

BAB V

PENUTUP

5.1 Kesimpulan

Berdasarkan hasil evaluasi performa dari tiga algoritma klasifikasi yang diterapkan dalam analisis sentimen, yaitu Naive Bayes, Support Vector Machine (SVM), dan Random Forest, dapat disimpulkan bahwa masing-masing model memiliki keunggulan dan keterbatasan yang berbeda. Model Naive Bayes menunjukkan performa terbaik secara keseluruhan dengan akurasi mencapai 88%, serta nilai presisi, *recall*, dan *F1-score* yang seimbang pada kisaran 88–89%. Hal ini menunjukkan bahwa Naive Bayes sangat efektif dalam memprediksi ketiga kelas sentimen, khususnya pada *dataset* yang relatif kecil dan seimbang. Di sisi lain, Random Forest menempati posisi kedua dengan akurasi 87% dan metrik evaluasi lain yang juga cukup tinggi dan konsisten. Keunggulan Random Forest terletak pada kemampuannya dalam menangkap pola-pola kompleks dan korelasi antar fitur, sehingga lebih stabil pada data yang memiliki kompleksitas lebih tinggi. Sementara itu, SVM meskipun memiliki performa sedikit lebih rendah dengan akurasi 84%, tetap menunjukkan hasil yang dapat diterima, khususnya pada aspek presisi. Namun, SVM lebih kompleks dan membutuhkan pengaturan parameter yang cermat agar performanya optimal.

Secara komprehensif, hasil ini menunjukkan bahwa Naive Bayes adalah pilihan utama untuk analisis sentimen dalam konteks *dataset* yang digunakan, terutama karena kemudahan dan efisiensinya, serta hasil klasifikasi yang seimbang. Jika dihadapkan pada tugas dengan pola data yang lebih rumit dan memerlukan kestabilan lebih, Random Forest bisa menjadi alternatif yang kuat. Sementara SVM, dengan tingkat kompleksitas dan kebutuhan tuning yang lebih tinggi, masih dapat digunakan dengan optimasi yang tepat. Pemilihan model sebaiknya disesuaikan dengan karakteristik data dan tujuan analisis.

5.2 Saran

Berdasarkan temuan dan pengalaman selama penelitian, beberapa saran dapat diberikan untuk pengembangan dan penelitian selanjutnya. Pertama, disarankan untuk memperbesar dan memperkaya *dataset* agar model-model yang digunakan dapat lebih baik dalam melakukan generalisasi terhadap data yang baru dan beragam, serta dapat meningkatkan performa khususnya pada model seperti SVM yang cenderung sensitif terhadap ukuran data. Kedua, optimasi model berupa *tuning hyperparameter* harus dilakukan secara lebih mendalam, terutama untuk model SVM dan Random Forest, dengan menggunakan teknik seperti *Grid Search* atau *Random Search* guna memperoleh konfigurasi yang optimal. Ketiga, eksplorasi terhadap model lain, khususnya yang berbasis *deep learning* seperti LSTM atau *transformer*, dapat dilakukan untuk mengatasi pola dan konteks teks yang lebih kompleks. Keempat, pengembangan fitur atau teknik representasi teks yang lebih canggih, seperti *feature engineering* atau penggunaan *word embeddings*, dapat meningkatkan kemampuan model dalam memahami nuansa dan makna teks. Kelima, evaluasi dan validasi model sebaiknya dilakukan secara menyeluruh menggunakan teknik seperti *cross-validation* serta pengujian pada dataset berbeda agar diperoleh kestabilan dan kemampuan generalisasi yang lebih baik. Terakhir, disarankan untuk mengimplementasikan model terbaik ke dalam aplikasi nyata, misalnya sistem analisis sentimen *online* atau *chatbot*, sehingga dapat memberikan manfaat praktis dan diuji secara langsung di lingkungan pengguna.