

BAB I PENDAHULUAN

1.1 Latar Belakang

Dalam beberapa tahun terakhir, perkembangan teknologi kecerdasan buatan *Artificial Intelligence (AI)* telah memungkinkan terciptanya gambar-gambar sintetis atau *deepfake* yang sangat realistis. Teknologi ini memberikan manfaat signifikan di berbagai bidang, seperti hiburan, seni digital, dan aplikasi kreatif lainnya. Namun, kemajuan ini juga memunculkan risiko besar, terutama dalam hal penyalahgunaan untuk tujuan yang merugikan, seperti penipuan, penyebaran informasi palsu, dan perundungan siber. Penyalahgunaan teknologi *deepfake* telah menjadi ancaman serius bagi masyarakat, khususnya generasi muda yang kerap terpapar dunia digital, seperti yang disampaikan oleh penelitian tentang pentingnya sosialisasi bahaya *deepfake* kepada pelajar [1]. Metode sosialisasi interaktif terbukti efektif meningkatkan pemahaman mereka, dengan peningkatan signifikan dalam hasil post-test sebesar 55%, yang menunjukkan keberhasilan dalam upaya membangun kesadaran terhadap bahaya teknologi ini. Hal ini menegaskan kebutuhan mendesak untuk mengedukasi masyarakat dalam membedakan antara gambar asli dan gambar buatan *AI*, terutama di sektor media, hukum, dan keamanan digital [2]. Secara khusus, arsitektur *DenseNet* memiliki sejumlah keunggulan penting yang menjadikannya unggul dalam tugas klasifikasi gambar, seperti kemampuan untuk mengurangi masalah vanishing gradient, memperkuat propagasi fitur, mendorong penggunaan ulang fitur dari lapisan sebelumnya, serta secara signifikan mengurangi jumlah parameter yang dibutuhkan [3]. Di sisi lain, *ResNet-50* menawarkan keseimbangan yang ideal antara kedalaman jaringan dan efisiensi komputasi melalui pendekatan residual learning, yang mempermudah proses optimasi jaringan dalam tanpa mengalami degradasi performa. Keunggulan ini semakin diperkuat oleh keberhasilan *ResNet* dalam berbagai kompetisi pengenalan citra berskala besar [4]. Berdasarkan hal tersebut, pemanfaatan *DenseNet-201* dan

ResNet-50 dinilai sebagai pilihan yang praktis dan efisien, terutama dalam konteks penelitian deteksi gambar buatan AI yang memerlukan akurasi tinggi dengan sumber daya komputasi yang terbatas.

Dengan memanfaatkan model *deep learning* seperti *DenseNet*, pengolahan data visual dapat dilakukan secara lebih efisien dan akurat [5]. *DenseNet* merupakan salah satu arsitektur *deep learning* yang telah terbukti efektif dalam berbagai tugas klasifikasi gambar, karena kemampuannya dalam memanfaatkan fitur dari seluruh lapisan jaringan. Dalam penelitian ini, *DenseNet* digunakan untuk mengidentifikasi gambar asli dan gambar buatan AI. Penelitian ini juga melibatkan langkah-langkah penting seperti *preprocessing* gambar (*resize*, *histogram equalization*, dan *augmentasi*), penyeimbangan data untuk menghindari bias, serta penerapan teknik seperti *early stopping* guna mengoptimalkan proses pelatihan [6]-[8]. Dengan pendekatan ini, diharapkan model dapat mencapai akurasi yang tinggi dalam membedakan gambar asli dan gambar buatan AI. Studi ini tidak hanya berkontribusi terhadap pengembangan teknologi deteksi *deepfake*, tetapi juga memberikan solusi praktis untuk mengatasi potensi penyalahgunaan gambar sintetis. Dengan demikian, penelitian ini diharapkan dapat menjadi pijakan awal untuk studi lanjutan dalam bidang deteksi gambar berbasis kecerdasan buatan, sekaligus mendukung pengembangan teknologi yang lebih bertanggung jawab.

1.2 Rumusan Masalah

Berdasar latar belakang yang telah dibuat diatas rumusan masalah diambil sebagai berikut:

1. Bagaimana model *DenseNet-201* dengan bobot *pre-trained* dapat diterapkan secara efektif untuk membedakan gambar asli dan gambar buatan AI?
2. Sejauh mana teknik augmentasi data seperti flipping, rotation, histogram equalization, dan Gaussian blur dapat meningkatkan akurasi dan kinerja

model klasifikasi?

3. Bagaimana penerapan teknik early stopping dapat mempengaruhi stabilitas pelatihan model serta meningkatkan kinerja validasi pada proses klasifikasi?

1.3 Batasan Masalah

Dalam penelitian Implementasi *DenseNet-201* dan *ResNet-50* untuk Identifikasi Gambar Asli dan Buatan AI ini diberikan beberapa batasan masalah, tujuannya agar pembahasan menjadi lebih terperinci dan tidak melebar. Antara lain:

1. Dataset yang digunakan hanya terdiri dari gambar asli dan buatan AI (deepfake), tanpa mencakup video atau jenis media lainnya.
2. Data dikumpulkan dari beberapa lokasi yaitu website kaggle.com serta gambar ai (FAKE) yang dihasilkan menggunakan *AI Image Generator* seperti *pixlr.com*, *dreamstudio.ai*, dan dataset tambahan dari kaggle.com.
3. Penelitian ini hanya membandingkan model *DenseNet-201* dan *ResNet-50*.
4. Evaluasi performa model hanya dilakukan berdasarkan metrik akurasi, presisi, recall, dan F1-score.
5. Penelitian ini tidak mencakup perancangan atau penerapan sistem yang dapat digunakan secara luas oleh publik.

1.4 Tujuan Penelitian

Penelitian ini mempunyai tujuan beberapa hal diantaranya sebagai berikut:

1. Untuk menginvestigasi metode yang tepat untuk mendeteksi gambar ai dan asli
2. Untuk membandingkan performa *DenseNet-201* dan *ResNet-50* pada klasifikasi gambar ai dan gambar asli.

1.5 Manfaat Penelitian

Penelitian ini diharapkan bisa memberikan manfaat kedepannya, baik secara teoritis maupun praktis, sebagai berikut:

1. Manfaat Teoritis

Manfaat teoritis, penelitian ini berkontribusi dalam pengembangan ilmu pengetahuan dengan menyediakan referensi bagi penelitian lebih lanjut terkait model deep learning untuk deteksi *deepfake*. Selain itu, penelitian ini menambah wawasan tentang efektivitas arsitektur *DenseNet-201* dalam klasifikasi gambar berbasis *deepfake* serta menguji penerapan teknik preprocessing seperti *gaussian Blur* dan *histogram equalization* dalam meningkatkan performa model. Analisis mendalam terhadap matrik evaluasi seperti akurasi, presisi, recall, dan F1-score juga memberikan pemahaman yang lebih luas mengenai kemampuan model dalam membedakan gambar asli dan gambar *ai*.

2. Manfaat Praktis

Manfaat praktis, penelitian ini memberikan manfaat bagi berbagai pihak, terutama organisasi dan individu yang memerlukan sistem deteksi *deepfake*. Lembaga keamanan siber dapat memanfaatkan penelitian ini untuk mengembangkan sistem deteksi yang lebih akurat dalam mengidentifikasi ancaman seperti penipuan digital, penyebaran berita palsu, dan pencurian identitas. Selain itu, industri media dan jurnalisme dapat menggunakan hasil penelitian ini untuk memverifikasi keaslian gambar guna mencegah penyebaran informasi palsu. Bagi pengguna internet dan masyarakat umum, penelitian ini berkontribusi dalam meningkatkan literasi digital serta kewaspadaan terhadap manipulasi gambar berbasis *AI*.

1.6 Sistematika Penulisan

Sistematika penulisan bertujuan untuk mempermudah pemahaman

mengenai penelitian ini. Berikut akan dijelaskan dalam beberapa bab:

BAB I PENDAHULUAN

Pada bab ini terdiri atas latar belakang, rumusan masalah, batasan masalah, tujuan penelitian, manfaat penelitian, tujuan penelitian, dan sistematika penulisan.

BAB II TINJAUAN PUSTAKA

Tinjauan pustaka dalam penelitian ini menyajikan ringkasan teori-teori dan pandangan para ahli yang relevan dengan topik yang diteliti. Bab ini mencakup kajian literatur atau jurnal penelitian sebelumnya yang mendukung penelitian, serta membahas landasan teori yang berkaitan dengan topik penelitian.

BAB III METODE PENELITIAN

Pada bab ini membahas terkait objek yang diteliti, teknik dari pengumpulan data dan juga langkah – langkah maupun metode yang telah diuraikan secara detail dipergunakan untuk memecahkan masalah yang terdapat pada penelitian ini.²⁴

BAB IV HASIL DAN PEMBAHASAN

Bab ini menyajikan hasil analisis sentimen yang diperoleh melalui pengolahan data. Selain itu, pembahasan mencakup perbandingan kinerja antara penerapan algoritma dengan dan tanpa metode SMOTE dan K-Fold Cross Validation, serta evaluasi metrik akurasi, presisi, dan recall yang diperoleh dari penerapan algoritma.

BAB V PENUTUP

Bab ini berisi kesimpulan yang dirumuskan berdasarkan hasil penelitian yang telah dilakukan, serta rekomendasi yang dapat diberikan oleh peneliti untuk pengembangan penelitian lebih lanjut di masa mendatang.