

HO3-SLAM: Human-Object Occlusion Ordering SLAM framework for Traversability Prediction

Jonathan Tay Yu Liang^[0009-0005-3985-0373] and Kanji Tanaka^[1111-2222-3333-4444]

¹ University of Fukui, Fukui, Japan
mf228029@g.u-fukui.ac.jp
tnkknj@u-fukui.ac.jp

Abstract. Detecting traversable areas from monocular vision is a fundamental problem for safe and efficient indoor robot navigation. Recent advances in deep learning have made it possible to detect floor regions stably and robustly. However, most existing methods rely on the direct visibility of the floor, which is often violated in crowded scenes like office environments, limiting their effectiveness. To address this limitation, we explore a novel approach that leverages pedestrian observation. By observing pedestrian behavior and assuming the floor region they occupy as traversable, we complement existing solutions. Recent advancements in computer vision techniques, including dynamic object detection (e.g., YOLO object detector), absolute localization of stationary objects (e.g., SLAM), and relative localization between dynamic and stationary objects (e.g., Occlusion ordering), enable the implementation of this approach. We developed a prototype system by extending state-of-the-art techniques in dynamic SLAM and object detection and conducted real-world experiments to demonstrate its basic performance. Our approach has significant potential for applications in challenging domains such as human-robot collaboration and long-horizon action planning.

Keywords: SLAM, Occlusion Ordering, Traversability Prediction

1 Introduction

Detecting traversable areas using monocular cameras is crucial for safe and efficient robot navigation. This problem has been extensively studied in computer vision, with various approaches such as obstacle detection [1], free space detection [2], and 3D reconstruction [3, 4, 5]. Many researchers have used deep learning models to predict pixel-level depth values from images [1], while others have developed vision-based algorithms that gradually expand into free space from the robot's feet [2]. Additionally, some studies have focused on reconstructing 3D geometric models using the robot's onboard camera view sequence [3, 4, 5]. However, most of the existing studies assume that the floor area is directly visible, which is not always the case in crowded scenarios. For instance, in crowded scenes like an office where humans and robots collaborate [6], obstacles (e.g., desks) often occlude the floor areas humans are standing on. In such situations, floor detection methods become ineffective, and neither obstacle detection nor 3D reconstruction methods can provide floor information.

To address this issue, we propose an innovative approach to learning traversable areas by tracking pedestrians, i.e., observing the behaviour of pedestrians assuming that the

floor region they are standing on is traversable. From our observation, it is recently becoming possible to implement this new approach by combining state-of-the-art of computer vision techniques, including (1) detection of dynamic objects (e.g., YOLO [7] object detector), (2) absolute localization of stationary objects (e.g., DynaSLAM [8]), and (3) relative localization between dynamic and stationary objects (e.g., Occlusion ordering). Moreover, typical objects in human environments like desks are not tall enough to occlude a pedestrian's entire body. Finally, we infer the occlusion order of stationary objects and pedestrians in the same image coordinate system to localize the pedestrians in terms of their position relative to the stationary objects, whose absolute coordinates are reconstructed. We will be using Detectron2 instance segmentation mask method [9] and bounding box method for occlusion order reasoning. After that, we make an ablation study on the combination of the methods to achieve more desirable results. We conducted preliminary experiments that showed how deep learning based YOLOv7 [7] detectors could easily detect regions in human images. Although the lower body and floor contact areas of humans are often occluded, the head and upper body are detectable and not occluded.

This research is closely related to the large-body research area of pedestrian observation. However, most of them are not intended for traversable region analysis. Furthermore, many of them assumed a fixed-point camera. To the best of our knowledge, the current work is the first to explore traversable area analysis from moving cameras [10].

In summary, our approach to learning traversable areas by tracking pedestrians represents a significant improvement over existing methods that rely on the assumption that the floor area is directly visible. By leveraging the maturity of human detection technology and the occlusion patterns of stationary objects and pedestrians, our approach can effectively detect traversable areas in crowded scenes, making robot navigation safer and more efficient.

2 Related Works

This section presents most relevant approaches and draws a relation to these previous works of the proposed method. Furthermore, we clearly state the application domain to which the proposed method contributes. In particular, we will focus on work that assumes monocular vision rather than other sensor modalities such as 3D point cloud scanners or stereo cameras. We extensively took advantage of the easy accessibility of monocular cameras for the usage of SLAM system and further explored the achievability of such system.

Traversable region detection problem is closely related to visual reconstruction problems such as structure-from-motion (SfM) [4], simultaneous localization and mapping (SLAM) [3], [11], [12] and visual odometry (VO) [5]. These approaches aim to reconstruct the robot's ego-motion and geometric environment models from the robot's view sequence. Among them, SLAM, a problem setup that sequentially updates environmental geometric models or maps from live view sequences, is most relevant to this work. Existing solutions of SLAM can be divided into filtering-based [13] and optimization-based methods [14], with the latter dominating the current state of the art. While most of existing optimization-based methods assume stationary or semi-stationary environments, new types of SLAM techniques have emerged in recent years that addresses a dynamic environment [12], as we focus on in this study. In particular, our

implementation is inspired by one of the state-of-the-art technologies, DynaSLAM [8], which will be detailed in the subsections described later. However, all of the existing reconstruction techniques only targeted specific traversable regions that are directly visible from the robot. In contrast, our approach aims to predict traversable regions that are not directly visible by using the indirect means of pedestrian observation.

Traversable region detection technology contributes to several important applications of visual navigation such as point goal navigation [15], object goal navigation [16] and coverage path planning [17]. In these applications, representations based on traversable regions are often used for state recognition and action planning. Compared with existing traversable region detection techniques that aim to reconstruct directly visible regions, our approach extends the scope of the reconstruction target from visual regions to even invisible regions which pedestrians are standing on. This brings a number of benefits. First, from the perspective of human-robot collaboration [6], the area where the pedestrian is standing often contains more contextual and semantic information that is important for the task compared to other areas. Second, from the perspective of state recognition and action planning, considering not only the visible area but also the invisible areas enable rich state representation and long-horizon exploration planning [18]. It should be noted that there are other approaches to traversable region detection [19], such as obstacle detection and floor detection [2] and monocular depth regression [1]. These approaches again aim to recognize obstacle-free regions or traversable floor regions that are directly visible to the robot. Therefore, they and our approach are in a complementary relationship.

DynaSLAM

DynaSLAM is a state-of-the-art visual simultaneous localization and mapping (SLAM) algorithm that builds on the foundation of ORB-SLAM2 [20]. One of the primary challenges in dynamic SLAM is distinguishing between camera motion and object motion within the scene. DynaSLAM overcomes this by using a probabilistic model to estimate the likelihood of an object's motion. If an object's motion is inconsistent with the predicted motion based on the camera's motion, it is identified as a dynamic object and removed from the map.

Our approach shares the design philosophy with DynaSLAM in two respects. First, our approach is also modular, with several modules such as static object modeling and dynamic object detection, each of which is connected to the ever-growing cutting-edge research field, including object detection [7], tracking [21], scene parsing [22], reconstruction [3], inpainting [23], occlusion reasoning, and action planning [17]. Second, the module group for stationary object recognition and the module group for dynamic object recognition are clearly separated, which allows us to focus on research and development of the latter module group without modifying the former module group.

However, our approach is different from DynaSLAM in many aspects. In particular, it introduces a new type of observation that is independent of the observations used by DynaSLAM. That is, relative observations of interactions between stationary and dynamic objects are used by our new observation model, in contrast to the absolute metric observations from stationary and dynamic objects used by DynaSLAM.

3 Approach

Our approach leverages the strengths of ORB-SLAM3 [24] and YOLOv7 [7] human detection model to address the challenge of reasoning the occlusion order of humans and objects in crowded environments. ORB-SLAM3 is a leading visual SLAM algorithm that provides real-time tracking of the camera pose and map features. Meanwhile, YOLOv7 is a highly accurate and efficient deep learning-based object detection model, which has superior performance in human detection.

Our framework has three sub-modules: static object modeling, moving object detection, and occlusion order inference. Given these modules, the system localizes the pedestrian region on the floor by sequentially performing static object modeling, moving object detection, and occlusion order inference. The subsections below provide details on each module used.

Static Object Modelling

The generation of 3D point clouds using ORB-SLAM3 [24] is a major advancement in vision and robotics. SLAM involves navigating and mapping unknown environments, and existing approaches can be categorized as filter-based (e.g., EKF [25]) or optimization-based (e.g., Graph-SLAM [26]). Filter-based methods provide real-time estimation updates, while optimization-based methods offer accurate and consistent maps. ORB-SLAM3, the latest version of the ORB-SLAM algorithms [27], is a powerful tool that can be used with various camera types and incorporates new features, including the robust DBoW3 feature descriptor [28]. Its output is a highly accurate point cloud representing the robot's environment, bringing significant progress to SLAM and applications in robotics and computer vision.

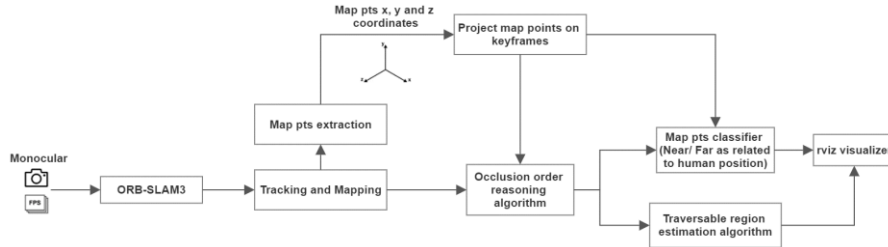


Figure 1: Block diagram of our proposed framework. A segmented video stream in rosbag format is pass into ORB-SLAM3 [24]. The feature points coordinates are extracted and projected onto the keyframes for visualization. By comparing the map pts 3D location/distance to the camera and the location of the human to the camera, we classify the map pts to two simple classes (Near and Far). We can obtain the human estimated location in the 3D map and hence estimate the traversable region. The final outcome can be visualized using rviz visualizer [30].

Dynamic Object Detection

In our study, we utilized two different approaches for pedestrian detection: YOLOv7 [7], a representative object detector and the state-of-the-art model at the time of this writing, and Detectron2, which we used to treat pedestrian detection as a semantic

segmentation problem. While YOLOv7 represents the human region as a bounding box, Detectron2 [9], represents it as a free-form region, allowing for a more accurate approximation of the human domain. Conversely, YOLOv7's bounding box representation offers a high frame rate, making it an attractive option in scenarios where speed is critical.

To compare the performance of these two methods, we conducted experiments to evaluate their accuracy and computational speed in the common field of real-time processing. Our results, which are reported in the experimental section, shed light on the relative strengths and weaknesses of the two approaches. Ultimately, our study provides valuable insights for researchers and practitioners seeking to optimize pedestrian detection performance for real-time applications.

Occlusion Order Inference

Pedestrian tracking is a fundamental task in computer vision that aims to locate and follow people's movement in visual scenes. However, occlusions caused by other humans or objects can pose significant challenges to accurate pedestrian localization. To address this issue, we have developed a technique called occlusion-ordering-based relative-depth estimation which aims to analyze depth values to determine the occlusion order of humans and objects. By leveraging occlusion-ordering information, it becomes possible to discern which person or object appears to be in front of another. Specifically, when one person's occlusion order is consistently greater than that of another person or object, it suggests that the latter is occluded by the former. This method is particularly useful in crowded environments where people and objects are likely to overlap, making it difficult to accurately track individual pedestrians.

The implementation of occlusion-ordering-based relative-depth estimation has practical applications for various fields, including indoor robot navigation, where accurate pedestrian localization is necessary for detecting traversable regions. By improving the accuracy of pedestrian tracking, this approach can enhance the ability of robots to navigate in crowded indoor environments and operate safely and efficiently in these settings.

Implementation Details

We developed a prototype of the proposed system by adding some changes to the public code of ORB-SLAM3 [24]. In particular, we modified some of the original code by adding some functionalities with some built-in ORB-SLAM3 APIs. We make use of ROS nodes to achieve concurrent process of ORB-SLAM3, extraction and projection of feature points and keyframes, occlusion ordering algorithm, and finally visualization on rviz visualizer. The framework is capable of processing data from a moving camera, perform human occlusion ordering algorithm and finally output the result on rviz visualizer concurrently. The traversable region will be clearly shown on the rviz visualizer as well. From Figure 1 above, we can see the flowchart block diagram of the framework. The entire framework consists of two main components: ORB-SLAM3 and occlusion mapping.

ORB-SLAM3 Components:

Feature points and keyframes extraction

The initial stage of this framework includes the ORB-SLAM3 process. We obtain feature points generated by ORB-SLAM3, which are shown as green feature points on the map in this context. In ORB-SLAM3, the feature points are obtained using the Oriented FAST and Rotated BRIEF (ORB) feature detector and descriptor. The ORB feature detector identifies key points in the image based on their high intensity gradients and corner-like structures. Once the key points are detected, the ORB descriptor is computed for each key point. These key points serve as potential feature points, each consisting of coordinate information that can be extracted for our purposes. In the original framework, ORB-SLAM3 also outputs the keyframes information.

To achieve this, we modified the original `ros_mono.cc`, `System.cc` and `System.h` files. In `ros_mono.cc`, a new function is added to save the keyframe trajectory and feature points coordinates. A logic is added to check if a new keyframe is added and save it as images in the image grabber function. Two functions are added to `System.cc` and `System.h`: `GetLastKeyFrameID()` , `SaveMapAndTrajectory()`

Modifications have been made to the image grabber class in `ros_mono.cc` to publish 3D points and 2D projections into a custom ROS topic. The 2D projection features the projection of feature points onto the keyframes.

Projection of feature points onto keyframes

While the coordinate information of the feature points is being extracted, the green feature points are projected onto the keyframes concurrently [10]. A Python code has been developed for this process. Initially, the keyframe pose is represented as a quaternion, so we need to convert it into a 3x3 matrix before passing it to OpenCV [29] which calculates it. We can visualize the projection with rviz [30] visualizer as shown in Figure 2 below. The trajectory and feature points are clearly visible in this visualizer. The green points represent the feature points that are being projected onto the current keyframe, while the red points represent the feature points generated by ORB-SLAM3. The green trajectory represents the camera position trajectory.

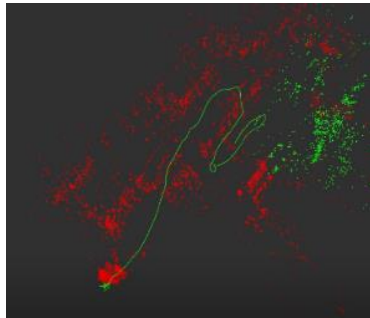


Figure 2: Projected feature points in rviz visualizer

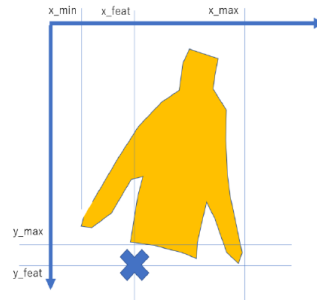


Figure 3: Occlusion ordering algorithm

Occlusion Mapping Components:

In this component, a custom message type called `TrackedFrame.msg` has been created. It will be used by ORB-SLAM3 to publish the required information.

Occlusion ordering algorithm:

A Python code `detect.py` which is modified from the original detectron2 code was used to provide human detection model, Detectron2 human masking and occlusion ordering algorithm. In this code, `TrackedFrame.msg` was read and near/ far points was calculated. The outcome was published to rviz visualizer for visualization purposes. The occlusion ordering algorithm makes use exclusively of the human area range and the feature points coordinates in 2D space. This is simply a 2-class classification problem where we classify the feature points as in front of or behind the human with higher probability.

As shown in Figure 3 above, the 'X' represents the feature points with coordinates (x_feat, y_feat). (x_min, x_max, y_min, y_max) represents the human area range. A simple if-else statement loop can determine and estimate the position of the feature points with higher probability. It could efficiently be integrated into SLAM systems and be used in environments with greatly degraded visual information. We label the feature points in our output video for a better visualization of the algorithm. The cases for the feature points labels are as follows:

d = '-1' when the feature points are in front of the human

d = '1' when the feature points are behind of the human

d = '0' when the feature points are not available/ not in range of the human area.



Figure 4: Feature points occlusion reasoning value annotations



Figure 5: Feature points annotations with occlusion ordering algorithm

As shown in Figure 4 above, the feature points which are in range of the human area are correctly classified as '-1' which means that the feature points are in front of the human. On the other hand, the feature points that are in the x range and out of y range can be classified as '1' which means the feature points are behind the human with a higher probability. In a '0' case, it means that the feature points are out of the human area range and are not available to be classified. We execute the occlusion ordering algorithm and then projects the feature points onto keyframes to find the estimated location of the human in 3D space. Once again, the outcome can be visualized with rviz visualizer with colored feature points. As shown in Figure 5 above, the board that

is behind the human has a lot of feature points built by ORB-SLAM3 [24]. The feature points 3D coordinates are used to be compared with the human location and hence are classified as '1' which means that is behind the human. We can ultimately visualize this with rviz [30] visualizer. For the scene in Figure 5, as shown in Figure 6 below, we can observe that there are points that are colored as Orange and Purple.

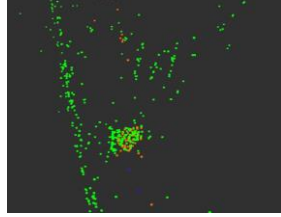


Figure 6: Feature points colored for visualization purpose

The orange cluster is the feature points which are behind the human in this scene. All the points that are labeled as '1' and observed at the current frame are colored as orange and classified as behind the human. On the other hand, feature points that are classified as purple are the points that are in front of the human. On the other hand, if the human is standing behind the board, as shown in Figure 7 and Figure 8 below, the purple cluster would simply represent the feature points of the board, which is in front of the human. The visualized results are as shown in Figure 10 below. By observation, we can conclude that the human is standing right in between the orange and purple cluster.



Figure 7: Human standing behind the board

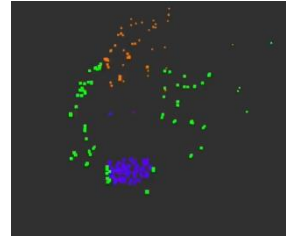


Figure 8: Human standing behind the board visualized

Occupancy Grid Map:



Figure 9: Occupancy Grid map as shown in rviz

We can concurrently display the map points using a 2D occupancy grid map using rviz visualizer as well as shown in Figure 9 above. The map shows a red box indicating the possible human position, and a map of obstacles being marked. The red arrow indicates the current camera point-of-view. The traversable region in the map is marked with grey contour in the visualizer. The grid map is obtained by using the built in OccupancyGrid message type in ROS. The occupancy grid map can hold values from 0 to 100. Close to 0 is considered to be free space and close to 100 is considered to be occupied. Every time an obstacle is marked the value will increase by 10. The concurrent process is achievable with the usage of ROS nodes, and we can continuously update the frames to the visualizer with a moving camera video.

4 Experimental Results

Experimental environment setup:

We first set up the test environment in a wide area to build a map with ORB-SLAM3 [24] as shown in Figure 10 below. A large object, such as a board, is positioned somewhere in the middle to serve as the occlusion object.



Figure 10: Example Scene



Figure 11: Human occluded scenario

The video stream is utilized to track and map the environment, aiming to generate a map with a sufficient number of sparse feature points using ORB-SLAM3. Consequently, the environment needs to be spacious enough to ensure the acquisition of a high-quality map with sparse feature points.

Preparing input video:

As the framework is capable of executing ORB-SLAM3 mapping and the occlusion ordering algorithm at the same time, a video stream featuring a human subject who for the occlusion ordering algorithm is present, as shown in Figure 11 above. The pedestrian performed a natural walking motion without understanding the intention of the experiment. We scan and map the environment as much as possible with the human moving in the scene. The video is then converted to rosbag format, `<video>.bag`.

Running with data obtained:

Once we have the rosbag file ready, we simply start up an ORB-SLAM3 process with the occlusion algorithm and rviz visualizer running concurrently. All the process runs concurrently at once as shown in Figure 12 below.

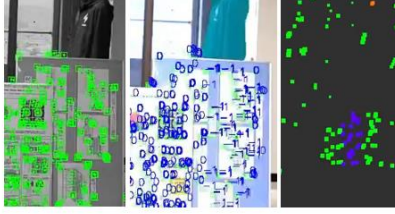


Figure 12: Processes running concurrently



Figure 13: An instance of ground-truth annotation

Studying with annotated ground truth:

We randomly sampled 50 feature points and manually annotated their correct classifications (Behind or in front). By comparing the ground truth with our system's runs, we investigated the algorithm's performance. We randomly selected 5 frames and calculated the average percentage of correctly labeled feature points. All the feature points are annotated relative to the human position, as shown in Figure 13 above. Observation by eyes and manual annotation is carried out to obtain an accurate ground truth. The framework classifies the feature points with different colors in rviz [30] for visualization purposes. As shown in Table 1 below, an outstanding correctness rate of 96.71% has been achieved.

The results are recorded in the Table 1 below:

No	Annotated feature points				%
	GT_Front	GT_Behind	R_Front	R_Behind	
1	2	48	2	48	100.00
2	4	46	4	46	100.00
3	7	43	7	43	100.00
4	35	15	32	18	91.43
5	38	12	35	15	92.11
Average = 96.71					

Table 1: Ground truth vs Actual run performance (Abbreviations: GT – Ground Truth, R – Run)

For ground truth assessment, we randomly selected 5 frames and manually annotated the green feature points. We compared the accuracy of our framework with the ground truth annotations to evaluate its performance. Based on our experiments, the algorithm demonstrates exceptional performance when the human is positioned in front of the occlusion. Our proposed framework exhibits robustness and is deemed to perform well.

Ablation Study:

Table 2 below presents the results of a series of ablation studies that combine the modules in the proposed framework. These studies utilize human-object occlusion ordering datasets to evaluate the performance of the methods employed. The occlusion ordering algorithm is executed using the YOLOv7 [7] human detection model. Traditionally, bounding boxes are applied around the detected objects when using such human detection models. In our case, we can directly use the bounding box edges (x_{min} , x_{max} , y_{min} , y_{max}) to determine the position of the feature points using the algorithm. A detailed explanation is provided in Figure 3 above. Alternatively, we can incorporate the Detectron2 instance segmentation mask for this purpose. It is worth noting that during the ORB-SLAM3 [24] process, unwanted map can occasionally be generated on dynamic objects, which in our case, is the human. To address this issue, we can apply a filter using the same Detectron2 mask on the human area. This filtering method is essentially what we have employed in our proposed framework. We processed 3 different datasets and evaluated their performances using different methods. The correctness ratio represents the percentage of correctly labeled feature points compared to manually annotated ground truths. The method used is the same as shown in Table 1. We randomly sampled 50 feature points from random frames and annotated them through observation. The limitations of the bounding box method in obtaining an accurate contour of the human pose present an obvious flaw in this study. For example, when the human is positioned behind an object, the bounding box may include points that are located at the side of the human contour, leading to a larger inaccuracy in the estimation. On the other hand, utilizing Detectron2 to acquire the contour of the human shape offers a significantly more reliable and accurate method. This approach allows us to filter out unwanted feature points associated with the human and obtain a precise boundary that corresponds to the actual human pose.

	<i>1st</i>	<i>2nd</i>	<i>3rd</i>
Mask with Mask filter (Proposed)	1.0	0.96	0.98
B.Box with Mask filter	0.87	0.79	0.78
B.Box	0.75	0.71	0.73

Table 2: Ablation studies with different module

5 Conclusions

In conclusion, our research introduces a simple yet robust human-object occlusion ordering framework for estimating traversable regions in crowded environments using SLAM systems. By leveraging the position of pedestrians, we can accurately predict traversable areas based on the occlusion order between humans and objects. The framework's simplicity makes it easy to implement and integrate into existing SLAM systems, while still demonstrating remarkable robustness in handling occlusions in complex and dense environments. This has significant implications for various applications, such as robot navigation and autonomous vehicles, where accurately predicting traversable areas is crucial for safe and efficient movement.

References

- [1] A. Bhoi, "Monocular Depth Estimation: A Survey," *arXiv:1901.09402*, p. 8, 2019.
- [2] Y. Xu and P. Ghamisi, "Consistency-Regularized Region-Growing Network for Semantic Segmentation of Urban Scenes with Point-Level Annotations," *IEEE Trans. Image Process*, vol. 31, pp. 5038-5051, 2022.
- [3] C. Cadena, L. Carlone, H. Carrillo, Y. Latif, D. Scaramuzza, J. Neira, I. Reid and J. J. Leonard, "Past, Present, and Future of Simultaneous Localization And Mapping: Towards the Robust-Perception Age," *IEEE Transactions on Robotics*, vol. 32, no. 6, pp. 1309-1332, 2017.
- [4] D. Moran, H. Koslowky, Y. Kasten, H. Maron, M. Galun and R. Basri, "Deep Permutation Equivariant Structure From Motion," *IEEE/CVF International Conference on Computer Vision (ICCV)*, pp. 5976-5986, October 2021.
- [5] S. Wang, R. Clark, H. Wen and N. Trigoni, "DeepVO: Towards End-to-End Visual Odometry with Deep Recurrent Convolutional Neural Networks," *In 2017 IEEE International Conference on Robotics and Automation (ICRA)*, pp. 2043-2050, 2017.
- [6] M. Selvaggio, M. Cagnetti, S. Nikolaidis, S. Ivaldi and B. Siciliano, "Autonomy in Physical Human-Robot Interaction: A Brief Survey," *IEEE Robotics and Automation Letters*, vol. 6, no. 4, pp. 7989-7996, 2021.
- [7] A. Bochkovskiy, H.-Y. M. Liao and C.-Y. Wang, "YOLOv7: Trainable bag-of-freebies sets new state-of-the-art for real-time object detectors," *arXiv*, 2022.
- [8] B. Bescos, J. M. F  cil, J. Civera and J. Neira, "Dynaslam: Tracking, mapping, and inpainting in dynamic scenes," *IEEE Robotics and Automation Letters*, vol. 3, no. 4, pp. 4076-4083, 2018.
- [9] Y. Wu, A. Kirillov, F. Massa, W.-Y. Lo and R. Girshick, "Detectron2," Facebook, 2019. [Online]. Available: <https://github.com/facebookresearch/detectron2>. [Accessed 19 03 2023].
- [10] T. K. Jonathan Tay, "HO3-SLAM: Human-Object Occlusion Ordering as Add-on for Enhancing Traversability Prediction in Dynamic SLAM," in *IEEE ICPRS2023*, Guayaquil, Ecuador, 2023.
- [11] A. J. Davison, I. D. Reid, N. D. Molton and O. Stasse, "MonoSLAM: Real-Time Single Camera SLAM," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 29, no. 6, pp. 1052-1067, 2007.
- [12] L. Zhao, B. Wei, L. Li and X. Li, "A Review of Visual SLAM for Dynamic Objects," *2022 IEEE 17th Conference on Industrial Electronics and Applications (ICIEA)*, pp. 1080-1085, 2022.
- [13] M. Lin, C. Yang and D. Li, "An Improved Transformed Unscented FastSLAM With Adaptive Genetic Resampling," *IEEE Transactions on Industrial*

Electronics, vol. 66, no. 5, pp. 3583-3594, 2019.

- [14] H. Zhou, Z. Hu, S. Liu and S. Khan, "Efficient 2d graph slam for sparse sensing," *iros*, pp. 6404-6411, 2022.
- [15] X. Zhao, H. Agrawal, D. Batra and A. G. Schwing, "The surprising effectiveness of visual odometry techniques for embodied pointgoal navigation," *In Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, pp. 16127-16136, October 2021.
- [16] Y. Zhu, R. Mottaghi, E. Kolve, J. J. Lim, A. Gupta, L. Fei-Fei and A. Farhadi, "Target-driven visual navigation in indoor scenes using deep reinforcement learning," *In 2017 IEEE International Conference on Robotics and Automation (ICRA)*, pp. 3357-3364, 2017.
- [17] D. S. Chaplot, D. Gandhi, S. Gupta, A. Gupta and R. Salakhutdinov, "Learning to explore using active neural slam," 2020.
- [18] C. Hoang, S. Sohn, J. Choi, W. Carvalho and H. Lee, "Successor feature landmarks for long-horizon goal-conditioned reinforcement learning," *Advances in Neural Information Processing Systems*, vol. 34, pp. 269633-26975, 2021.
- [19] D. Stadler and J. Beyerer, "Improving multiple pedestrian tracking by track management and occlusion handling," *In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 10958-10967, June 2021.
- [20] R. Mur-Artal and J. D. Tardos, "ORB-SLAM2: An open-source slam system for monocular, stereo, and rgb-d cameras," *IEEE Transactions on Robotics*, vol. 33, no. 5, pp. 1255-1262, 2017.
- [21] X. Li, C. Ma, B. Wu, Z. He and M.-H. Yang, "Target-aware deep tracking," *In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2019.
- [22] R. Singh, P. Gupta, P. Shenoy and R. Sarvadevabhatla, "Float: Factorized learning of object attributes for improved multi-object multi-part scene parsing," *In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 1445-1455, June 2022.
- [23] D. Kim, S. Woo, J.-Y. Lee and I. S. Kweon, "Deep video inpainting," *In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2019.
- [24] C. Campos, R. Elvira, J. J. G. Rodriguez, J. M. M. Montiel and J. D. Tardos, "ORB-SLAM3: An Accurate Open-Source Library for Visual, Visual-Inertial and Multi-Map SLAM," *IEEE Transactions on Robotics*, vol. 37, no. 6, pp. 1874-1890, 2021.
- [25] R. Agramonte and C. Urrea, "Kalman Filter: Historical Overview and Review of Its Use in Robotics 60 Years after Its Creation," vol. 2021, 2021.
- [26] S. Thrun and M. Montemerlo, "The GraphSLAM Algorithm with Applications to Large-Scale Mapping of Urban Structures," 2006.

- [27] R. Mur-Artal, J. M. M. Montiel and J. D. Tardos, "ORB-SLAM: A Versatile and Accurate Monocular SLAM System," *IEEE Transactions on Robotics*, vol. 31, no. 5, pp. 1147-1163, October 2015.
- [28] R. Muñoz-Salinas, "DBow3," 17 02 2017. [Online]. Available: <https://github.com/rmsalinas/DBow3>. [Accessed 19 03 2023].
- [29] I. Culjak, D. Abram, H. Dzapo, T. Pribanic and M. Cifrek, "A brief introduction to OpenCV," *2012 Proceedings of the 35th International Convention MIPRO, Opatija, Croatia*, pp. 1725-1730, 2012.
- [30] W. Woodal, "rviz," ROS, 16 05 2018. [Online]. Available: <http://wiki.ros.org/rviz>. [Accessed 19 03 2023].