

BAB V

PENUTUP

5.1. Kesimpulan

1. Berdasarkan hasil penelitian yang dilakukan, dapat disimpulkan bahwa metode *Naive Bayes* mampu melakukan klasifikasi tingkat kepuasan pelanggan dengan sangat baik. Proses klasifikasi diawali dengan tahap pengumpulan dan pemrosesan data, kemudian dilanjutkan dengan perancangan model menggunakan komponen *File*, *Data Table*, *Naive Bayes*, *Test and Score*, dan *Confusion Matrix* pada aplikasi Orange. Model ini dievaluasi untuk melihat seberapa baik performanya dalam memprediksi kategori kepuasan pelanggan, apakah termasuk kategori “Tinggi” atau “Rendah”.
2. Hasil evaluasi model menunjukkan bahwa algoritma *Naive Bayes* memiliki tingkat akurasi, presisi, dan recall sebesar 100%. Hal ini menunjukkan bahwa model mampu mengklasifikasikan semua data dengan benar tanpa kesalahan prediksi. *Confusion Matrix* menunjukkan bahwa dari 100 data uji, seluruh data diklasifikasikan dengan tepat ke dalam kelas yang sesuai. Tidak ditemukan data yang diklasifikasikan secara salah, yang menandakan bahwa data yang digunakan sangat cocok dengan karakteristik algoritma *Naive Bayes*.
3. Dengan hasil tersebut, dapat disimpulkan bahwa metode *Naive Bayes* sangat efektif dan akurat untuk digunakan dalam klasifikasi tingkat kepuasan pelanggan pada penelitian ini. Meskipun demikian, hasil ini perlu diuji

kembali dengan data yang lebih beragam atau skenario dunia nyata untuk mengetahui kestabilan model secara umum. Penggunaan metode ini dapat menjadi solusi sederhana namun kuat untuk analisis klasifikasi berbasis data kategorikal di berbagai bidang penelitian.

5.2. Saran

Berdasarkan hasil klasifikasi yang menunjukkan tingkat akurasi yang sangat tinggi, peneliti menyarankan agar metode *Naive Bayes* dapat diterapkan secara lebih luas pada jenis data serupa, terutama yang memiliki karakteristik kategorikal. Metode ini sangat cocok digunakan oleh peneliti pemula atau praktisi yang membutuhkan hasil klasifikasi yang cepat dan akurat dengan proses yang sederhana. Aplikasi Orange sebagai alat bantu visual juga sangat mendukung proses pemodelan yang intuitif tanpa memerlukan banyak kode pemrograman.

Namun demikian, untuk penelitian selanjutnya, disarankan agar dilakukan pengujian model dengan data yang lebih bervariasi dan jumlah yang lebih besar. Hal ini bertujuan untuk menghindari hasil yang terlalu sempurna yang mungkin disebabkan oleh data yang terlalu ideal atau homogen. Penggunaan data yang kompleks dan tidak seimbang dapat membantu menguji sejauh mana ketahanan model *Naive Bayes* dalam menghadapi tantangan nyata pada sistem klasifikasi.

Selain itu, peneliti juga disarankan untuk melakukan perbandingan metode *Naive Bayes* dengan algoritma klasifikasi lainnya seperti *Decision Tree*, *K-Nearest Neighbor*, atau *Support Vector Machine* guna mengetahui algoritma mana yang paling optimal untuk masalah klasifikasi tertentu.