

Pred-NBV: Prediction-guided Next-Best-View Planning for 3D Object Reconstruction

Harnaik Dhami*

Vishnu D. Sharma*

Pratap Tokekar

Abstract—Prediction-based active perception has shown the potential to improve the navigation efficiency and safety of the robot by anticipating the uncertainty in the unknown environment. The existing works for 3D shape prediction make an implicit assumption about the partial observations and therefore cannot be used for real-world planning and do not consider the control effort for next-best-view planning. We present *Pred-NBV*, a realistic object shape reconstruction method consisting of PoinTr-C, an enhanced 3D prediction model trained on the ShapeNet dataset, and an information and control effort-based next-best-view method to address these issues. *Pred-NBV* shows an improvement of 25.46% in object coverage over the traditional methods in the AirSim simulator, and performs better shape completion than PoinTr, the state-of-the-art shape completion model, even on real data obtained from a Velodyne 3D LiDAR mounted on DJI M600 Pro.

I. INTRODUCTION

The goal of this paper is to improve the efficiency of mapping and reconstructing an object of interest with a mobile robot. This is a long studied and fundamental problem in the field of robotics [1]. In particular, the commonly used approach is Next-Best-View (NBV) planning. In NBV planning, the robot seeks to find the best location to go to next and obtain sensory information that will aid in reconstructing the object of interest. A number of approaches for NBV planning have been proposed over the years [2]. In this paper, we show how to leverage the recent improvements in perception due to deep learning to improve the efficiency of 3D object reconstruction with NBV planning. In particular, we present a 3D shape prediction technique that can predict a full 3D model based on the partial views of the object seen so far by the robot to find the NBV. Notably, our method works “in the wild” by eschewing some common assumptions made in 3D shape prediction, namely, assuming that the partial views are still centered at the full object center.

There are several applications where robots are being used for visual data collection. Some examples include inspection for visual defect identification of civil infrastructure such as bridges [3], [4], ship hulls [5] and aeroplanes [6], digital mapping for real estate [7], [8], and precision agriculture [9]. The key reasons why robots are used in such applications is that they can reach regions that are not easily accessible for humans and we can precisely control where the images are taken from. However, existing practices for the most part require humans to specify a nominal trajectory for the robots

that will visually cover the object of interest. Our goal in this paper is to automate this process.

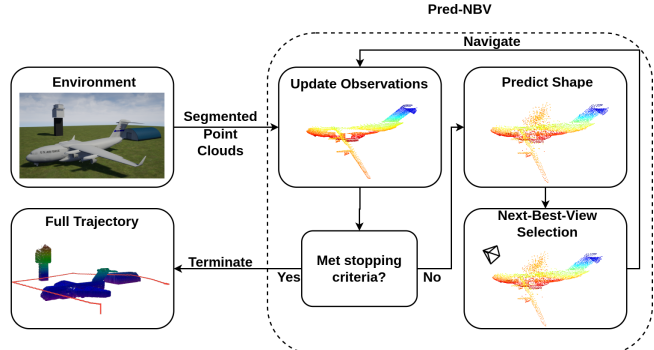


Fig. 1: Overview of the proposed approach

The NBV planning method is the commonly used approach to autonomously decide where to obtain the next measurement from. NBV planning typically uses geometric cues such as symmetry [10] or prior information [11] for deciding the next best location. In this work, we do not rely on such assumptions but instead leverage the predictive power of deep neural networks for 3D shape reconstruction.

Recent works have explored predictions as a way of improving these systems by anticipating the unknowns with prediction and guiding robots’ motion accordingly. This approach has been studied for robot navigation, exploration, and manipulation [12]–[15] with the help of neural network-based methods that learn from datasets. While 2D map presentation works have shown the benefits of robotic tasks in simulation as well as the real-world, similar methods for 3D prediction have been limited to simplistic simulations. The latter approaches generally rely on synthetic datasets due to the lack of realistic counterparts for learning.

Strong reliance on data results in the neural network learning implicit biases. Some models can make predictions only for specific objects [16]. Others require implicit knowledge of the object’s center, despite being partially visible [17]. These situations are invalid in real-world, mapless scenarios and may result in inaccurate shape estimation. Many shape prediction works assume the effortless motion of the robot [18], whereas an optimal path for a robot should include the effort required to reach a position and the potential information gain due to time and power constraints. Monolithic neural networks replacing both the perception and planning components present an alternative, but they tend to be specific to the datasets used for training and may require extensive finetuning for real-world deployment. Lack

*Equal contribution. Names are listed alphabetically.

Authors are with the Department of Computer Science, University of Maryland, U.S.A. {dhami, vishnu, tokekar}@umd.edu.

This work is supported by the National Science Foundation under grant number 1943368 and ONR under grant number N00014-18-1-2829.

of transparency in such networks poses another challenge that could be critical for the safe operation of a robot when working alongside humans.

To make 3D predictive planning more realistic, efficient, and safe, we propose a method consisting of a 3D point cloud completion model, that relaxes the assumption about implicit knowledge of the object’s center using a curriculum learning framework [19], and an NBV framework, that maximizes the information gain from image rendering and minimizes the distance traveled by the robot. Furthermore, our approach is modular making it interpretable and easy to upgrade.

We make the following contributions to this work:

- We use curriculum learning to build an improved 3D point cloud completion model, which does not require the partial point cloud to be centered at the full point cloud’s center and is more robust to perturbations than earlier models. We show that this model, termed *PoinTr-C*, outperforms the base model, *PoinTr* [17], by at least **23.06%** and show qualitative comparison on ShapeNet [20] dataset, and real point cloud obtained with a Velodyne 3D LiDAR mounted on DJI M600 Pro.
- We propose a next-best-view planning approach that performs object reconstruction without any prior information about the geometry, using predictions to optimize information gain and control effort over a range of objects in a model-agnostic fashion.
- We show that our method covers on average **25.46%** more points on all models evaluated for object reconstruction in AirSim [21] simulations compared to the non-predictive baseline approach, *Basic-Next-Best-View* [22] and performs even better for complex structures such as airplanes.

We share the qualitative results, project code, and visualization from our method on our project website¹.

II. RELATED WORK

Active reconstruction in an unknown environment can be accomplished through next-best-view (NBV) planning, which has been studied by the robotics community for a long time [23]. In this approach, the robot builds a partial model of the environment based on observations and then moves to a new location to maximize the cumulative information gained. The NBV approaches can be broadly classified into information-theoretic and geometric methods. The former builds a probabilistic occupancy map from the observations and uses the information-theoretic measure [2] to select the NBV. The latter assumes the partial information to be exact and determines the NBV based on geometric measures [24].

The existing works on NBV with robots focus heavily on information-theoretic approaches for exploration in 2D and 3D environments [25], [26]. Subsequent development for NBV with frontier and tree-based approaches was also designed for exploration by moving the robot towards unknown regions [27]–[30]. Prior works on NBV for object reconstruction also rely heavily on information-theoretic approaches

to reduce uncertainty in pre-defined closed spaces [31], [32]. Geometric approaches require knowing the model of the object in some form and thus have not been explored to a similar extent. Such existing works try to infer the object geometry from a database or as an unknown closed shape [33], [34], and thus may be limited in application.

In recent years, prediction-based approaches have emerged as another solution. One body of these approaches works in conjunction with other exploration techniques to improve exploration efficiency by learning to predict structures in the environment from a partial observation. This is accomplished by learning the common structures in the environment (buildings and furniture, for example) from large datasets. This approach has recently gained traction and has been shown to work well for mobile robots navigation with 2D occupancy map representations [12], [14], [35], exploration [13], high-speed maneuvers [36], and elevation mapping [37].

Similar works on 3D representations have focused mainly on prediction modules. Works along this line have proposed generating 3D models from novel views using single RGB image input [38], depth images [39], normalized digital surface models (nDSM) [40], point clouds [17], [41], [42], etc. The focus of these works is solely on inferring shapes based on huge datasets of 3D point clouds [20]. They do not discuss the downstream task of planning. A key gap in these works is that they assume a canonical representation of the object, such as the center of the whole object, to be provided either explicitly or implicitly. Relaxing this assumption does not work well in the real-world where the center of the object may not be estimated accurately, discouraging the adoption of 3D prediction models for prediction-driven planning.

Another school of work using 3D predictions combines the perception and planning modules as a neural network. These works, aimed at predicting the NBV to guide the robot from partial observations, were developed for simple objects [43], 3D house models [44], and a variety of objects [45] ranging from remotes to rockets. Peralta et al. [44] propose a reinforcement-learning framework, which can be difficult to implement due to sampling complexity issues. The supervised-learning approach proposed by Zeng et al. [45] predicts the NBV using a partial point cloud, but the candidate locations must lie on a sphere around the object, restricting the robot’s planning space. Moreover, monolithic neural networks suffer from a lack of transparency and real-world deployment may require extensive fine-tuning of the hyperparameters. Prediction-based modular approaches solve these problems as the intermediate outputs are available for interpretation and the prediction model can be plugged in with the preferred planning method for a real environment.

A significant contribution of our work is to relax the implicit assumption used in many works that the center and the canonical orientation of the object under consideration are known beforehand, even if the 3D shape completion framework uses partial information as the input. A realistic inspection system may not know this information and thus the existing works may not be practically deployable.

¹Project webpage: <http://raaslab.org/projects/PredNBV/>

III. PROBLEM FORMULATION

We are given a robot with a 3D sensor onboard that explores a closed object with volume $\mathcal{V} \in \mathbb{R}^3$. The set of points on the surface of the object is denoted by $\mathcal{S} \in \mathbb{R}^3$. The robot can move in free space around the object and observe its surface. The surface of the object $s_i \subset \mathcal{S}$ perceived by the 3D sensor from the pose $\phi_i \subset \Phi$

is represented as a voxel-filtered point cloud. We define the relationship between the set of points observed from a view-point ϕ_i with a function f , i.e., $s_i = f(\phi_i)$. The robot can traverse a trajectory ξ that consists of view-points $\{\phi_1, \phi_2, \dots, \phi_m\}$. The surface observed over a trajectory is the union of surface points observed from the consisting viewpoints, i.e. $s_\xi = \bigcup_{\phi \in \xi} f(\phi)$. The distance traversed by the robot between two view-points ϕ_i and ϕ_j is denoted by $d(\phi_i, \phi_j)$.

Our objective is to find a trajectory ξ_i from the set of all possible trajectories Ξ , such that it observes the whole voxel-filtered surface of the object while minimizing the distance traversed.

$$\xi^* = \arg \min_{\xi \in \Xi} \sum_{i=1}^{|\xi|-1} d(\phi_i, \phi_{i+1}), \text{ such that } \bigcup_{\phi_i \in \xi} f(\phi_i) = \mathcal{S}. \quad (1)$$

In unseen environments, $\mathcal{S}$ is not known apriori, hence the optimal trajectory can not be determined. We assume that the robot starts with a view of the object. If not, we can always first explore the environment until the object of interest becomes visible.

IV. PROPOSED APPROACH

We propose *Pred-NBV*, a prediction-guided NBV method for 3D object reconstruction highlighted in Fig. 1. Our method consists of two key modules: (1) *PoinTr-C*, a robust 3D prediction model that completes the point cloud using only partial observations, and (2) an NBV framework that uses prediction-based information gain to reduce the control effort for active object reconstruction. We provide the details in the following subsection.

A. PoinTr-C: 3D Shape Completion Network

Given the current set of observations $v_o \in \mathcal{V}$, we predict the complete volume using a learning-based predictor g , i.e., $\hat{\mathcal{V}} = g(v_o)$.

To obtain $\hat{\mathcal{V}}$, we use PoinTr [17], a transformer-based architecture that uses 3D point clouds as the input and output.

PoinTr uses multiple types of machine learning methods to perform shape completion. It first identifies the geometric relationship in low resolution between points in the cloud by clustering. Then it generates features around the cluster centers, which are then fed to a transformer [46] to capture the long-range relationships and predict the centers for the missing point cloud. Finally, a coarse-to-fine transformation over the predicted centers using a neural network outputs the missing point cloud.

This model was trained on the ShapeNet [20] dataset and outperforms the previous methods on a range of objects.

However, PoinTr was trained with implicit knowledge of the center of the object. Moving the partially observed point cloud to its center results in incorrect prediction from PoinTr.

To improve predictions, we fine-tune PoinTr using a curriculum framework, which dictates training the network over easy to hard tasks by increasing the difficulty in steps during learning [19]. Specifically, we fine-tune PoinTr over increasing perturbations in rotation and translation to the canonical representation of the object to relax the assumption about implicit knowledge of the object's center. We use successive rotation-translation pairs of $(25^\circ, 0.0)$, $(25^\circ, 0.1)$, $(45^\circ, 0.1)$, $(45^\circ, 0.25)$, $(45^\circ, 0.5)$, $(90^\circ, 0.5)$, $(180^\circ, 0.5)$, and $(360^\circ, 0.5)$ for curriculum training. We assume that the object point cloud is segmented well, which can be achieved using distance-based filters or segmentation networks.

B. Next-Best View Planner

Given the predicted point cloud $\hat{\mathcal{V}}$ for the robot after traversing the trajectory ξ_t , we generate a set of candidate poses $\mathcal{C} = \{\phi_1, \phi_2, \dots, \phi_m\}$ around the object observed so far. Given v_o , the observations so far, we define the objective to select the shortest path that results in observing at least $\tau\%$ of the maximum possible information gain over all the candidate poses. Considering $\hat{\mathcal{V}}$ as an exact model, we use a geometric measure to quantify the information gained from the candidate poses. Specifically, we define a projection function $I(\xi)$, over the trajectory ξ , which first identifies the predicted points distinct from the observed point cloud over the trajectory, then apply a hidden point removal operator on them [47], without reconstructing a surface or estimating a normal, and lastly, find the number of points that will be observed if we render an image on the robot's camera. Thus, we find the NBV from the candidate set $\mathcal{C}$ as follows:

$$\phi_{t+1} = \arg \min_{\phi \in \mathcal{C}} d(\phi, \phi_t), \text{ such that } \frac{I(\xi_t \cup \phi)}{\max_{\phi \in \mathcal{C}} I(\xi_t \cup \phi)} \geq \tau.$$

We find the $d(\phi_i, \phi_j)$, using RRT-Connect [48], which incrementally builds two rapidly-exploring random trees rooted at ϕ_i and ϕ_j through the observed space to provide a safe trajectory. After selecting the NBV, the robot follows this trajectory to reach the prescribed view-point. We repeat the prediction and planning process until the ratio of observations in the previous step and current step is 0.95 or higher.

To generate the candidate set $\mathcal{C}$, we first find the distance d_{max} of the point farthest from the center of the predicted point cloud $\hat{\mathcal{V}}$ and z-range. Then, we generate candidate poses on three concentric circles: one centered at $\hat{\mathcal{V}}$ with radius $1.5 \times d_{max}$ at steps of 30° , and one $0.25 \times$ z-range above and below with radius $1.2 \times d_{max}$ at steps of 30° . We use $\tau = 0.95$ for all our experiments.

V. EVALUATION

In this section we evaluate the *Pred-NBV* pipeline. We start with a qualitative example followed by a comparison of the individual modules against respective baseline methods. The results show that *Pred-NBV* is able to outperform the baselines significantly using large-scale models from the Shapenet [20] dataset and with real-world 3D LiDAR data.

A. Qualitative Example

Fig. 2 show the reconstructed point cloud of a C17 airplane and the path followed by a UAV in AirSim [21]. We create candidate poses on three concentric rings at different heights around the center of the partially observed point cloud. The candidate poses change as more of the object is visible.

As shown in Fig. 4, *Pred-NBV* observes more points than the NBV planner without prediction in the same time budget.

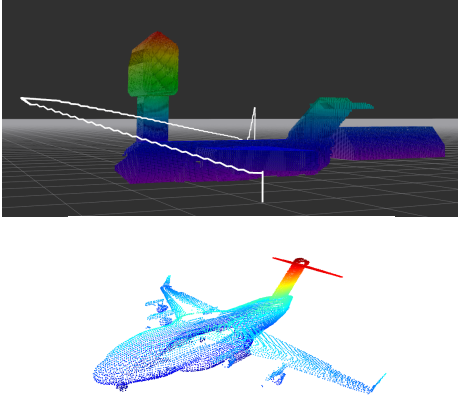


Fig. 2: Flight path and total observations of C17 Airplane after running our NBV planner in AirSim simulation.

B. 3D Shape Prediction

1) *Setup*: We train *PoinTr-C* on a 32-core, 2.10Ghz Xeon Silver-4208 CPU and Nvidia GeForce RTX 2080Ti GPU with 11GB of memory. The network is fine-tuned over the ShapeNet [20] dataset, trained with perturbation as described in Section IV. Similar to PoinTr [17], we use Chamfer distance (CD) and Earth Mover’s Distance (EMD), permutation-invariant metrics suggested by Fan et al. [49], as the loss function for training *PoinTr-C*. For evaluation we use two versions of Chamfer distance: $CD-l_1$ and $CD-l_2$, which use L1 and L2-norm, respectively, to calculate the distance between two sets of points, and F-score which quantifies the percentage of points reconstructed correctly.

2) *Results*: Table I summarizes our findings regarding the effect of perturbations. *PoinTr-C* outperforms the baseline in both scenarios. It only falters in $CD-l_2$ in the ideal condition, i.e., no augmentation. Furthermore, *PoinTr-C* doesn’t undergo large changes in the presence of augmentations, making it more robust than PoinTr. The relative improvement for *PoinTr-C* at least 23.05% (F-Score).

TABLE I: Comparison between the baseline model (PoinTr) and *PoinTr-C* over test data with and without perturbation. Arrows show if a higher ($\uparrow$) or a lower ($\downarrow$) value is better.

Perturbation	Approach	F-Score $\uparrow$	$CD-l_1$ $\downarrow$	$CD-l_2$ $\downarrow$
$\times$	PoinTr [17]	0.497	11.621	0.577
	<i>PoinTr-C</i>	0.550	10.024	0.651
$\checkmark$	PoinTr [17]	0.436	16.464	1.717
	<i>PoinTr-C</i>	0.550	10.236	0.717

We provide a qualitative comparison of the predictions from the two models for various objects under perturbations on our project webpage. Fig. 3 shows the results for a real

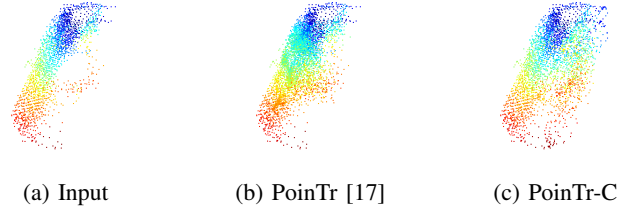


Fig. 3: Results over the real-world point cloud of a car obtained using LiDAR (Interactive figure available on our webpage).

point cloud of a car (visualized from the camera above the left headlight) obtained with a Velodyne LiDAR sensor. The results show the predictions from PoinTr are scattered around the center of the visible point cloud, whereas *PoinTr-C* makes more realistic predictions. Our webpage provides interactive visualizations of these point clouds for further inspection.

C. Next-Best-View Planning

1) *Setup*: We use Robot Operating System (ROS) Melodic and AirSim [21] on Ubuntu 18.04 to carry out the simulations. We equipped the virtual UAV with a RGBD camera. AirSim’s built-in image segmentation is used to segment out the target object from the rest of the environment. We created a ROS package to publish a segmented depth image containing pixels only belonging to the bridge based on the RGB camera segmentation. This segmented depth image was then converted to a point cloud. We use the MoveIt [50] software package based on the work done by Köse [51] to implement the RRT connect algorithm. MoveIt uses RRT connect and the environmental 3D occupancy grid to find collision-free paths for point-to-point navigation.

2) *Qualitative Example*: We evaluate *Pred-NBV* on 20 objects from 5 ShapeNet classes: airplane, rocket, tower, train, and watercraft. We selected these classes as they represent larger shapes suitable for inspection. Fig. 1 shows the path followed by the UAV using *Pred-NBV* for the C-17 airplane simulation. There are non-target obstacles in the environment, such as a hangar and air traffic control tower. *Pred-NBV* finds a collision-free path that selects viewpoints targeting maximum coverage of the airplane.

3) *Comparison with Baseline*: We compare the performance of *Pred-NBV* with a baseline NBV method [22]. The baseline selects poses based on frontiers in the observed space using occupancy grids. We modified the baseline to improve it for our application and make it comparable to *Pred-NBV*. The modifications include using our segmentation for the occupancy grid so that frontiers are weighted towards the target object. We also set the orientation of the selected poses towards the center of the target object similar to how *Pred-NBV* works. The algorithms had the same stopping criteria as *Pred-NBV*.

We see in Table II that our method observes on average 25.46% more points than the baseline for object reconstruction across multiple models from various classes. In Fig. 4, we show that *Pred-NBV* observes more points per step than the baseline while not flying further per each step.

TABLE II: Points observed by *Pred-NBV* and the baseline NBV method [22] for all models in AirSim.

Class	Model	Points Seen <i>Pred-NBV</i>	Points Seen Baseline	Improvement
Airplane	747	11922	9657	20.99%
	A340	8603	5238	48.62%
	C-17	12916	7277	55.85%
	C-130	9900	7929	22.11%
	Fokker 100	10192	9100	11.32%
Rocket	Atlas	1822	1722	5.64%
	Maverick	2873	2643	8.34%
	Saturn V	1111	807	31.70%
	Sparrow	1785	1639	8.53%
	V2	1264	1086	15.15%
Tower	Big Ben	4119	3340	20.89%
	Church	2965	2588	13.58%
	Clock	2660	1947	30.95%
	Pylon	3181	2479	24.80%
	Silo	5674	3459	48.51%
Train	Diesel	3421	3161	7.90%
	Mountain	4545	4222	7.37%
Watercraft	Cruise	4733	3522	29.34%
	Patrol	3957	2306	52.72%
	Yacht	9499	6016	44.90%

VI. CONCLUSION

We propose a realistic and efficient planning approach for robotic inspection using learning-based predictions. Our approach fills the gap between the existing works and the realistic setting by proposing a curriculum-learning-based point cloud prediction model, and a distance and information gain aware inspection planner for efficient operation. Our analysis shows that our approach is able to outperform the baseline approach in observing the object surface by 25.46%. Furthermore, we show that our predictive model is able to provide satisfactory results for real-world point cloud data. We believe the modular design paves the path to further improvement by enhancement of the constituents.

In this work we use noise free observations, but show that *Pred-NBV* has the potential to work well on real, noisy inputs with pre-processing. In future work, we will explore making the prediction network robust to noisy inputs and with implicit filtering capabilities. We used a geometric measure for NBV in this work, and will extend it to information-theoretic measures using ensemble of predictions and uncertainty extraction techniques in future work.

REFERENCES

- [1] R. Bajcsy, Y. Aloimonos, and J. K. Tsotsos, "Revisiting active perception," *Autonomous Robots*, vol. 42, pp. 177–196, 2018.
- [2] J. Delmerico, S. Isler, R. Sabzevari, and D. Scaramuzza, "A comparison of volumetric information gain metrics for active 3d object reconstruction," *Autonomous Robots*, vol. 42, no. 2, pp. 197–208, 2018.
- [3] P. Shanthakumar, K. Yu, M. Singh, J. Orevillo, E. Bianchi, M. Hebdon, and P. Tokekar, "View planning and navigation algorithms for autonomous bridge inspection with uavs," in *International Symposium on Experimental Robotics (ISER)*, 2018, accepted.
- [4] H. Dhami, K. Yu, T. Williams, V. Vajipey, and P. Tokekar, "GATSBI: An online GTSP-based algorithm for targeted surface bridge inspection," in *2023 International Conference on Unmanned Aircraft Systems (ICUAS)*. IEEE, jun 2023. [Online]. Available: <https://doi.org/10.1109/2Ficuas57906.2023.10156013>

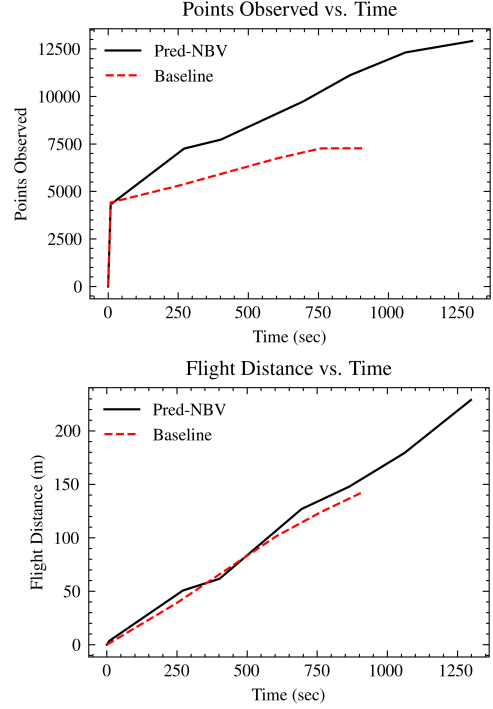


Fig. 4: Comparison between **Pred-NBV** and the **baseline NBV algorithm** [22] for a C-17 airplane.

- [5] A. Kim and R. Eustice, "Pose-graph visual slam with geometric model selection for autonomous underwater ship hull inspection," in *2009 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2009, pp. 1559–1565.
- [6] L. Ropek, "This startup is using drones to conduct aircraft inspections," Mar 2021. [Online]. Available: <https://gizmodo.com/this-startup-is-using-drones-to-conduct-airplane-inspec-1846527002>
- [7] T. Zhou, R. Tucker, J. Flynn, G. Fyffe, and N. Snavely, "Stereo magnification: Learning view synthesis using multiplane images," *ACM Trans. Graph. (Proc. SIGGRAPH)*, vol. 37, 2018. [Online]. Available: <https://arxiv.org/abs/1805.09817>
- [8] S. K. Ramakrishnan, A. Gokaslan, E. Wijmans, O. Maksymets, A. Clegg, J. M. Turner, E. Undersander, W. Galuba, A. Westbury, A. X. Chang, M. Savva, Y. Zhao, and D. Batra, "Habitat-matterport 3d dataset (HM3d): 1000 large-scale 3d environments for embodied AI," in *Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks Track (Round 2)*, 2021. [Online]. Available: <https://openreview.net/forum?id=v4OuqNs5P>
- [9] H. Dhami, K. Yu, T. Xu, Q. Zhu, K. Dhakal, J. Friel, S. Li, and P. Tokekar, "Crop height and plot estimation for phenotyping from unmanned aerial vehicles using 3d lidar," in *Proceedings of the IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2020.
- [10] P. E. Debevec, C. J. Taylor, and J. Malik, "Modeling and rendering architecture from photographs: A hybrid geometry-and image-based approach," in *Proceedings of the 23rd annual conference on Computer graphics and interactive techniques*, 1996, pp. 11–20.
- [11] M. Breyer, L. Ott, R. Siegwart, and J. J. Chung, "Closed-loop next-best-view planning for target-driven grasping," in *2022 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. IEEE, 2022, pp. 1411–1416.
- [12] S. K. Ramakrishnan, Z. Al-Halah, and K. Grauman, "Occupancy anticipation for efficient exploration and navigation," in *Computer Vision—ECCV 2020: 16th European Conference, Glasgow, UK, August 23–28, 2020, Proceedings, Part V 16*. Springer, 2020, pp. 400–418.
- [13] G. Georgakis, B. Bucher, A. Arapin, K. Schmeckpeper, N. Matni, and K. Daniilidis, "Uncertainty-driven planner for exploration and navigation," in *2022 International Conference on Robotics and Automation (ICRA)*, 2022, pp. 11 295–11 302.
- [14] M. Wei, D. Lee, V. Isler, and D. Lee, "Occupancy map inpainting for online robot navigation," in *2021 IEEE International Conference on Robotics and Automation (ICRA)*, 2021, pp. 8551–8557.

- [15] X. Yan, M. Khansari, J. Hsu, Y. Gong, Y. Bai, S. Pirk, and H. Lee, "Data-efficient learning for sim-to-real robotic grasping using deep point cloud prediction networks," *arXiv preprint arXiv:1906.08989*, 2019.
- [16] B. Yang, S. Rosa, A. Markham, N. Trigoni, and H. Wen, "Dense 3d object reconstruction from a single depth view," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 41, no. 12, pp. 2820–2834, 2019.
- [17] X. Yu, Y. Rao, Z. Wang, Z. Liu, J. Lu, and J. Zhou, "PointR: Diverse point cloud completion with geometry-aware transformers," in *Proceedings of the IEEE/CVF international conference on computer vision*, 2021, pp. 12498–12507.
- [18] Z. Wu, S. Song, A. Khosla, F. Yu, L. Zhang, X. Tang, and J. Xiao, "3d shapenets: A deep representation for volumetric shapes," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2015, pp. 1912–1920.
- [19] Y. Bengio, J. Louradour, R. Collobert, and J. Weston, "Curriculum learning," in *Proceedings of the 26th annual international conference on machine learning*, 2009, pp. 41–48.
- [20] A. X. Chang, T. Funkhouser, L. Guibas, P. Hanrahan, Q. Huang, Z. Li, S. Savarese, M. Savva, S. Song, H. Su, J. Xiao, L. Yi, and F. Yu, "ShapeNet: An Information-Rich 3D Model Repository," Stanford University — Princeton University — Toyota Technological Institute at Chicago, Tech. Rep. arXiv:1512.03012 [cs.GR], 2015.
- [21] S. Shah, D. Dey, C. Lovett, and A. Kapoor, "Airsim: High-fidelity visual and physical simulation for autonomous vehicles," in *Field and Service Robotics*, M. Hutter and R. Siegwart, Eds. Cham: Springer International Publishing, 2018, pp. 621–635.
- [22] J. Aleotti, D. L. Rizzini, R. Monica, and S. Caselli, "Global registration of mid-range 3d observations and short range next best views," in *2014 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. IEEE, 2014, pp. 3668–3675.
- [23] W. R. Scott, G. Roth, and J.-F. Rivest, "View planning for automated three-dimensional object reconstruction and inspection," *ACM Computing Surveys (CSUR)*, vol. 35, no. 1, pp. 64–96, 2003.
- [24] K. A. Tarabanis, P. K. Allen, and R. Y. Tsai, "A survey of sensor planning in computer vision," *IEEE transactions on Robotics and Automation*, vol. 11, no. 1, pp. 86–104, 1995.
- [25] B. Kuipers and Y.-T. Byun, "A robot exploration and mapping strategy based on a semantic hierarchy of spatial representations," *Robotics and autonomous systems*, vol. 8, no. 1–2, pp. 47–63, 1991.
- [26] J. I. Vázquez-Gómez, L. E. Sucar, R. Murrieta-Cid, and E. López-Damian, "Volumetric next-best-view planning for 3d object reconstruction with positioning error," *International Journal of Advanced Robotic Systems*, vol. 11, no. 10, p. 159, 2014.
- [27] B. Yamauchi, "A frontier-based approach for autonomous exploration," in *Proceedings 1997 IEEE International Symposium on Computational Intelligence in Robotics and Automation CIRA'97. Towards New Computational Principles for Robotics and Automation*. IEEE, 1997, pp. 146–151.
- [28] H. H. González-Banos and J.-C. Latombe, "Navigation strategies for exploring indoor environments," *The International Journal of Robotics Research*, vol. 21, no. 10–11, pp. 829–848, 2002.
- [29] B. Adler, J. Xiao, and J. Zhang, "Autonomous exploration of urban environments using unmanned aerial vehicles," *Journal of Field Robotics*, vol. 31, no. 6, pp. 912–939, 2014.
- [30] A. Bircher, M. Kamel, K. Alexis, O. Oleynikova, and R. Siegwart, "Receding horizon path planning for 3d exploration and surface inspection," *Autonomous Robots*, vol. 42, no. 2, pp. 291–306, 2018.
- [31] K. Morooka, H. Zha, and T. Hasegawa, "Next best viewpoint (nbv) planning for active object modeling based on a learning-by-showing approach," in *Proceedings. Fourteenth International Conference on Pattern Recognition (Cat. No. 98EX170)*, vol. 1. IEEE, 1998, pp. 677–681.
- [32] J. I. Vázquez-Gómez, E. López-Damian, and L. E. Sucar, "View planning for 3d object reconstruction," in *2009 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. IEEE, 2009, pp. 4015–4020.
- [33] J. E. Banta, L. Wong, C. Dumont, and M. A. Abidi, "A next-best-view system for autonomous 3-d object reconstruction," *IEEE Transactions on Systems, Man, and Cybernetics-Part A: Systems and Humans*, vol. 30, no. 5, pp. 589–598, 2000.
- [34] S. Kriegel, M. Brucker, Z.-C. Marton, T. Bodenmüller, and M. Suppa, "Combining object modeling and recognition for active scene exploration," in *2013 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2013, pp. 2384–2391.
- [35] V. D. Sharma, J. Chen, and P. Tokekar, "Proxmap: Proximal occupancy map prediction for efficient indoor robot navigation," in *2023 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. IEEE, 2023, (in press).
- [36] K. D. Katyal, A. Poleyov, J. Moore, C. Knuth, and K. M. Popek, "High-speed robot navigation using predicted occupancy maps," in *2021 IEEE International Conference on Robotics and Automation (ICRA)*, 2021, pp. 5476–5482.
- [37] B. Yang, Q. Zhang, R. Geng, L. Wang, and M. Liu, "Real-time neural dense elevation mapping for urban terrain with uncertainty estimations," *IEEE Robotics and Automation Letters*, vol. 8, no. 2, pp. 696–703, 2022.
- [38] N. Häni, S. Engin, J.-J. Chao, and V. Isler, "Continuous object representation networks: Novel view synthesis without target view supervision," in *Proceedings of the 34th International Conference on Neural Information Processing Systems*, ser. NIPS'20. Red Hook, NY, USA: Curran Associates Inc., 2020.
- [39] B. Yang, S. Rosa, A. Markham, N. Trigoni, and H. Wen, "Dense 3d object reconstruction from a single depth view," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 41, no. 12, pp. 2820–2834, dec 2019. [Online]. Available: <https://doi.org/10.1109/2Ftpami.2018.2868195>
- [40] F. Alidoost, H. Arefi, and F. Tombari, "2d image-to-3d model: Knowledge-based 3d building reconstruction (3dbr) using single aerial images and convolutional neural networks (cnns)," *Remote Sensing*, vol. 11, no. 19, 2019. [Online]. Available: <https://www.mdpi.com/2072-4292/11/19/2219>
- [41] H. Xie, H. Yao, S. Zhou, J. Mao, S. Zhang, and W. Sun, "Grnet: Gridding residual network for dense point cloud completion," in *Computer Vision—ECCV 2020: 16th European Conference, Glasgow, UK, August 23–28, 2020, Proceedings, Part IX*. Springer, 2020, pp. 365–381.
- [42] W. Yuan, T. Khot, D. Held, C. Mertz, and M. Hebert, "Pcn: Point completion network," in *2018 international conference on 3D vision (3DV)*. IEEE, 2018, pp. 728–737.
- [43] A. Pop and L. Tamas, "Next best view estimation for volumetric information gain," *IFAC-PapersOnLine*, vol. 55, no. 15, pp. 160–165, 2022, 6th IFAC Conference on Intelligent Control and Automation Sciences/CONS 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S2405896322010369>
- [44] D. Peralta, J. Casimiro, A. M. Nilles, J. A. Aguilar, R. Atienza, and R. Cajote, "Next-best view policy for 3d reconstruction," in *Computer Vision – ECCV 2020 Workshops: Glasgow, UK, August 23–28, 2020, Proceedings, Part IV*. Berlin, Heidelberg: Springer-Verlag, 2020, p. 558–573. [Online]. Available: https://doi.org/10.1007/978-3-030-66823-5_33
- [45] R. Zeng, W. Zhao, and Y.-J. Liu, "Pc-nbv: A point cloud based deep network for efficient next best view planning," in *2020 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2020, pp. 7050–7057.
- [46] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, Ł. Kaiser, and I. Polosukhin, "Attention is all you need," *Advances in neural information processing systems*, vol. 30, 2017.
- [47] S. Katz, A. Tal, and R. Basri, "Direct visibility of point sets," *ACM Trans. Graph.*, vol. 26, no. 3, p. 24-es, jul 2007. [Online]. Available: <https://doi.org/10.1145/1276377.1276407>
- [48] J. J. Kuffner and S. M. LaValle, "Rrt-connect: An efficient approach to single-query path planning," in *Proceedings 2000 ICRA. Millennium Conference. IEEE International Conference on Robotics and Automation. Symposia Proceedings (Cat. No. 00CH37065)*, vol. 2. IEEE, 2000, pp. 995–1001.
- [49] H. Fan, H. Su, and L. J. Guibas, "A point set generation network for 3d object reconstruction from a single image," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2017, pp. 605–613.
- [50] D. Coleman, I. Sucan, S. Chitta, and N. Correll, "Reducing the barrier to entry of complex robotic software: a moveit! case study," *arXiv preprint arXiv:1404.3785*, 2014.
- [51] T. Köse, "tahsinkose/hector-moveit: Hector quadrotor with moveit! motion planning framework," <https://github.com/tahsinkose/hector-moveit>, 2018.