

BAB V

KESIMPULAN DAN SARAN

5.1 Kesimpulan

Berikut merupakan kesimpulan dari hasil dan temuan dari penelitian yang telah dilakukan:

- 1) Pengembangan model deteksi emosi dan *arousal-valence* menggunakan model EfficientNetB0 dan EfficientNetB7 memiliki hasil model terbaik pada model EfficientNetB0 dengan metode *fine-tuning* dengan *learning rate* 0.00001 tanpa *image augmentation*. Akurasi yang didapatkan dari model terbaik tersebut memiliki akurasi pada set pengujian untuk *output* klasifikasi emosi sebesar 56%, untuk *output* nilai valensi MAE 0.2889, MSE 0.1333, RMSE 0.3651, dan untuk *output* nilai arousal MAE 0.2795, MSE 0.1244, RMSE 0.3527.
- 2) Hasil perbandingan kinerja model deteksi emosi dan *arousal-valence* menggunakan model EfficientNetB0 dan EfficientNetB7 menunjukkan bahwa meskipun EfficientNetB7 mampu mencapai akurasi dan *f1-score* yang lebih tinggi dalam klasifikasi emosi, serta menghasilkan nilai *RMSE* yang lebih rendah pada *output* linier nilai valensi dan *arousal*, model ini cenderung *overfit* setelah iterasi ke-25. Selain itu, EfficientNetB7 memiliki nilai *loss* yang lebih tinggi jika dibandingkan dengan EfficientNetB0. Sebaliknya, EfficientNetB0 menunjukkan kinerja yang lebih stabil dan konsisten sepanjang pelatihan, dengan *loss* yang lebih rendah dan riwayat pelatihan yang lebih normal. Oleh karena itu, meskipun EfficientNetB7 memiliki keunggulan dalam beberapa metrik, EfficientNetB0 lebih stabil dan lebih cocok untuk penggunaan jangka panjang dalam deteksi emosi dan *arousal-valence* tanpa risiko *overfitting*.
- 3) Hasil implementasi model terbaik pada API dengan penerapan ke layanan *cloud* memiliki hasil akurasi 99.55% akurat dengan hasil yang didapatkan sebelum dilakukan implementasi. Pengujian layanan API di *cloud* juga dilakukan dengan hasil waktu *request* rata-rata 281.61 milidetik, waktu *request* terkecil 166 milidetik, total *request* sebanyak 21.728, *request* per

detik yang tidak memiliki perbandingan jauh yaitu sebanyak 28.1 *request* dari total 35 pengguna virtual dengan setiap pengguna mengakses API di *cloud* bersamaan setiap satu detik sekali, dan waktu respons 220 milidetik di persentil ke-50 yang berarti setengah dari semua *request* yang dilakukan oleh pengguna virtual selesai dalam waktu yang lebih singkat dari nilai tersebut atau median dan setengah lainnya memerlukan waktu lebih lama.

5.2 Saran

Berikut merupakan saran dan rekomendasi dari penelitian yang telah dilakukan:

- 1) Pengembangan model deteksi emosi dan *arousal-valence* pada penelitian ini masih belum mencapai akurasi yang maksimal, peningkatan akurasi model masih diperlukan dengan teknik pengolahan data dan pengembangan model yang berbeda.
- 2) Penerapan atau implementasi model deteksi emosi dan *arousal-valence* pada penelitian ini dilakukan pada API yang diterapkan ke layanan *cloud* Google Cloud. Saran untuk penelitian selanjutnya untuk dapat menerapkan hasil implementasi ke aplikasi untuk digunakan pada aplikasi nyata.
- 3) Menerapkan teknik-teknik untuk otomatisasi *pipeline machine learning* guna melakukan monitor, mengotomatiskan proses pengembangan, pelatihan, pengujian, dan *deployment* model *machine learning*. Dengan demikian, setiap perubahan pada model dapat diuji dan di-*deploy* secara otomatis, mempercepat proses waktu produksi, dan memastikan model yang dihasilkan selalu beradaptasi dengan perubahan data serta dalam kondisi optimal, dan dapat dipantau secara *real-time*.