

BAB V

PENUTUP

5.1. Kesimpulan

Berdasarkan perancangan dan implementasi dari BAB sebelumnya, maka penelitian ini dapat diambil kesimpulan sebagai berikut:

1. Metode *pruning* dan *vector quantization* berhasil mengurangi ukuran model AlexNet dan SqueezeNet tanpa mengurangi akurasi secara signifikan terutama untuk ukuran data yang relatif kecil dan tidak banyak kelas.
2. Semakin *sparse* bobot pada model maka akurasi model juga akan semakin berkurang dan ukuran memori pada model setelah di kompres juga akan menjadi lebih kecil.
3. K-means *clustering* dapat digunakan untuk melakukan *vector quantization* pada bobot yang ada pada tiap layer dalam model.
4. Semakin kecil jumlah bits untuk *vector quantization* maka semakin kecil juga ukuran model setelah di kompres tetapi akurasi model juga akan semakin berkurang.

5.2. Saran

Sistem ini masih memiliki banyak kekurangan. Adapun saran yang dapat dijadikan sebagai acuan untuk penelitian selanjutnya adalah:

1. Melakukan percobaan menggunakan metode ini untuk data yang lain atau data yang lebih besar dan lebih banyak kelas.

2. Melakukan percobaan *pruning* dengan metode lain misalnya: *structured pruning*, *magnitude pruning*.
3. Melakukan percobaan *vector quantization* dengan menggunakan berbagai macam jumlah bits dan juga menggunakan metode lain untuk melakukan *clustering* dan melihat pengaruhnya terhadap model.
4. Melakukan optimisasi untuk kecepatan training data di karenakan pada penelitian ini memerlukan waktu training data yang cukup lama.