

## **BAB V**

### **PENUTUP**

#### **5.1 KESIMPULAN**

Berdasarkan penelitian yang dilaksanakan mengenai klasifikasi tumor otak dengan metode *Support Vector Machine* (SVM) dan *K-Nearest Neighbors* (KNN), dapat disimpulkan bahwa:

1. Berdasarkan hasil pengujian, algoritma KNN menunjukkan performa yang lebih stabil dan konsisten dalam seluruh skenario evaluasi. Pada pembagian data 80:20, KNN menunjukkan kemampuan klasifikasi yang optimal pada data latih, sementara pada pembagian 70:30, KNN terbukti lebih unggul dibandingkan SVM dalam klasifikasi data uji, karena mampu memberikan hasil yang lebih baik pada hampir seluruh metrik evaluasi, termasuk kesalahan prediksi negatif (FPR) yang lebih rendah.
2. Hasil keseluruhan dari penelitian ini menunjukkan bahwa algoritma *K-Nearest Neighbor* (KNN) memiliki kinerja yang lebih unggul dibandingkan *Support Vector Machine* (SVM) dalam mengidentifikasi penyakit tumor otak. KNN mencatatkan nilai yang lebih tinggi pada aspek akurasi, presisi, recall, dan F1-score, serta menunjukkan kestabilan model yang lebih baik pada kedua skenario pembagian data yang diuji.
3. Selain membandingkan algoritma SVM dan KNN, penelitian ini juga mengevaluasi dampak penggunaan feature selection (FS) terhadap performa model. Hasil menunjukkan bahwa algoritma KNN tanpa FS

memberikan performa paling optimal, khususnya pada data latih dengan skenario pembagian 80:20. Sementara itu, penggunaan FS tidak menunjukkan peningkatan performa yang signifikan baik pada algoritma KNN maupun SVM. Dengan demikian, pada konteks dataset yang digunakan, penghapusan fitur tidak selalu menjamin peningkatan akurasi atau efektivitas model, sehingga pemilihan fitur perlu disesuaikan secara lebih selektif berdasarkan analisis karakteristik data.

## 5.2 SARAN

Berdasarkan hasil penelitian yang dilakukan, dapat diberikan beberapa saran sebagai berikut:

1. Penting untuk menggunakan dataset yang lengkap, representatif, dan relevan terhadap permasalahan yang ingin diselesaikan. Dataset yang berkualitas akan membantu mengurangi potensi bias dalam proses pelatihan model dan meningkatkan kemampuan generalisasi model terhadap data baru. Selain itu, data yang memiliki distribusi kelas yang seimbang serta jumlah atribut yang bermakna dapat meningkatkan performa algoritma klasifikasi.
2. Pemilihan fitur (feature selection) yang tepat dapat meningkatkan efisiensi dan efektivitas model. Meskipun dalam penelitian ini algoritma KNN menunjukkan performa tinggi, baik dengan feature selection maupun tanpa feature selection, namun penerapan feature selection tetap disarankan untuk mengurangi kompleksitas model dan waktu komputasi, serta untuk menghindari overfitting akibat fitur yang tidak relevan.

3. Disarankan untuk menguji kembali data atau penelitian ini dengan algoritma lain, seperti Random Forest, Naïve bayes, Decetion Tree ,atau metode berbasis *deep learning* untuk meningkatkan akurasi dan stabilitas hasil.