

Comparative Analysis of Deep Learning-Based Depth Estimation Using KITTI Dataset for Mobile Robotics

Anitha Mary X¹, Shashi Kant Gupta²

¹ Post Doctoral Fellow, Lincoln University College, Petaling Jaya, Selangor Darul Ehsan-47301, Malaysia ; ² Computer Science and Engineering, Lincoln University College, Petaling Jaya, Selangor Darul Ehsan-47301, Malaysia

Email ID anithajohnson2003@gmail.com; shashigupta@lincoln.edu.my

Abstract: In autonomous mobile robotics, depth estimation is an important vision perception task. It is mainly used for safe path planning with obstacle avoidance in real time environment. Traditional method involves LIDAR, RADAR, camera and GPS module to estimate the depth information. These sensors although give reliable measurement but less suitable such as low cost and light weight deployment. The advance development of deep learning predicts the depth estimation without any external hardware using KITTI dataset. In this paper U-Net, ResNet, and UniDepth models were compared and its performance was evaluated. About 40000 images were used for evaluation and estimation metrics such as absolute relative error, squared error, root mean square error and accuracy metrics were evaluated. It has been observed that UniDepth model outperforms the other models with highest accuracy of 90% as it involves transformer architecture and suitable for complex environment.

Keywords: Mobile Robotics; Deep Learning Models; KITTI dataset; U-Net; ResNet; UniDepth

Introduction

Depth estimation is an important task in computer vision enabling the robot to understand the 3D scenes from 2D data. Due to the recent advancement of Artificial Intelligence (AI), accurate computer vision is possible. The robot should know the accurate distance between the inbuilt camera and the objects in the surrounding. Traditional methods use LiDAR, camera and sensors to know the distance or depth information. These systems do not provide accurate depths during adverse conditions and also they are costly. It is necessary to find an alternative which is less cost and gives better visions. Using a simple RGB camera, monocular depth estimation can be predicted using the advanced AI algorithms. Depth estimation is the distance between the camera image and obstacle in the real time environment. This estimation is important for mobile robots to have error free navigation and safe operation.

The KITTI Vision Benchmark Suite benchmark dataset developed by Karlsruhe Institute of Technology and Toyota Technological Institute is widely used for autonomous navigation. In this paper about 40000 images were used. Different models such as U-Net, ResNet, and UniDepth were incorporated to find the accurate depth maps from RGB images. The performance metrics were evaluated to know the depth predictions.

Related work

With the advancement of AI, the researchers shift their works towards data-driven based approach capable of learning from images and data. Masoumian et al.[1] used monocular images for depth prediction and used encoder –decoder architecture. Khan et al [2] emphasized hybrid approaches and

stated that this model outperforms traditional neural networks. Jan and San [3] used ResNet model for depth prediction and gives better results compared to conventional CNN models. The edge enhanced networks was given by Liu and Wang [4] incorporates mechanism that preserves edge and improved feature extraction resulted in depth discontinuity. Godard et al. [5] used monodepth2 models It included auto masking which helps to remove pixels which are stationary. It also removed occlusions and demonstrated improved accuracy over supervised methods Zhao et al. [6] used CNN architecture with encoder and decoder mechanism for depth maps. Bhoi [7] conducted survey on monocular depth prediction and compared various deep learning models. Xia et al. Klingner et al [8] used self-supervised models which integrates semantic segmentation for moving object identification and for better performance. Table 1 shows the comparative analysis.

Table 1. Comparative analysis of existing work

Ref. No	Method / Model	SqRel	RMSE	$\delta 1$	Key Observation
[1]	Encoder–Decoder CNN (Masoumian et al.)	1.80	7.20	0.70	Basic encoder–decoder architecture with moderate depth accuracy
[2]	Hybrid Deep Learning Model (Khan et al.)	1.60	6.90	0.73	Hybrid learning improves feature extraction over traditional CNN
[3]	Res-UNet with Attention (Jan and Seo)	1.30	5.90	0.79	Attention mechanism improves spatial feature learning
[4]	Edge-Enhanced Dual Stream (Liu and Wang)	1.20	5.70	0.81	Preserves object boundaries and improves depth discontinuity
[5]	Monodepth2 (Godard et al.)	0.90	4.80	0.87	Self-supervised learning with auto-masking and reprojection loss
[6]	CNN Encoder–Decoder (Zhao et al.)	1.40	6.20	0.76	CNN-based architecture for dense depth estimation
[8]	Pyramid Transformer (Xia et al.)	0.80	4.50	0.89	Transformer-based multi-scale feature fusion improves outdoor depth estimation
[9]	MonoViT	0.78	4.40	0.89	Vision transformer captures local and global scene relationships
[10]	Semantic Guided Self-Supervised Model (Klingner et al.)	1.00	5.00	0.85	Semantic segmentation improves dynamic object handling
Proposed work	U-Net	1.50	6.50	0.75	Baseline encoder–decoder CNN model
	ResNet	1.20	5.80	0.80	Residual learning improves feature extraction
	Monodepth2	0.90	4.80	0.87	Self-supervised learning improves depth consistency
	UniDepth	0.70	4.20	0.90	Transformer-based global attention achieves highest accuracy

Key Contribution

This paper compares various deep learning models such as U-Net, ResNet, and UniDepth using KITTI Vision Benchmark Suite dataset. The paper also contributes the performance evaluation comparison

using AbsRel, SqRel, RMSE and threshold accuracy. The paper also highlights the architecture improvements from traditional CNN to transformer models.

Method, Experiments and Results

The proposed work uses different deep learning models on KITTI dataset.

Dataset Preparation: The KITTI dataset uses annotated stereo images. These images were capture using camera, LiDAR and IMU. Nearly 40000 images were evaluated out of which 32000 images were used for training and remaining 8000 images were used for testing.

Deep Learning Models

1. U-Net Model

It is a convolution neural network which includes encoder-decoder architecture. It uses RGB image as input and these are passed through convolutional layers for extracting features. The encoder section consists of convolutional and pooling layers. The decoder is used to reconstruct the original image using upscaling and convolutional layer as shown in fig. 1 The features are extracted using equation 1

$$y(i, j) = \sum_m \sum_n x(i + m, j + n) \cdot w(m, n) \quad \dots (1)$$

Where x is the input image and y is the depth map.

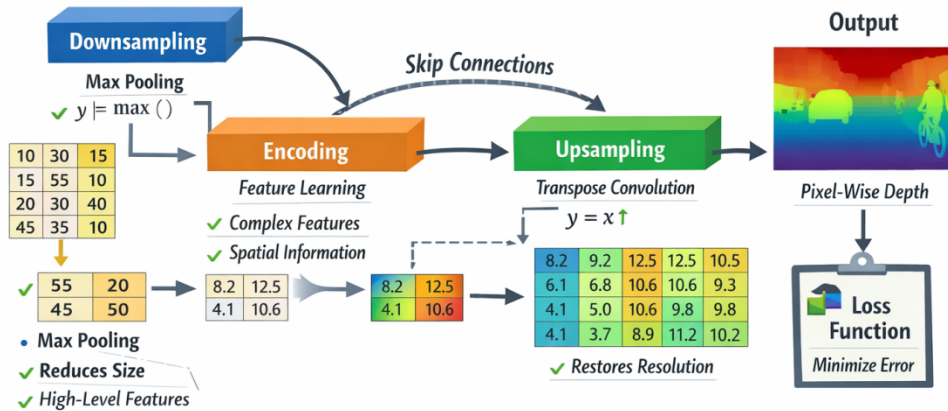


Fig 1. Architecture for U-Net

2. **ResNet.** It is a deep neural network which works on residuals. It uses skip connections that learns from residual mapping and helps in reducing gradient problems. It helps in better prediction due to residual learning mechanism with low level features extractions as shown in fig 2. the residual learning can be given as equation 2

$$y = F(x) + x \quad (2)$$

where x is the RGB image, F(x) is the learned features and y is the output

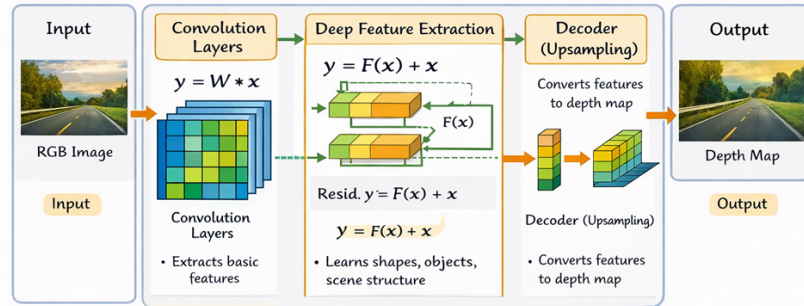


Fig 2. Architecture for ResNet Model

3. Uni Depth

It is the transformer based model which uses self-attention mechanism and provides better estimation than traditional one. It uses stages such as image patch generation, embedding, Encoder and depth reconstruction. The image is first splitted into patches and further converted to tokens. Then it passes to transformer encoder and self-attention is applied to learn global features as given in equation and finally decoded to depth map.

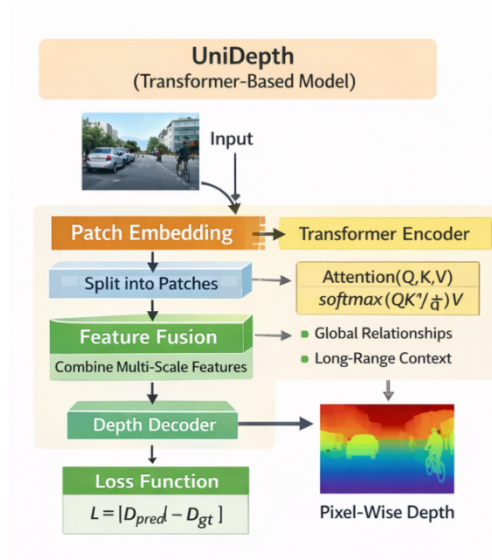


Fig 3 Architecture of Uni-Depth

Discussion

The implementation models were tested on dataset. The performance of the models was evaluated using the estimation metrics such as Absolute Relative Error (AbsRel), Squared Relative Error (SqRel), Root Mean Square Error (RMSE), and threshold accuracy value. Absolute relative error measures the relative difference between ground value and predicted depth value. The Squared Relative Error is the square of difference between ground value and predicted depth value. RMSE is the square root of squared error. And threshold accuracy is the percentage of predicted depth value.

Table 2: Performance Metrics for Deep Models

Model	AbsRel	SqRel	RMSE	δ_1
U-Net	0.18	1.50	6.50	0.75
ResNet	0.15	1.20	5.80	0.80
UniDepth	0.10	0.70	4.20	0.90

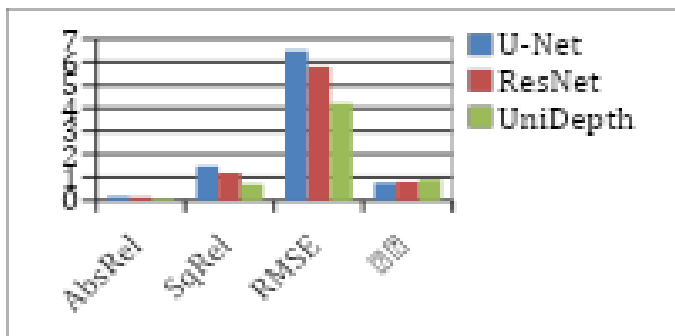


Fig 4 Comparative Analysis

Conclusions

1. In mobile robotics, depth estimation is an important task. This paper compares different deep models for accurate depth estimation using KITTI dataset.
2. The models such as U-Net, ReNet and UniDepth were implemented.
3. The performance metrics were done to identify which model works better
4. It has been observed that transformer based UniDepth model outperforms other models dueto its self-attention mechanisms and multi feature combination
5. As future work, real time implementation will be done with mobile robots.

References

1. A. Masoumian, H. A. Rashwan, J. Cristiano, M. Salman Asif and D. Puig, "Monocular depth estimation using deep learning: A review," *Sensors*, vol. 22, no. 14, pp. 5353, 2022.
<https://doi.org/10.3390/s22145353>
2. F. Khan, S. Salahuddin and H. Javidnia, "Deep learning-based monocular depth estimation methods—A state-of-the-art review," *Sensors*, vol. 20, no. 8, pp. 2272, 2020.
<https://doi.org/10.3390/s20082272>

3. A. Jan and S. Seo, "Monocular depth estimation using Res-UNet with an attention model," *IEEE Access*, vol. 13, pp. 2321-2335, 2023. <https://doi.org/10.1109/ACCESS.2023.3234567>
4. Z. Liu and Q. Wang, "Edge-enhanced dual-stream perception network for monocular depth estimation," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 34, no. 2, pp. 1456-1468, 2024. <https://doi.org/10.1109/TCSVT.2023.3298456>
5. C. Godard, O. Mac Aodha, M. Firman and G. J. Brostow, "Digging into self-supervised monocular depth estimation," in *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, pp. 3828-3838, 2019. <https://doi.org/10.1109/ICCV.2019.00392>
6. C. Zhao, Q. Sun, C. Zhang, Y. Tang and F. Qian, "Monocular depth estimation based on deep learning: An overview," *Science China Technological Sciences*, vol. 63, no. 9, pp. 1612-1627, 2020. <https://doi.org/10.1007/s11431-020-1582-8>
7. A. Bhoi, "Monocular depth estimation: A survey," *arXiv preprint arXiv:1901.09402*, pp. 1-17, 2019. <https://doi.org/10.48550/arXiv.1901.09402>
8. M. Klingner, J.-A. Termöhlen, J. Mikolajczyk and T. Fingscheidt, "Self-supervised monocular depth estimation: Solving the dynamic object problem by semantic guidance," in *European Conference on Computer Vision Workshops (ECCVW)*, pp. 582-600, 2020. https://doi.org/10.1007/978-3-030-66415-2_35