

Abstrak

Deteksi objek menggunakan *deep learning* menghadapi tantangan signifikan pada skenario dunia nyata dengan objek teroklusi. Arsitektur YOLOv8n yang populer untuk deteksi *real-time* dapat ditingkatkan kinerjanya menggunakan modul perhatian (*attention*). Penelitian ini bertujuan mengevaluasi pengaruh lokasi penyisipan dua modul perhatian, CBAM (*Convolutional Block Attention Module*) dan CoordAtt (*Coordinate Attention*), pada kinerja YOLOv8n dalam mendeteksi objek teroklusi. Fokus evaluasi adalah pada dataset KITTI yang dimodifikasi untuk merepresentasikan tingkat oklusi berbeda pada enam kelas objek gabungan (kendaraan dan pejalan kaki). Permasalahan utama adalah bagaimana lokasi penyisipan (*Neck*, *Backbone*, atau Keduanya) dan jenis modul perhatian (CBAM vs CoordAtt) mempengaruhi kemampuan YOLOv8n mengatasi oklusi. Hasil eksperimen secara konsisten menunjukkan penyisipan modul perhatian pada bagian *Neck* memberikan peningkatan kinerja paling signifikan. Konfigurasi CBAM-*Neck* mencapai mAP@0.5 tertinggi (0.683) dan mAP@0.5:0.95 (0.476), menunjukkan peningkatan substansial dari *baseline* (0.659 dan 0.449). Analisis kualitatif mengonfirmasi kemampuan model dengan perhatian di *Neck* untuk lebih baik mendeteksi objek teroklusi sedang, meskipun oklusi berat tetap menjadi tantangan. Penyisipan pada *Backbone* atau *Both* terbukti kurang efektif.

Kata kunci : YOLOv8, CBAM, CoordAtt, deteksi objek, oklusi, perhatian (*attention*), KITTI.