

Autonomous Visual Assistance for Robot Operations Using a Tethered UAV



Xuesu Xiao, Jan Dufek, and Robin R. Murphy

1 Introduction

Tele-operated robots are still widely used in Dangerous, Dirty, and Dull (DDD) environments where human presence is extremely difficult or impossible, due to those environments' mission-critical task execution and current technological limitations. Projecting human presence in remote environments is still an effective approach to leverage current technologies and actual field demand. However, one major challenge of tele-operation is the insufficient situational awareness of the remote field, caused by the onboard sensing limitations, such as relatively stationary and limited field of view and lack of depth perception from the robot's onboard camera. Therefore, the emerging state of the practice for nuclear operations, bomb squad, disaster robots, and other domains with novel tasks or highly occluded environments is to use two robots, a primary and a secondary that acts as a visual assistant to overcome the perceptual limitations of the sensors by providing an external viewpoint.

However, the usage of tele-operated visual assistants also causes problems: it requires an extra human operator, or even an operating crew, to tele-operate the secondary visual assistant. Human operators also tend to choose suboptimal viewpoints based on experience only. Most importantly, communication between the two operators of the primary and secondary robots requires extra teamwork demand, in

X. Xiao (✉)

Department of Computer Science, University of Texas at Austin,
Austin, USA

e-mail: xiao@cs.utexas.ed

J. Dufek · R. R. Murphy

Department of Computer Science and Engineering, Texas A&M University,
College Station, USA

e-mail: dufek@tamu.edu

R. R. Murphy

e-mail: robin.r.murphy@tamu.edu

© Springer Nature Singapore Pte Ltd. 2021

G. Ishigami and K. Yoshida (eds.), *Field and Service Robotics*, Springer Proceedings
in Advanced Robotics 16, https://doi.org/10.1007/978-981-15-9460-1_2

addition to the task and perceptual demands of the tele-operation. In this research, an autonomous tethered aerial visual assistant is developed to replace the secondary robot and its operator, reducing the human-robot ratio from 2:2 to 1:2. The co-robots team will then consist of one tele-operated primary ground robot, one autonomous aerial visual assistant, and one human operator, whose situational awareness is maintained by the visual feedback streamed from a series of optimal viewpoints for the particular tele-operation task.

Unmanned Ground Vehicles (UGVs) are stable, reliable, durable, and can thus represent humans to actuate upon the real world, while Unmanned Aerial Vehicles (UAVs) have superior mobility and workspace coverage and therefore are capable of providing enhanced situational awareness [7]. Researchers have looked into utilizing the advantages and avoiding the disadvantages by teaming up the two types of robots [1, 3]. A more relevant area was to use a UAV to augment the UGV's perception or assist UGV's task execution, such as "an eye in the sky" for UGV localization [2], providing stationary third-person view for construction machines [6], improving navigation in case of GPS loss [5], and UGV control with UAV's visual feedback [10, 16] using differential flatness [9]. However, instead of prior works' flight path execution in wide-open space or hovering at a stationary and elevated viewpoint, our aerial visual assistant needs to navigate through unstructured or confined spaces in order to provide visual assistance to the UGV operator from a series of good viewpoints. It is able to reason about the motion execution risk in complex environments and plan a path that provides good visual assistance. In particular, this work uses a tethered UAV, with the purpose of matching its battery duration with UGV's and as a fail-safe in case of malfunction in mission-critical tasks. Our tethered visual assistant utilizes the advantages and mitigates the disadvantages of the tether, in terms of tether-based indoor localization, motion primitives, and environment contact planning.

The remainder of this article is organized as follows: Sect. 2 presents the heterogeneous co-robots team. The high-level visual assistance components are described in Sect. 3, while low-level tether-based motion implementation in Sect. 4. System demonstrations are provided in Sect. 5. Section 6 concludes the paper.

2 Co-Robots Team

This section presents the co-robots team: a tele-operated ground primary robot, an autonomous tethered aerial visual assistant, and a human operator of the primary robot under the visual assistance of the aerial vehicle (Fig. 1).

2.1 Tele-Operated Ground Primary Robot

In the co-robots team, the primary robot is a tele-operated Endeavor PackBot 510 (Fig. 1 upper left). PackBot has a chassis with two main differential treads that allow

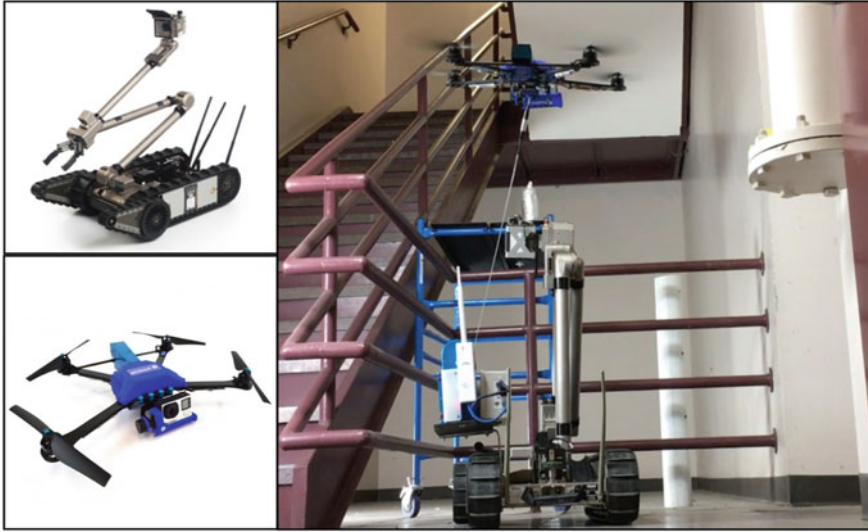


Fig. 1 The Co-Robots Team: tele-operated primary robot, Endeavor PackBot 510 (upper left), and autonomous tethered visual assistant, Fotokite Pro (lower left), picking up a sensor and dropping it into a radiation pipe in a confined staircase (right)

zero radius turn and maximum speed up to 9.3 km/h. Two articulated flippers with treads are used to climb over obstacles or stairs (up to 40°). PackBot's three-link manipulator locates on the top of the chassis, with an articulated gripper on the second link and an onboard camera on the third. The manipulator can lift 5 kg at full extension and 20 kg close in. Motor encoders on the arm provide a precise position of the articulated joints, including the gripper, the default visual assistance point of interest. Four onboard cameras provide first-person-views, but are all limited to the robot body. On the chassis, a Velodyne Puck LiDAR constantly scans the 3-D environments, providing the map for the co-robots team to navigate through. The map does not necessarily need to be global, with the unknown parts being assumed as obstacles. Four BB-2590 batteries provide up to 8 hrs run time.

2.2 Autonomous Aerial Visual Assistant

A tethered UAV, Fotokite Pro, is used as the autonomous aerial visual assistant (Fig. 1 lower left). It could be deployed from a landing platform mounted on the ground robot's chassis. The UAV is equipped with an onboard camera with a 2-DoF gimbal (pitch and roll). The camera's yaw is controlled dependently by the vehicular yaw. The main purpose of the tether is to match the run time of the aerial visual assistant with that of the ground primary robot, since flight power could be transmitted via the

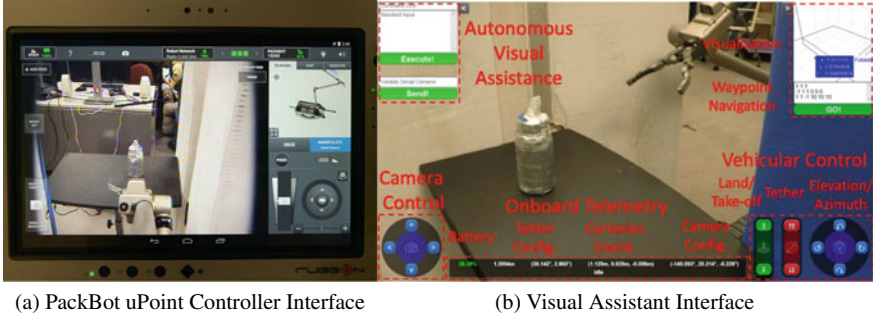


Fig. 2 Interfaces with the human operator

tether. Additionally, the tether serves as a fail-safe in mission-critical environments. The UAV's flight controller is based on the tether sensory feedback, including the tether length, azimuth, and elevation angles. The six-dimensional coverage of the workspace makes the UAV suitable for the visual assistance purpose.

2.3 Human Operator

The human operator tele-operations the primary ground robot with the visual assistance of the UAV. In addition to the default PackBot uPoint controller with onboard first-person-view, the visual feedback from the visual assistant's onboard camera is also available for enhanced situational awareness. For example, the visual assistant could move to a location perpendicular to the tele-operation action, providing extra depth perception to the operator. The visual assistant could be either manually controlled or automated. For the focus of this research, autonomous visual assistance, a 3-D map is provided by the primary robot's LiDAR, and a risk-aware path is planned using a pre-established viewpoint quality map (discussed in the following sections). The uPoint tele-operation and visual assistance interfaces are shown in Fig. 2.

3 Visual Assistance Components

This section introduces the key components of autonomous visual assistance, including a viewpoint quality map based on the cognitive science concept of affordances, an explicit path risk representation with a focus on unstructured or confined environments, and a planning framework to balance the trade-off between reward (viewpoint quality) and risk (motion execution).

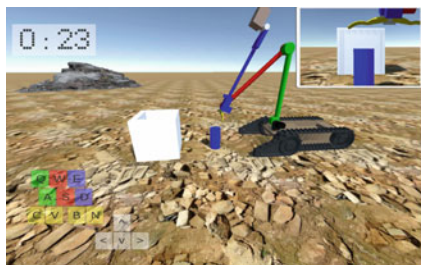
3.1 Viewpoint Quality Reward

The cognitive science concept of affordances is used to determine viewpoint quality, where the potential for an action can be directly perceived without knowing intent or models, and thus is universal to all robots and tasks [8]. For this work, four affordances are used: manipulability, passability, reachability, and traversability (Fig. 3). In order to determine the viewpoint quality (reward) for each affordance, we use a computer simulation to collect performance data with 30 professional PackBot operators. A hemisphere centered at each affordance is created, with 30 viewpoints evenly distributed on it. The 30 viewpoints are divided into five groups: left, right, front, back, and above the affordance. For each affordance, every test subject is randomly given one viewpoint within each of the five groups and tries to finish the tele-operation task in an error-free and fast manner. The number of errors, such as colliding with the wall for passability or falling off the ledge for traversability, and the completion time are recorded. The average value of the product of error and time collected by all subjects is the metric to reflect the viewpoint quality. Given any point in the 3-D space, its viewpoint reward is assigned to be the viewpoint quality of the closest point on the hemisphere. This study is still ongoing and its results will be reported in a separate paper.

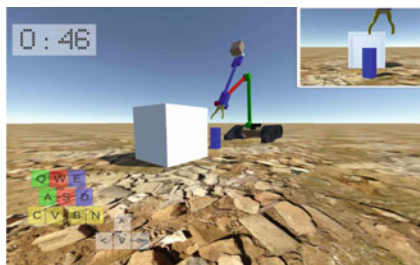
3.2 Explicit Risk Representation

In contrast to the traditional state-dependent risk representation or probabilistic uncertainty modeling, this work uses an explicit risk representation as a function of the entire path. The workspace $\mathbb{W}$ of the robot could be constructed by an occupancy grid from the LiDAR, excluding a set of obstacles $\mathbb{OB} = \{ob_i | i = 1, 2, \dots, o\}$, where o is the number of obstacles. Given a start location of the visual assistant S , a simple path P could be defined as $P = \{s_0, s_1, \dots, s_n\}$ where s_i denotes the i th state on the path P while $s_0 = S, \forall 1 \leq i \leq n, 1 \leq j \leq o, s_i \cap ob_j = \emptyset$, and $\forall i \neq j, 1 \leq i, j \leq n, s_i \neq s_j$. In a conventional state-dependent risk representation, risk at state s_i is defined based on a function mapping from one state to a risk index $r : s_i \mapsto ri$ and the risk of executing a path P is a simple summation of all individual states $risk(P) = \sum_{i=1}^n r(s_i)$. In the proposed path-dependent risk representation, however, risk at state s_i can be evaluated by not only the state, but also the path leading to s_i , $P_i = \{s_0, s_1, \dots, s_i\}$ and the risk at s_i is computed through the mapping $R : (s_0, s_2, \dots, s_i) \mapsto ri$. The path-level risk is relaxed from the simple summation to a more general representation $risk(P) = risk(s_0, s_1, \dots, s_i)$.

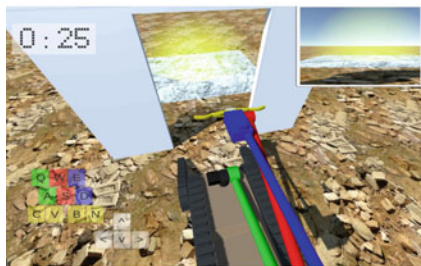
Our explicit path risk representation does not exclude traditional state-dependent risk elements. Those risk functions $r : s_i \mapsto ri$ include risk elements caused by the distance to the closest obstacles $r_{di} = dis(s_i)$, visibility from isovist lines (Fig. 4 upper left) $r_{vi} = vis(s_i)$, altitude due to propeller ceiling or ground effect $r_{ai} = alt(s_i)$, and tether singularity $r_{si} = sig(s_i)$. Those risk elements are additive along the path. Path-



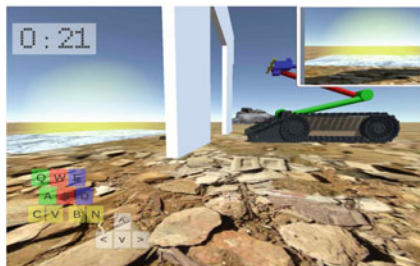
(a) Good Manipulability Viewpoint



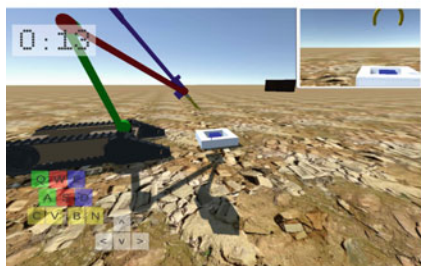
(b) Bad Manipulability Viewpoint



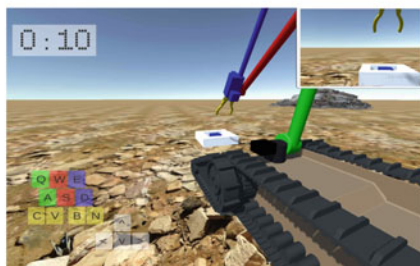
(c) Good Passability Viewpoint



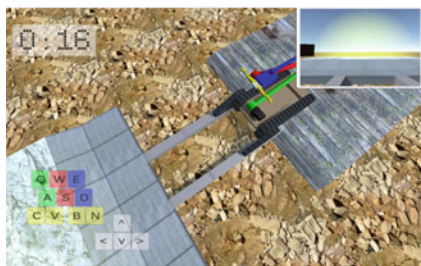
(d) Bad Passability Viewpoint



(e) Good Reachability Viewpoint



(f) Bad Reachability Viewpoint



(g) Good Traversability Viewpoint



(h) Bad Traversability Viewpoint

Fig. 3 Ongoing viewpoint quality study in simulation with professional PackBot operators

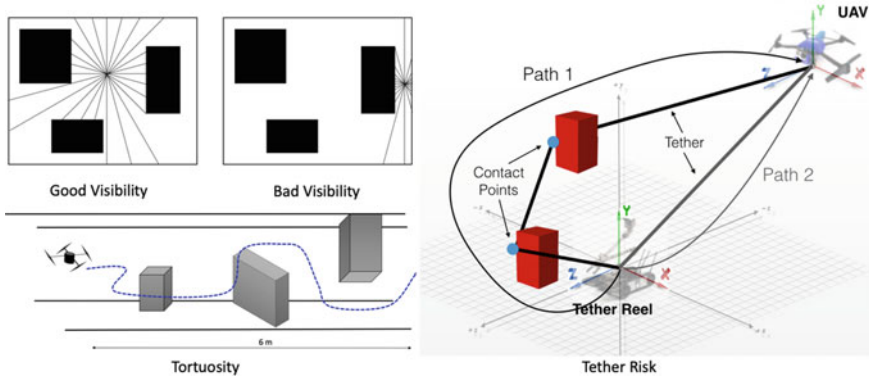


Fig. 4 Examples of risk elements: visibility (upper left), tortuosity (lower left), tether length, number of contacts, and azimuth (right)

dependent risk elements include action length, access element, tortuosity, number of tether contacts, tether length, and azimuth (Table 1). In order to define these risk elements, we further define the action between two consecutive states s_{i-1} and s_i to be A_i . So the whole sequence of actions to execute the entire path P is defined as $\mathbb{A} = \{A_1, A_2, \dots, A_n\}$. For action length, $\|A_i\|$ denotes the length of the executing action A_i . For access elements, the function AE evaluates the difficulty of entering from the void where s_{i-1} is located to the void of s_i . The only positive difficulty is added to the risk index. Tortuosity characterizes the number of “turns” necessary to reach the state. More generally speaking, this is the difference by some measurement between two consecutive actions $\|A_i - A_{i-1}\|$. Tether length is a function of the entire path, e.g., taking path 1 and path 2 in Fig. 4 right will have a completely different tether length. The number of contact points and azimuth angle are also different. Risk index should never decrease with the execution of a path, which is guaranteed by the norm and max operations for the first three elements in Table 1. Thus, they only need to be evaluated once for each path (unitary). For the last three, however, the risk associated with each state may decrease, e.g., contact points may be relaxed [15] and tether length may decrease. But this does not cancel the previous risk. Therefore, those three elements need to be added to all states. Given a path P , its execution risk could be evaluated based on each risk element. Weighted sum or fuzzy logic could be used to combine all elements into one total risk index, quantifying the difficulty of executing that path. Detailed information on the explicit risk representation could be found in [13].

Table 1 Path-dependent risk elements

Risk element	Risk index	Property
Action length	$R_{AL}(P) = \sum_{i=1}^n \ A_i\ $	Unitary
Access element	$R_{AE}(P) = \sum_{i=2}^n \max(AE(void(s_{i-1}), void(s_i)), 0)$	Unitary
Tortuosity	$R_T(P) = \sum_{i=2}^n \ A_i - A_{i-1}\ $	Unitary
Tether length	$R_{TL}(P) = f(s_0, s_1, \dots, s_n)$	Additive
Number of contacts	$R_{NC}(P) = g(s_0, s_1, \dots, s_n)$	Additive
Azimuth	$R_A(P) = h(s_0, s_1, \dots, s_n)$	Additive

3.3 Risk-Aware Reward-Maximizing Planning

Given a viewpoint quality map as a reward and motion execution risk as a function of the path, the risk-aware reward-maximizing planner plans a minimum-risk path to each state [14], evaluates the reward collected, and then picks the one with optimal utility value, defined as the ratio between the reward and the risk. Executing the optimal utility path approximates the optimal visual assistance behavior.

4 Tethered Motion

With a high-level risk-aware optimal utility path, this section presents a low-level motion suite to realize the path on the tethered aerial visual assistant.

4.1 Tether-Based Indoor Localization

Our aerial visual assistant uses its tether to localize in GPS-denied environments. The sensory input is the tether length L , elevation angle θ , and azimuth angle ϕ . The mechanics model M in [17] corrects the preliminary localization model under taut and straight tether assumption (Fig. 5a) using the Free Body Diagram (FBD) of the UAV (Fig. 5b) and tether (Fig. 5c) in order to achieve accurate localization of the airframe $M : (\theta, \phi, L) \mapsto (x, y, z)$, from tether sensory input to Cartesian space location.

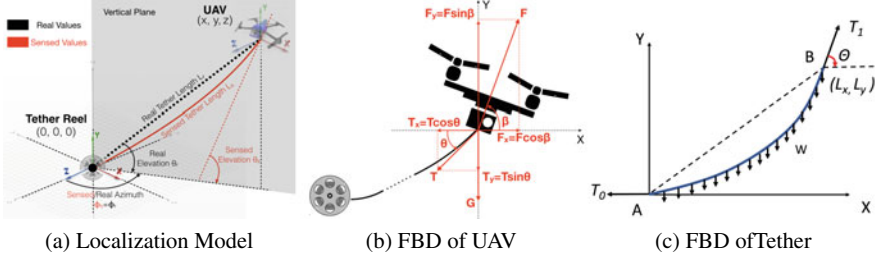


Fig. 5 Tether-based localization [17]

4.2 Motion Primitives

Two types of motion primitives are used, which map the waypoints in Cartesian space into tether-based motion commands: position control uses the inverse transformation from polar to Cartesian coordinates and three independent PID controllers to drive the position of L , θ , and ϕ to their desired values (Eq. 1). On the other hand, velocity control based on the system's inverse Jacobian matrix computes velocity commands L' , θ' , and ϕ' using an instantaneous velocity vector pointing from the current location to the next waypoint $d\vec{X}/dt$ (Eq. 2). The vehicular yaw and camera pitch and roll reactively point at the center of the affordance along the entire path.

$$\begin{cases} L = \sqrt{x^2 + y^2 + z^2} \\ \theta = \arcsin \frac{y}{\sqrt{x^2 + y^2 + z^2}} \\ \phi = a \tan 2\left(\frac{x}{z}\right) \end{cases} \quad (1)$$

$$\begin{pmatrix} \frac{dx}{dt} \\ \frac{dy}{dt} \\ \frac{dz}{dt} \end{pmatrix} = \begin{pmatrix} \cos \theta \sin \phi & -L \sin \theta \sin \phi & L \cos \theta \cos \phi \\ \sin \theta & L \cos \theta & 0 \\ \cos \theta \cos \phi & -L \sin \theta \cos \phi & -L \cos \theta \sin \phi \end{pmatrix} \begin{pmatrix} L' \\ \theta' \\ \phi' \end{pmatrix} \quad (2)$$

The vehicular yaw and camera gimbal pitch are controlled using the 3-D vehicular position localization and the 3-D Cartesian coordinates of the visual assistance point of interest. The camera gimbal roll is passively controlled to align with gravity so that the operator's viewpoint is in level with the ground. Therefore, the visual assistant's camera is pointing toward the point of interest along the entire motion sequence [4, 11]. Xiao et al. [12] report detailed benchmarking results of the motion primitives.

4.3 Tether Contact Planning

In the case when some good viewpoints are located behind an obstacle and the UAV cannot reach with a straight tether, contact points of the tether with the environment

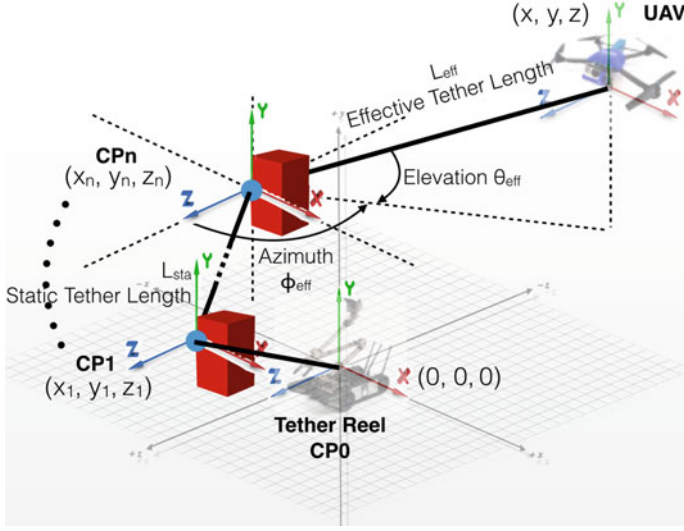


Fig. 6 Motion execution with contact(s) planning and relaxation: given multiple contact points along the tether, static tether length denotes the portion of the tether that wraps around the obstacles (Eq. 3), while the effective length is the last moving segment (Eq. 4). Effective elevation and azimuth angles (Eq. 4) are with respect to the last contact point (CP_n), instead of the tether reel (CP_0)

are necessary. The tether contact point(s) planning and relaxation framework in [15], which allows the UAV to fly as if it were tetherless, is implemented on the tethered visual assistant. A new contact is planned when the current contact is no longer within the line-of-sight of the UAV, while the current contact is relaxed when the last contact becomes visible again. Figure 6 shows the motion execution with multiple contact points (CPs).

$$L_{sta} = \sum_{i=0}^{n-1} \sqrt{(x_{i+1} - x_i)^2 + (y_{i+1} - y_i)^2 + (z_{i+1} - z_i)^2} \quad (3)$$

$$\begin{cases} L_{eff} = \sqrt{(x - x_n)^2 + (y - y_n)^2 + (z - z_n)^2} \\ \theta_{eff} = \arcsin\left(\frac{y - y_n}{\sqrt{(x - x_n)^2 + (y - y_n)^2 + (z - z_n)^2}}\right) \\ \phi_{eff} = \arctan 2\left(\frac{x - x_n}{z - z_n}\right) \end{cases} \quad (4)$$

5 System Demonstration

This section presents two system demonstrations in both indoor and outdoor environments and shows the enhanced situational awareness of the operator achieved by the visual assistance.

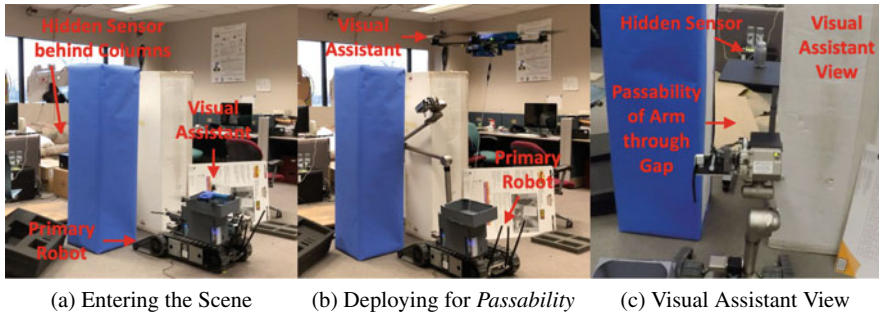


Fig. 7 Visual assistance for *Passability*

5.1 Indoor Test

In this demonstration, the co-robots team drives into a cluttered indoor environment, with the aim of retrieving a hidden sensor. The ground robot is tele-operated and creates a map of the environment. The entry points to the sensor are all blocked by the clutter, leaving the only retrieval option through the narrow gap between the two columns (shown in blue and white in Fig. 7a). Based on the viewpoint quality for *passability*, the visual assistant takes off and deploys to a viewpoint from behind and above to help perceive arm *passability* through a gap (Fig. 7b). The visual assistant view is shown in Fig. 7c, where the relative location of the arm to the narrow gap along with the hidden sensor is well perceived.

After the arm passes through the gap, the visual assistant switches to assist with *manipulability*. Good viewpoints for *manipulability* are located at the side of the gripper. After balancing the viewpoint quality reward and motion execution risk, the planner finds a goal and a path leading to it, which contains two tether contact points with the obstacles. Figure 8a shows the obstacles (red), inflated space for UAV flight tolerance (yellow), waypoints on the planned path (purple), and two contact points on the obstacles (green). The tether configuration is illustrated with black lines. The actual deployment is shown in Fig. 8b. The onboard camera view on the left of Fig. 8c completely misses the depth perception. With this onboard view alone, the risk of not reaching or even knocking off the sensor is high. This lack of depth perception is compensated by the visual assistant view (right).

5.2 Outdoor Test

The co-robots team is also deployed in an outdoor disaster environment, Disaster City[®] Prop 133 in College Station, Texas (Fig. 9). The environment simulates a collapsed multi-story building and the mission for the co-robots team is to navigate into the building and search for victims and threats in two stranded cars.

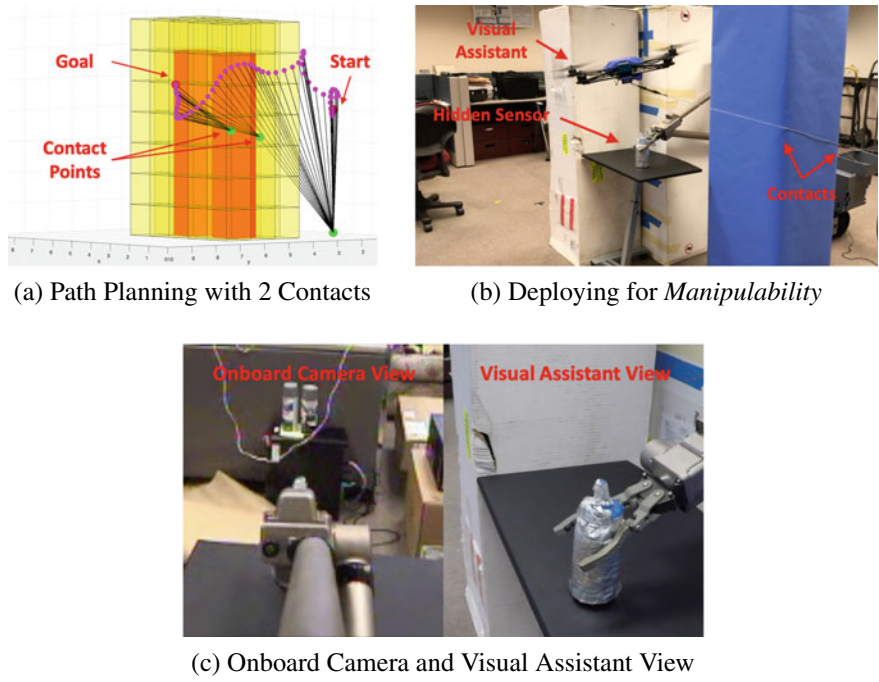


Fig. 8 Visual assistance for *Manipulability*

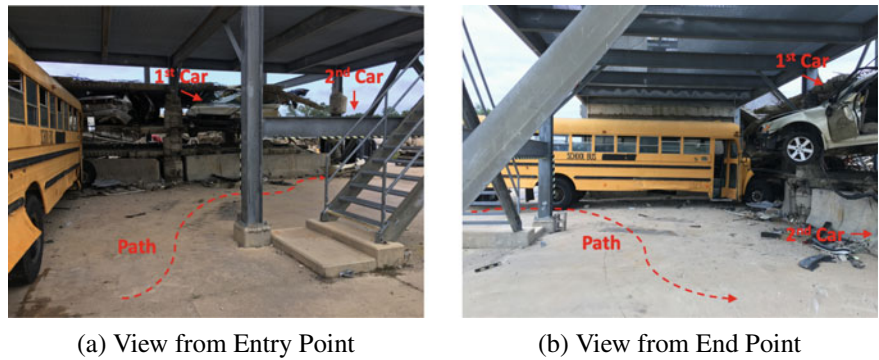
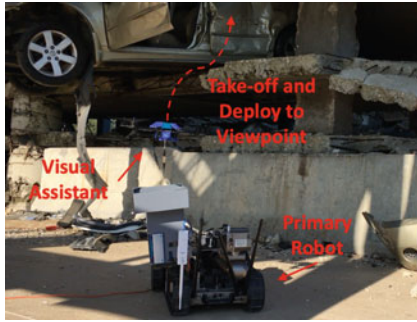


Fig. 9 Disaster City® Prop 133



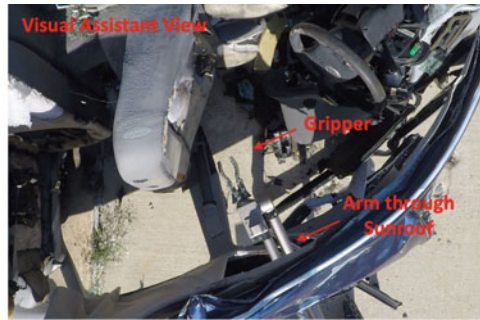
(a) Take-off and Deployment



(b) No Victim in 1st Car

Fig. 10 Enhanced coverage through visual assistance for the first car

(a) Inspection for 2nd Car



(b) Assisting Insertion Depth Perception

Fig. 11 Car inspection through open sunroof for the second car

After reaching the first stranded car, which was on the second floor but is now squeezed down by the collapse, the primary robot's onboard camera is not able to reach the height to search for victims inside the car. The visual assistant takes off from the landing platform and autonomously navigates to a manually specified viewpoint to look inside the car (Fig. 10a). Through the elevated viewpoint provided by the visual assistant, it is confirmed that no victim is trapped in the first car. The visual assistant lands back on the primary robot and the team is tele-operated to the second car.

The second car was tipped over during the collapse, with its sunroof open on the side. The operator intends to insert the manipulator arm into the interior for a thorough search and retrieval if necessary. For safety reasons, the goal is not automatically selected, but manually specified above the side window (Fig. 11a). Looking down through the side window, the depth of the arm insertion into the car interior is clearly visible. No victim or hazardous material exists in the car. The visual assistant lands, the co-robots team finishes the mission and navigates back.

6 Conclusions

We present a co-robots team equipped with autonomous visual assistance for robot tele-operations in unstructured or confined environments using a tethered UAV. The tele-operated primary ground robot projects human presence to remote environments, while the autonomous visual assistant provides enhanced situational awareness to the human operator. The autonomy is realized through a formal study on viewpoint quality, an explicit risk representation to quantify the difficulty of path execution, and a planner that balances the trade-off between viewpoint quality reward and motion execution risk. With the help of a low-level motion suite, including tether-based localization, motion primitives, and contact(s) planning, the high-level path is implemented on a tethered aerial visual assistant given the existence of obstacles. The co-robots team is deployed in both indoor and outdoor search and rescue scenarios, as a proof of the concept of the system. Future work will focus on quantitatively measuring the performance of the co-robots team, including the reward collected, risk encountered, flight accuracy of the autonomous visual assistant, and the improvement in the tele-operation of the primary robot.

Acknowledgements This work is supported by NSF 1637955, NRI: A Collaborative Visual Assistant for Robot Operations in Unstructured or Confined Environments.

References

1. Chaimowicz, L., Cowley, A., Gomez-Ibanez, D., Grocholsky, B., Hsieh, M., Hsu, H., Keller, J., Kumar, V., Swaminathan, R., Taylor, C.: Deploying air-ground multi-robot teams in urban environments. In: *Multi-Robot Systems. From Swarms to Intelligent Automata*, vol. III, pp. 223–234. Springer (2005)
2. Chaimowicz, L., Grocholsky, B., Keller, J.F., Kumar, V., Taylor, C.J.: Experiments in multirobot air-ground coordination. In: *IEEE International Conference on Robotics and Automation*, 2004, Proceedings, ICRA'04, vol. 4, pp. 4053–4058. IEEE (2004)
3. Cheung, C., Grocholsky, B.: UAV-UGV collaboration with a PackBot UGV and Raven SUAV for pursuit and tracking of a dynamic target. In: *Unmanned systems technology X*, vol. 6962, p. 696216. International Society for Optics and Photonics (2008)
4. Dufek, J., Xiao, X., Murphy, R.: Visual pose stabilization of tethered small unmanned aerial system to assist drowning victim recovery. In: *2017 IEEE International Symposium on Safety, Security and Rescue Robotics (SSRR)*, pp. 116–122. IEEE (2017)
5. Frietsch, N., Meister, O., Schlaile, C., Trommer, G.: Teaming of an UGV with a VTOL-UAV in urban environments. In: *2008 IEEE/ION Position, Location and Navigation Symposium*, pp. 1278–1285. IEEE (2008)
6. Kiribayashi, S., Yakushigawa, K., Nagatani, K.: Design and development of tether-powered multirotor micro unmanned aerial vehicle system for remote-controlled construction machine. In: *Field and Service Robotics*, pp. 637–648. Springer (2018)
7. Murphy, R., Dufek, J., Sarmiento, T., Wilde, G., Xiao, X., Braun, J., Mullen, L., Smith, R., Allred, S., Adams, J., et al.: Two case studies and gaps analysis of flood assessment for emergency management with small unmanned aerial systems. In: *2016 IEEE International Symposium on Safety, Security, and Rescue Robotics (SSRR)*, pp. 54–61. IEEE (2016)

8. Murphy, T.B.: Apprehending remote affordances: assessing human sensor systems and their ability to understand a distant environment. Ph.D. Thesis, The Ohio State University (2013)
9. Rao, R., Kumar, V., Taylor, C.: Visual servoing of a UGV from a UAV using differential flatness. In: Proceedings 2003 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS 2003) (Cat. No. 03CH37453), vol. 1, pp. 743–748. IEEE (2003)
10. Xiao, X., Cappel, E., Zhen, W., Dai, J., Sun, K., Gong, C., Travers, M.J., Choset, H.: Locomotive reduction for snake robots. In: 2015 IEEE International Conference on Robotics and Automation (ICRA), pp. 3735–3740. IEEE (2015)
11. Xiao, X., Dufek, J., Murphy, R.: Visual servoing for teleoperation using a tethered UAV. In: 2017 IEEE International Symposium on Safety, Security and Rescue Robotics (SSRR), pp. 147–152. IEEE (2017)
12. Xiao, X., Dufek, J., Murphy, R.: Benchmarking tether-based UAV motion primitives. In: 2019 IEEE International Symposium on Safety, Security, and Rescue Robotics (SSRR). IEEE (2019)
13. Xiao, X., Dufek, J., Murphy, R.: Explicit motion risk representation. In: 2019 IEEE International Symposium on Safety, Security, and Rescue Robotics (SSRR). IEEE (2019)
14. Xiao, X., Dufek, J., Murphy, R.: Explicit-risk-aware path planning with reward maximization (2019). [arXiv:1903.03187](https://arxiv.org/abs/1903.03187)
15. Xiao, X., Dufek, J., Suhail, M., Murphy, R.: Motion planning for a UAV with a straight or kinked tether. In: 2018 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS), pp. 8486–8492. IEEE (2018)
16. Xiao, X., Dufek, J., Woodbury, T., Murphy, R.: UAV assisted USV visual navigation for marine mass casualty incident response. In: 2017 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS), pp. 6105–6110. IEEE (2017)
17. Xiao, X., Fan, Y., Dufek, J., Murphy, R.: Indoor UAV localization using a tether. In: 2018 IEEE International Symposium on Safety, Security, and Rescue Robotics (SSRR), pp. 1–6. IEEE (2018)