

Depth Coefficients for Depth Completion: Supplementary Material

Anonymous CVPR submission

Paper ID 6770

1. Estimated Depth-Map and Point-Cloud in KITTI Dataset

In this supplementary section, we show some more visual examples of predicted depth-maps and point-clouds at different scenes for comparison with various methods. We added the point-cloud of ground-truth with these visual examples along with the 3D bounding boxes of objects on top of the 3D point-clouds generated by various methods.

Figure 1 show the visual example of a complex scene where several cars are parked in parallel and a pedestrian is walking very closely by one of the cars (recognized in the color image and the 3D bounding boxes of the objects in the ground-truth). Both the predicted depth-maps and 3D point-clouds show that ours with 3 coefficient can recover the shapes of the cars and the pedestrians reasonably well. The point-cloud estimated by using 3-coeff. also matches closely with ground-truth point-cloud generated by KITTI.

Figure 2 show the visual example of a challenging scene where a pedestrian (obscured by shadows and low-light) is walking on the pavement close to a building (recognized by the color image and the 3D bounding boxes of the pedestrian in the ground-truth). This time, our method with 3 coefficient is creating more artifacts (since there are much less cues from color image to build on), but is doing reasonably well compared to other methods in order to recover the shapes of the pedestrian. The point-cloud estimated by using 3-coeff. also has more artifacts and distortion in this case, but still was able to separate itself out from the side-by building compared to other methods, and they still appear to be closer to the ground-truth point-cloud.

References

- [1] F. Ma, G. V. Cavalheiro, and S. Karaman. Self-supervised sparse-to-dense: Self-supervised depth completion from lidar and monocular camera. *arXiv preprint arXiv:1807.00275*, 2018. 2, 3

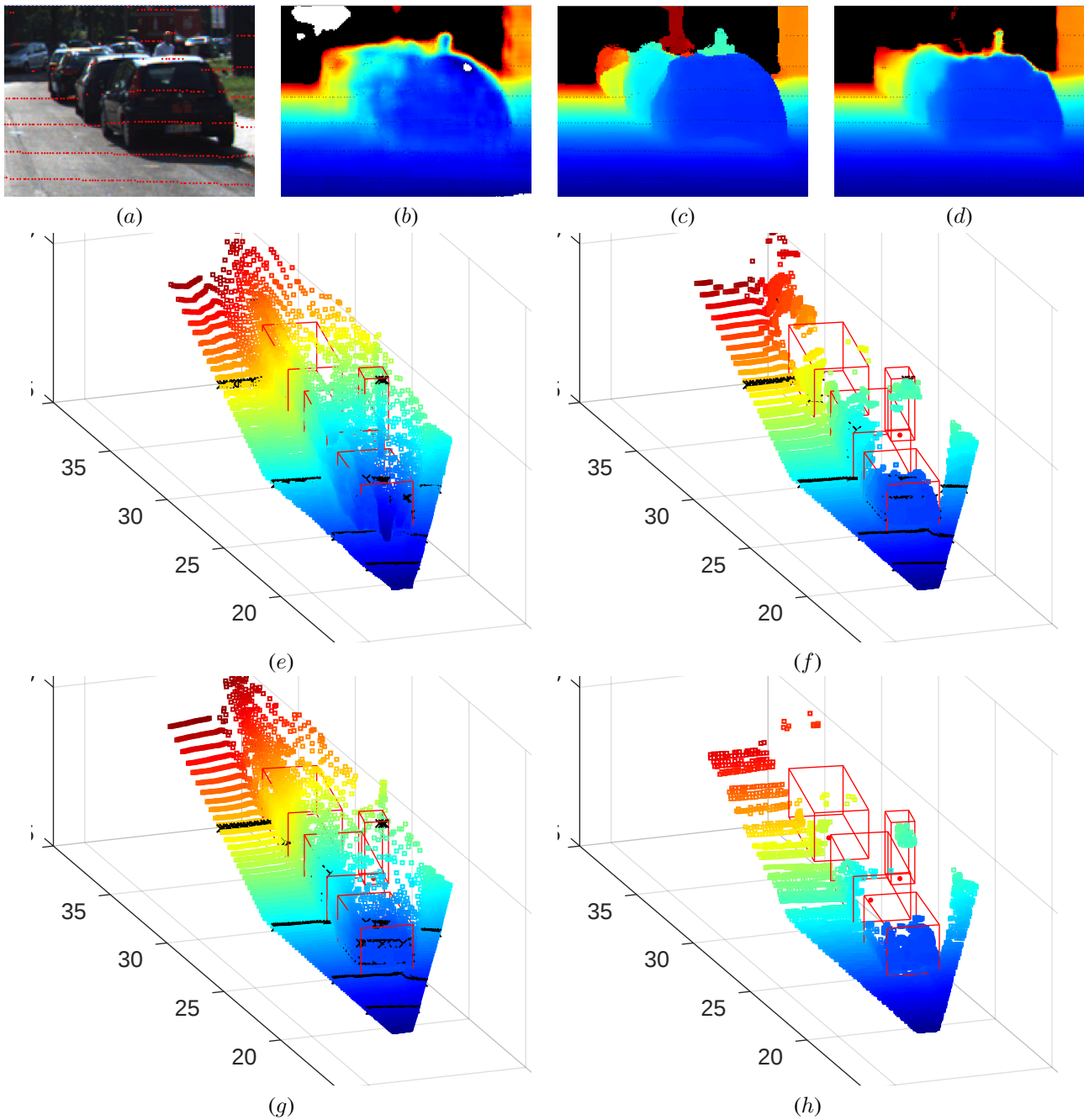


Figure 1. Estimated Depth-Map and Point-Clouds. (a) Color Image with Subsampled Lidar points (in red), (b), (c) and (d), (e), (f), (g) show predicted depth-maps and point-cloud of [1], Ours-3coeff, and Ours-allcoeff respectively; and (h) show point-cloud of Ground-truth (semi-dense annotated depth-map provided by KITTI). It shows (e), being state of the art, is introducing significant depth-mixing(g), being our method with all coefficients, has also introduced some depth mixing. (f), being our method with 3 coefficients, is more visually closer to the ground-truth.

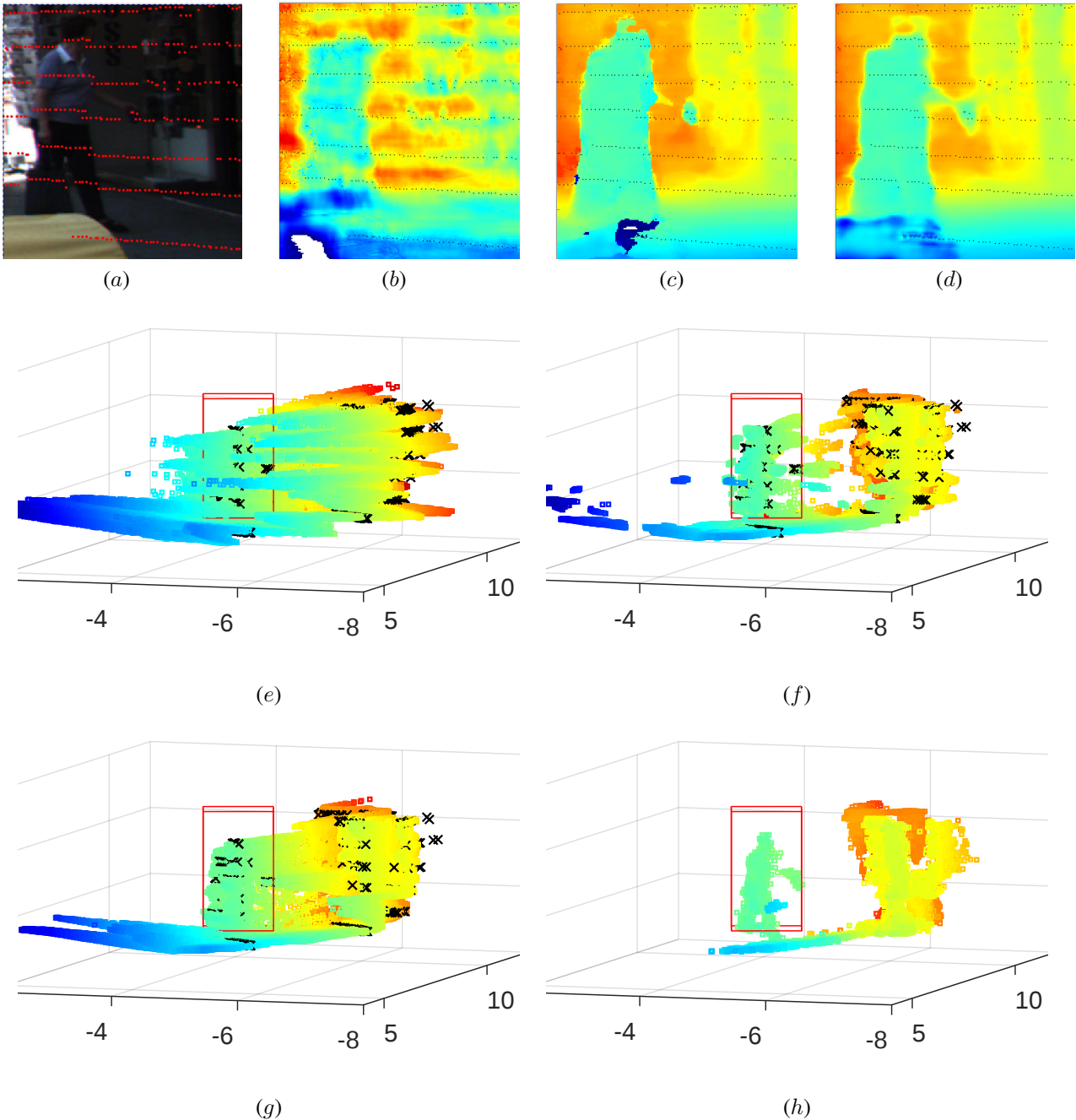


Figure 2. Estimated Depth-Map and Point-Clouds. (a) Color Image with Subsampled Lidar points (in red), (b), (c) and (d), (e), (f), (g) show predicted depth-maps and point-cloud of [1], Ours-3coeff, and Ours-allcoeff respectively; and (h) show point-cloud of Ground-truth (semi-dense annotated depth-map provided by KITTI). It shows (e), being state of the art, is introducing significant depth-mixing(g), being our method with all coefficients, has also introduced some depth mixing. (f), being our method with 3 coefficients, is more visually closer to the ground-truth.