

EVALUASI PERFORMA ALGORITMA YOLOV8 DALAM DETEKSI DAN KLASIFIKASI POSISI TIDUR

Abstrak

Deteksi objek menggunakan teknologi computer vision memungkinkan komputer mengenali dan menentukan lokasi serta jenis objek dalam gambar atau video. Metode tradisional seperti HOG (*Histogram of Oriented Gradien*) dan SIFT (*Scale-Invariant Feature Transform*) memiliki keterbatasan dalam menangani variasi objek, sementara pendekatan deep learning seperti CNN dan algoritma YOLO (*You Only Look Once*) menunjukkan kinerja yang lebih baik. YOLO memproses gambar secara real-time dan memprediksi kotak pembatas serta kelas objek dalam satu jaringan saraf. Penelitian ini bertujuan untuk menguji efektivitas algoritma YOLOv8 dalam mendeteksi orang dan mengklasifikasikan posisi tidur, seperti telentang, miring kiri, miring kanan, dan posisi lainnya. Hasil pengujian menunjukkan bahwa YOLO berhasil mendeteksi manusia di tempat tidur dengan akurasi di atas 90% pada dataset yang mirip dengan data pelatihan, namun performa menurun signifikan pada dataset yang berbeda dengan akurasi 25%-40%. Pengujian pada dataset campuran tetap menunjukkan akurasi di atas 90%. Dalam klasifikasi gambar, YOLO juga menunjukkan hasil baik dengan akurasi di atas 85% pada dataset serupa, tetapi menurun menjadi 20% pada dataset yang berbeda. Penggunaan dataset pelatihan yang seimbang dan proporsional sangat penting untuk performa optimal. Perbedaan antara YOLOv8n dan YOLOv8s tidak terlalu signifikan, namun YOLOv8n lebih cocok untuk perangkat dengan komputasi rendah dan dataset kecil, sedangkan YOLOv8s memiliki performa keseluruhan lebih baik tetapi membutuhkan komputasi lebih tinggi dan dataset lebih besar.

Kata Kunci: deteksi, klasifikasi, posisi tidur, yolov8.

Abstract

Object detection using computer vision technology enables computers to recognize and determine the location and type of objects in images or videos. Traditional methods such as HOG and SIFT have limitations in handling object variations, while deep learning approaches such as CNN and YOLO (You Only Look Once) algorithms demonstrate better performance. YOLO processes images in real-time and predicts bounding boxes and object classes in a single neural network. This study aims to evaluate the effectiveness of the YOLOv8 algorithm in detecting humans and classifying sleeping positions, such as supine, left lateral, right lateral, and others. The test results show that YOLO successfully detects humans in bed with over 90% accuracy on a similar dataset to the training data, but performance drops significantly to 25%-40% accuracy on different datasets. Testing on mixed datasets still shows accuracy above 90%. In image classification, YOLO also performs well with over 85% accuracy on similar datasets, but drops to 20% on different datasets. Balanced and proportional training datasets are crucial for optimal performance. The difference between YOLOv8n and YOLOv8s is not significant, but YOLOv8n is more suitable for devices with low computation and small datasets, whereas YOLOv8s has better overall performance but requires higher computation and larger datasets.

Keywords: *classification, detection, sleep pose, yolov8.*

1. PENDAHULUAN

Deteksi objek adalah teknologi computer vision yang bisa mengenali objek tertentu dalam gambar atau video, lalu memberikan informasi mengenai lokasi dan jenis objek yang terdeteksi. Teknologi ini memungkinkan komputer untuk mengidentifikasi objek dalam sebuah gambar (Lee & Hwang, 2021). Metode machine learning tradisional seperti *Histogram of Oriented Gradients (HOG)* dan *Scale-Invariant Feature Transform (SIFT)* telah lama digunakan untuk deteksi objek. Meskipun metode-metode ini efektif, metode tersebut sangat bergantung pada fitur tertentu dan kurang fleksibel dalam menangani variasi dalam penampilan objek (Bahri et al., 2019). Pendekatan berbasis deep learning seperti *Convolutional Neural Networks (CNN)* telah menunjukkan kinerja yang lebih baik dalam deteksi objek dan telah terbukti sangat efektif dalam pengenalan gambar. Hal ini dikarenakan CNN berusaha meniru cara sistem visual pada korteks manusia dalam mengenali gambar, sehingga memberikannya kemampuan dalam menganalisis gambar (Eka Putra, 2016). Selain itu, YOLO (*You Only Look Once*) juga merupakan algoritma dalam *deep learning* yang digunakan untuk deteksi objek dan klasifikasi gambar yang berbeda dari algoritma lain karena menggunakan jaringan saraf tunggal untuk memeriksa keseluruhan gambar sekaligus (Redmon et al., 2016).

Dalam pemantauan pose tidur secara tradisional, seringkali metode ini dapat mengganggu kenyamanan pasien dan membutuhkan upaya manual yang intensif. Salah satu solusi alternatif yang telah diajukan sebelumnya di antaranya adalah deteksi dan klasifikasi posisi pasien yang terbaring di tempat tidur tanpa mengganggu kenyamanan pasien (Ostadabbas et al., 2012). Karena itu, tujuan dari penelitian ini berfokus pada pelatihan dan pengujian model serta evaluasi performa algoritma *deep learning* untuk deteksi dan klasifikasi posisi tidur.

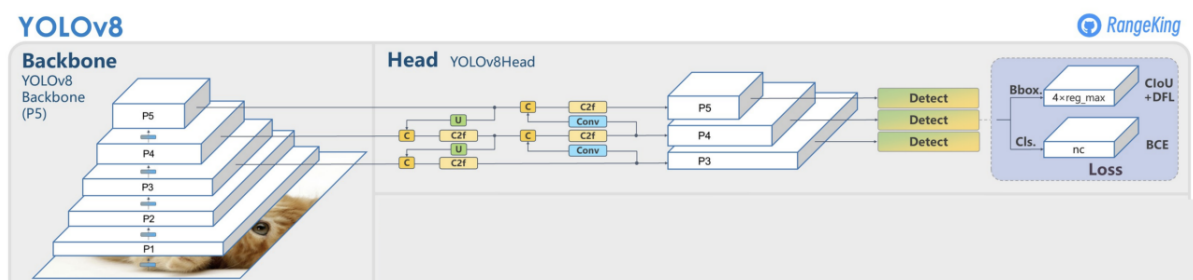
YOLO diperkenalkan pada tahun 2015 oleh Joseph Redmon dan Ali Farhadi. Algoritma ini mengoperasikan gambar secara real-time melalui jaringan saraf konvolusi (CNN), langsung memprediksi kotak pembatas objek serta probabilitas kelasnya. Seiring waktu, YOLO telah mengalami serangkaian pengembangan. YOLOv2 (2016) memperkenalkan kotak jangkar untuk meningkatkan akurasi deteksi dengan mengatasi masalah *bounding box* yang kurang akurat di versi sebelumnya (Gao et al., 2019). YOLOv3 (2018) memperkenalkan pemrosesan multi-skala dan meningkatkan performa secara keseluruhan (Hong et al., 2021). Kemudian, YOLOv4 (2020) yang diperbaharui lagi di tahun 2021 dan membawa berbagai pengoptimalan dan peningkatan akurasi (Du et al., 2021). Sementara YOLOv5 (2021) memperkenalkan arsitektur model baru dan metode pelatihan yang lebih baik, dengan akurasi yang lebih tinggi terutama untuk objek kecil, dan tersedia dalam bahasa *Python* untuk kemudahan implementasi dalam sistem (Ji et al., 2022). YOLOv6

(2022) memperkenalkan model yang lebih ringan untuk perangkat *edge computing* (Yanto et al., 2023), sementara YOLOv7 (2022) membawa pengoptimalan lebih lanjut dan peningkatan akurasi (Zhou et al., 2022) dan yang terbaru, YOLOv8 (2023) dengan pengoptimalan tambahan dan modul baru untuk meningkatkan akurasi dan efisiensi deteksi objek (Wang et al., 2023).

Banyak penelitian telah menggunakan metode *You Only Look Once* (YOLO) dalam deteksi objek. Sebuah penelitian sebelumnya menggunakan arsitektur YOLOv5 untuk mengklasifikasikan simbol tangan secara real-time, dan hasil kinerja model untuk mengklasifikasikan simbol tangan mencapai akurasi 80% (Wibowo & Sugiarto, 2023). Penelitian lebih lanjut tentang klasifikasi penyakit mata berdasarkan citra fundus menggunakan algoritma YOLOv8 memperoleh hasil akurasi terbaik sebesar 92% (Muhammad Nur Ihsan Muhlashin & Stefanie, 2023). Dengan adanya penelitian yang membuktikan performa yang baik pada algoritma YOLO dalam kasus tertentu, maka dalam penelitian ini penulis akan menggunakan algoritma YOLOv8 yang lebih baru untuk mengevaluasi kinerjanya dalam mendeteksi dan mengklasifikasi posisi tidur. Oleh karena itu, penelitian ini berjudul **“Evaluasi Performa Algoritma YOLOv8 Dalam Deteksi dan Klasifikasi Posisi Tidur”** dengan studi kasus klasifikasi posisi tidur telentang (*supine*), miring kiri (*left*), miring kanan (*right*), dan negatif (pose selain 3 kelas sebelumnya atau yang tidak dikenali).

2. METODE

YOLOv8 adalah versi baru dari algoritma YOLO yang telah mengalami pengembangan dalam deteksi objek secara real-time dengan memperhatikan arsitektur lapisan-lapisan yang dikembangkan sebelumnya. Untuk arsitektur YOLOv8 dapat dilihat pada Gambar 1.

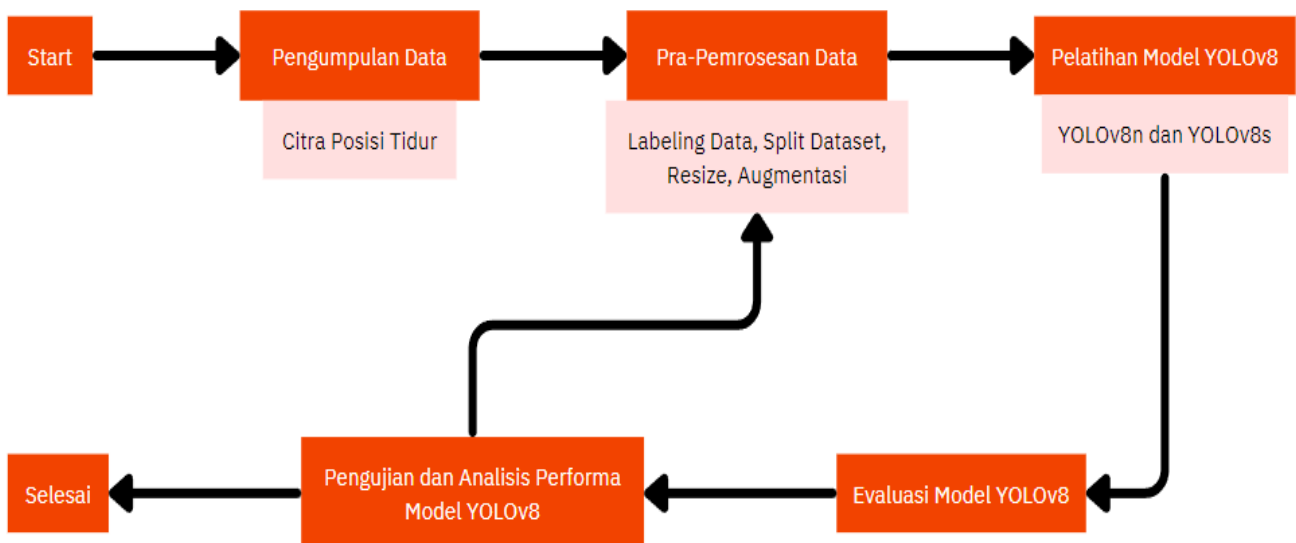


Gambar 1. Arsitektur YOLOv8

1. *Input Layer*: Arsitektur menerima gambar sebagai input, yang kemudian melalui pra-pemrosesan untuk dinormalisasi dan diubah ukurannya jika diperlukan.
2. *Backbone Network* (Jaringan Punggung) : Bagian dari jaringan ini mengekstrak fitur dari gambar input. Biasanya terdiri dari beberapa lapisan konvolusi dengan berbagai ukuran kernel. Jaringan punggung YOLOv8 menggunakan arsitektur yang dimodifikasi yang kemungkinan mencakup aspek dari versi sebelumnya dengan peningkatan untuk ekstraksi fitur yang lebih baik.

3. *Neck Network* (Jaringan Leher) : Setelah fitur diekstraksi oleh jaringan punggung, mereka melewati jaringan leher. YOLOv8 menggunakan kombinasi *Feature Pyramid Network* (FPN) dan *Path Aggregation Network* (PAN). FPN membantu dalam mendeteksi objek pada skala yang berbeda dengan menghasilkan peta fitur pada beberapa skala dan resolusi. PAN menggabungkan fitur dari berbagai tingkat jaringan untuk mempertahankan fitur bergranulasi halus pada semua skala.
4. *Head Network* (Jaringan Kepala) : Bagian dari arsitektur di mana deteksi berlangsung. Terdiri dari beberapa kepala deteksi, masing-masing bertanggung jawab untuk menghasilkan kotak pembatas, probabilitas kelas, dan skor kepercayaan. Kepala YOLOv8 akan menghasilkan prediksi ini untuk peta fitur mereka masing-masing.
5. *Output Layer*: YOLOv8 akan menghasilkan tiga jenis output untuk setiap objek yang terdeteksi: koordinat kotak pembatas, probabilitas kelas untuk setiap kotak yang menunjukkan kemungkinan objek tersebut termasuk dalam kelas tertentu, dan skor kepercayaan yang menandakan seberapa yakin model bahwa suatu objek ada dalam kotak pembatas yang diprediksi.
6. *Post-Processing*: Setelah prediksi mentah dihasilkan, dilakukan penekanan non-maksimum. Proses ini menyaring kotak pembatas yang tumpang tindih, memastikan bahwa setiap objek yang terdeteksi dihitung hanya sekali. (Reis et al., 2023)

Penelitian ini bertujuan untuk menguji seberapa efektif algoritma YOLOv8 dalam mendeteksi orang dan mengklasifikasikan posisi tidur. Dalam penelitian ini, metode yang digunakan adalah *machine learning life cycle* yang terdiri dari beberapa langkah. Langkah-langkah tersebut dapat dilihat pada Gambar 2.



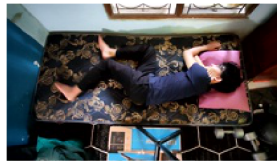
Gambar 2. Machine Learning Life Cycle

A. Pengumpulan Data

Jenis data yang akan digunakan untuk melatih dan mengevaluasi model adalah kumpulan data citra posisi tidur yang diperoleh dari perekaman video yang melibatkan 3 orang dan terdiri dari empat kelas, yaitu telentang (supine), miring kiri (left), miring kanan (right), dan negatif (pose selain 3 kelas sebelumnya atau yang tidak dikenali). Setiap orang melakukan perekaman sesuai kelas yang ada kecuali kelas negatif, sehingga setiap orang memiliki 3 rekaman video tetapi, untuk kelas negatif hanya ada 1 rekaman video. Setiap video memiliki perekaman video dengan durasi 90 detik dan beresolusi 1080p 30fps yang direkam menggunakan kamera handphone. Kemudian video tersebut diambil 1 gambar setiap 3 detik sehingga menghasilkan 90 gambar untuk masing-masing kelas telentang, miring kiri dan miring kanan. Sedangkan, kelas negatif memiliki 30 gambar. Jadi total data jika digabungkan yaitu $90 \times 3 \text{ kelas} + 30 \text{ gambar negatif} = 300 \text{ gambar}$ untuk setiap dataset A (dataset A1, dataset A2 dan dataset A3). Dataset A memiliki jumlah total gambar sebanyak 900 dengan pembagian untuk pelatihan model sebanyak 600 gambar, 180 gambar untuk setiap kelas *left*, *right* dan *supine*. Sedangkan, 60 gambar untuk kelas negatif. Dataset A yang digunakan untuk validasi dan pengujian model masing-masing sebanyak 150 gambar, 45 gambar untuk setiap kelas *left*, *right* dan *supine*. Sedangkan, 15 gambar untuk kelas negatif. Pengumpulan data dilakukan sebanyak 3 kali dengan variasi seperti yang tertera pada Tabel 1 dan sampel dari dataset A dapat dilihat pada Gambar 3, Gambar 4 dan gambar 5.

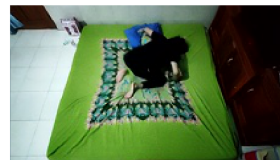
Tabel 1. Variasi Dataset

Dataset	Tempat	Tipe Kasur	Warna Sprei	Warna Bantal	Pakaian Orang 1	Pakaian Orang 2	Pakaian Orang 3	Posisi kamera
A1	Kamar 1	Single Bed	Pink	Hitam	Kaos kuning dan celana hitam	Kaos putih dan celana hitam	Kaos ungu dan celana hitam	Landscape
A2	Kamar 1	Single Bed	Biru	Pink	Kaos biru dan celana hitam	Kaos hijau dan celana putih	Kaos merah dan celana hitam	Landscape
A3	Kamar 2	Twin Bed	Hijau	Hijau	Kaos hitam dan celana hitam	Kaos biru tua dan celana hitam	Kaos hijau dan kain sarung coklat	Potrait



Gambar 3. Sampel Dataset A1

Gambar 4. Sampel Dataset A2

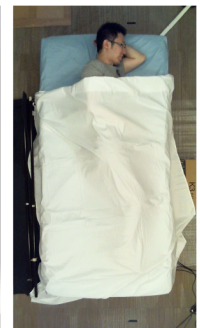


Gambar 5. Sampel Dataset A3

Selain dataset A, akan digunakan dataset lain yang tersedia dan dapat diunduh di internet yaitu dari IEEE VIP CUP 2021 Dataset “Privacy-Preserving In-Bed Human Pose Estimation”. Dataset ini dibagi menjadi 2 yaitu B1 dan B2 dimana B1 merupakan dataset tanpa selimut (uncover) dan B2 merupakan dataset dengan selimut (cover). Total dataset B yang digunakan sebanyak 450 gambar, untuk pelatihan model sebanyak 360 gambar, 120 gambar untuk setiap kelasnya (*left*, *right*, dan *supine*). Sedangkan untuk validasi dan pengujian model masing-masing sebanyak 90 gambar, 30 gambar untuk setiap kelasnya (*left*, *right*, dan *supine*). Untuk sampel dari dataset B1 dan B2 dapat dilihat pada Gambar 6 dan gambar 7.



Gambar 6. Sampel Dataset B1



Gambar 7. Sampel Dataset B2

B. Pra-pemrosesan Data

Ini merupakan tahap pemrosesan data sebelum melalui proses pelatihan dan evaluasi.

Langkah awal dalam tahap pra-pemrosesan data adalah memberi label data untuk klasifikasi yang terbagi menjadi 4 kelas, yaitu telentang (supine), miring kiri (left), miring kanan (right), dan negatif (pose selain 3 kelas sebelumnya atau yang tidak dikenali). Sedangkan, untuk deteksi gambar hanya memiliki 1 label saja yaitu orang (person). Kemudian, membagi data menjadi dua bagian, yaitu data latih yang digunakan untuk melatih model sebanyak 80% (data train), dan data validasi yang digunakan untuk mengevaluasi model sebanyak 20% (data validation) dari total data. Setelah itu, gambar-gambar diubah ukurannya agar memiliki dimensi yang seragam untuk mempercepat proses pelatihan model.

C. Pelatihan Model YOLOv8

Model yang akan digunakan adalah YOLOv8n dan YOLOv8s. Pada tahap ini dilakukan pelatihan model menggunakan data pelatihan yang sudah melalui pra-pemrosesan data yaitu sebanyak 80% dari total data. Pertama, model akan dilatih untuk mendeteksi orang (person detection) kemudian akan dilakukan pemotongan (cropping) gambar pada dataset sesuai dengan bounding box yang dihasilkan selama labeling. Setelah pemotongan gambar dataset selesai, dilakukan pelatihan untuk klasifikasi pose tidur dengan dataset tersebut. Proses pemotongan gambar ini juga berlaku pada saat melakukan evaluasi dan pengujian model. Selama pelatihan model deteksi orang, akan dihitung nilai box(precision), recall dan mAP50 (mean average precision). Kemudian dalam klasifikasi akan dihitung nilai precision, recall, F1-score dan akurasi. Hasil dari pelatihan deteksi dan klasifikasi akan dijadikan sebagai indikator kinerja hasil pelatihan.

D. Evaluasi Model YOLOv8

Ini adalah tahap di mana akan dilakukan evaluasi seberapa baik model dalam mendeteksi orang dan mengklasifikasi posisi tidur dengan data yang berbeda dari data yang digunakan dalam pelatihan model. Dalam tahap ini, model akan dievaluasi menggunakan data validasi yang telah ditentukan sebelumnya yaitu sebanyak 20% dari total data. Kemudian, akan dihitung nilai box(precision), recall dan mAP50 (mean average precision) untuk model deteksi dan nilai precision, recall, F1-score dan akurasi untuk model klasifikasi. Hasil dari evaluasi deteksi dan klasifikasi akan dijadikan sebagai indikator kinerja hasil evaluasi model.

E. Pengujian dan Analisis Performa Model YOLOv8

Ini adalah proses penghitungan yang bertujuan untuk mengevaluasi hasil dan seberapa baik algoritma YOLOv8 dalam mendeteksi orang dan mengklasifikasikan citra posisi tidur telentang (supine), miring kiri (left), miring kanan (right), dan negatif (pose selain 3 kelas sebelumnya atau yang tidak dikenali). Parameter-parameter yang dihitung untuk mengetahui

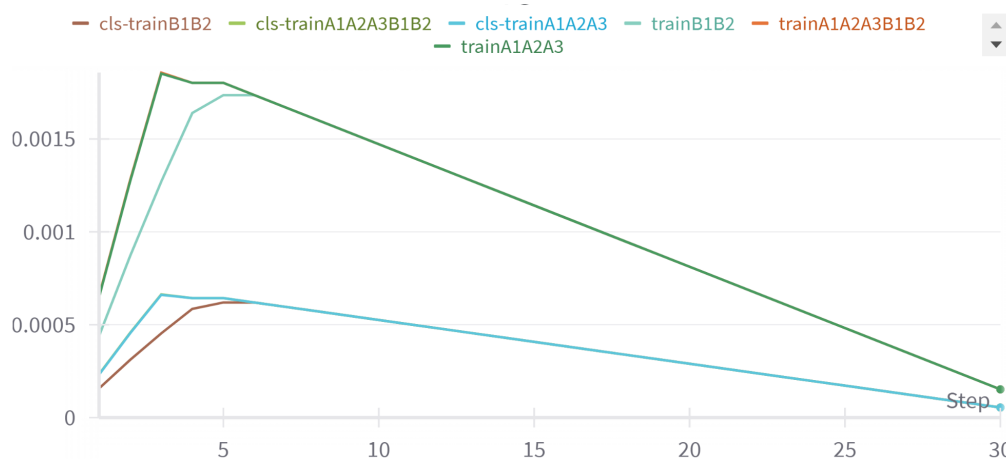
performa model tersebut antara lain adalah nilai precision, recall, F1-score dan akurasi untuk model klasifikasi dan box(precision), recall dan map50 untuk model deteksi.

3. HASIL DAN PEMBAHASAN

Pelatihan model YOLO melibatkan beberapa langkah untuk mengajarkan komputer mengenali, mendeteksi orang dan mengklasifikasikan berbagai pose tidur. Pertama, akan disiapkan kumpulan data gambar pose tidur. Kemudian, data ini dibagi menjadi beberapa bagian kecil yang disebut batch. Model akan melihat setiap batch ini satu per satu untuk belajar dari data. Selama pelatihan, kita menggunakan metode khusus untuk mengukur seberapa baik model membuat prediksi yang benar, dan kita menyesuaikan pengaturan model berdasarkan hasil pengukuran ini. Untuk *hyperparameter* dan perangkat yang digunakan selama pelatihan dapat dilihat dalam Tabel 2 dan Gambar 8.

Tabel 2. Hyperparameter dan perangkat yang digunakan.

Hyperparameter	Value	
	Deteksi	Klasifikasi
Batch size	16	16
Epoch	30	30
Input image size	640	224
Spesifikasi perangkat yang digunakan		
Server	Google colab	
CPU	Intel Xeon 2.20 GHz	
GPU	T4 15 GB	
RAM	12.7 GB	



Gambar 8. *Learning rate* yang digunakan selama pelatihan model deteksi

A. Pelatihan Model Deteksi

Tabel 3. Hasil pelatihan model deteksi

	Pelatihan dataset A1A2A3		Pelatihan dataset B1B2		Pelatihan dataset A1A2A3B1B2	
	YOLO v8n	YOLO v8s	YOLO v8n	YOLO v8s	YOLO v8n	YOLO v8s
Box(p)	1	0.999	0.997	0.926	0.995	0.99
Recall	0.993	0.987	1	0.837	0.996	0.837
mAP50	0.995	0.995	0.995	0.907	0.995	0.92



Gambar 9. Train loss dan val loss selama pelatihan model deteksi

Berdasarkan Tabel 3 dan Gambar 9, nilai mAP50 (mean average precision) pada setiap model yang dilatih hampir mencapai 100%. Ini menunjukkan bahwa model bisa mendeteksi dengan akurat pada dataset yang mirip. Namun, belum diketahui apakah model tersebut bisa mendeteksi secara akurat pada dataset yang berbeda seperti di lingkungan yang berbeda atau dengan orang yang berbeda.

B. Pelatihan Model Klasifikasi

Tabel 4. Hasil pelatihan model klasifikasi

	Pelatihan dataset A1A2A3		Pelatihan dataset B1B2		Pelatihan dataset A1A2A3B1B2	
	YOLO v8n	YOLO v8s	YOLO v8n	YOLO v8s	YOLO v8n	YOLO v8s
Presisi	0.92	0.96	0.8	0.82	0.9	0.88
Recall	0.93	0.95	0.8	0.82	0.92	0.88
F1-Score	0.92	0.95	0.8	0.82	0.91	0.88
Akurasi	0.933	0.967	0.8	0.82	0.896	0.871



Gambar 10. Train loss dan val loss selama pelatihan model klasifikasi.

Dari hasil pelatihan model klasifikasi yang ditunjukkan pada Tabel 4 dan Gambar 10, nilai matriks yang dihasilkan pada setiap model yang dilatih rata-rata mencapai 90%. Ini menunjukkan bahwa model bisa mengklasifikasi dengan akurat pada dataset yang mirip. Namun, belum diketahui apakah model tersebut bisa mengklasifikasi secara akurat pada dataset yang berbeda seperti di lingkungan yang berbeda atau dengan orang yang berbeda.

C. Pengujian Model Deteksi

Tabel 5. A1A2A3 model, diuji dengan A1A2A3 tes dataset

	YOLOv8n	YOLOv8s
Box(p)	1	1
Recall	1	0.92
mAP50	0.995	0.995

Tabel 6. A1A2A3 model, diuji dengan B1B2 tes dataset

	YOLOv8n	YOLOv8s
Box(p)	0.445	0.3
Recall	0.3	0.4
mAP50	0.283	0.281

Mengacu pada Tabel 5 dan Tabel 6, apabila model yang dilatih menggunakan dataset A1A2A3 dan diuji dengan dataset A1A2A3 yang memiliki kemiripan tinggi dari segi lingkungan (background) maka akan menghasilkan performa yang baik. Namun ketika diuji dengan dataset B1B2 performa menurun secara signifikan. Hal itu disebabkan karena besarnya perbedaan antara dataset A1A2A3 dan B1B2. Perbedaan performa model YOLOv8n dan YOLOv8s pun tidak terlalu signifikan.

Tabel 7. B1B2 model, diuji dengan A1A2A3 tes dataset

	YOLOv8n	YOLOv8s
Box(p)	0.392	0.57
Recall	0.279	0.633
mAP50	0.283	0.513

Tabel 8. B1B2 model, diuji dengan B1B2 tes dataset

	YOLOv8n	YOLOv8s
Box(p)	0.999	0.991
Recall	1	0.8
mAP50	0.995	0.885

Hampir sama dengan pengujian sebelumnya, Tabel 7 dan Tabel 8 menunjukkan bahwa model yang dilatih menggunakan dataset B1B2 dan diuji dengan dataset B1B2 yang memiliki kemiripan tinggi dari segi lingkungan menghasilkan performa yang baik. Namun, ketika diuji dengan dataset A1A2A3, performanya menurun secara signifikan karena perbedaan yang besar antara dataset A1A2A3 dan B1B2. Meski begitu, performa model yang dilatih dengan YOLOv8s terlihat lebih baik dalam mendeteksi gambar dengan lingkungan yang berbeda, meskipun performanya hanya sekitar 50-60% saja.

Tabel 9. A1A2A3B1B2 model, diuji dengan A1A2A3B1B2 tes dataset

	YOLOv8n	YOLOv8s
Box(p)	0.996	0.983
Recall	1	0.838
mAP50	0.995	0.916

Dari hasil pengujian yang terlihat pada Tabel 9, setelah dataset digabung dan

dilakukan pengujian, menghasilkan performa yang cukup baik. Perbedaan antara performa model yang dilatih dengan YOLOv8n dan YOLOv8s pun cukup signifikan. Ini menunjukkan bahwa meskipun YOLOv8s memiliki performa yang lebih baik secara keseluruhan, YOLOv8n juga mampu memberikan hasil yang memadai, terutama ketika dilatih dengan dataset yang lebih kecil.

D. Pengujian Model Klasifikasi

Tabel 10. A1A2A3 model, diuji dengan A1A2A3 tes dataset

	YOLOv8n	YOLOv8s
Presisi	0.995	0.995
Recall	0.994	0.994
F1-Score	0.994	0.994
Akurasi	0.994	0.953

Tabel 11. A1A2A3 model, diuji dengan B1B2 tes dataset

	YOLOv8n	YOLOv8s
Presisi	0.25	0.32
Recall	0.33	0.41
F1-Score	0.13	0.21
Akurasi	0.121	0.22

Selama pengujian, informasi dari Tabel 10 dan Tabel 11, model yang dilatih menggunakan dataset A1A2A3 dan diuji dengan dataset A1A2A3 yang memiliki kemiripan tinggi dari segi lingkungan maka akan menghasilkan performa yang cukup baik. Namun ketika diuji dengan dataset B1B2 performa menurun sangat signifikan. Hal itu disebabkan karena besarnya perbedaan antara dataset A1A2A3 dan B1B2. Perbedaan performa model YOLOv8n dan YOLOv8s pun tidak terlalu signifikan.

Tabel 12. B1B2 model, diuji dengan A1A2A3 tes dataset

	YOLOv8n	YOLOv8s
Presisi	0.83	0.8
Recall	0.8	0.78
F1-Score	0.79	0.78
Akurasi	0.8	0.78

Tabel 13. B1B2 model, diuji dengan B1B2 tes dataset

	YOLOv8n	YOLOv8s
Presisi	0.73	0.8
Recall	0.73	0.8
F1-Score	0.73	0.8
Akurasi	0.73	0.8

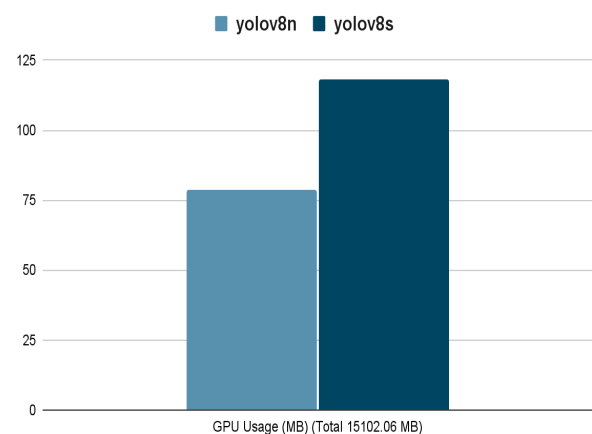
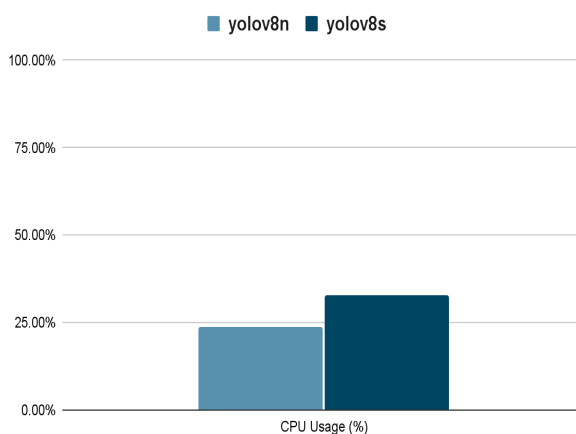
Tabel 12 dan Tabel 13 menunjukkan performa model yang dilatih menggunakan dataset B1B2 dan diuji dengan dataset B1B2 maupun A1A2A3 hanya mencapai sekitar 70% - 80%. Meskipun demikian, model yang dilatih dengan YOLOv8s menunjukkan peningkatan dalam klasifikasi gambar. Namun, peningkatan tersebut tidak terlalu signifikan.

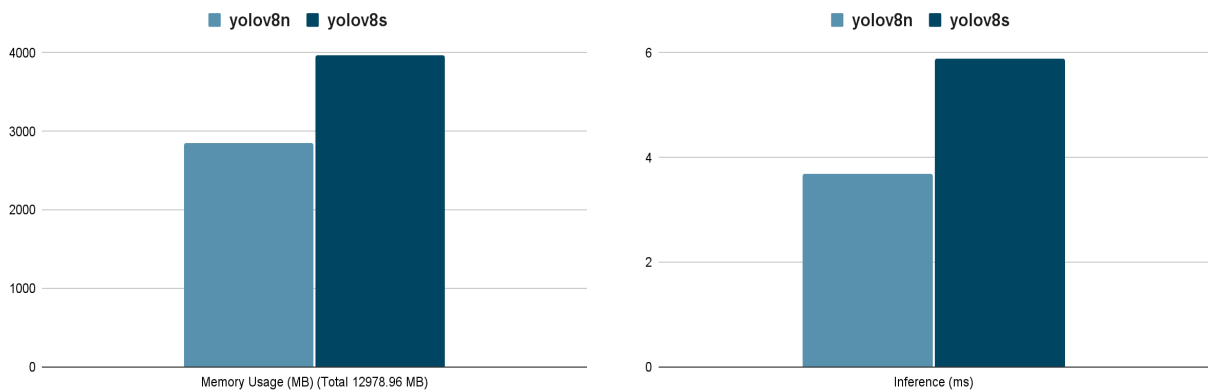
Tabel 14. A1A2A3B1B2 model, diuji dengan A1A2A3B1B2 tes dataset

	YOLOv8n	YOLOv8s
Presisi	0.85	0.91
Recall	0.86	0.89
F1-Score	0.85	0.9
Akurasi	0.846	0.879

Berdasarkan Tabel 14, setelah dataset digabung dan dilakukan pengujian, menghasilkan performa yang cukup baik. Perbedaan antara performa model yang dilatih dengan YOLOv8n dan YOLOv8s pun cukup terlihat walau tidak terlalu signifikan. Ini menunjukkan bahwa YOLOv8s memiliki performa yang lebih baik dalam klasifikasi gambar secara keseluruhan.

E. Penggunaan Sumber Daya Komputasi





Gambar 11. Penggunaan Sumber Daya Komputasi

Ketika dilakukan pengujian A1A2A3B1B2 model dengan A1A2A3B1B2 tes dataset, penggunaan sumber daya komputasi model YOLOv8s lebih tinggi. Namun, memiliki performa yang lebih baik dalam deteksi dan klasifikasi gambar secara keseluruhan. YOLOv8n juga memiliki kelebihan, yaitu waktu yang diperlukan untuk mendeteksi gambar (inference) lebih singkat yaitu 3.7ms dimana dalam 1 detik dapat mendeteksi kurang lebih 270 gambar.

4. PENUTUP

Pengujian menunjukkan bahwa model YOLO berhasil melakukan human detection (deteksi orang) pada tempat tidur dengan cukup baik pada dataset yang memiliki kemiripan tinggi dengan data pelatihan dengan rata-rata akurasi di 90% keatas. Namun, ketika diuji dengan dataset yang memiliki kemiripan rendah, model mengalami penurunan performa yang sangat signifikan dengan rata-rata akurasi di 25% hingga 40%. Ketika dilakukan pengujian pada dataset campuran, performa yang dihasilkan cukup baik, rata-rata di 90% keatas. Dengan begitu, dataset pelatihan perlu disesuaikan pada lingkungan yang akan dilakukan deteksi agar menghasilkan performa yang baik.

Pengujian model YOLO dalam klasifikasi gambar pun cukup baik pada dataset yang memiliki kemiripan tinggi dengan akurasi rata-rata di 85% keatas. Model A1A2A3 ketika diuji dengan dataset B1B2, performanya menurun drastis dengan rata-rata akurasi 20%. Namun, ketika model B1B2 diuji dengan dataset A1A2A3 performa yang dihasilkan cukup walaupun kurang maksimal dengan rata-rata akurasi di 75%. Dengan begitu bisa dikatakan bahwa dataset B1 dan B2 dalam klasifikasi memiliki variasi yang lebih banyak sehingga bisa mengklasifikasikan gambar dengan lebih baik dalam berbagai kondisi. Pengujian pada dataset campuran pun juga memiliki performa yang cukup baik dengan rata-rata 85%. Maka dari itu, performa deteksi dan klasifikasi sangat dipengaruhi oleh variasi dataset yang ada dan seberapa general dataset yang digunakan.

Perbedaan antara YOLOv8n dan YOLOv8s juga tidak terlalu signifikan. Namun memiliki kelebihan-nya masing-masing dimana YOLOv8n lebih cocok digunakan pada perangkat edge yang tidak memerlukan proses komputasi yang berat dan juga dataset yang relatif kurang banyak, YOLOv8s juga memiliki kelebihan dalam keseluruhan performa modelnya. Namun, memerlukan proses komputasi yang lumayan tinggi dan dibutuhkan datasets yang relatif banyak.

DAFTAR PUSTAKA

- Black, J., Baharestani, M. M., Cuddigan, J., Dorner, B., Edsberg, L., Langemo, D., Posthauer, M. E., Ratliff, C., Taler, G., & National Pressure Ulcer Advisory Panel. (2007). National Pressure Ulcer Advisory Panel's updated pressure ulcer staging system. *Advances in Skin & Wound Care*, 20(5), 269–274. <https://doi.org/10.1097/01.ASW.0000269314.23015.e9>
- Du, S., Zhang, P., Zhang, B., & Xu, H. (2021). Weak and Occluded Vehicle Detection in Complex Infrared Environment Based on Improved YOLOv4. *IEEE Access*, PP, 1–1. <https://doi.org/10.1109/ACCESS.2021.3057723>
- Eka Putra, W. S. (2016). Klasifikasi Citra Menggunakan Convolutional Neural Network (CNN) pada Caltech 101. *Jurnal Teknik ITS*, 5(1). <https://doi.org/10.12962/j23373539.v5i1.15696>
- Gao, M., Du, Y., Yang, Y., & Zhang, J. (2019). Adaptive anchor box mechanism to improve the accuracy in the object detection system. *Multimedia Tools and Applications*, 78(19), 27383–27402. <https://doi.org/10.1007/s11042-019-07858-w>
- Hong, Z., Yang, T., Tong, X., Zhang, Y., Jiang, S., Zhou, R., Han, Y., Wang, J., Yang, S., & Liu, S. (2021). Multi-Scale Ship Detection From SAR and Optical Imagery Via A More Accurate YOLOv3. *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*, 14, 6083–6101. <https://doi.org/10.1109/JSTARS.2021.3087555>
- Huang, T., Mariani, S., & Redline, S. (2020). Sleep Irregularity and Risk of Cardiovascular Events: The Multi-Ethnic Study of Atherosclerosis. *Journal of the American College of Cardiology*, 75(9), 991–999. <https://doi.org/10.1016/j.jacc.2019.12.054>
- Ji, W., Pan, Y., Xu, B., & Wang, J. (2022). A Real-Time Apple Targets Detection Method for Picking Robot Based on ShufflenetV2-YOLOX. *Agriculture*, 12(6), Article 6. <https://doi.org/10.3390/agriculture12060856>
- Liu, S., & Ostadabbas, S. (2017). A Vision-Based System for In-Bed Posture Tracking. *2017 IEEE International Conference on Computer Vision Workshops (ICCVW)*, 1373–1382. <https://doi.org/10.1109/ICCVW.2017.163>
- Muhammad Nur Ihsan Muhlashin, & Stefanie, A. (2023). Klasifikasi Penyakit Mata Berdasarkan Citra Fundus Menggunakan YOLO V8. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 7(2), 1363–1368. <https://doi.org/10.36040/jati.v7i2.6927>
- Ostadabbas, S., Yousefi, R., Nourani, M., Faezipour, M., Tamil, L., & Pompeo, M. Q. (2012). A resource-efficient planning for pressure ulcer prevention. *IEEE Transactions on Information Technology in Biomedicine: A Publication of the IEEE Engineering in Medicine and Biology Society*, 16(6), 1265–1273. <https://doi.org/10.1109/TITB.2012.2214443>
- Redmon, J., Divvala, S., Girshick, R., & Farhadi, A. (2016). *You Only Look Once: Unified, Real-Time Object Detection*. 779–788. <https://doi.org/10.1109/CVPR.2016.91>
- Reis, D., Kupec, J., Hong, J., & Daoudi, A. (2023). *Real-Time Flying Object Detection with YOLOv8* (arXiv:2305.09972). arXiv. <https://doi.org/10.48550/arXiv.2305.09972>

- Smith, D. M. (1995). Pressure ulcers in the nursing home. *Annals of Internal Medicine*, 123(6), 433–442. <https://doi.org/10.7326/0003-4819-123-6-199509150-00008>
- Wang, H., Xu, X., Liu, Y., Lu, D., Liang, B., & Tang, Y. (2023). Real-Time Defect Detection for Metal Components: A Fusion of Enhanced Canny–Devernay and YOLOv6 Algorithms. *Applied Sciences*, 13(12), Article 12. <https://doi.org/10.3390/app13126898>
- Wibowo, S., & Sugiarto, I. (2023). Hand Symbol Classification for Human-Computer Interaction Using the Fifth Version of YOLO Object Detection. *CommIT (Communication and Information Technology) Journal*, 17(1), 43–50. <https://doi.org/10.21512/commit.v17i1.8520>
- Yanto, Y., Aziz, F., & Irmawati, I. (2023). YOLO-V8 Peningkatan Algoritma Untuk Deteksi Pemakaian Masker Wajah. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 7(3), Article 3. <https://doi.org/10.36040/jati.v7i3.7047>
- Zhou, Y., Tang, Y., Zou, X., Wu, M., Tang, W., Meng, F., Zhang, Y., & Kang, H. (2022). Adaptive Active Positioning of Camellia oleifera Fruit Picking Points: Classical Image Processing and YOLOv7 Fusion Algorithm. *Applied Sciences*, 12(24), Article 24. <https://doi.org/10.3390/app122412959>