

Single-Image Face Recognition For Student Identification Using Facenet512 And Yolov8 In Academic Environtment With Limited Dataset

**Almas Najiib Imam Muttaqin¹, Ardytha Luthfiarta^{*2}, Adhitya Nugraha³,
Pramesya Mutia Salsabila⁴**

^{1,2,3,4}Informatics Engineering, Faculty of Computer Science, Dian Nuswantoro University, Indonesia

Email: ²ardytha.luthfiarta@dsn.dinus.ac.id

Received : Oct 7, 2025; Revised : Sep 23, 2025; Accepted : Sep 23, 2025; Published : Oct 16, 2025

Abstract

Face recognition has become one of the most significant research areas in image processing and computer vision, mainly due to its wide applications in security, identity verification, and human and machine interaction. In this study, FaceNet512 and YOLOv8 models are used to overcome the challenges in face recognition with a limited dataset, which is only one formal photo per individual. The application of image augmentation to the model achieved 90% accuracy and ROC curve of 0.82, while the model without augmentation achieved 89% accuracy and ROC curve of 0.79. FaceNet512 showed superiority in producing more accurate and detailed facial representations compared to other models, such as ArcFace and FaceNet, especially in handling minimal facial variations. Meanwhile, YOLOv8 provides efficient face detection across various lighting conditions and viewing angles. The main challenge in this research is the low quality of the original image, which can reduce the accuracy of face recognition. These results show the great potential of using deep learning-based face recognition systems in the real world, especially for automatic attendance applications in academic environments.

Keywords: *Face Detection, Face Recognition, FaceNet512, Image Augmentation, Limited Dataset, YOLOv8*

This work is an open access article and licensed under a Creative Commons Attribution-Non Commercial 4.0 International License



1. PENDAHULUAN

Pengenalan wajah (face recognition) telah menjadi salah satu bidang utama dalam pengolahan citra dan visi komputer karena aplikasinya yang luas, terutama di sektor keamanan, verifikasi identitas, dan interaksi antara manusia dan mesin[1], [2]. Pada lingkungan akademik, teknologi ini semakin banyak digunakan, terutama untuk sistem presensi otomatis, kontrol akses, dan identifikasi mahasiswa. Meskipun teknologi pengenalan wajah terus berkembang, penerapannya di dunia nyata masih menghadapi sejumlah tantangan, terutama ketika dataset terbatas hanya pada satu gambar formal per individu[3], [4], [5].

Tantangan utama dalam menggunakan dataset yang hanya terdiri dari satu gambar adalah minimnya variasi gambar untuk melatih model pengenalan wajah. Model deep learning umumnya membutuhkan variasi yang signifikan dalam sudut pandang, pencahayaan, dan ekspresi wajah agar dapat melakukan generalisasi yang baik. Ketika variasi ini tidak ada, model cenderung mengalami kesulitan untuk mengenali wajah secara akurat dalam kondisi yang berbeda dari gambar latih, seperti saat pencahayaan berubah atau sudut wajah berbeda[6], [7]. Minimnya variasi ini menyebabkan model lebih rentan terhadap kesalahan dalam proses pengenalan, yang dapat berdampak pada akurasi keseluruhan sistem[3], [8].

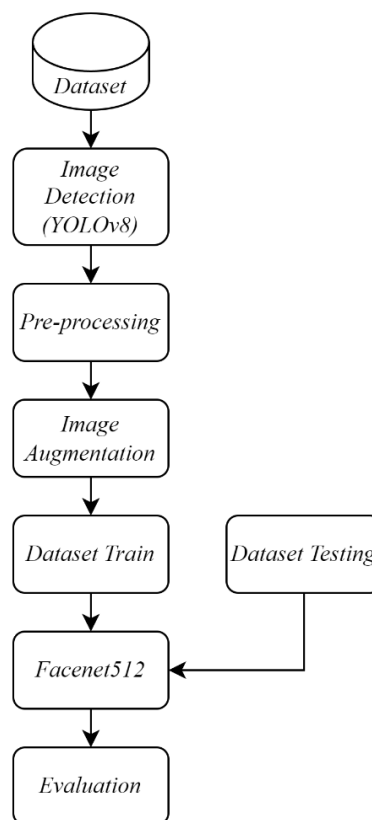
Selain itu, kualitas gambar yang rendah, seperti gambar dengan pencahayaan yang buruk, resolusi rendah, atau fokus yang tidak tajam, dan gambar yang telah lampau dapat semakin memperburuk kinerja model pengenalan wajah. Pada gambar berkualitas rendah, model sering kali kesulitan mengekstraksi

fitur-fitur penting dari wajah, seperti bentuk mata, hidung, dan mulut, yang diperlukan untuk membedakan satu wajah dengan wajah lainnya[9], [10]. Keterbatasan kualitas ini sangat relevan dalam lingkungan akademik, di mana pengenalan wajah sering kali diterapkan dalam skala besar namun dengan sumber data yang terbatas, seperti foto formal mahasiswa yang diambil pada saat pendaftaran[11], [12].

Penelitian terdahulu cenderung berhasil dalam mengatasi tantangan pengenalan wajah dengan menggunakan dataset besar yang memiliki banyak variasi gambar. Namun, pendekatan ini tidak dapat diterapkan dalam kasus di mana dataset sangat terbatas, seperti hanya menggunakan satu gambar per individu. Dalam skenario ini, ada gap penelitian yang signifikan terkait dengan cara mengoptimalkan model pengenalan wajah untuk bekerja dengan data yang sangat terbatas[13]. Sementara beberapa studi telah menyarankan bahwa augmentasi gambar dapat membantu meningkatkan variasi data latih, eksplorasi tentang bagaimana teknik augmentasi dapat diterapkan secara efektif pada dataset gambar tunggal masih sangat terbatas[2], [12], [14], [15].

Penelitian ini bertujuan untuk mengatasi masalah tersebut dengan menggunakan teknik augmentasi gambar guna mensimulasikan variasi gambar yang lebih luas, seperti perubahan sudut pandang, skala, pencahayaan, dan noise. Teknik augmentasi ini diharapkan dapat memperkaya dataset latih secara artifisial, sehingga meningkatkan kemampuan model pengenalan wajah untuk melakukan generalisasi dengan lebih baik meskipun hanya memiliki satu gambar per individu[5], [11]. Dengan demikian, penelitian ini diharapkan dapat mengisi gap penelitian yang ada terkait penggunaan augmentasi gambar untuk meningkatkan akurasi pengenalan wajah dalam kondisi dataset terbatas, terutama di lingkungan akademik untuk penggunaan sistem absensi kehadiran[12].

2. METODE PENELITIAN



Gambar 1. Alur Penelitian

Seperti pada gambar 1, Alur penelitian ini dimulai dengan pengumpulan dataset[16] yang terdiri dari satu foto formal per mahasiswa. Setelah itu, dilakukan deteksi wajah menggunakan YOLOv8. Selanjutnya, dilakukan *pre-processing* untuk mempersiapkan data seperti memotong bagian wajah yang sebelumnya telah dideteksi. Augmentasi gambar kemudian diterapkan untuk memperbesar variasi data yang terbatas. Setelah augmentasi, dataset *train* disiapkan untuk melatih model. Selanjutnya, model dilatih menggunakan FaceNet512 untuk mengekstraksi fitur wajah. Pada tahap akhir, dataset testing digunakan untuk mengevaluasi performa model menggunakan metrik seperti akurasi, *precision*, *recall*, dan *f1-score* untuk mengukur kemampuan model dalam mengenali wajah mahasiswa. Pada penelitian ini juga menggunakan *hardware* yaitu AMD Ryzen 9 4900HS dengan *base clockspeed* 3.00GHz, NVIDIA GTX 1660-Ti (Mobile GPU) dengan *VRAM* 6GB, dan RAM yang digunakan yaitu 24GB. *Libraries* dan *tools* yang digunakan meliputi *DeepFace Framework* sebagai platform utama untuk model pengenalan wajah, TensorFlow 2.10.1 untuk melatih dan menjalankan model berbasis GPU serta membuat Image Generator menggunakan modul *ImageDataGenerator*. *Cupy* menggantikan *numpy* dalam penghitungan vektor untuk operasi berbasis GPU dan digunakan juga untuk melatih model, sedangkan FaceNet512 merupakan model dasar pengenalan wajah yang lebih akurat. YOLOv8 digunakan untuk deteksi wajah yang cepat dan efisien. Untuk pemrosesan gambar lebih lanjut, OpenCV digunakan untuk melakukan pemrosesan gambar, sementara ImgAug membantu dalam membuat augmentasi gambar yang bervariasi untuk meningkatkan variasi data pelatihan.

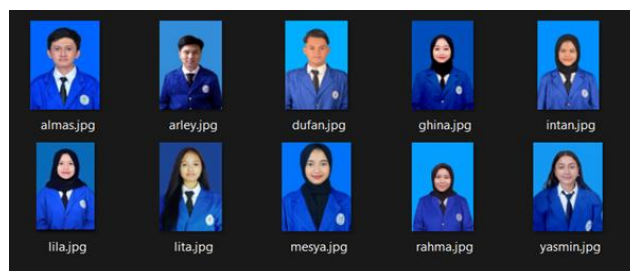
2.1.1 Dataset

Dataset yang digunakan dalam penelitian ini diperoleh dari foto formal mahasiswa Universitas Dian Nuswantoro yang diambil pada tahun pertama pada saat pendaftaran. Setiap individu dalam dataset direpresentasikan oleh satu foto formal, yang menunjukkan penampilan formal dengan seragam universitas. Pada penelitian ini, terdapat total 10 foto yang diambil dari 10 mahasiswa yang berbeda yang ditampilkan pada gambar 2.

2.1.2 Pre-processing

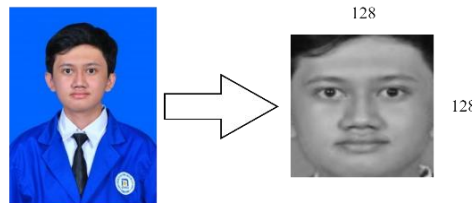
Pada tahap *pre-processing*, beberapa langkah dilakukan untuk mempersiapkan data gambar agar lebih seragam dan siap digunakan dalam pelatihan model. Pada tahap sebelumnya setiap gambar foto formal diproses menggunakan algoritma deteksi wajah YOLOv8 untuk mendeteksi area wajah. Setelah wajah terdeteksi, gambar dipotong sesuai dengan area yang teridentifikasi dan diubah ukurannya menjadi resolusi standar 128x128 piksel. Proses pemotongan dan penyeragaman ukuran ini bertujuan untuk menjaga konsistensi input yang diberikan ke model selama pelatihan.

Selanjutnya, gambar hasil pemotongan diubah menjadi *grayscale* yang ada pada gambar 3. Pemrosesan *grayscale* dipilih karena gambar hitam-putih hanya membutuhkan satu channel warna, berbeda dengan gambar berwarna yang memiliki tiga channel (RGB). Selain itu, dalam konteks deteksi wajah, informasi warna seringkali tidak terlalu penting, sementara fitur-fitur penting seperti tepi, kontur, dan bentuk tetap terjaga serta cukup untuk kebutuhan klasifikasi atau deteksi oleh model [17].



Gambar 2. Foto formal mahasiswa Universitas Dian Nuswantoro

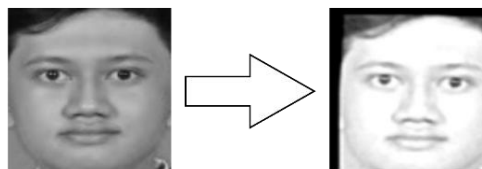
Pada tahap berikutnya, penambahan jumlah gambar ditentukan terlebih dahulu untuk memperbesar variasi data latih. Jumlah gambar diperbanyak dengan menghasilkan beberapa salinan dari gambar asli, misalnya 5, 10, atau 15 gambar tambahan per gambar. Penambahan ini bertujuan untuk memperluas dataset sehingga model dapat menjadi lebih tangguh dan mampu melakukan generalisasi yang lebih baik pada data baru [12], [15].



Gambar 3. Gambar wajah yang telah di-*grayscale*

2.1.3 Augmentasi Gambar

Dalam proses augmentasi data pada gambar 4, berbagai teknik diterapkan untuk memperbanyak variasi data latih guna meningkatkan ketahanan model terhadap berbagai kondisi. Teknik augmentasi yang digunakan meliputi penambahan *noise* acak, rotasi gambar secara acak, serta perubahan ukuran gambar melalui pembesaran dan pengecilan. Selain itu, dilakukan translasi gambar dengan memindahkan posisi objek secara acak dari titik tengah, serta penambahan detail gambar. Pembalikan gambar secara horizontal juga diterapkan, diikuti dengan penyesuaian kontras dan blur untuk mensimulasikan variasi kondisi gambar. Intensitas piksel dan kecerahan gambar diubah, serta dilakukan penajaman gambar untuk memastikan model dapat menangkap fitur penting. Kombinasi dari berbagai teknik ini membantu model dalam menghadapi variasi kondisi dengan lebih baik, sehingga meningkatkan performanya secara keseluruhan [12], [15], [18].



Gambar 4. Gambar wajah grayscale yang telah diaugmentasi

2.1.4 Deepface Framework

DeepFace Framework[3]-[5] adalah sebuah platform *open-source* yang dirancang untuk memudahkan integrasi berbagai model pengenalan wajah. **DeepFace Framework** adalah sebuah platform *open-source* yang dirancang untuk memudahkan integrasi berbagai model pengenalan wajah berbasis *deep learning*. Framework ini mendukung beberapa model unggulan seperti FaceNet[20], FaceNet512[21], ArcFace[8], VGG-Face, OpenFace[22], SFace,[23] dan DeepID[13], yang memungkinkan penggunaan teknologi pengenalan wajah dalam berbagai aplikasi tanpa perlu membangun algoritma dari awal. **DeepFace** dilatih menggunakan dataset seperti *Labelled Face in the Wild* (LFW)[24] dan *Youtube Face* (YTF)[25], yang membuatnya mampu bekerja dalam berbagai kondisi pengujian dengan performa yang dapat diandalkan.

Framework ini dipilih dalam penelitian ini karena fleksibilitas dan kemampuannya dalam mendukung berbagai model pengenalan wajah dengan antarmuka yang sederhana. Hal ini memungkinkan integrasi yang cepat ke dalam sistem yang dikembangkan, terutama dalam kondisi keterbatasan data, seperti hanya satu foto per individu. **DeepFace Framework** juga mendukung deteksi wajah dengan menggunakan beberapa model yang efisien, seperti YOLOv8, untuk memastikan wajah

dapat dideteksi dengan baik sebelum dilakukan ekstraksi fitur oleh model seperti FaceNet512. Kombinasi ini memastikan sistem tetap dapat bekerja secara efektif meskipun dengan data yang terbatas, yang merupakan tantangan utama dalam penelitian pengenalan wajah di lingkungan akademik.

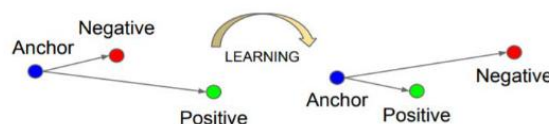
2.1.5 FaceNet512

FaceNet512 adalah model deep learning untuk pengenalan wajah yang merupakan pengembangan dari FaceNet asli yang dikembangkan oleh Google[20]. FaceNet512 dipilih karena representasi vektor yang lebih tinggi dari 128 dimensi yang dimiliki oleh FaceNet menjadi 512 dimensi[21], [26]. Peningkatan ini memungkinkan FaceNet512 menghasilkan representasi fitur wajah yang lebih detail dan akurat, terutama dalam kondisi di mana wajah-wajah yang dibandingkan memiliki tingkat kemiripan yang sangat tinggi.



Gambar 5. Struktur model FaceNet512

Secara arsitektural pada gambar 5, FaceNet512 menggunakan jaringan *Convolutional Neural Network* (CNN) yang dalam proses pelatihannya menerapkan triplet loss. Pada tahap pelatihan, model dilatih menggunakan triplet gambar yang terdiri dari tiga elemen: *anchor* (gambar wajah yang ingin dikenali), *positive* (gambar wajah yang sama dengan *anchor*), dan *negative* (gambar wajah yang berbeda). Tujuan dari *triplet loss* pada gambar 6, adalah untuk meminimalkan jarak antara *anchor* dan *positive*, serta memaksimalkan jarak antara *anchor* dan *negative*. Proses ini menghasilkan vektor fitur yang mampu secara efektif membedakan antara wajah-wajah yang berbeda, bahkan ketika perbedaannya sangat halus[34]-[36].



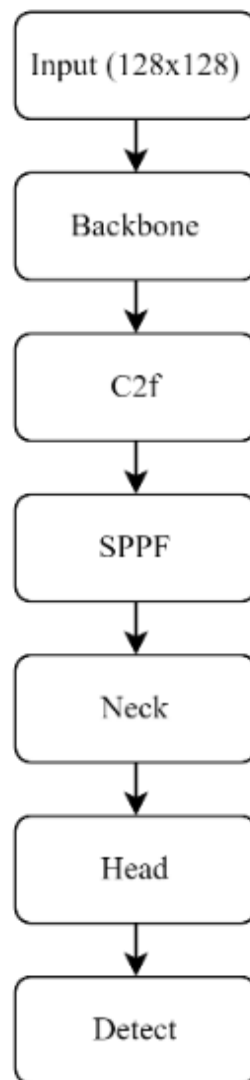
Gambar 6. The Triplet Loss Training

Dalam penerapannya, FaceNet512 memproses gambar wajah melalui beberapa lapisan konvolusi yang bertujuan untuk mengekstraksi fitur-fitur penting, seperti posisi relatif mata, hidung, dan mulut. Setelah melalui lapisan konvolusi, gambar wajah dikonversi menjadi **embedding** atau representasi vektor berdimensi 512. Vektor ini menyimpan informasi unik dari wajah yang dapat digunakan untuk proses verifikasi atau pengenalan wajah, dengan cara membandingkan embedding wajah baru dengan embedding wajah yang sudah ada.

Keunggulan dari arsitektur FaceNet512 terletak pada fleksibilitasnya dalam menangani variasi wajah yang luas, baik dalam hal ekspresi, pencahayaan, maupun sudut pengambilan gambar. Model ini juga dapat bekerja dengan baik meskipun hanya ada satu gambar formal untuk setiap individu, karena vektor berdimensi tinggi yang dihasilkan memiliki detail yang cukup untuk memisahkan wajah yang mirip dengan akurasi tinggi. FaceNet512 juga mengandalkan *L2 normalization*, yang memastikan bahwa panjang vektor fitur tetap konsisten, sehingga perbandingan antar vektor menjadi lebih stabil dan akurat.

2.1.6 YOLOv8

YOLOv8[11]-[16] model deteksi yang dirancang untuk mendeteksi objek secara efisien dan akurat yang dirancang untuk mendeteksi objek dengan cepat dan akurat dalam satu tahap. Model ini menggunakan arsitektur anchor-free, yang lebih sederhana dibandingkan dengan pendahulunya, dan memproses gambar menggunakan *Convolutional Neural Network* (CNN). YOLOv8 dipilih karena kemampuannya mendeteksi berbagai ukuran objek dengan efisiensi tinggi, sehingga dapat bekerja dengan optimal dalam beragam kondisi pencahayaan dan sudut pandang[10]. YOLOv8 dirancang untuk menangani berbagai ukuran objek dengan efisiensi tinggi, membuatnya mampu bekerja dengan baik dalam situasi dengan beragam kondisi pencahayaan dan sudut pandang.



Gambar 7. Arsitektur model deteksi YOLOv8

Alur pada gambar 7 dimulai dengan gambar input beresolusi 128x128 yang dimasukkan ke dalam model YOLOv8[10]. Gambar tersebut mewakili objek yang akan dideteksi oleh sistem. Tahap pertama adalah *Backbone*, di mana fitur awal dari gambar diekstraksi melalui serangkaian lapisan konvolusi. *Backbone* bertugas untuk mendeteksi pola dasar dalam gambar, seperti tepi, bentuk, dan warna, dengan menghasilkan representasi fitur yang lebih abstrak dan kaya akan informasi untuk diproses lebih lanjut.

Selanjutnya, fitur yang telah diekstraksi diproses melalui modul C2f (*Cross Stage Partial Networks*). C2f bertindak sebagai bottleneck, meningkatkan efisiensi dengan mengurangi kompleksitas komputasi tanpa mengorbankan performa. Modul ini membagi aliran data ke dalam beberapa bagian yang diproses secara paralel dan kemudian digabungkan kembali untuk menciptakan representasi fitur yang lebih kuat.

Setelah C2f, fitur gambar diproses oleh SPPF (*Spatial Pyramid Pooling Fast*). SPPF bertujuan untuk menangkap informasi multi-skala dari gambar, yang sangat penting untuk mendeteksi objek dalam berbagai ukuran. Dengan *pooling* di berbagai tingkat skala, SPPF memastikan bahwa model dapat mendeteksi objek secara konsisten, bahkan jika ukuran objek berubah dalam gambar.

Berikutnya, fitur yang telah diproses oleh Backbone dan SPPF diteruskan ke bagian Neck, yang bertugas menggabungkan fitur dari berbagai skala tersebut. Operasi *Upsample* dan *Concat* digunakan untuk memastikan deteksi objek dapat dilakukan pada berbagai ukuran dan tingkat skala. Proses di Neck memperkuat kemampuan model dalam menghadapi variasi ukuran objek yang berbeda.

Pada tahap Head, model memulai proses prediksi. Head melakukan tugas deteksi dengan memprediksi lokasi objek (*bounding box*) dan kelas objek berdasarkan fitur yang telah dikumpulkan dari Backbone dan Neck. Head dirancang untuk memberikan prediksi akhir dengan tingkat akurasi yang tinggi.

Akhirnya, dalam tahap Detect, model mengidentifikasi objek dalam gambar berdasarkan semua fitur yang telah diproses. Model menghasilkan hasil berupa posisi (*bounding box*) dan kelas objek yang terdeteksi. Pada penelitian ini, YOLOv8 digunakan sebagai detektor wajah untuk mendeteksi wajah sebelum proses ekstraksi fitur dilakukan oleh model FaceNet512. YOLOv8 dipilih karena kemampuannya yang tinggi dalam mendeteksi wajah secara efisien dengan tetap mempertahankan akurasi tinggi, meskipun bekerja dengan data yang minimal, seperti penggunaan hanya satu foto formal per individu.

2.1.7 Evaluasi

Untuk mengevaluasi performa model pengenalan wajah, empat metrik utama yang sering digunakan adalah akurasi, *precision*, *recall*, dan *F1 score*. Metrik-metrik ini memberikan penilaian yang menyeluruh mengenai kemampuan model dalam mengenali wajah dengan tepat serta mengukur bagaimana model menangani kesalahan selama proses prediksi.

Akurasi mengukur persentase prediksi yang benar dibandingkan dengan jumlah total gambar yang diuji, memberikan gambaran umum tentang kinerja model secara keseluruhan. Dalam akurasi manual, wajah yang dikenali salah, dikategorikan sebagai "Unknown", atau tidak dapat dideteksi atau dikenali dianggap sebagai kesalahan (1). *Precision* mengevaluasi seberapa tepat model dalam menghasilkan prediksi positif, yaitu dari semua wajah yang dikenali, berapa banyak yang benar-benar sesuai dengan wajah sebenarnya (2). *Recall* atau sensitivitas mengukur kemampuan model dalam mendeteksi semua wajah yang seharusnya dikenali, menunjukkan seberapa baik model dapat mengenali wajah-wajah yang sebenarnya ada dalam gambar (3). *F1 score* menggabungkan *precision* dan *recall* menjadi satu metrik yang menyeimbangkan keduanya, yang sangat berguna ketika kita ingin mengukur keseimbangan antara kemampuan model dalam menghindari kesalahan positif dan mendeteksi wajah dengan akurat (4).

True Positive (TP) merujuk pada prediksi yang tepat di mana model berhasil mengenali wajah yang benar, sedangkan *True Negative* (TN) mengacu pada kasus di mana model dengan benar tidak mengenali objek yang bukan wajah atau wajah yang tidak relevan. *False Positive* (FP) terjadi ketika model salah mengidentifikasi objek yang bukan wajah atau wajah yang tidak sesuai sebagai wajah yang benar. *False Negative* (FN) muncul ketika model gagal mendeteksi wajah yang sebenarnya ada dalam gambar.

$$Akurasi = \frac{Jumlah\ Prediksi\ Benar}{Total\ Gambar} \quad (1)$$

$$Precision = \frac{TP}{TP+FP} \quad (2)$$

$$Recall = \frac{TP}{TP+FN} \quad (3)$$

$$F1\ Score = \frac{2 \times Precision \times Recall}{Precision + Recall} \quad (4)$$

2.1.8 Threshold Cosine Similarity

Dalam penelitian ini, kesamaan wajah diukur menggunakan metrik Cosine Similarity[28], [33], yang menghitung sudut antara dua vektor representasi wajah yang diekstraksi oleh model FaceNet512. Cosine Similarity (5) digunakan untuk menilai kemiripan dua wajah berdasarkan hubungan antar vektor tersebut. Semakin mendekati nilai maksimal (1), semakin mirip kedua wajah.

Threshold sebesar 0.75 berdasarkan studi[33] yang digunakan untuk menentukan apakah dua wajah dianggap mirip. Jika nilai cosine similarity lebih kecil atau sama dengan 0.75, wajah tersebut dikenali, sedangkan jika lebih besar dari 0.75, atau dengan kata lain dengan nilai error lebih dari 0.25 maka wajah dianggap tidak dikenali. Pemilihan threshold ini mengacu pada studi sebelumnya yang menunjukkan bahwa nilai ini mampu mengurangi false positives (kesalahan positif) sambil mempertahankan tingkat recall yang baik, yaitu kemampuan model dalam mendeteksi wajah yang benar[34], [35] dan juga didasarkan pada batas minimal di mana model dapat mengenali wajah mahasiswa dengan tingkat keakuratan yang tinggi[35].

Penggunaan threshold ini membantu menjaga keseimbangan antara menghindari kesalahan pengenalan wajah yang salah dan memastikan bahwa wajah yang benar dapat dikenali secara akurat. Nilai threshold ini sesuai untuk skenario dengan data terbatas, seperti pengenalan wajah dalam sistem absensi kehadiran.

Rumus Cosine Similarity yang digunakan untuk menghitung kemiripan antara dua vektor adalah sebagai berikut :

$$Cosine = \frac{x \cdot y}{\|x\| \|y\|} \quad (5)$$

di mana X dan Y merupakan vektor representasi dua wajah yang dibandingkan.

3. HASIL DAN PEMBAHASAN

Pada bagian ini, disajikan analisis mendalam mengenai proses pengujian, penerapan teknik augmentasi, dan evaluasi kinerja model dalam tugas pengenalan wajah yang menggunakan FaceNet512 serta deteksi wajah berbasis YOLOv8, dan YOLOv8 yang digunakan yaitu YOLOv8n-face.

Pembahasan mencakup bagaimana penerapan teknik augmentasi gambar berkontribusi dalam meningkatkan variasi data pelatihan dan memperkuat ketahanan model terhadap kondisi variatif di dunia nyata. Kinerja model dianalisis melalui metrik evaluasi, yaitu akurasi, *precision*, *recall*, dan *F1 score*, dengan jumlah augmentasi gambar yang bervariasi. Setiap eksperimen dibandingkan untuk mengidentifikasi pengaruh variasi data yang dihasilkan dari augmentasi terhadap kemampuan model dalam mengenali wajah secara konsisten. Visualisasi hasil pengenalan wajah juga disajikan sebagai representasi kinerja model setelah penerapan augmentasi pada data pelatihan.

Tabel 1. Teknik Augmentasi

Teknik Augmentasi		
Jenis	Nilai	Efek
Noise	Intensitas: 0 - 15	Menambah noise acak untuk mensimulasikan efek noise pada gambar
Poisson		
Rotasi	Sudut: -10 hingga 10	Mengubah orientasi gambar
Skala	Faktor Skala:0.5-1.0	Mengubah ukuran objek dalam gambar, mensimulasikan variasi jarak subjek terhadap kamera.
Translasi	Persentase Translasi: $x = -0.1$ hingga 0.1 , $y = -0.1$ hingga 0.1	Menggeser posisi objek dalam gambar, membantu model mengenali objek di posisi berbeda.
Penambahan Detail	Diterapkan	Meningkatkan kontras mikro, menonjolkan fitur-fitur kecil seperti tekstur atau garis halus
Flip Horizontal	Diterapkan	Membalik gambar secara horizontal
Kontras Gamma	Nilai Gamma 0.5-2.0	Mengubah kontras gambar
Blur Gaussian	Sigma : 0 - 1	Membuat efek blur pada gambar
Perubahan Intensitas	Faktor:0.5-1.5	Mengubah kecerahan gambar, mensimulasikan variasi pencahayaan.
Penajaman Gambar	Alpha 0.0, Tingkat Kecerahan:0.5-1.5	Mempertajam detail gambar

3.1.1 Augmentasi Gambar

Dalam proses augmentasi gambar[18] pada tabel 1, berbagai teknik diterapkan untuk memperkaya variasi data latih dan meningkatkan ketahanan model terhadap kondisi dunia nyata yang bervariasi, seperti pencahayaan, posisi, orientasi, dan skala gambar. Augmentasi ini dilakukan untuk meningkatkan kemampuan generalisasi model, terutama dalam kondisi di mana hanya terdapat sedikit data pelatihan.

Salah satu teknik augmentasi yang digunakan adalah penambahan *noise Poisson* dengan intensitas antara 0 hingga 15, yang mensimulasikan kondisi pencahayaan rendah atau penggunaan sensor berkualitas rendah. Teknik ini membantu model menjadi lebih tahan terhadap noise acak yang muncul pada gambar.

Gambar wajah dirotasi secara acak dengan sudut antara -10 hingga 10 derajat, sehingga model dilatih untuk mengenali wajah meskipun terdapat sedikit rotasi dalam gambar. Selain itu, gambar diperbesar atau diperkecil dengan skala antara 50% hingga 100%, yang mensimulasikan variasi jarak antara subjek dan kamera. Teknik ini membantu model mengenali wajah dengan ukuran yang berbeda.

Translasi gambar dilakukan dengan menggeser posisi wajah dalam gambar sebesar -10% hingga 10% dari lebar dan tinggi gambar asli. Ini membantu model untuk tetap mengenali wajah meskipun posisinya tidak berada di tengah gambar. Teknik lain seperti penambahan detail meningkatkan kontras mikro dalam gambar, yang memperjelas fitur-fitur kecil seperti tekstur wajah atau garis-garis halus.

Pembalikan gambar secara horizontal dilakukan untuk melatih model untuk mengenali wajah dari kedua sisi. Ini meningkatkan adaptabilitas model terhadap perubahan orientasi wajah dalam aplikasi dunia nyata. Perubahan kontras gamma dengan nilai antara 0.5 hingga 2.0 mensimulasikan berbagai kondisi pencahayaan, yang melatih model untuk tetap mengenali wajah dalam situasi pencahayaan yang beragam.

Blur Gaussian diterapkan dengan tingkat blur bervariasi antara 0 hingga 1.0 untuk mensimulasikan gambar yang buram, melatih model untuk tetap mampu mengenali wajah meskipun kualitas gambar menurun. Selain itu, perubahan intensitas piksel diterapkan dengan nilai antara 0.5 hingga 1.5, mensimulasikan variasi pencahayaan yang dapat terjadi pada gambar. Penajaman gambar dan perubahan kecerahan juga diterapkan untuk memperjelas detail gambar, sehingga model dapat menangkap fitur wajah dengan lebih baik.

Teknik augmentasi ini secara keseluruhan membantu model dalam menghadapi berbagai variasi kondisi gambar yang mungkin muncul di dunia nyata, seperti perubahan pencahayaan, orientasi, skala, dan noise. Dengan adanya variasi yang lebih besar dalam data latih, model menjadi lebih tangguh dalam mengenali wajah di berbagai situasi.

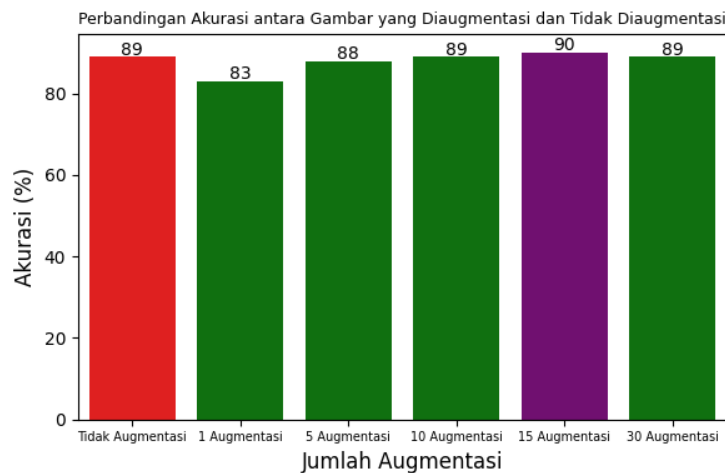
3.1.2 Hasil Pengujian

Dalam pengujian ini berdasarkan tabel 2, model pengenalan wajah diuji menggunakan 100 foto dari 10 orang dengan pose yang berbeda, dengan masing-masing orang memiliki 10 foto. Ketika augmentasi tidak diterapkan, model berhasil mencapai akurasi sebesar 89%, dengan precision 89%, recall 100%, dan F1 score 94% dalam waktu 3 detik. Ketika augmentasi diterapkan pada satu gambar per mahasiswa, akurasi model turun menjadi 83%, disertai penurunan precision menjadi 83%, sementara recall tetap 100%, namun F1 score turun menjadi 88%. Penurunan akurasi ini dapat dikaitkan dengan meningkatnya variasi yang dihasilkan dari augmentasi, yang berpotensi menyebabkan overfitting terhadap beberapa transformasi tertentu dan mengakibatkan model membutuhkan lebih banyak penyesuaian dalam mempelajari pola-pola baru yang muncul. Dalam konteks aplikasi dunia nyata, penurunan performa ini menunjukkan bahwa model pengenalan wajah yang digunakan mungkin kurang optimal jika diterapkan pada situasi yang memerlukan konsistensi tinggi dalam berbagai kondisi pencahayaan dan pose.

Tabel 2. Hasil Pengujian

Jumlah Foto Augmentasi yang dihasilkan per foto formal	Penerapan Augmentasi	Akurasi	Precision	Recall	F1 Score	Waktu Pelatihan
1x1	Tidak	89	89	100	94	3 detik
1x1	Diterapkan	83	83	100	88	3 detik
1x5	Diterapkan	88	88	100	93	12 detik
1x10	Diterapkan	85	85	100	92	25 detik
1x15	Diterapkan	90	90	100	95	38 detik
1x30	Diterapkan	89	89	100	94	1m 14s

Ketika jumlah gambar augmentasi yang dihasilkan ditingkatkan menjadi 5 gambar per mahasiswa, akurasi kembali meningkat menjadi 88%, dengan *precision*, *recall*, dan *F1 score* masing-masing 88%, 100%, dan 93% dalam waktu 12 detik. Ini menunjukkan bahwa variasi augmentasi yang lebih banyak membantu model dalam mengenali berbagai variasi pada wajah, meskipun meningkatkan kompleksitas data. Namun, saat jumlah augmentasi meningkat menjadi 10 gambar per mahasiswa, akurasi turun menjadi 85%, dengan *precision* dan *F1-score* juga menurun, meskipun recall tetap tinggi di angka 100%, dan waktu pelatihannya bertambah menjadi 12 detik.

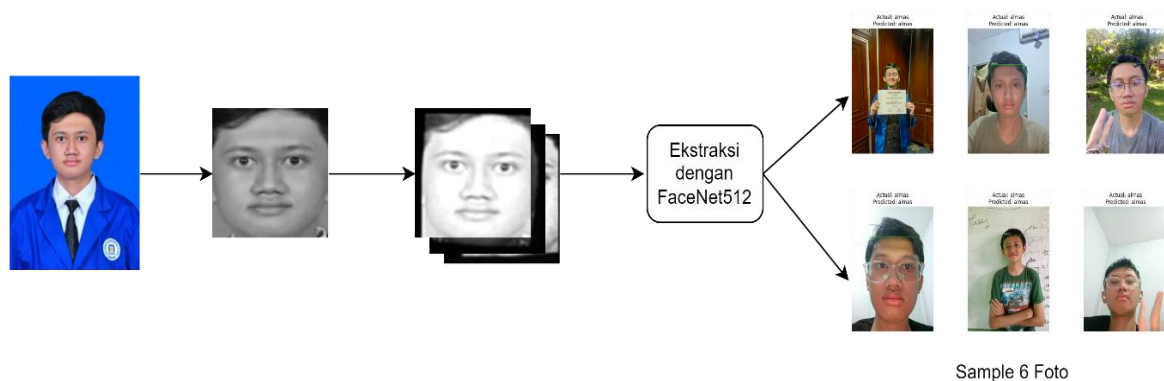


Gambar 8. Perbandingan Akurasi

Pada augmentasi 15 gambar per mahasiswa, akurasi meningkat menjadi 90% pada gambar 8, dan *F1 score* mencapai 95% dalam waktu 38 detik. Hal ini menunjukkan bahwa augmentasi yang lebih besar dapat membantu model beradaptasi dengan variasi data yang lebih luas. *F1 score* sebesar 95% ini mengindikasikan bahwa model mampu mempertahankan keseimbangan yang baik antara precision dan recall, yang sangat penting dalam skenario aplikasi dunia nyata, seperti keamanan dan identifikasi, di mana kesalahan deteksi dan false positives harus diminimalkan.

Terakhir, peningkatan jumlah augmentasi hingga 30 gambar per mahasiswa justru menyebabkan penurunan performa model, di mana akurasi menurun menjadi 89% dengan *F1 score* sebesar 94% dalam waktu 1 menit 14 detik. Hal ini menunjukkan bahwa jumlah gambar augmentasi yang terlalu banyak dapat menurunkan performa model dalam mengenali wajah. Jumlah gambar yang berlebihan ini menyebabkan model kesulitan dalam mengidentifikasi pola-pola yang relevan dan signifikan dalam pengenalan wajah.

Pada gambar 9, visualisasi hasil pengenalan wajah menggunakan model FaceNet512 setelah diterapkan proses augmentasi dapat dilihat. Proses ekstraksi fitur dari gambar mahasiswa yang diekstraksi menjadi vektor embedding membantu model dalam melakukan prediksi dengan berbagai tingkat augmentasi yang diuji. Dengan confidence interval sebesar 95% berdasarkan studi[36] pada akurasi $90\% \pm 0.06\%$, peningkatan performa model ini signifikan secara statistik dan menunjukkan dampak penggunaan augmentasi terhadap hasil yang dicapai.



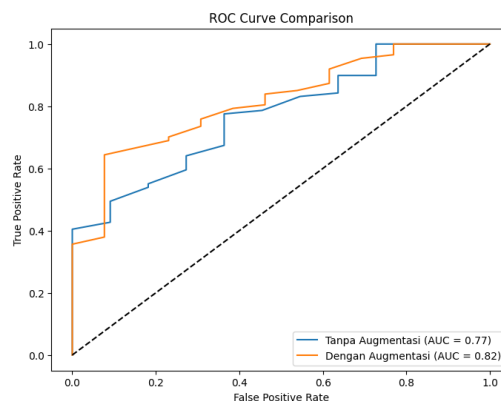
Gambar 9. Hasil Pengenalan Wajah

4. DISKUSI

Hasil penelitian ini menunjukkan bahwa penggunaan model FaceNet512 dan deteksi wajah dengan YOLOv8 pada dataset terbatas menghasilkan performa yang kompetitif. Penerapan augmentasi gambar terbukti efektif dalam meningkatkan akurasi model, meskipun waktu pelatihan juga meningkat seiring dengan penambahan jumlah augmentasi. Dengan menggunakan hanya satu foto formal per mahasiswa, penelitian ini berhasil mencapai akurasi maksimal sebesar 90%, yang sejalan dengan studi-studi terdahulu yang menunjukkan bahwa augmentasi data berperan penting dalam memperbaiki kinerja model deep learning dalam kondisi keterbatasan data[1], [13].

Namun, ada beberapa bias dalam dataset yang digunakan yang dapat memengaruhi performa model. Teknik augmentasi yang diterapkan berperan dalam meningkatkan ketahanan model terhadap variasi tersebut, tetapi mungkin tidak sepenuhnya menangani variasi dunia nyata secara menyeluruh. Dalam konteks ini, augmentasi membantu dalam memperluas cakupan variasi, namun beberapa variasi lingkungan yang nyata masih mungkin belum tercakup.

Jika dibandingkan dengan penelitian sebelumnya, kombinasi FaceNet512 dengan YOLOv8 menunjukkan performa yang lebih baik dibandingkan model FaceNet standar dengan Haar Cascade berdasarkan studi[37]. Kombinasi ini lebih efektif dalam mendeteksi wajah pada berbagai kondisi lingkungan[10], terutama dalam hal variasi pencahayaan dan sudut pandang yang berbeda-beda. Ini menunjukkan bahwa penambahan YOLOv8 sebagai detektor wajah membantu model dalam menghadapi skenario yang lebih kompleks.



Gambar 10. ROC Curve

Sebagai tambahan untuk memperkuat analisis hasil, kami juga menghasilkan visualisasi ROC curve untuk mengevaluasi keseimbangan antara true positive rate dan false positive rate. Seperti yang ditunjukkan pada Gambar 10, ROC AUC model tanpa augmentasi adalah 0,77, sementara dengan augmentasi meningkat menjadi 0,82. Ini menunjukkan bahwa augmentasi berperan dalam memperbaiki performa model dalam mendeteksi wajah secara akurat di bawah variasi yang lebih luas.

Penelitian ini juga mengidentifikasi tantangan terkait kualitas foto formal mahasiswa, yang sangat mempengaruhi akurasi pengenalan wajah. Meskipun augmentasi gambar mampu meningkatkan akurasi, kualitas gambar asli yang rendah tetap menjadi faktor yang berpotensi menurunkan performa model secara signifikan. Penambahan teknik peningkatan kualitas gambar sebelum augmentasi dapat dipertimbangkan untuk riset masa depan.

Secara keseluruhan, penelitian ini memberikan kontribusi terhadap pengembangan sistem pengenalan wajah yang efisien dan akurat meskipun dengan keterbatasan dataset. Penggunaan model FaceNet512 dan YOLOv8 menunjukkan hasil yang positif dan memiliki potensi untuk diterapkan dalam skenario dunia nyata, seperti presensi otomatis di lingkungan akademik.

5. KESIMPULAN

Kesimpulan dari penelitian ini menunjukkan bahwa penggunaan model FaceNet512 dan deteksi wajah dengan YOLOv8 pada dataset terbatas memberikan hasil yang baik, mendukung hipotesis bahwa augmentasi gambar dapat meningkatkan akurasi pengenalan wajah dalam skenario dataset terbatas. Dengan hanya menggunakan satu foto formal per mahasiswa dan penerapan augmentasi gambar, sistem pengenalan wajah ini mampu mencapai akurasi hingga 90%, yang menunjukkan potensi aplikasi praktis dalam sistem kehadiran akademik. FaceNet512 menunjukkan keunggulan dalam menangani variasi wajah dibandingkan model seperti ArcFace dan FaceNet standar, terutama ketika dataset memiliki jumlah gambar yang minim. Di sisi lain, YOLOv8 juga efisien dalam mendeteksi wajah dalam berbagai kondisi pencahayaan dan sudut pandang, menjadikannya pilihan efektif sebagai detektor wajah.

Tantangan utama yang diidentifikasi dalam penelitian ini adalah kualitas gambar asli yang bervariasi, yang dapat memengaruhi akurasi pengenalan wajah secara signifikan. Untuk penelitian selanjutnya, disarankan untuk menggunakan dataset yang lebih besar dan menerapkan teknik augmentasi yang lebih beragam. Selain itu, implementasi sistem ini di lingkungan pendidikan nyata untuk sistem absensi perlu dieksplorasi lebih lanjut guna validasi tambahan dan pengembangan kualitas pengenalan wajah, seperti teknik upscaling gambar, untuk meningkatkan ketajaman dan detail dalam ekstraksi fitur.

DAFTAR PUSTAKA

- [1] I. Adjabi, A. Ouahabi, A. Benzaoui, and A. Taleb-Ahmed, "Past, Present, and Future of Face Recognition: A Review," *Electronics* 2020, Vol. 9, Page 1188, vol. 9, no. 8, p. 1188, Jul. 2020, doi: 10.3390/ELECTRONICS9081188.
- [2] M. O. Oloyede, G. P. Hancke, and H. C. Myburgh, "A review on face recognition systems: recent approaches and challenges," *Multimed Tools Appl*, vol. 79, no. 37–38, pp. 27891–27922, Oct. 2020, doi: 10.1007/S11042-020-09261-2/TABLES/15.
- [3] R. Min, S. Xu, and Z. Cui, "Single-sample face recognition based on feature expansion," *IEEE Access*, vol. 7, pp. 45219–45229, 2019, doi: 10.1109/access.2019.2909039.
- [4] N. Kumar and V. Garg, "Single sample face recognition in the last decade: A survey," *Intern J Pattern Recognit Artif Intell*, vol. 33, no. 13, p. 1956009, Dec. 2019, doi: 10.1142/s0218001419560093.
- [5] V. Tomar, N. Kumar, and A. R. Srivastava, "Single sample face recognition using deep learning: a survey," *Artificial Intelligence Review* 2023 56:1, vol. 56, no. 1, pp. 1063–1111, Jul. 2023, doi: 10.1007/S10462-023-10551-Y.
- [6] Y. Wen, H. Yi, Z. Fan, Z. Xu, Y. Xue, and Y. Li, "Gallery-sensitive single sample face recognition based on domain adaptation," *Neurocomputing*, vol. 458, pp. 626–638, Oct. 2021, doi: 10.1016/j.neucom.2020.06.136.
- [7] F. Liu, F. Wang, Y. Wang, J. Zhou, and F. Xu, "Cycle-autoencoder based block-sparse joint representation for single sample face recognition," *Computers and Electrical Engineering*, vol. 101, Jul. 2022, doi: 10.1016/j.compeleceng.2022.108003.
- [8] J. Deng, J. Guo, J. Yang, N. Xue, I. Kotsia, and S. Zafeiriou, "ArcFace: Additive Angular Margin Loss for Deep Face Recognition," *IEEE Trans Pattern Anal Mach Intell*, vol. 44, no. 10, pp. 5962–5979, Oct. 2022, doi: 10.1109/TPAMI.2021.3087709.
- [9] J. Terven, D. M. Córdova-Esparza, and J. A. Romero-González, "A Comprehensive Review of YOLO Architectures in Computer Vision: From YOLOv1 to YOLOv8 and YOLO-NAS," *Mach Learn Knowl Extr*, vol. 5, no. 4, pp. 1680–1716, 2023, doi: 10.3390/make5040083.
- [10] Y. Zhao, F. Sun, and X. Wu, "FEB-YOLOv8: A multi-scale lightweight detection model for underwater object detection," *PLoS One*, vol. 19, no. 9, p. e0311173, Sep. 2024, doi: 10.1371/JOURNAL.PONE.0311173.
- [11] B. Ríos-Sánchez, D. Costa-da-Silva, N. Martín-Yuste, and C. Sánchez-Ávila, "Deep learning for facial recognition on single sample per person scenarios with varied capturing conditions," *Applied Sciences*, vol. 9, no. 24, p. 5474, Dec. 2019, doi: 10.3390/app9245474.

-
- [12] C. Shorten and T. M. Khoshgoftaar, "A survey on Image Data Augmentation for Deep Learning," *J Big Data*, vol. 6, no. 1, pp. 1–48, Dec. 2019, doi: 10.1186/S40537-019-0197-0/FIGURES/33.
- [13] S. I. Serengil and A. Ozpinar, "LightFace: A Hybrid Deep Face Recognition Framework," *Proceedings - 2020 Innovations in Intelligent Systems and Applications Conference, ASYU 2020*, Oct. 2020, doi: 10.1109/ASYU50717.2020.9259802.
- [14] A. Makalesi, R. Article Şefik İlkin SERENGİL, and A. Özpinar, "A Benchmark of Facial Recognition Pipelines and Co-Usability Performances of Modules," *Journal of Information Technologies*, vol. 17, no. 2, pp. 95–107, Apr. 2024, doi: 10.17671/GAZIBTD.1399077.
- [15] P. Chlap, H. Min, N. Vandenberg, J. Dowling, L. Holloway, and A. Haworth, "A review of medical image data augmentation techniques for deep learning applications," *J Med Imaging Radiat Oncol*, vol. 65, no. 5, pp. 545–563, Aug. 2021, doi: 10.1111/1754-9485.13261.
- [16] N. Kadek, D. A. Putri, A. Luthfiarta, P. Langgeng, and W. E. Putra, "OPTIMIZING BUTTERFLY CLASSIFICATION THROUGH TRANSFER LEARNING: FINE-TUNING APPROACH WITH NASNETMOBILE AND MOBILENETV2," *Jurnal Teknik Informatika (JUTIF)*, vol. 5, no. 3, pp. 685–692, 2024, doi: 10.52436/1.jutif.2024.5.3.1583.
- [17] Y. Xu, H. Kan, and G. Han, "Fourier-Reflexive Partitions and Group of Linear Isometries with Respect to Weighted Poset Metric," in *IEEE International Symposium on Information Theory - Proceedings*, Institute of Electrical and Electronics Engineers Inc., 2022, pp. 1975–1980. doi: 10.1109/ISIT50566.2022.9834567.
- [18] A. Jung, "imgaug — imgaug 0.4.0 documentation." Accessed: Sep. 13, 2024. [Online]. Available: <https://imgaug.readthedocs.io/en/latest/index.html>
- [19] S. I. Serengil and A. Ozpinar, "HyperExtended LightFace: A Facial Attribute Analysis Framework," *7th International Conference on Engineering and Emerging Technologies, ICEET 2021*, 2021, doi: 10.1109/ICEET53442.2021.9659697.
- [20] F. Schroff, D. Kalenichenko, and J. Philbin, "FaceNet: A unified embedding for face recognition and clustering," *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, vol. 07-12-June-2015, pp. 815–823, Oct. 2015, doi: 10.1109/CVPR.2015.7298682.
- [21] "GitHub - davidsandberg/facenet: Face recognition using Tensorflow." Accessed: Sep. 13, 2024. [Online]. Available: <https://github.com/davidsandberg/facenet>
- [22] T. Baltrusaitis, P. Robinson, and L. P. Morency, "OpenFace: An open source facial behavior analysis toolkit," *2016 IEEE Winter Conference on Applications of Computer Vision, WACV 2016*, May 2016, doi: 10.1109/WACV.2016.7477553.
- [23] Y. Zhong, W. Deng, J. Hu, D. Zhao, X. Li, and D. Wen, "SFace: Sigmoid-Constrained Hypersphere Loss for Robust Face Recognition," *IEEE Transactions on Image Processing*, vol. 30, pp. 2587–2598, 2021, doi: 10.1109/TIP.2020.3048632.
- [24] G. B. Huang, M. Ramesh, T. Berg, and E. Learned-Miller, "Labeled Faces in the Wild: A Database for Studying Face Recognition in Unconstrained Environments." [Online]. Available: <http://vis-www.cs.umass.edu/lfw/>.
- [25] C. Ferrari, S. Berretti, and A. Del Bimbo, "Extended YouTube Faces: A Dataset for Heterogeneous Open-Set Face Identification," *Proceedings - International Conference on Pattern Recognition*, vol. 2018-August, pp. 3408–3413, Nov. 2018, doi: 10.1109/ICPR.2018.8545642.
- [26] A. Firmansyah, T. F. Kusumasari, and E. N. Alam, "Comparison of Face Recognition Accuracy of ArcFace, Facenet and Facenet512 Models on Deepface Framework," in *ICCoSITE 2023 - International Conference on Computer Science, Information Technology and Engineering: Digital Transformation Strategy in Facing the VUCA and TUNA Era*, Institute of Electrical and Electronics Engineers Inc., 2023, pp. 535–539. doi: 10.1109/ICCoSITE57641.2023.10127799.
- [27] I. William, D. R. Ignatius Moses Setiadi, E. H. Rachmawanto, H. A. Santoso, and C. A. Sari, "Face Recognition using FaceNet (Survey, Performance Test, and Comparison)," *Proceedings of 2019 4th International Conference on Informatics and Computing, ICIC 2019*, Oct. 2019, doi: 10.1109/ICIC47613.2019.8985786.
- [28] N. Mardiana, R. D. Dana, Faisal, I. Farida, A. G. Azwar, and Nurwathi, "Similarity Measures Implementation on Face Authentication using Indonesian Citizen ID Card," *Proceeding of 2023*
-

- 17th International Conference on Telecommunication Systems, Services, and Applications, TSSA 2023*, 2023, doi: 10.1109/TSSA59948.2023.10366880.
- [29] T. Wu and Y. Dong, "YOLO-SE: Improved YOLOv8 for Remote Sensing Object Detection and Recognition," *Applied Sciences* 2023, Vol. 13, Page 12977, vol. 13, no. 24, p. 12977, Dec. 2023, doi: 10.3390/AP132412977.
- [30] J. Torres, "YOLOv8 Documentation: A Deep Dive into the Documentation - YOLOv8." Accessed: Sep. 13, 2024. [Online]. Available: <https://yolov8.org/yolov8-documentation/>
- [31] J. Straka and I. Gruber, "Object Detection Pipeline Using YOLOv8 for Document Information Extraction," 2023. [Online]. Available: <https://github.com/strakaj/YOLOv8-for-document-understanding.git>.
- [32] "GitHub - derronqi/yolov8-face: yolov8 face detection with landmark." Accessed: Sep. 19, 2024. [Online]. Available: <https://github.com/derronqi/yolov8-face?tab=readme-ov-file>
- [33] Y. Li, J. Wang, B. Pullman, N. Bandeira, and Y. Papakonstantinou, "Index-based, High-dimensional, Cosine Threshold Querying with Optimality Guarantees," *Theory of Computing Systems ∨ Mathematical Systems Theory*, vol. 65, no. 1, pp. 42–83, Jan. 2021, doi: 10.1007/S00224-020-10009-6.
- [34] J. Deng, J. Guo, X. An, Z. Zhu, and S. Zafeiriou, "Masked Face Recognition Challenge: The InsightFace Track Report," *Proceedings of the IEEE International Conference on Computer Vision*, vol. 2021-October, pp. 1437–1444, Aug. 2021, doi: 10.1109/ICCVW54120.2021.00165.
- [35] Y. Shi and A. Jain, "Probabilistic face embeddings," in *Proceedings of the IEEE International Conference on Computer Vision*, Institute of Electrical and Electronics Engineers Inc., Oct. 2019, pp. 6901–6910. doi: 10.1109/ICCV.2019.00700.
- [36] J. H. Grabman and C. S. Dodson, "Stark Individual Differences: Face Recognition Ability Influences the Relationship Between Confidence and Accuracy in a Recognition Test of Game of Thrones Actors," *J Appl Res Mem Cogn*, vol. 9, no. 2, pp. 254–269, Jun. 2020, doi: 10.1016/j.jarmac.2020.02.007.
- [37] Jamal Rosid, "Face Recognition Dengan Metode Haar Cascade dan Facenet," *Indonesian Journal of Data and Science*, vol. 3, no. 1, pp. 30–34, 2022, doi: 10.56705/ijodas.v3i1.38.