

Forensik Citra Digital Berbasis XceptionNet dengan Kerangka Kerja DFRWS untuk Deteksi Deepfake

Muh. Hajar Akbar^{*1}, Jimsan², Yahya³, Nasrullah⁴, Ilcham⁵

^{1,2,3,4,5}Prodi Sistem Informasi, Universitas Sembilanbelas November Kolaka

E-mail: ^{*1}hajarakbar16@gmail.com, ²jimsan190894@gmail.com, ³yahya@usn.ac.id,

⁴nasrullah@usn.ac.id, ⁵ilch4m@gmail.com

Abstrak

Penelitian ini menyajikan pendekatan baru (*novel approach*) untuk deteksi *deepfake* dengan mengintegrasikan kerangka kerja DFRWS (*Digital Forensics Research Workshop*) dengan arsitektur *deep learning* berbasis XceptionNet. Kemajuan pesat teknologi *deepfake* menimbulkan ancaman signifikan terhadap autentisitas media digital, sehingga memerlukan metode deteksi yang tangguh (*robust*). Penelitian kami mengimplementasikan model XceptionNet yang telah ditala halus (*fine-tuned*) dengan teknik regularisasi tambahan dan berfokus pada analisis fitur wajah. Model ini dilatih pada dataset seimbang yang terdiri dari 2.000 gambar, yang terbagi rata antara sampel autentik dan *deepfake*. Hasil eksperimental menunjukkan kinerja yang luar biasa (*exceptional*), dengan pencapaian akurasi 91.25%, presisi 88.73%, recall 94.50%, dan skor AUC 0.9710. Model yang diusulkan menunjukkan peningkatan signifikan dalam mendeteksi artefak manipulasi yang tetap mempertahankan efisiensi komputasi.

Kata kunci— *Deepfake, XceptionNet, Digital Forensics, DFRWS*

1. PENDAHULUAN

Perkembangan teknologi kecerdasan buatan telah menghadirkan berbagai inovasi yang mengagumkan sekaligus mengkhawatirkan, salah satunya adalah teknologi *deepfake* [1] [2]. Kemampuan untuk memanipulasi gambar dan video secara digital dengan hasil yang sangat realistis telah menciptakan tantangan baru dalam aspek keamanan digital dan forensik [3]. Teknologi *deepfake* yang semakin canggih memungkinkan pembuatan konten manipulatif yang sulit dibedakan dari konten asli, bahkan oleh mata manusia sekalipun [4][5].

Beberapa penelitian sebelumnya telah mengusulkan berbagai pendekatan untuk mendeteksi *deepfake*, mulai dari analisis tekstur wajah hingga penggunaan model *deep learning* [6][7][8]. Pendekatan berbasis Convolutional Neural Network (CNN) telah menunjukkan hasil yang menjanjikan [9], namun masih menghadapi tantangan dalam hal akurasi dan efisiensi komputasi [10][11]. XceptionNet, sebagai arsitektur yang menggunakan *depthwise separable convolution*, telah terbukti efektif dalam tugas klasifikasi gambar [12][13], namun belum banyak dieksplorasi dalam konteks forensik digital dengan framework yang terstruktur.

Meskipun telah banyak penelitian tentang deteksi *deepfake*, masih terdapat kesenjangan dalam integrasi metode *deep learning* dengan framework forensik digital yang sistematis [14]. Penelitian ini mengusulkan pendekatan baru dengan mengintegrasikan model XceptionNet yang dioptimalkan dengan framework DFRWS untuk analisis forensik digital. Tujuan utama penelitian ini adalah mengembangkan sistem deteksi *deepfake* yang tidak hanya akurat dalam klasifikasi, tetapi juga memenuhi standar investigasi forensik digital yang dapat dipertanggungjawabkan secara ilmiah.

2. METODE PENELITIAN

Penelitian ini menggunakan pendekatan eksperimental dengan mengikuti framework DFRWS dan mengintegrasikannya dengan deep learning. Implementasi metode penelitian dibagi menjadi beberapa tahapan utama.

2.1 Dataset dan Preprocessing

Dataset yang digunakan dalam penelitian ini terdiri dari 2000 gambar yang terbagi menjadi dua kategori: gambar asli dan gambar deepfake. Dataset dibagi menjadi tiga bagian dengan proporsi 60% data training (1200 gambar), 20% data validasi (400 gambar), dan 20% data testing (400 gambar). Setiap gambar diproses melalui tahapan preprocessing yang meliputi resize ke ukuran 196x196 pixel, normalisasi nilai pixel, dan augmentasi data untuk set training. Sumber dataset: <https://www.kaggle.com/api/v1/datasets/download/manjilkarki/deepfake-and-real-images>

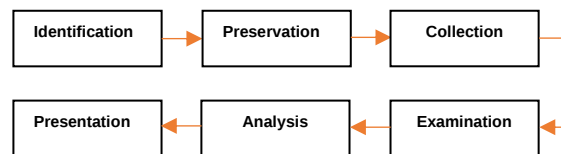
2.2 Arsitektur Model

Model yang digunakan adalah XceptionNet dengan modifikasi pada layer terakhir untuk tugas klasifikasi binary. Arsitektur model terdiri dari:

1. Base model XceptionNet dengan pre-trained weights dari ImageNet;
2. Global Average Pooling layer;
3. Dense layer dengan 1024 unit dan aktivasi ReLU;
4. Dropout layer dengan rate 0.5;
5. Output layer dengan aktivasi sigmoid.

2.3 Framework DFRWS

Implementasi framework DFRWS dalam penelitian ini dapat dilihat pada Gambar 1.



Gambar 1. Framework DFRWS [15]

Penjelasan pada Gambar 1 di atas adalah sebagai berikut:

1. Identification: Mengidentifikasi dan mengumpulkan dataset gambar
2. Preservation: Memastikan integritas data melalui proses preprocessing
3. Collection: Pengorganisasian dataset untuk training, validasi, dan testing
4. Examination: Penerapan model XceptionNet untuk analisis
5. Analysis: Evaluasi hasil deteksi dan pengukuran performa
6. Presentation: Dokumentasi hasil dan visualisasi performa model

2.4 Proses Training dan Evaluasi

Model dilatih menggunakan *optimizer Adam* dengan *learning rate* 0.00001. Proses training dilakukan selama 20 epoch dengan implementasi *early stopping* untuk mencegah overfitting. Evaluasi performa model menggunakan metrik *Accuracy*, *Precision*, *Recall*, *AUC score*, *Loss value*.

1. Accuracy

Akurasi merupakan proses perhitungan dengan membandingkan jumlah data yang diprediksi dengan tepat terhadap total keseluruhan data. Persamaan akurasi dapat ditunjukkan pada persamaan (1).

$$Accuracy = \frac{TP + TN}{N} \quad (1)$$

2. Precision

Presisi merupakan proses perhitungan dengan membandingkan jumlah data positif yang terklasifikasi dengan tepat terhadap total data yang terklasifikasi, baik positif maupun negatif. Rumus presisi dapat dilihat seperti ditunjukkan pada persamaan (2).

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

3. Recall

Recall merupakan proses perhitungan dengan membandingkan jumlah data positif yang berhasil diidentifikasi dengan tepat oleh sistem terhadap total keseluruhan data positif, baik yang terklasifikasi benar maupun salah. Rumus recall dapat dilihat seperti ditunjukkan pada persamaan (3).

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

4. AUC Score

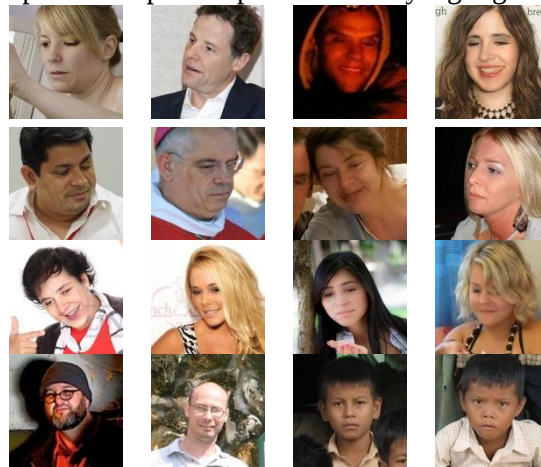
AUC Score merupakan perhitungan dari gambaran tentang keseluruhan pengukuran atas kesesuaian dari model yang digunakan seperti ditunjukkan pada persamaan (4).

$$AUC\ Score = \frac{1 + TP\ rate - FP\ rate}{2} \quad (4)$$

3. HASIL DAN PEMBAHASAN

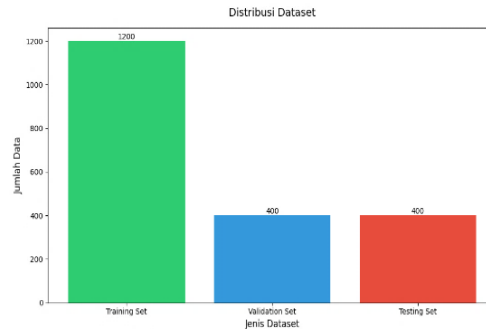
3.1 Identification

Pada tahap ini, dilakukan identifikasi terhadap karakteristik deepfake yang akan dideteksi menggunakan pendekatan berbasis deep learning, khususnya arsitektur XceptionNet. Dataset yang digunakan dalam penelitian ini bersumber dari situs open source Kaggle dengan judul "Deepfake and Real Images" yang dikumpulkan oleh Manjil Karki pada tahun 2021. Total dataset yang digunakan berjumlah 2000 gambar dengan pembagian yang seimbang antara gambar asli dan palsu. Gambar 1 merupakan sample Deepfake Dataset yang digunakan.



Gambar 1 Sampel dataset penelitian

Distribusi dataset dilakukan dengan proporsi 70-15-15: data training menggunakan 700 gambar asli dan 700 gambar palsu, data validasi menggunakan 150 gambar asli dan 150 gambar palsu, serta data testing yang terdiri dari 150 gambar asli dan 150 gambar palsu seperti yang ditunjukkan pada Gambar 2. Pembagian ini penting untuk memverifikasi kemampuan model dalam menggeneralisasi pada data yang belum pernah dilihat sebelumnya. Proses identifikasi juga mencakup penentuan preprocessing yang diperlukan, termasuk normalisasi citra, deteksi wajah, dan ekstraksi facial landmarks yang akan digunakan dalam tahap analisis selanjutnya.



Gambar 2 Distribusi dataset

3.2 Preservation

Pada tahap ini, dataset disimpan dan diorganisir dalam struktur folder yang sistematis di Google Drive dengan pembagian tiga folder utama yaitu Train, Validation, dan Test. Setiap folder utama memiliki dua subfolder yang terdiri dari 'real' untuk gambar asli dan 'fake' untuk gambar deepfake. Untuk memastikan integritas data, setiap gambar diberikan nilai hash menggunakan algoritma MD5 dan SHA-256. Proses hashing ini penting untuk memverifikasi bahwa tidak ada perubahan yang tidak diinginkan pada dataset selama proses penelitian, seperti ditunjukkan pada Tabel 1.

Tabel 1 Hash Dataset

Category	Image Name	SHA-256 Hash
Real	Real_100.jpg	d76e17aa744ad2579fa34e28f40f3621e04c0d5db4a5a92cb444a7fa48afc33e
Real	Real_10.jpg	f0cd2cecefa920afcda2032a9bcfca57d4529a11eb3e104f025f2ceb0baa7c
Real	Real_105.jpg	2d827821dc4bc62a588a86a26ebc49f7e311e162ffd2de263a28bf6a8ac5bf92
Fake	Fake_109.jpg	2361003b79d0a3382b23a83520d064e63b0fd22d33725d5f7166b284cbb14b8d
Fake	Fake_106.jpg	b044d877022eb22dc9c44c6607e33c7300b7f9458905bb125c829aa888bf7628

Berdasarkan Tabel 1, nilai hash (seperti d76e17... untuk 'Real_100.jpg') berfungsi sebagai "sidik jari digital" yang unik untuk setiap berkas gambar. Dengan mencatat nilai hash ini di awal, dapat menjamin integritas data di seluruh proses penelitian.

3.3 Collection

Tahap Pengumpulan dalam proses forensik digital untuk deepfake berfokus pada pengorganisasian dataset yang telah melalui pra-pemrosesan. Dataset tersebut, yang terdiri dari total 2000 gambar, dibagi menjadi tiga subset menggunakan metode `train_test_split` dengan proporsi 60-20-20. Hal ini menghasilkan himpunan data latih (training set) yang terdiri dari 1200 gambar (600 palsu dan 600 asli), himpunan data validasi (validation set) yang berisi 400 gambar (200 palsu dan 200 asli), dan himpunan data uji (testing set) yang mencakup 400 gambar (200 palsu dan 200 asli). Pemisahan dataset dilakukan secara acak dengan mempertahankan `random_state=42` untuk memastikan reproduktibilitas hasil eksperimen.

Teknik augmentasi data diterapkan pada data latih untuk memperkaya keragamannya dan meningkatkan kemampuan generalisasi model. Hal ini dicapai dengan menggunakan `ImageDataGenerator` dengan berbagai parameter transformasi, termasuk rotasi gambar hingga 20 derajat, pergeseran horizontal dan vertikal hingga 20%, transformasi shearing hingga 20%, rentang zoom hingga 20%, dan pembalikan horizontal. Proses augmentasi ini sangat penting untuk meningkatkan kemampuan model dalam menggeneralisasi data yang belum pernah dilihat dan mencegah overfitting (penyesuaian berlebihan), terutama karena model deep learning memerlukan data yang beragam untuk pembelajaran yang efektif. Dataset validasi dan uji tidak dikenai augmentasi untuk memastikan evaluasi yang objektif terhadap kinerja model.

3.4 Examination

Arsitektur model dirancang dengan memanfaatkan bobot pra-latih (pre-trained weights) dari ImageNet sebagai model dasar, yang selanjutnya dimodifikasi dengan menambahkan beberapa lapisan kustom (custom layers) untuk mengoptimalkan proses deteksi. Lapisan-lapisan ini terdiri dari lapisan Global Average Pooling untuk reduksi dimensi fitur, lapisan Dense dengan 1024 unit dan aktivasi ReLU untuk ekstraksi fitur yang kompleks, serta lapisan Dropout dengan laju 0.5 untuk mencegah overfitting (penyesuaian berlebih). Terakhir, sebuah lapisan keluaran (output layer) dengan aktivasi sigmoid ditambahkan untuk klasifikasi biner.

Fine-tuning diterapkan pada 30 lapisan terakhir XceptionNet agar model dapat beradaptasi dengan karakteristik spesifik dari dataset deepfake, sementara lapisan-lapisan awal tetap dibekukan (frozen) untuk mempertahankan fitur-fitur fundamental yang telah dipelajari dari ImageNet. Optimasi model memanfaatkan optimizer Adam dengan laju pembelajaran (learning rate) yang kecil (0.00001) untuk memastikan pembelajaran yang stabil, serta penerapan regularisasi L2 dengan nilai 0.01 pada lapisan dense untuk mencegah overfitting. Model dikompilasi dengan binary cross-entropy sebagai loss function (fungsi kerugian) dan menggunakan metrik-metrik seperti akurasi, presisi, recall, dan AUC untuk evaluasi kinerja yang komprehensif.

3.5 Analysis

Model dievaluasi menggunakan dataset pengujian yang telah disiapkan dan mencapai kinerja yang signifikan dengan akurasi 91.25%, presisi 88.73%, recall 94.50%, dan skor AUC 0.9710. Analisis ini juga mencakup evaluasi terhadap nilai loss (nilai kerugian) dan proses pelatihan model untuk memahami perilaku pembelajarannya (learning behavior) selama masa pelatihan. Kinerja model disajikan pada Tabel 2 dan 3.

Tabel 2 Kinerja Model pada Dataset Latih dan Validasi

Model Metrics	Training	Validation
Accuracy	95.54%	90.50%
Precision	94.93%	91.75%
Recall	96.15%	89.00%
AUC	0.9907	0.9654
Loss	9.3098	9.2581

Tabel 2 menyajikan perbandingan rinci antara metrik kinerja model pada dataset latih (training) dan dataset validasi (validation). Data pada kolom 'Training' menunjukkan bahwa model berhasil mempelajari pola dari data latih dengan sangat baik, yang dibuktikan dengan pencapaian akurasi 95.54%, presisi 94.93%, recall 96.15%, dan skor AUC 0.9907. Di sisi lain, kolom 'Validation' mengukur kemampuan generalisasi model terhadap data yang tidak terlihat selama pelatihan, di mana model mencapai akurasi 90.50% dan AUC 0.9654. Meskipun terdapat selisih (sekitar 5%) antara akurasi latih dan validasi, hal ini masih dalam batas wajar dan tidak mengindikasikan overfitting yang serius. Bukti terkuat dari generalisasi yang baik ini terlihat pada nilai Loss, di mana nilai Loss latih (9.3098) dan Loss validasi (9.2581) memiliki nilai yang sangat berdekatan. Perbedaan yang relatif kecil ini, yang juga tergambar pada grafik di Gambar 5 dan 6, menunjukkan bahwa model memiliki kemampuan generalisasi yang baik dan tidak hanya "menghafal" data latih.

Table 3. Kinerja Model pada Dataset Pengujian

Metrics	Value
Test Accuracy	91.25%
Test Precision	88.73%
Test Recall	94.50%
Test AUC	0.9710
Test Loss	9.6008

Tabel 3 menyajikan hasil evaluasi akhir yang merupakan tolok ukur kinerja model yang paling objektif, karena diuji menggunakan dataset pengujian yang sama sekali belum pernah dilihat oleh model selama proses pelatihan maupun validasi. Model mencapai kinerja yang signifikan dengan akurasi keseluruhan sebesar 91.25%, yang menunjukkan kemampuannya mengklasifikasikan 91.25% dari 400 data uji dengan benar. Lebih lanjut, nilai presisi 88.73% mengindikasikan tingginya kebenaran model saat memprediksi sebuah gambar sebagai deepfake.

Nilai recall yang tinggi sebesar 94.50% sangat penting dalam konteks ini, karena menunjukkan bahwa model berhasil mengidentifikasi 94.50% dari seluruh gambar deepfake yang ada di dalam dataset pengujian. Kinerja komprehensif ini dilengkapi dengan skor AUC 0.9710, yang menegaskan kemampuan diskriminatif (daya pembeda) model yang sangat baik antara kelas asli dan palsu. Hasil pada tabel ini, yang dianalisis lebih lanjut melalui confusion matrix pada Gambar 7, mengonfirmasi bahwa model XceptionNet yang dimodifikasi memiliki kemampuan generalisasi yang kuat dan efektif untuk deteksi deepfake pada data baru.

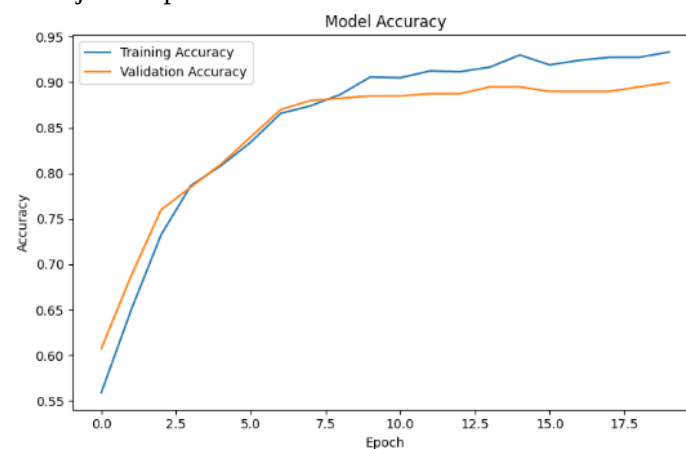
3.6 Presentation

Tahap ini menyajikan hasil analisis dan evaluasi model XceptionNet dalam format yang komprehensif dan mudah dipahami. Visualisasi hasil disajikan dalam bentuk grafik kinerja model dan confusion matrix (matriks kebingungan) untuk memberikan gambaran yang jelas mengenai kemampuan model dalam mendeteksi deepfake. Rancangan arsitektur model XceptionNet yang digunakan dalam penelitian ini ditunjukkan pada Gambar 4. Model ini dirancang dengan mempertimbangkan aspek-aspek penting dalam deteksi deepfake, mulai dari pra-pemrosesan data hingga proses klasifikasi akhir.

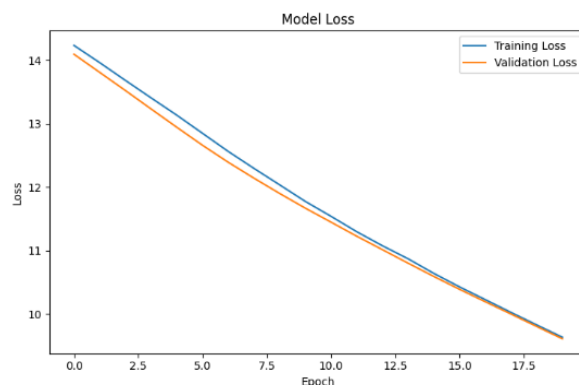
Layer (type)	Output Shape	Param #	Connected to
input_layer_3 (InputLayer)	(None, 196, 196, 3)	0	-
block1_conv1 (Conv2D)	(None, 97, 97, 32)	964	input_layer_3[][]
block1_conv1_bn (BatchNormalization)	(None, 97, 97, 32)	128	block1_conv1[][]
block1_conv1_act (Activation)	(None, 97, 97, 32)	0	block1_conv1_bn[][]
block1_conv2 (Conv2D)	(None, 97, 97, 64)	18,432	block1_conv1_act[][]
block1_conv2_bn (BatchNormalization)	(None, 97, 97, 64)	256	block1_conv2[][]
block1_conv2_act (Activation)	(None, 97, 97, 64)	0	block1_conv2_bn[][]
block2_sepconv1 (SeparableConv2D)	(None, 97, 97, 128)	6,768	block1_conv2_act[][]
block2_sepconv1_bn (BatchNormalization)	(None, 97, 97, 128)	512	block2_sepconv1[][]

Gambar 4 Arsitektur Model XceptionNet untuk Deteksi Deepfake

Hasil pelatihan model XceptionNet selama 20 epoch menunjukkan kemajuan yang signifikan dalam pembelajaran model. Proses pelatihan menunjukkan peningkatan kinerja yang stabil, dengan akurasi pelatihan (training accuracy) mencapai 95.54% dan akurasi validasi (validation accuracy) mencapai 90.50%. Meskipun terdapat selisih sekitar 5% antara akurasi pelatihan dan validasi, hal ini masih dalam rentang yang dapat diterima dan tidak mengindikasikan overfitting (penyesuaian berlebihan) yang serius. Nilai loss (nilai kerugian) pada akhir pelatihan menunjukkan 9.3098 untuk pelatihan dan 9.2581 untuk validasi, dengan perbedaan yang relatif kecil menunjukkan bahwa model memiliki kemampuan generalisasi yang baik, sebagaimana ditunjukkan pada Gambar 5 dan 6.

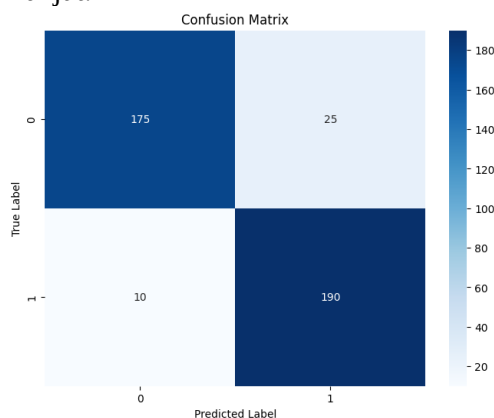


Gambar 5 Model Accuracy



Gambar 6 Model Loss

Kinerja model dievaluasi lebih lanjut dengan menganalisis confusion matrix (matriks kebingungan) untuk mendapatkan pemahaman yang lebih mendalam mengenai kapabilitas klasifikasi model, sebagaimana diilustrasikan pada Gambar 7. Confusion matrix menyajikan visualisasi yang jelas mengenai distribusi prediksi model dalam mengklasifikasikan gambar asli (genuine) dan deepfake. Dari total 400 gambar pengujian, model berhasil mengidentifikasi dengan benar mayoritas kasus, baik untuk gambar asli (True Negatives) maupun gambar deepfake (True Positives). Meskipun demikian, terdapat beberapa kasus misklasifikasi, yang ditunjukkan oleh False Positives dan False Negatives, yang menawarkan wawasan berharga untuk pengembangan model lebih lanjut.



Gambar 7 Confusion matrix

4. KESIMPULAN

Penelitian ini berfokus pada pengembangan sistem deteksi deepfake menggunakan model XceptionNet yang dimodifikasi, yang diimplementasikan dengan TensorFlow dan Keras. Dataset yang digunakan terdiri dari gambar wajah asli (genuine) dan palsu (deepfake), yang telah dibagi menjadi himpunan data latih (training), validasi (validation), dan uji (testing). Teknik augmentasi data diterapkan untuk meningkatkan keragaman data latih dan mencegah overfitting (penyesuaian berlebih). Model yang dikembangkan telah diuji dan dievaluasi menggunakan metrik akurasi, presisi, recall, dan AUC. Hasil evaluasi menunjukkan bahwa model XceptionNet yang dimodifikasi mencapai kinerja yang baik dalam mendeteksi deepfake. Pemanfaatan teknik regularisasi dan fine-tuning (penalaan halus) secara efektif meningkatkan akurasi dan kemampuan generalisasi model. Meskipun demikian, masih terdapat ruang untuk peningkatan lebih lanjut, seperti mengeksplorasi arsitektur model yang lebih kompleks, mengoptimalkan hyperparameter, serta menggunakan dataset yang lebih besar dan lebih beragam. Proyek ini berkontribusi pada pengembangan sistem deteksi deepfake yang dapat membantu memerangi penyebaran misinformasi dan melindungi integritas media digital.

DAFTAR PUSTAKA

- [1] A. Saxena et al., "Detecting Deepfakes: A Novel Framework Employing XceptionNet-Based Convolutional Neural Networks," *Trait. du Signal*, vol. 40, no. 3, pp. 835–846, 2023, doi: 10.18280/ts.400301.
- [2] M. Karaköse, İ. İlhan, H. Yetiş, and S. Ataş, "A New Approach for Deepfake Detection with the Choquet Fuzzy Integral," *Appl. Sci.*, vol. 14, no. 16, 2024, doi: 10.3390/app14167216.
- [3] L. Verdoliva, "Media Forensics and DeepFakes: An Overview," *IEEE J. Sel. Top. Signal Process.*, vol. 14, no. 5, pp. 910–932, 2020, doi: 10.1109/JSTSP.2020.3002101.
- [4] M. Elpeltagy, A. Ismail, M. S. Zaki, and K. Eldahshan, "A Novel Smart Deepfake Video Detection System," *Int. J. Adv. Comput. Sci. Appl.*, vol. 14, no. 1, pp. 407–419, 2023, doi: 10.14569/IJACSA.2023.0140144.
- [5] D. Gong, Y. J. Kumar, O. S. G. Z. Ye, and W. Chi, "DeepfakeNet, an Efficient Deepfake Detection Method," *Int. J. Adv. Comput. Sci. Appl.*, vol. 12, no. 6, pp. 201–207, 2021, doi: 10.14569/IJACSA.2021.0120622.
- [6] Z. Guo, G. Yang, J. Chen, and X. Sun, "Fake face detection via adaptive manipulation traces extraction network," *Comput. Vis. Image Underst.*, vol. 204, no. January, p. 103170, 2021, doi: 10.1016/j.cviu.2021.103170.
- [7] R. Caldelli, L. Galteri, I. Amerini, and A. Del Bimbo, "Optical Flow based CNN for detection of unlearned deepfake manipulations," *Pattern Recognit. Lett.*, vol. 146, pp. 31–37, 2021, doi: 10.1016/j.patrec.2021.03.005.
- [8] A. Kohli and A. Gupta, "Detecting DeepFake, FaceSwap and Face2Face facial forgeries using frequency CNN," *Multimed. Tools Appl.*, vol. 80, no. 12, pp. 18461–18478, 2021, doi: 10.1007/s11042-020-10420-8.
- [9] J. Liu, K. Zhu, W. Lu, X. Luo, and X. Zhao, "A lightweight 3D convolutional neural network for deepfake detection," *Int. J. Intell. Syst.*, vol. 36, no. 9, pp. 4990–5004, 2021, doi: 10.1002/int.22499.
- [10] Y. Patel et al., "An Improved Dense CNN Architecture for Deepfake Image Detection," *IEEE Access*, vol. 11, no. January, pp. 22081–22095, 2023, doi: 10.1109/ACCESS.2023.3251417.
- [11] O. A. H. H. Al-Dulaimi and S. Kurnaz, "A Hybrid CNN-LSTM Approach for Precision Deepfake Image Detection Based on Transfer Learning," *Electronics*, vol. 13, no. 9, 2024, doi: 10.3390/electronics13091662.
- [12] A. Chintha, A. Rao, S. Sohrawardi, K. Bhatt, M. Wright, and R. Ptucha, "Leveraging edges and optical flow on faces for deepfake detection," *IJCB 2020 - IEEE/IAPR Int. Jt. Conf. Biometrics*, 2020, doi: 10.1109/IJCB48548.2020.9304936.
- [13] A. Saxena et al., "Detecting Deepfakes: A Novel Framework Employing XceptionNet-Based Convolutional Neural Networks," *Trait. du Signal*, vol. 40, no. 3, pp. 835–846, 2023, doi: 10.18280/ts.400301.
- [14] T. M. Jawad Abbas and A. S. Abdulmajeed, "Identifying digital forensic frameworks based on processes models," *Iraqi J. Sci.*, vol. 2021, pp. 249–258, 2021, doi: 10.24996/ij.2021.SI.1.35.
- [15] Pollitt MM. 2007. An ad hoc review of digital forensic models. Conference: Second .IEEE International Workshop on Systematic Approaches to Digital Forensic Engineering, SADFE 2007, 12-Seattle, Washington, USA, April 10, 2007.