processing Data

Data yang dikumpulkan akan melewati tahap *pre-processing*, yang mencakup beberapa langkah penting untuk menjamin kualitas *data* yang digunakan dalam pelatihan *model*. Proses ini bertujuan untuk menyiapkan *data* agar dapat diproses secara optimal oleh *model*. Berikut adalah langkah-langkah yang akan dilakukan.

3.5.1 Labeling

Labeling merupakan tahap penting dalam proses *pre-processing data*, di mana setiap gambar diberikan anotasi yang sesuai untuk membantu *model* dalam mengenali objek dengan akurat. Pada penelitian ini, proses *labeling* dilakukan menggunakan Roboflow, sebuah platform yang mempermudah pembuatan anotasi *bounding box* secara efisien. Dengan menggunakan Roboflow, setiap gambar dikategorikan ke dalam kelas tertentu, seperti Kacamata, SarungTangan, Tanpa_Kacamata, dan Tanpa_SarungTangan, sehingga *data* siap digunakan untuk pelatihan model. Contoh *data* yang sudah di *labeling* pada Gambar 3.5



Gambar 3. 5 Data Labeling

3.5.2 Augmentasi

Teknik augmentasi *data* seperti *Flipping*, *Mosaic*, *Rotation*, dan *Motion Blur* digunakan untuk meningkatkan variasi gambar dalam *dataset*. Proses ini membantu *model* dalam mengenali berbagai karakteristik Kelengkapan Penggunaan APD dengan lebih baik, sehingga meningkatkan akurasi dan kemampuan generalisasi *model* selama pelatihan :

a. Mosaic

Teknik augmentasi *data* yang menggabungkan empat gambar berbeda ke dalam satu gambar baru. Keempat gambar tersebut dipotong dan disusun menjadi satu bingkai (*frame*) besar. Teknik ini memungkinkan *model* melihat objek dari berbagai konteks, latar belakang, dan ukuran dalam satu kali pelatihan, sehingga dapat meningkatkan kemampuan generalisasi *model* deteksi objek seperti *YOLO*.

b. Motion Blur

Motion Blur adalah efek buram yang terjadi akibat pergerakan kamera atau objek saat pengambilan gambar. Teknik ini digunakan dalam augmentasi untuk membuat *model* lebih mampu mengenali objek meskipun dalam kondisi gambar yang sedikit kabur, seperti saat mendeteksi objek dalam video atau gambar yang diambil dengan kecepatan tinggi.

c. *Rotation*

Teknik augmentasi yang memutar gambar input sejauh sudut tertentu, misalnya -15° hingga $+15^\circ$. Tujuannya adalah agar *model* dapat belajar mengenali objek dalam posisi yang berbeda-beda. Dengan melakukan rotasi, *model* tidak hanya belajar dari objek yang selalu berada pada orientasi *horizontal* atau tegak lurus, tetapi juga bisa mengenali jika objek sedikit miring atau terbalik.

d. *Flip*

adalah teknik augmentasi yang membalik gambar secara *horizontal* atau *vertikal*. Teknik ini meningkatkan variasi *dataset*, sehingga *model* dapat mengenali objek dalam berbagai orientasi dan sudut pandang.

e. *Resize*

adalah proses mengubah ukuran gambar agar sesuai dengan dimensi yang dibutuhkan oleh model. Teknik ini memastikan bahwa semua gambar dalam *dataset* memiliki ukuran yang seragam, misalnya 640x640 piksel pada *model YOLO*, sehingga *model* dapat bekerja dengan lebih optimal.

3.6 Implementasi YOLOv11

Pada tahap ini, *YOLOv11* digunakan untuk membangun *model* yang mampu mendeteksi kelengkapan penggunaan Alat Pelindung Diri (APD) pada siswa saat praktikum pengelasan. *YOLOv11* dipilih karena terkenal cepat dan akurat dalam mendeteksi objek secara langsung dari gambar. Sebelum pelatihan dimulai, semua gambar telah melalui proses pelabelan menggunakan *Roboflow*. Dalam proses ini, setiap objek pada gambar diberi kotak dan *label* sesuai dengan kelasnya, yaitu: Kacamata, SarungTangan, Tanpa_Kacamata, dan Tanpa_SarungTangan. Untuk menambah variasi *data*, *Roboflow* juga secara otomatis melakukan augmentasi seperti *flip*, *rotation*, *blur*, dan *mosaic*. Semua gambar kemudian diubah ukurannya menjadi 640x640 piksel agar sesuai dengan kebutuhan input *model YOLO*.

Pelatihan *model* dilakukan menggunakan *YOLOv11* dari awal tanpa menggunakan bobot dari *dataset* umum seperti *COCO*. *Model* dilatih menggunakan GPU untuk mempercepat proses dan menggunakan *optimizer AdamW* agar proses belajar *model* lebih stabil. Tujuan akhirnya adalah agar *model* ini mampu

mendeteksi apakah siswa menggunakan APD secara lengkap, sebagian, atau tidak lengkap dengan akurasi yang baik.

3.7 Implementasi *MobileNetV2*

Setelah *model YOLOv11* berhasil mendeteksi keberadaan Alat Pelindung Diri (APD) dalam gambar, tahap selanjutnya adalah melakukan klasifikasi hasil deteksi tersebut menggunakan *model MobileNetV2*. Tujuan dari tahap ini adalah untuk mengelompokkan kondisi penggunaan APD oleh siswa ke dalam tiga kategori, yaitu APD Lengkap, Sebagian, dan Tidak Lengkap. *Model* klasifikasi ini tidak melakukan klasifikasi ulang terhadap setiap objek terdeteksi, tetapi bekerja berdasarkan hasil akhir deteksi *YOLO*, dengan mempertimbangkan kombinasi objek yang terdeteksi. Sebagai contoh, jika *YOLO* mendeteksi kacamata dan sarung tangan dalam satu gambar, maka gambar tersebut diklasifikasikan sebagai Lengkap. Jika hanya salah satu APD yang terdeteksi, maka masuk kategori Sebagian. Sedangkan jika tidak ada APD terdeteksi, maka dikategorikan sebagai Tidak Lengkap. *Model MobileNetV2* dilatih menggunakan *dataset* hasil pengolahan citra dari deteksi *YOLOv11* yang telah diberi *label* sesuai tiga kategori tersebut. *MobileNetV2* dipilih karena merupakan arsitektur *CNN* yang ringan dan efisien, sehingga mampu memberikan hasil klasifikasi yang cepat dan akurat. Hasil pengujian menunjukkan bahwa *model* ini mampu mengenali kondisi penggunaan APD dengan akurasi yang tinggi dan performa yang stabil.

3.8 Evaluasi Model

Setelah proses pelatihan *model* selesai, langkah berikutnya adalah melakukan evaluasi untuk mengetahui seberapa baik *model* bekerja dalam mendeteksi dan mengklasifikasikan penggunaan Alat Pelindung Diri (APD) pada siswa. Evaluasi dilakukan dalam dua tahap, yaitu untuk *model* deteksi objek dengan *YOLOv11* dan untuk *model* klasifikasi *MobileNetV2*. Beberapa ukuran yang digunakan dalam evaluasi ini antara lain adalah *precision*, *recall*, *F1-score*, dan *mean Average Precision (mAP)*:

1. *Mean Average Precision (mAP)* adalah metrik utama dalam evaluasi *model* deteksi objek seperti *YOLOv11*. Metrik ini menghitung rata-rata *precision* pada berbagai tingkat *recall*. Semakin tinggi nilai *mAP*, semakin baik *model* dalam mendeteksi objek secara konsisten dan akurat.

2. *Recall* merupakan metode yang mengukur seberapa banyak objek yang benar terdeteksi dari total objek yang seharusnya terdeteksi. Berikut adalah rumusnya :

$$Recall = \frac{True\ Positive\ (TP)}{True\ Positive(TP) + False\ Positive(FP)} \quad (3.1)$$

- a. *Recall* adalah metrik evaluasi dalam *machine learning*, khususnya pada tugas klasifikasi dan deteksi objek, yang mengukur kemampuan *model* dalam menemukan semua objek yang seharusnya terdeteksi.
 - b. *True Positive (TP)*: Jumlah objek yang benar-benar ada dan berhasil dideteksi oleh *model*.
 - c. *False Negative (FN)*: Jumlah objek yang sebenarnya ada tetapi tidak berhasil dideteksi oleh *model*.
3. *Precision* menunjukkan proporsi dari semua deteksi yang benar terhadap total deteksi yang dilakukan oleh *model*. Semakin tinggi *precision*, semakin sedikit kesalahan *model* dalam mengidentifikasi objek yang tidak relevan. Berikut adalah rumusnya :

$$Precision = \frac{True\ Positive\ (TP)}{True\ Positive(TP) + False\ Positive(FP)} \quad (3.2)$$

- a. *Precision* adalah metrik evaluasi dalam *machine learning* yang mengukur seberapa akurat *model* dalam mendeteksi objek yang benar dibandingkan dengan semua deteksi yang dibuat.
4. *F1 Score* adalah metrik evaluasi dalam *machine learning* yang merupakan rata-rata harmonik dari *Precision* dan *Recall*. Metrik ini digunakan untuk mencari keseimbangan antara *precision* dan *recall*, terutama saat kita ingin mempertimbangkan baik ketepatan (*precision*) maupun kelengkapan (*recall*) dari prediksi *model*.

$$F1\ Score = 2x \frac{Precision \times Recall}{Precision + Recall} \quad (3.3)$$