

RESEARCH ARTICLE

A Targeted Region-based Optical Flow and Motion Vector Deriving Approach for Self-Driving Vehicles using Monocular Vision

P.A.D.S.N. Wijesekara ^{1*}

¹ Department of Electrical and Information Engineering, Faculty of Engineering, University of Ruhuna, Galle 80000, Sri Lanka

Abstract: In autonomous driving, object dynamics are crucial for improved driving performance and fewer collisions, in conjunction with three-dimensional awareness of the environment. Using an aerial perspective, several previous efforts have followed the course of automobiles and examined optical flow for self-driving transportation. Because they track objects without priority, are susceptible to object obstructions, and experience a lack of situational awareness, these works lack emphasis on regions of interest (RoIs). In this work, we therefore propose to address this research gap by first obtaining the RoIs from the driver's perspective by identifying moving objects, such as vehicles and pedestrians, using an innovative divide and conquer strategy to produce various hierarchical 2D bounding boxes utilizing a segmented image, addressing overlying segments leveraging pixel matching. In order to prioritize regions of interest above other areas and ensure data accuracy and smoothness in the areas of interest, we then created the optical flow for the chosen ROIs by taking into account a scaled energy formulation. Last but not least, two segmented images are used for motion vector generation. To mitigate the effects of object deformation (deploying augmentation), erroneous positives (deploying neighborhood scan), and partial template (deploying template splitting) issues, we conduct the enhanced template standardized cross-correlation within the vicinity of RoIs to assess the movement of entities. The motion vectors, which show the relative velocity of objects in reference to the moving vehicle, are obtained by using the centroids of the old and new objects after template matching. To assess how well the suggested method performs, we choose a dataset from the KITTI dataset and the CARLA simulator. The findings demonstrate the effectiveness of the suggested RoI identification and optical flow determination models, and that motion vectors closely match the actual object relative velocities.

Keywords: *Image processing, optical flow, motion vector, self-driving, dual TV, monocular vision*

Introduction

Autonomous driving is a core component in an intelligent transportation system (Zhang and Lu, 2020). In self-driving, automobiles are operated by themselves with little assistance from people (Liu *et al.*, 2021). In the early days, autonomous driving was more heuristic, meaning that rules governed the driving rather than autonomous inference, like using a PID control system (Gottschalk *et al.*, 2010). However, nowadays, such heuristic-based decisions are rarely made, and now, the decisions are taken using more modern approaches like machine learning and optimization (Guo *et al.*, 2025). They have been

heavily utilized in trajectory prediction and path planning tasks in autonomous driving by using specific techniques like deep learning or reinforcement learning (Yu *et al.*, 2024).

Since most autonomous driving works rely on their judgments on the current timestep, they mostly lack a suitable method to get the moving patterns of key entities (moving cars and people) (Wijesekara, 2022). Studies done towards the end of the 2010s have proposed a modular pipeline architecture to derive driving actions using monocular vision of a single RGB image (Li and Qin, 2018). There exist numerous studies on imitation learning,

*corresponding author: nilmantha@eie.ruh.ac.lk

 <https://orcid.org/0000-0002-3045-1596>



This article is published under the Creative Commons CC-BY-ND License (<http://creativecommons.org/licenses/by-nd/4.0/>). This license permits commercial and non-commercial reuse, distribution, and reproduction in any medium, provided the original work is not changed in any way and is properly cited.

where the autonomous driving model learns to imitate the actions performed by an expert driver (Le Mero *et al.*, 2022). These studies infer actions from the observations in the present time and lack understanding about object movements in the vicinity.

Region-of-interest-aware visual flow and movement indicators for inference are not generated by the few known research studies that combine temporal characteristics. To be more precise, they have inspected the vehicular motion from an aerial perspective (Hu *et al.*, 2022, Jiang *et al.*, 2023). Similarly, BEV Former interacts with spatial and temporal space using previously known grid-like BEV queries that aggregate spatial input using point clouds and cameras, while temporal knowledge is utilized by using temporal self-attention (Li *et al.*, 2025). For instance, in the study (Kang and Lee, 2021), a driver's visual attention is captured using motion information estimation using deep neural networks to provide locations and levels with the aid of optical flow maps. On the other hand, some researchers have utilized spatial and temporal cues utilizing vehicular monocular observations and vehicle historical states using LSTM (Chi and Mu, 2017). This lacks situational awareness because the view may be obstructed by objects, areas of interest are not prioritized, and the motion with respect to the driving vehicle is not taken into account.

Incorporating dynamic characteristics into autonomous driving scenarios has been explored in a few studies. By employing a distinct speed prediction branch, some have tried to integrate speed prediction within the model (Codevilla *et al.*, 2019). To ascertain temporal linkages in the environment, Olier *et al.* (2017) use observational learning in conjunction with Bayesian filtering. To drive perception by filtering out space-time aspects of the map, some have taken into consideration previous motions in an aerial perspective of the area (Hu *et al.*, 2022). Likewise, other studies propose a method to recognize 3D items from the self-driving vehicle and a self-motion monitoring system that measures the speed of onboard cameras (Li and Qin, 2018, Wang, H. *et al.*, 2021). In a similar manner, Jiang *et al.* (2023) also monitor vehicle speeds to provide an aerial perspective of the traffic landscape. An optical flow energy potential for the entire environment is incorporated in works in Wang, Hengli *et al.* (2021) and Yuan *et al.* (2023). In a similar manner, another study helps with trajectory planning by using optical flow from six cameras. Some have recognized and separated the moving item in the image by attempting to employ optical flow for moving object recognition (Siam *et al.*, 2018, Rashed *et al.*, 2019). Lastly, some have determined the optical

flow of the environment using semantic characteristics and camera pictures (Yuan *et al.*, 2023, Wijesekara, 2025a).

Even though some have proposed optical flow to be utilized in autonomous driving, optical flow alone cannot be utilized to learn the relative motion of the objects in the surroundings. In order to cater to this research void, in this work, we use image processing techniques and use two novel algorithms to produce motion vectors when two segmented images are given. Note that, different from traditional RGB images for object detection, we use segmented images, as the latter contain two classes of possible dynamic items of interest: pedestrians and vehicles. The first algorithm extracts those regions of interest by drawing bounding boxes around possible dynamic objects and splitting the bounding boxes in case overlapping classes exist. Next, we derive dual TV-L1 optical flow for the regions of interest by using a prioritized energy expression that emphasizes the optical flows of regions of interest to indicate whether there exists motion or not. Finally, we generate motion vectors, which are the core of this work, by using template matching, mitigating inherent problems in traditional template matching by using template augmentation, splitting, and neighborhood search.

The novelty of this research can be stated as follows: 1) Different from existing work, we use segmented images derived from a machine learning model to detect dynamic items of interest. 2) Use a novel divide-and-conquer algorithm to extract regions of interest, considering object overlaps. 3) Based on the derived RoIs, dual TV-LV1 optical flow for RoIs is derived using a prioritized energy expression. 4) A motion vector extraction algorithm is presented that utilizes augmented template matching with standardized cross-correlation and neighborhood search.

A review of existing literature (Codevilla *et al.*, 2019, Li and Qin, 2018, Siam *et al.*, 2018) indicates that we are the frontline researchers to derive motion vectors that resemble relative motion with respect to the driver's point of view for moving regions of interest using optical flow, normalized cross-correlation, and template matching for segmented images. As suggested by the simulation experiments, the proposed framework is better than state-of-the-art techniques in terms of different metrics. Thus, this work significantly improves the current literature on optical flow and motion estimation. The model can be readily used in an autonomous driving model to improve the driving

performance; thus, it can contribute to improving autonomous driving.

Materials and Methods

First, the segmented images required for the proposed model were obtained using machine learning, as per previous studies (Wijesekara, 2022).

Region of interest finding

Initially, two masks (binary pictures) containing segmented objects exclusive to that class are obtained using the segmented picture. Equation (1) may be utilized to express a binary picture for a certain class.

$$Img_{mask} = \begin{cases} 255; & \text{if } Img_{segment}(h, v) == color \\ 0; & \text{otherwise} \end{cases} \quad (1)$$

Note that $Img_{segment}$ is the input segmented picture in Equation (1).

Next, using Equations (2) to (3), we identify the contour's minimum and maximum points that may be encapsulated by a rectangle to establish the first-level bounding box.

$$h_{min} = \min(h_1, h_2, h_3, \dots, h_n) \quad (2)$$

$$h_{max} = \max(h_1, h_2, h_3, \dots, h_n) \quad (3)$$

$(h_1, h_2, h_3, \dots, h_n)$ are the horizontal coordinates of every point on the contour in Equations (2) and (3). Keep in mind that matching vertical values need to be calculated similarly to how they are in Equations (2) and (3).

Algorithm 1 Region of interest finding (ROIF)

```

1: procedure ROIF( $Img, color, thresh1, thresh2, min\_area$ )
2:    $Bin\_img \leftarrow SEARCH\_BIN(Img, color)$ 
3:    $1L\_cont \leftarrow SEARCH\_CONT(Bin\_img)$ 
4:   for each  $cont$  in  $1L\_cont$  do
5:     if  $Area(cont) > min\_area$  then
6:        $2L\_cont1 = SPLIT\_CONT(cont, color, thresh1, vertical)$ 
7:        $2L\_cont2 = CONQUER\_CONT(2L\_cont1, color, thresh1, 1)$ 
8:       if  $len(2L\_cont2) > 0$  then
9:          $3L\_cont1 = SPLIT\_CONT(2L\_cont2, color, thresh2, horiz.)$ 
10:         $3L\_cont2 = CONQUER\_CONT(3L\_cont1, color, thresh2, 2)$ 
11:        if  $len(3L\_cont2) > 0$  then
12:           $cont\_set.PUSH\_BACK(3L\_cont2)$ 
13:        else
14:           $cont\_set.PUSH\_BACK(2L\_cont2)$ 
15:        end if
16:      else
17:         $cont\_set.PUSH\_BACK(cont)$ 
18:      end if
19:    end for
20:  return  $cont\_set$ 
21: end procedure

```

We use a divide-and-conquer strategy to separate a first-level bounding box into many bounding boxes and contemplate conquering them according to a

threshold in order to avoid prior behavior (see Algorithm 1).

Optical flow creation

Equation (4) provides the function $R(h, v)$, which we use to describe the region of interest.

$$R(h, v) = \begin{cases} 1; & \text{if } (h, v) \text{ within } RoI \\ 0.9; & \text{if } (h, v) \text{ in neighborhood} \\ 0.8; & \text{otherwise} \end{cases} \quad (4)$$

It should be noted that the neighborhood is defined as the magnitude of the height in both plus and minus v (vertical) directions and the width of the ROI in both plus and minus h (horizontal) directions. The image coordinates are (h, v) in Equation (4).

According to Equation (5), we calculate the weight function ($P(h, v)$) for dual TV-L1 optical flow computation based on the region of interest.

$$P(h, v) = \lambda R(h, v) + \mu(1 - R(h, v)) \quad (5)$$

The constants $0 \leq \lambda \leq 1$ and $0 \leq \mu \leq 1$ determine the significance of the region of interest in Equation (5).

Motion vector creation

We use standardized cross-correlation to match augmented templates in the original template's neighborhood with optional splitting if a satisfactory match is not found. After identifying the best match position, we use Equation (6) to get the centroid (C_{opt}) of the best match region.

$$C_{opt} = \{h_{opt} + \frac{W}{2}, v_{opt} + \frac{H}{2}\} \quad (6)$$

In Equation (6), h_{opt} , v_{opt} are the top left coordinates of the best match location, and W and H are the width and height of the template. Let B be the template area, $|F|$ and α be the motion vector magnitude and direction, and $C_{initial} = \{h_{initial}, v_{initial}\}$ be the centroid of the template in the original picture. Lastly, Equation (7) provides the motion vector (E).

$$E = \{C_{initial}, B, |F|, \alpha\} \quad (7)$$

Algorithm 2 displays the motion vector creation technique.

Algorithm 2 Motion Vector Creation (MVC)

```

1: procedure MVC( $Img_1, Img_2, Img_{of}, RoIs, Thr$ )
2:   for each  $R$  in  $RoIs$  do
3:      $Got, Centr \leftarrow MATCH\_TEMPL(Img_1, Img_2, R, Thr)$ 
4:      $Vec \leftarrow CALCULATE\_MOT\_VEC(Centr, R)$ 
5:      $\{Img_{of\_mv}, Iseg\_mv\} \leftarrow NOTE\_MOT\_VEC(Img_1, Img_{of}, Vec)$ 
6:   end for
7:   return  $Img_{of\_mv}, Img_{seg\_mv}, Vec$ 
8: end procedure

```

Results and Discussion

In order to gather a segmented picture dataset for assessing the suggested models, we utilize CARLA to emulate automobiles and pedestrians in the self-driving scenarios. This is done in client-server mode with a sufficient frame rate of 30 fps for video perception (Wijesekara, 2025b). In front of the car, the first-person camera is placed at random points with different yaw, roll, and pitch. The KITTI dataset is used to acquire ground truth optical flow pictures in order to assess the optical flow's performance.

Performance assessment metrics

Equation (8) provides the Bounding-box Error (BE), which we employ.

$$BE = \frac{1}{N} \sum_{r=1}^N \frac{1}{M} \sum_{s=1}^M \sqrt{(a_{prs}^{min} - a_{grs}^{min})^2 + (b_{prs}^{min} - b_{grs}^{min})^2 + (a_{prs}^{max} - a_{grs}^{max})^2 + (b_{prs}^{max} - b_{grs}^{max})^2} \quad (8)$$

The subscript p in Equation (8) represents the predicted value, g the ground truth, a the x coordinate, b the y coordinate, min and max abbreviate the lowest and highest values, subscript s the s^{th} bounding box, M the total number of bounding contours in a picture, N the number of images, and r the r^{th} image.

We utilize the End-point Error (EE), which is computed using Equation (9), to assess the efficacy of the ROI-aware optical flow.

$$EE = \frac{1}{N} \sum_{r=1}^N \frac{1}{M} \sum_{s=1}^M \frac{1}{|R_s|} \sum_{x,y \in R_s} \sqrt{(e_{prs}(x,y) - e_{grs}(x,y))^2 + (f_{prs}(x,y) - f_{grs}(x,y))^2} \quad (9)$$

In Equation (9), e and f represent the lateral and perpendicular elements of optical flow, respectively, and R_s represents the s^{th} area of interest. As previously stated, other symbols have appropriate meaning.

We employ standardized Relative Velocity Error (RVE), as given in Equation (10), to assess the efficacy of the suggested motion vector estimation scheme (Algorithm 2).

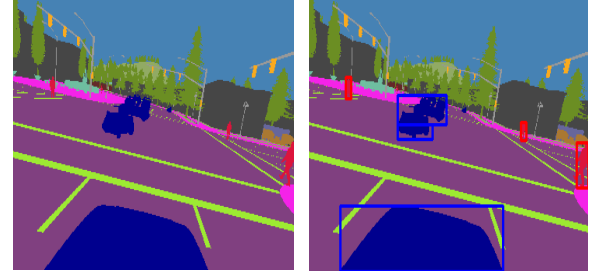
$$RVE = \frac{1}{N} \sum_{r=1}^N \frac{1}{M} \sum_{s=1}^M \frac{|\frac{m_{prs}}{m_{prs}^{max}} - \frac{m_{grs}}{m_{grs}^{max}}|}{|\frac{m_{prs}}{m_{prs}^{max}}| + |\frac{m_{grs}}{m_{grs}^{max}}|} \quad (10)$$

The velocity is represented by m in Equation (10), and all other symbols are the same as those found in Equation (9).

Performance evaluation

Figure 1 displays the suggested ROIs derived from Algorithm 1.

Instead of wrongly allocating a single motion vector to both automobiles, the algorithm was able to recognize two overlapping cars individually (Figure 1), allowing for the generation of distinct motion vectors for each car. The optical flow produced for the provided segmented pictures with ROIs is displayed in Figure 2.



(a) Input: Segmented picture (left column) (b) Output: Picture with ROIs (right column)

Figure 1: Algorithm 1 results

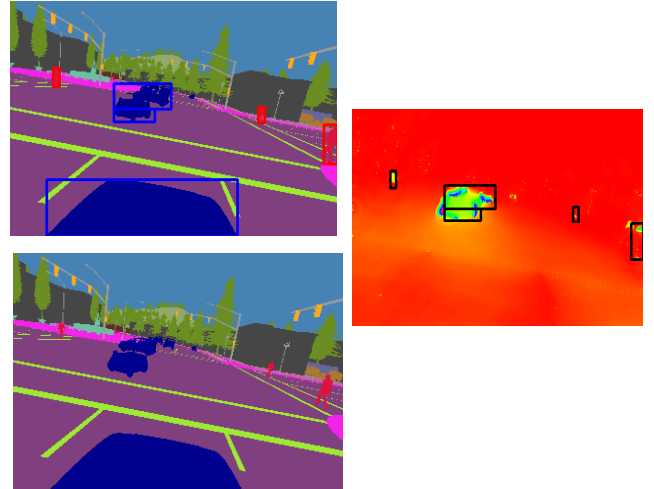


Figure 2: ROI-based optical flow creation. (a) Inputs: Segmented pictures (left column) (b) Output: Optical flow having ROIs (right column)

The non-dark red color patches in Figure 2(b) indicate the detection of optical flow for the pedestrians and two autos. Please take note that the color map we utilized is magenta for elevated optical flow and dark red for low or no optical flow. The hues change from red to orange, orange to yellow, yellow to green, green to cyan, cyan to blue, and blue to magenta as the optical flow rises. The regions of interest in Figure 2(b) are shown by black rectangles.

The input and output pictures for motion vector creation are displayed in Figure 3.

The output photos demonstrate that Algorithm 2 was effective in producing the motion vectors, as seen in Figure 3. It should be noted that motion vectors are represented by black pointers in the optical flow picture and white pointers in the segmented picture.

The effectiveness of Algorithm 1 in comparison to current methods is displayed in Figure 4.

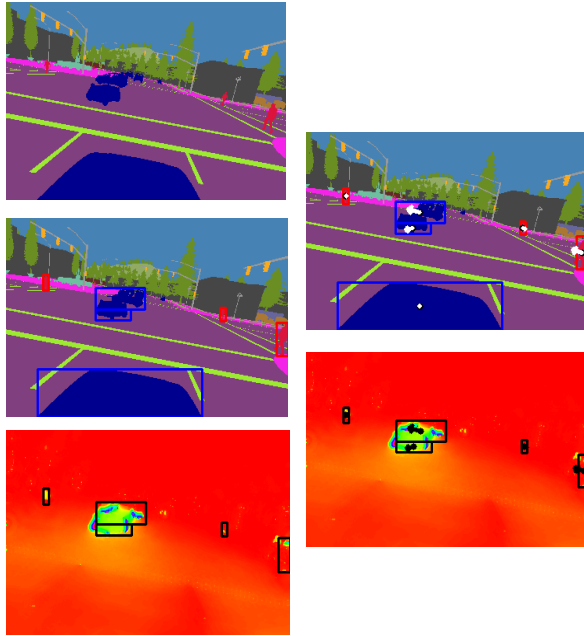


Figure 3: Motion vector creation (a) Inputs: segmented & optical flow pictures (left column) (b) Outputs: pictures with motion vectors (right column)

Figure 4 makes it clear that when there are few ROIs per picture, the suggested ROI determination technique performs somewhat worse than Faster-RCNN. However, the performance of both YOLO and Faster-RCNN starts to deteriorate as the image's ROIs increase. However, when the number of objects increases, the BBE of the suggested ROI determination approach only minimally increases. The fundamental cause of the previous variance is that Algorithm 1 divides the bounding boxes according to matching pixel proportions and a threshold using a divide and merge strategy since it is aware of overlapping items of the same class in the segmented picture.

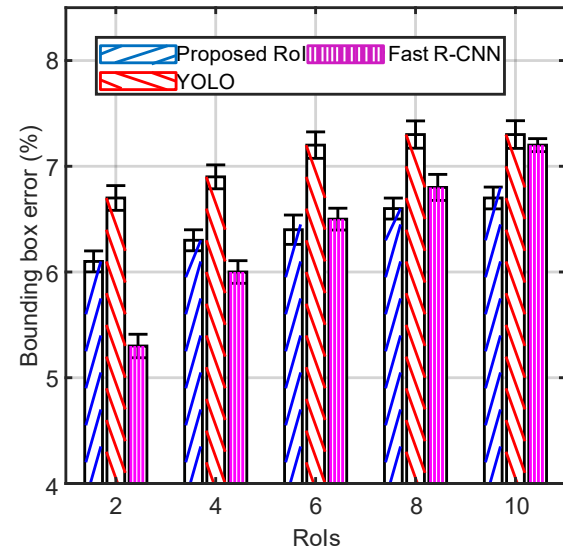


Figure 4: Performance assessment of ROI determination

To verify the ROI determination statistically, we assumed two null hypotheses: H_{0R1} : "The performance of the proposed ROI is worse than YOLO" and H_{0R2} : "The performance of the proposed ROI is worse than Fast-RCNN." We conducted paired t-tests to statistically compare the performance of the proposed method with Fast R-CNN and YOLO, using the reported bounding box error values. For evaluation of H_{0R1} : $t = -16.5$ and $p = 3.95e-9$; thus, H_{0R1} can be rejected, and the alternative hypothesis, which is that the bounding box error of the proposed method is lower than YOLO, is true. However, for H_{0R2} , $t = -0.27$ and $p = 0.4$, suggesting that the null hypothesis cannot be rejected. This is owing to the fact that Fast R-CNN has better performance than the proposed ROI determination under a lower number of ROIs, and vice versa.

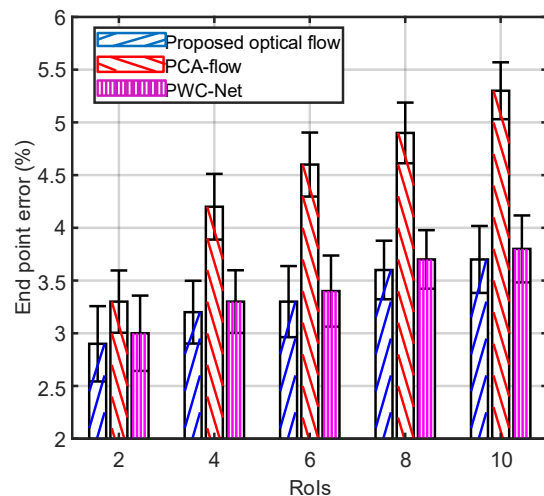


Figure 5: Assessment of ROI-driven dynamic vision

Figure 5 compares the suggested RoI-driven dynamic vision to PCA-Flow (Chen *et al.*, 2021) and PWC-Net (Sun *et al.*, 2018) for assessment. Figure 5: Assessment of RoI-driven dynamic vision.

The EE of the suggested RoI-aware optical flow creating module is much lower than that of PCA-flow under any number of RoIs, as shown in Figure 5. However, the EE of the proposed technique is slightly better than PWC-Net owing to the fact that PWC-Net does not cater to optical flow in regions of interest, unlike the proposed one, despite it being one of the top-performing optical flow generating models. In order to verify the RoI determination statistically, we assumed two null hypotheses: H_{0Q1} : “The performance of the proposed optical flow is worse than PCA flow,” and H_{0Q2} : “The performance of the proposed optical flow is worse than PWC-Net.” We conducted paired Wilcoxon signed-rank tests to statistically compare the performance of the proposed method with PCA flow and PWC-Net, using the reported endpoint error values. Here, we used the Wilcoxon signed-rank test, as the differences among the compared values were too small for a paired t-test. For evaluation of H_{0Q1} : $v = 0$ and $p = 0.026$; thus, H_{0Q1} can be rejected, and the alternative hypothesis, which is that the end point error of the proposed optical flow is lower than PCA-flow, is true. Similarly, for H_{0Q2} , $v = 0$ and $p = 0.031$, suggesting that the null hypothesis can be rejected and the alternative hypothesis, H_{0Q2} , which is that the endpoint error of the proposed optical flow is less than PWC-Net, is true. The results indicate that the proposed method achieves statistically significant reductions in endpoint error compared to both baselines.

Therefore, the suggested approach works well for supplying optical flow to support autonomous driving. Please take note that we have utilized 40 photos for evaluation for each data value in Figure 5.

Figure 6 shows how NRVE is changed with the increment of RoIs.

The relative velocity inaccuracy rises as the number of ROIs in the picture does, as seen in Figure 6. This is because motion prediction gets difficult when there are more objects around the vehicle because of the increase in dynamic object area in the segmented picture, which might lead to some interference between optical flow and ROI determination. The fact that the standardized RVE is much lower, however, is significant since it shows that motion vectors properly anticipate the real relative velocities between the ego-driving automobile and other dynamic objects.

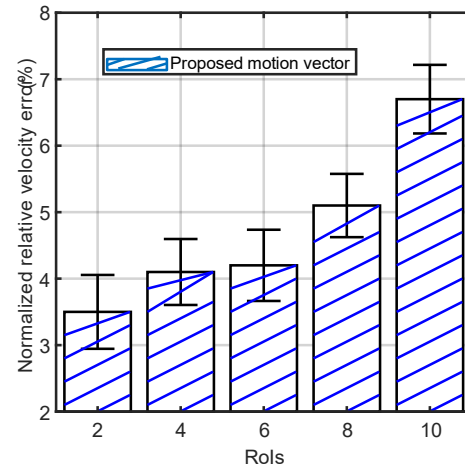


Figure 6: Performance assessment of motion vector creation

Since this is the first study on relative motion vector generation using monocular segmented images using a region of interest-based optical flow model, we don't have any other technique to compare our results with. In order to substantiate the motion vector prediction with at least close work, we refer to the study (Rill, 2019) that shows that incorporating depth information along with optical flow improves speed estimation and has reported a speed error of less than 1 m/s. As our normalized relative velocity error is under 7% under dense vehicular scenarios, that error can map to an absolute speed error of less than 1 m/s under dense network scenarios, suggesting that our work achieves similar performance.

Conclusion

In order to facilitate autonomous driving decision-making, this research proposed a motion vector and optical flow creation model that is aware of the region of interest. After segmented picture identification that accommodated overlapping objects, we employed enhanced template matching in the neighborhood with potential splitting to create motion vectors and optical flow for the regions of interest. The findings demonstrated that, in contrast to current methods, motion vectors closely approximate the true relative velocity between the moving vehicle and the objects of interest, optical flow production is efficient, and region of interest identification yields low bounding box inaccuracy.

Acknowledgement

We acknowledge the original contributors of the KITTI dataset.

References

- Chen, H., Li, H., Xu, Z., Zhao, Y., and He, T. (2021). Real-time action feature extraction via fast PCA-Flow. *Concurrency and Computation: Practice and Experience* 33:e5507. <https://doi.org/10.1002/cpe.5507>.
- Chi, L., and Mu, Y. (2017). Deep steering: Learning end-to-end driving model from spatial and temporal visual cues. *arXiv preprint arXiv:1708.03798:1708.03798*. <https://doi.org/10.48550/arXiv.1708.03798>.
- Codevilla, F., Santana, E., López, A. M., and Gaidon, A. (2019). Exploring the limitations of behavior cloning for autonomous driving *Proceedings of the IEEE/CVF international conference on computer vision*,9329-9338
- Gottschalk, R., Burgos-Artizzu, X. P., Ribeiro, A., Pajares, G., and , A. S. M. (2010). Real-time image processing for the guidance of a small agricultural field inspection vehicle. *International Journal of Intelligent Systems Technologies and Applications* 8:434-443. <https://doi.org/10.1504/ijista.2010.030222>.
- Guo, X., Jiang, F., Chen, Q., Wang, Y., Sha, K., and Chen, J. (2025). Deep learning-enhanced environment perception for autonomous driving: MDNet with CSP-DarkNet53. *Pattern Recognition* 160:111174. <https://doi.org/10.1016/j.patcog.2024.111174>.
- Hu, S., Chen, L., Wu, P., Li, H., Yan, J., and Tao, D. (2022). ST-P3: End-to-End Vision-Based Autonomous Driving via Spatial-Temporal Feature Learning. *Springer Nature Switzerland, Cham*,533-549
- Jiang, B., Chen, S., Xu, Q., Liao, B., Chen, J., Zhou, H., Zhang, Q., Liu, W., Huang, C., and Wang, X. (2023). Vad: Vectorized scene representation for efficient autonomous driving *Proceedings of the IEEE/CVF International Conference on Computer Vision*,8340-8350
- Kang, B., and Lee, Y. (2021). A Driver's Visual Attention Prediction Using Optical Flow. *Sensors* 21:3722. <https://doi.org/10.3390/s21113722>.
- Le Mero, L., Yi, D., Dianati, M., and Mouzakitis, A. (2022). A Survey on Imitation Learning Techniques for End-to-End Autonomous Vehicles. *IEEE Transactions on Intelligent Transportation Systems* 23:14128-14147. [10.1109/TITS.2022.3144867](https://doi.org/10.1109/TITS.2022.3144867).
- Li, P., and Qin, T. (2018). Stereo vision-based semantic 3d object and ego-motion tracking for autonomous driving *Proceedings of the European Conference on Computer Vision (ECCV)*,646-661
- Li, Z., Wang, W., Li, H., Xie, E., Sima, C., Lu, T., Yu, Q., and Dai, J. (2025). BEVFormer: Learning Bird's-Eye-View Representation From LiDAR-Camera via Spatiotemporal Transformers. *IEEE Transactions on Pattern Analysis and Machine Intelligence* 47:2020-2036. <https://doi.org/10.1109/TPAMI.2024.3515454>.
- Liu, S., Yu, B., Tang, J., and Zhu, Q. (2021). Towards Fully Intelligent Transportation through Infrastructure-Vehicle Cooperative Autonomous Driving: Challenges and Opportunities 2021 58th ACM/IEEE Design Automation Conference (DAC),1323-1326
- Olier, J. S., Marín-Plaza, P., Martín, D., Marcenaro, L., Barakova, E., Rauterberg, M., and Regazzoni, C. (2017). Dynamic representations for autonomous driving 2017 14th IEEE International Conference on Advanced Video and Signal Based Surveillance (AVSS),1-6
- Rashed, H., Ramzy, M., Vaquero, V., El Sallab, A., Sistu, G., and Yogamani, S. (2019). Fusemodnet: Real-time camera and lidar based moving object detection for robust low-light autonomous driving *Proceedings of the IEEE/CVF international conference on computer vision workshops*,1-10
- Rill, R. A. (2019). Speed estimation evaluation on the KITTI benchmark based on motion and monocular depth information. *arXiv preprint arXiv:1907.06989*.
- Siam, M., Mahgoub, H., Zahran, M., Yogamani, S., Jagersand, M., and El-Sallab, A. (2018). MODNet: Motion and Appearance based Moving Object Detection Network for Autonomous Driving 21st International

- Conference on Intelligent Transportation Systems (ITSC),2859-2864
- Sun, D., Yang, X., Liu, M.-Y., and Kautz, J. (2018). Pwc-net: Cnns for optical flow using pyramid, warping, and cost volume Proceedings of the IEEE conference on computer vision and pattern recognition,8934-8943
- Wang, H., Cai, P., Fan, R., Sun, Y., and Liu, M. (2021). End-to-end interactive prediction and planning with optical flow distillation for autonomous driving Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition,2229-2238
- Wang, H., Cai, P., Sun, Y., Wang, L., and Liu, M. (2021). Learning Interpretable End-to-End Vision-Based Motion Planning for Autonomous Driving with Optical Flow Distillation 2021 IEEE International Conference on Robotics and Automation (ICRA),13731-13737
- Wijesekara, P. A. D. S. N. (2022). Deep 3D dynamic object detection towards successful and safe navigation for full autonomous driving. The Open Transportation Journal 16:1-15. <https://doi.org/10.2174/18744478-v16-e2208191>.
- Wijesekara, P. A. D. S. N. (2025a). A Region of Interest-aware Motion Vector and Optical Flow Generation Model for Autonomous Driving 22nd Academic Sessions, University of Ruhuna. University of Ruhuna, Sri Lanka,32
- Wijesekara, P. A. D. S. N. (2025b). Towards efficient and reliable video communication: A survey on scalability, error protection, and multicasting. The Indonesian Journal of Computer Science 14:34-67.
- Yu, H., Huo, S., Zhu, M., Gong, Y., and Xiang, Y. (2024). Machine learning-based vehicle intention trajectory recognition and prediction for autonomous driving 7th International Conference on Advanced Algorithms and Control Engineering (ICAACE),771-775
- Yuan, S., Yu, S., Kim, H., and Tomasi, C. (2023). Semarflow: Injecting semantics into unsupervised optical flow estimation for autonomous driving Proceedings of the IEEE/CVF International Conference on Computer Vision,9566-9577
- Zhang, H., and Lu, X. (2020). Vehicle communication network in intelligent transportation system based on Internet of Things. Computer Communications 160:799-806. <https://doi.org/10.1016/j.comcom.2020.03.041>.