

Control and Feedback Design for Teleoperation of Continuum Robots in Obstacle-Constrained and Partially Observable Environments

Tonia Mielke^{1,2}, Chengnan Shentu¹, Jessica Burgner-Kahrs¹

Abstract—Effective teleoperation of continuum robots requires control and feedback strategies that remain robust under obstacle constraints and limited visibility. However, most prior studies evaluate simplified free-space scenarios with full visual access to the robot. This work presents a user study (N = 21) comparing joint-space and task-space control, an additional external view conveying robot shape and position, and visual and vibrotactile contact force feedback under obstacle-constrained and partially observable task conditions. Results show that providing explicit visual information about robot shape and position yields the largest improvement, significantly reducing task completion time and workload while increasing performance. Joint-space control enables faster and less demanding interaction than task-space control in constrained environments. Although contact force feedback does not significantly improve objective performance, it is preferred by most participants and perceived as supportive for motion understanding. These findings highlight the critical role of configuration observability in continuum robot teleoperation and inform the design of control and visualization strategies for realistic applications.

Index Terms—Soft Robot Applications, Modeling, Control, and Learning for Soft Robots.

I. INTRODUCTION

CONTINUUM robots bend continuously along their length, enabling navigation through narrow and curved environments, facilitating tasks that are challenging or infeasible for rigid-link robots [1], [2]. Control of continuum robots is commonly achieved through teleoperation, where operators directly command the robot's motion [1], [2]. Therefore, effective robot control requires user interfaces that enable intuitive and efficient control. However, interaction remains challenging due to the complex kinematics and the frequent lack of direct view of robot shape and environment [3], [4].

Prior work has shown that teleoperation performance depends on factors such as the control method [5], [6] and the input modality [5], [7]. In addition, feedback strategies have



Fig. 1: User controlling the continuum robot simulator with a gamepad controller. The screen displays the user interface, including two views of the application task.

been explored, including visualization of the robot's shape [4], feedback on contact forces [8], or auditory cues [9]. However, most studies consider simplified free space scenarios without environmental interaction and assume that the entire robot body remains continuously visible to the teleoperator [5], [10].

In practical applications, the robot is often partially occluded, and contact with the environment alters its kinematics, thereby complicating the mapping between operator input and robot motion. These factors limit the transferability of existing findings to realistic tasks and highlight the need for evaluation under more representative conditions. To address this gap, this work presents a systematic investigation of control methods and feedback designs for teleoperating continuum robots under realistic constraints (see Fig. 1).

II. RELATED WORK

Control methods for continuum robot teleoperation can broadly be categorized into task-space and configuration-space approaches. Task-space methods command desired end-effector motion, typically using haptic devices, joysticks, 3D mice, or gamepads [5], [7], [11], relying on inverse kinematics or local Jacobian mappings to generate the corresponding actuation commands. In contrast, configuration-space approaches aim to provide more direct control over the robot's shape. These include drawing desired backbone curves [12], mapping the robot configuration to arm or finger bending [13], [14], or employing kinematically similar physical master devices [10], [15]–[17]. Gesture-based techniques further enable users to manipulate feature points or individual segments through pre-defined interactions [18]–[20].

Several studies have directly compared task-space and joint-space control. Majewicz et al. [6] evaluated both approaches

Manuscript received: March 9, 2026; Revised: June 2026; Accepted: August 9, 2026. This article was recommended for publication by Editor Y.-L. Park upon evaluation of the Associate Editor and Reviewers' comments. The work was supported by the European Regional Development Fund under the operation number 'ZS/2023/12/182010' as part of the initiative "Sachsen-Anhalt WISSENSCHAFT Schwerpunkte". The authors acknowledge the funding support of the Canada Foundation for Innovation (CFI) and the Ontario Research Fund (ORF).

¹ The authors are with the Continuum Robotics Laboratory, Department of Mathematical and Computational Sciences, University of Toronto, Mississauga, ON L5L 1C6, Canada
jessica.burgnerkahrs@utoronto.ca

² Tonia Mielke is with the Faculty of Computer Science, Otto von Guericke University, and with the Forschungscampus STIMULATE, Magdeburg, 39104 Magdeburg, Germany

Digital Object Identifier (DOI): see top of this page.

for a planar needle steering task and reported faster and more accurate movements with task-space control, although some participants preferred joint-space control due to a greater perceived sense of control. El-Hussieny et al. [5] compared multiple input mappings for a vine robot in simulated path-following and wrapping tasks, observing faster completion times for position control variants, higher success rates for rate and shape control, but lower perceived workload for shape-based control. While these studies provide valuable comparative insights, they evaluate steering tasks in simplified or free-space environments with full visibility of the robot and rely on custom or expensive input devices. It therefore remains unclear whether the reported advantages of task-space or joint-space control persist when navigation requires contact-rich interaction and the robot is only partially observable.

To enhance spatial awareness during teleoperation, several approaches provide additional information about the continuum robot's shape. Lin et al. [4] reconstruct and visualize continuum robot shape using augmented reality. Bern et al. [20] and Martín-Barrio et al. [18] employ virtual or mixed reality environments to visualize soft or hyper-redundant robots, while Bhattacharjee et al. [11] present desktop-based shape visualizations. Beyond shape visualization, feedback has also been used to convey motion intent or guidance. El-Hussieny et al. [5] visualize predicted future paths of a soft robot, and Majewicz et al. [6] provide force feedback to indicate the position error. Xie et al. [3] similarly use haptic cues to guide users along a path and warn of potential collisions. Finally, contact force feedback has been explored through haptic interfaces [8] and cutaneous fingertip feedback [21] to communicate contact forces during manipulation. Although providing additional feedback on robot shape and contact forces aims to improve spatial understanding and interaction quality, their effectiveness has not been systematically evaluated under realistic teleoperation constraints such as obstacle-rich environments, environmental contact, and limited visibility. Consequently, their practical benefits in realistic settings remain largely unexplored.

To address these gaps, this work contributes: (1) a systematic comparison of task-space and joint-space control strategies for continuum robot teleoperation, (2) an evaluation of the effect of additional visual feedback on robot shape and position, and (3) an assessment of visual and vibrotactile contact force feedback, all investigated in a controlled user study under obstacle-rich, contact-aware, and partially observable task conditions, focusing on usability aspects such as efficiency and user experience. While this work does not propose a new control algorithm, it provides the first systematic evaluation, to our knowledge, of control and feedback design choices for continuum robot teleoperation under obstacle-rich, contact-aware, and partially observable conditions.

III. METHODOLOGY

To investigate different control and feedback modes, a simulated tendon-driven continuum robot (TDCR) with four segments is used, including a shorter distal segment as a wrist for orienting a tip-mounted camera mimicking typical

inspection tasks. To extend the workspace and facilitate base positioning and insertion, the TDCR is mounted on a robotic arm controlled in task-space. A standard gamepad controller is used as the control interface, as these devices are widely available, inexpensive, and familiar to most users, making them a practical baseline. While alternative interfaces, such as 3D mice, haptic devices [7], or custom-built joysticks [3], may offer more precise or natural control, they are often more expensive, less accessible, or require specialized setup. Moreover, previous work has shown that gamepads can achieve comparable efficiency and workload to more natural interfaces, including gestures or speech [18].

A. Control Modes

1) *Joint-space Control*: The user commands the bending of individual TDCR segments. Each segment has 2 DOF of bending, controlled independently, with an additional insertion command that advances the base robot further into the workspace. This mode offers direct manipulation of the robot's shape but requires the operator to mentally coordinate multiple segments to achieve a desired tip motion.

2) *Task-space Control*: The user commands the desired Cartesian velocity of a control point located at the endpoint of the active segment; the segment-selection input switches which endpoint serves as the controlled point. For a control point with k proximal segments, the Jacobian $\mathbf{J} \in \mathbb{R}^{3 \times (2k+1)}$ relates its 3D position to the bending DOFs of each proximal segment and one axial insertion input. $\mathbf{J}$ is estimated online by central finite differences: each input is perturbed in a cloned simulation state invisible to the operator, the clone is simulated forward over a short horizon, and the resulting control-point displacement forms the corresponding Jacobian column. Each column thus reflects the short-horizon response of the contact-affected system, implicitly accounting for the effects of environmental contacts on the TDCR [22]. The commanded velocity is resolved into actuation through the Moore–Penrose pseudoinverse of $\mathbf{J}$, with singular values below a relative threshold of 0.01 truncated for conditioning. No additional filtering is applied and $\mathbf{J}$ is held constant between re-estimations, with contact-induced changes in the local mapping handled implicitly through continuous re-estimation. This allows the operator to focus on tip placement rather than segment shapes. The shorter distal segment remains under joint-space control for independent camera orientation control in both modes; the two therefore differ only in how the proximal segments are commanded.

3) *Mappings*: The gamepad input is mapped to both the base robot and the TDCR motion, as shown in Fig. 2. For base control, the three DOFs are operated using the right joystick, trigger, and bumper. The same controls are reassigned for TDCR operation to provide 3-DOF task-space control, 2-DOF joint-space control, and 1-DOF insertion.

B. View Modes

This work focuses on a teleoperation scenario in which the robot's surroundings are perceived through a camera mounted at the tip of the TDCR, with the live video stream being

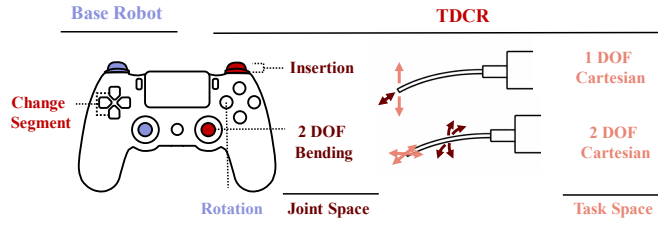


Fig. 2: Mappings of gamepad input onto robot control.

displayed on a screen. Therefore, information about the robot's shape is also provided on the same screen by an additional *top-view* camera (see Fig. 3). This simulated camera is positioned above the continuum robot and provides a cross-sectional view at the height of the robot tip, allowing the user to simultaneously observe the robot shape and the locations of nearby obstacles. This simulated top-view provides near-perfect robot shape and position information, whereas real-world performance would depend on accurate shape estimation and localization.

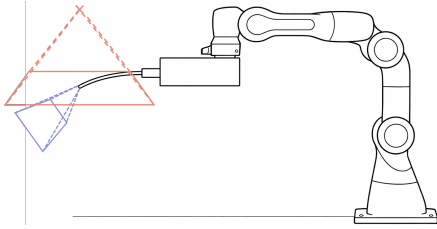


Fig. 3: Camera placements for different view modes. ■ shows front view camera and ■ shows top view camera.

C. Force Display Modes

To convey contact forces, two sensory substitution approaches are employed: visual and vibrotactile feedback. Both modalities are designed to communicate force magnitude and direction relative to the TDCR.

To localize contacts, each contact point $\mathbf{p}_c$ expressed in world coordinates is transformed into the local coordinate frame of the TDCR link with index i , corresponding to the link with the smallest Euclidean distance to the contact point. The rigid-body transformation $\mathbf{T}_w^i$ represents the pose of link i in the world frame and maps points from the link frame to the world frame. The relative position is computed as $\mathbf{p}_c^{(i)} = (\mathbf{T}_w^i)^{-1}\mathbf{p}_c$. The resulting local coordinates are thresholded along the principal axes, after normalization using the predefined threshold of 0.6 to classify the contact direction. Feedback for a direction is activated when at least one contact occurs on the corresponding side of the TDCR. The displayed magnitude is defined as the maximum force among all contacts within that direction.

For the visual feedback, this information is visualized as a colored border at the edge of the display's field of view, where both border width and color scale with force magnitude (see Fig. 4). Additionally, contact points are shown directly in the top view using the same color coding.

For vibrotactile feedback, two motors located on the left and right sides of the controller are used. Consequently, only lateral directional information can be conveyed independently, while

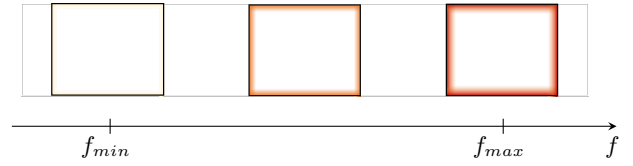


Fig. 4: Concept of visual force feedback.

contacts in other directions result in simultaneous activation of both motors. Based on previous work indicating that vibrotactile feedback is more effective at conveying force information via pulse frequency than via varying amplitude [23], force magnitude is encoded in vibration frequency, with higher frequencies corresponding to larger forces.

D. Implementation

The user interface is implemented in Unity (Unity Software Inc., USA), which captures input from a PlayStation DualShock 4 controller (Sony, Japan) and renders the study environment. Input is converted into velocity commands and sent via UDP to MuJoCo, where the TDCR is modeled as a series of rigid links connected by elastic universal joints following the pseudo-rigid body method [24], with four tendons per segment. MuJoCo simulates dynamics and contacts, and returns joint angles to Unity for rendering the camera views in real time. For joint-space control, the 2-DOF bending input is analytically mapped to tendon tension increments [25] for the active segment. For task-space control, each bending DOF is perturbed by ± 1 mN and the insertion input by ± 2 mm, magnitudes corresponding to approximately 0.1 s of sustained user input. Each perturbed clone state is simulated forward by 0.2 s to allow the transient response to settle. The Jacobian is refreshed every fifth control cycle (5 Hz), while operator commands are sampled at 25 Hz.

IV. EXPERIMENT

A. Study Design

To investigate the different control modes, view modes, and force display modes, a within-subjects repeated-measures user study with a simulated robot setup was conducted. Each participant tested each combination of modes across two repetitions for an abstract task, resulting in 16 trials per participant (2 control modes $\times$ 2 view modes $\times$ 2 force display modes $\times$ 2 trials). Additionally, an application task, designed to resemble real-world scenarios, was conducted to gain more in-depth insights. This study was approved by the University of Toronto Research Ethics Board (Protocol #00049459).

B. Tasks

As one promising application of continuum robots is the inspection of spaces with complex geometries, such as in aviation, this study investigates an inspection task in a cluttered environment. To this end, two tasks were designed: an *abstract task* and an *application task*, both aiming to visualize targets using a camera attached to the TDCR tip. The *abstract task* is intended to cover varying task complexities in a simplified environment, whereas the *application task* examines a more

realistic scenario, investigating the feasibility of the control methods in such conditions.

1) *Abstract Task*: In the abstract task, targets are occluded by obstacles at three different difficulty levels (see Fig. 5). At the easy level, the target is occluded by a single obstacle, and visibility can be achieved either by translating and rotating the base robot or by combining base robot translation with TDCR actuation. At the medium level, the target is occluded by two obstacles arranged such that visualization through base robot motion alone is not possible, as the obstacles limit the reachable viewing angles. Therefore, TDCR actuation is required to look behind the obstacles. At the difficult level, the target lies behind a set of obstacles that requires a more complex TDCR configuration, as the robot must weave through the occlusions to visualize the target. To reduce learning effects, four different task configurations with comparable difficulties are created, each containing all three difficulty levels.

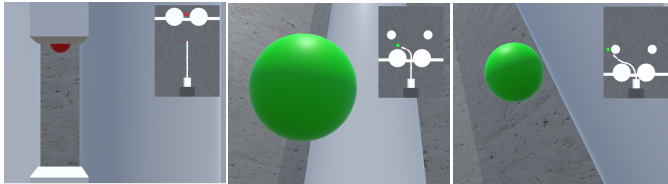


Fig. 5: Front and top view during abstract task at three difficulty levels (from left to right: easy, medium, hard).

2) *Application Task*: In the application task, a CAD model of a jet engine is used as the environment, with targets strategically placed between the rotor blades such that successful visualization requires both TDCR actuation and base robot motion, representing a narrow and complex inspection scenario. Additionally, the lighting is designed to mimic real-world conditions more closely, with a light source attached to the robot tip (see Fig. 6).

C. Measures

1) *Abstract Task*: To assess the efficiency of the proposed system, both objective and subjective measures are collected. The first objective measure is the *performance*, defined as the proportion of targets that participants are able to visualize using the camera attached to the TDCR end-effector. A target is considered successfully visualized when it lies within the camera frustum and is not occluded. The frustum condition is evaluated by checking that the target center lies within the camera viewing volume. Occlusion is evaluated by sampling equidistant points on the surface of the target sphere and verifying that all rays from these sample points to the camera center do not intersect any obstacles. The second objective measure is the *task completion time* (TCT), which is defined as the time from the first user interaction until the last target is successfully visualized. A cutoff time of 120 s per task is used. Targets that are not visualized within this limit are classified as failures and are assigned a TCT of 120 s.

As a subjective measure, the NASA TLX questionnaire is used to assess perceived workload across six dimensions: mental demand, physical demand, temporal demand, performance, effort, and frustration. The NASA TLX score is reported as

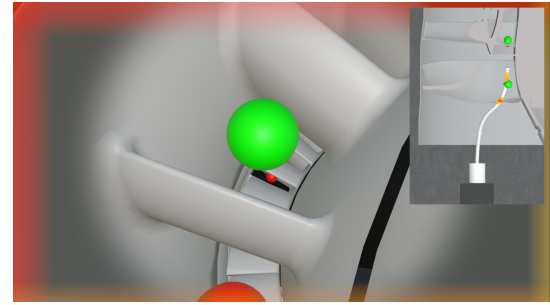


Fig. 6: Application task with visual force display: force magnitude is encoded by border width and color (front view) and by color (top view), while contact locations are shown by the border side (front view) and point position (top view).

the raw TLX, computed by averaging the ratings across all dimensions and multiplying by 5 to obtain a percentage-like score ranging from 0 to 100.

2) *Application Task*: For the application task, *performance*, defined as the proportion of targets participants were able to visualize, and TCT were again recorded. Additionally, participants were free to choose which control mode, view mode, and force display mode they preferred. Therefore, their *user preference* for the different modes was recorded.

D. Procedure

Upon arrival, participants were verbally informed about the study design, the control and feedback methods, and the task. They were then provided with the consent form and invited to ask any questions before signing. Afterward, they completed a demographic questionnaire including questions about age, gender, background, and task-relevant experience.

The experimental phase began with the abstract tasks. In this phase, different combinations of control mode (task-space vs. joint-space), viewpoint (front view only vs. front and top view), and force display (no force display vs. force display information) were tested for the abstract task. The order of conditions was counterbalanced across participants using a Latin square design to reduce potential learning effects. To further minimize learning effects and ensure task proficiency, each trial block began with a training session during which participants could familiarize themselves with the interface using one randomly selected task configuration and ask questions. Once participants indicated that they felt confident performing the task, two evaluation trials were conducted in which they completed two randomly selected configurations of the abstract task. After these evaluation trials, participants completed the NASA TLX questionnaire. This procedure was repeated for all conditions. After completing all trial blocks, a semi-structured interview was conducted to gather in-depth qualitative feedback on the user experience.

Following the abstract tasks, participants performed the application task, in which they were free to choose any combination of the previously introduced parameters. After completing this task, a second semi-structured interview was conducted to obtain additional qualitative feedback. The experiment lasted approximately 60 minutes per participant.

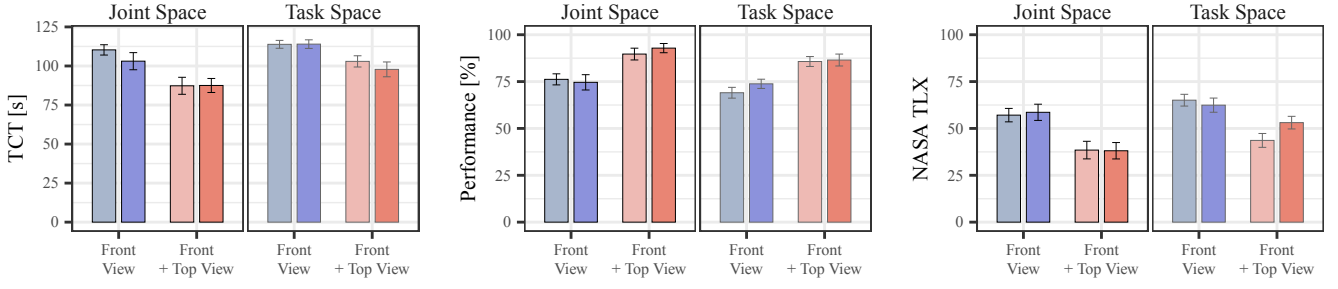


Fig. 7: Descriptive results. ■ & ■ show 'force display', while ■ & ■ show 'no force display' conditions.

E. Participants

A total of 21 participants took part in the study (9 female, 11 male, and 1 non-binary). They were aged between 20 and 32 years (median (m) = 27) and had a technical background (14 graduate students, 3 postdoctoral researchers, 2 engineers, and 2 undergraduate students). All participants had normal or corrected-to-normal vision. Participants rated their technical affinity between 3 and 5 ($m = 4$; 1 = low, 5 = high). Their task-relevant experience, rated on a scale from 1 ("no experience") to 5 ("a lot of experience"), was 1–5 ($m = 4$) for Continuum robots, 1–5 ($m = 2$) for Human–Robot Interaction, and 1–5 ($m = 3$) for video games.

F. Statistical Analysis

1) *Abstract Task*: Since two trials are conducted for each combination of the three factors, TCT and Performance are averaged across identical experimental conditions for each participant. For each dependent variable assumptions of normality and homogeneity of variance are assessed using the Shapiro-Wilk and Levene's tests. When these assumptions are satisfied, three-way repeated-measures ANOVAs are performed. If the assumptions are violated, robust three-way ANOVAs for within-subject designs based on trimmed means are applied instead (see [26]). For conventional ANOVAs, the effect size η^2 is reported. For robust ANOVAs, the effect sizes for main effects are estimated using δ_t (see [27]). For interaction effects, the conventional effect size η^2 is used.

2) *Application Task*: The application task is only evaluated descriptively, as its aim is not to compare different methods formally, but to illustrate whether participants can successfully complete this more complex and realistic task using one of the investigated interaction modes, and to determine whether the interaction and feedback approaches are also suitable in this more complex scenario.

V. RESULTS

A. Quantitative Results

1) *Abstract Task*: The descriptive results are depicted in Fig. 7 and the statistical results are summarized in Table I.

For TCT, the assumptions of normality and homogeneity were violated; robust ANOVAs were conducted. Significant main effects were found for *Control Mode* ($Q = 7.45$, $p = 0.01$, $\delta_t = -0.42$, small effect) and *View Mode* ($Q = 0.98$, $p < 0.001$, $\delta_t = 0.87$, large effect). No other significant main or interaction effects were found. For Performance, normality

and homogeneity were again violated, so robust ANOVAs were used. A significant main effect of *View Mode* with a large effect size was observed ($Q = 58.64$, $p < 0.001$, $\delta_t = -1.20$), while no other significant main or interaction effects were found. For the NASA TLX, the assumptions of normality and homogeneity were satisfied; conventional ANOVAs were conducted. Significant main effects were found for *Control Mode* ($F = 6.24$, $p = 0.02$, $\eta^2 = 0.04$, small effect) and *View Mode* ($F = 63.51$, $p < 0.001$, $\eta^2 = 0.19$, large effect). Additionally, a significant three-way interaction was observed ($F = 7.02$, $p = 0.02$, $\eta^2 = 0.008$, very small effect). No further significant main or interaction effects were detected.

TABLE I: Summary of all statistical analyses ($\alpha < .05$). For repeated measures ANOVAs, test statistic F and effect size η^2 are reported. For robust repeated measures ANOVAs, test statistic Q and effect size δ_t are given.

Variable / Effect type	Factor	Test Statistic	P	Effect Size
TCT				
Main	<i>Control</i>	7.45	0.01	-0.42
	<i>View</i>	0.98	<0.001	0.87
	<i>Force Display</i>	0.47	0.49	-
Two-way	<i>Control</i> $\times$ <i>View</i>	3.16	0.08	-
	<i>Control</i> $\times$ <i>Force Display</i>	<0.001	1.00	-
	<i>View</i> $\times$ <i>Force Display</i>	0.08	0.77	-
Three-way		0.54	0.46	-
Performance				
Main	<i>Control</i>	1.82	0.18	-
	<i>View</i>	58.64	<0.001	-1.20
	<i>Force Display</i>	0.59	0.44	-
Two-way	<i>Control</i> $\times$ <i>View</i>	0.72	0.40	-
	<i>Control</i> $\times$ <i>Force Display</i>	0.23	0.63	-
	<i>View</i> $\times$ <i>Force Display</i>	0.02	0.89	-
Three-way		0.02	0.89	-
TLX				
Main	<i>Control</i>	6.24	0.02	0.04
	<i>View</i>	63.51	<0.001	0.19
	<i>Force Display</i>	0.59	0.45	-
Two-way	<i>Control</i> $\times$ <i>View</i>	0.67	0.42	-
	<i>Control</i> $\times$ <i>Force Display</i>	1.25	0.28	-
	<i>View</i> $\times$ <i>Force Display</i>	3.67	0.07	-
Three-way		7.02	0.02	0.008

2) *Application Task*: The modes selected by participants to complete the application task are depicted in Fig. 8. Descriptively, these results show that participants predominantly preferred using joint-space control, having the additional top view, and receiving feedback on the contact forces. The

TABLE II: Summary and frequency of statements (#) received during the interview, listing reported advantages (+), disadvantages (-), and general comments (*).

Joint Space Control	Additional Top View	Force Display
+ Predictable (8), intuitive (8), easy (7) + Direct control of shape (4) + Motion feels controlled (3) + Separation between Base and TDCR (3) - Requires feedback (6), is unintuitive (2) * Future versions should allow simultaneous base and TDCR control (4)	+ Seeing TDCR shape is helpful (14) + Top view generally helpful (13) + Seeing TDCR position is helpful (10) + Supports motion comprehension (7) + Enables perception of contact points (7) + Makes control easier (5) + Seeing base position is helpful (3) + Helps to correct contact (2) - Depth perception difficult (7) - Less helpful for engine task (6) - Tip orientation hard to see (2) * When not available mental map needed (3) * Additional cameras e.g., at base would be helpful (3)	+ Haptic feedback is helpful (15) + Supports motion comprehension (9) + Enhances perception of collision (7) + Helpful when not having top view (5) + Visual feedback easy to understand (4) + Supports perception of TDCR position (3) + Directional information in visual helpful (3) + Brings attention to collisions (3) + Allows to reduce collisions (2) + Helps learning control (2) - Vibration is stressful (7), disruptive (6), too urgent (5), and annoying (2), especially when contact is required (4) - Vibratory directional information unclear (3) - Not necessary when having top view (3) * Should be less intense (3)
Task Space Control		
+ Intuitive (3) + Simultaneous TDCR and Base movement (3) - Base motion confusing (10) - Unpredictable shape (5) and motion (2) especially when in contact (2) - Requires training (2) - Not intuitive (2), challenging (2)		

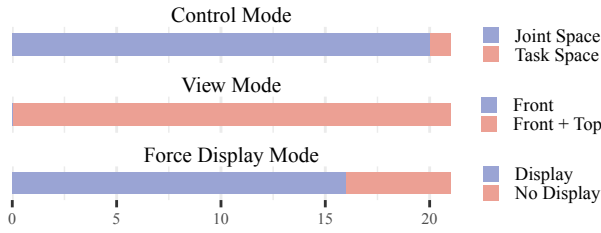


Fig. 8: User preference for modes in the application task.

performance of participants for this task was $84.92 \pm 21.02\%$, with an average TCT of 83.58 ± 22.71 s.

B. Qualitative Results

For the qualitative feedback analysis, participant statements from the semi-structured interviews were paraphrased and organized into thematic clusters. Only statements reported by at least two participants were included in the clustering process. Of the 305 individual statements collected, 260 met this criterion and were subsequently condensed into 56 summarizing statements (see Table II).

C. Interpretation of Results

1) *Control Mode*: Significant main effects of *Control Mode* are shown in Fig. 9. Joint-space control resulted in significantly lower TCT and perceived workload compared to task-space control. Qualitative feedback indicates that participants perceived joint-space control as predictable, intuitive, and easier to manage, whereas task-space control was often described as less predictable, particularly in the presence of obstacle

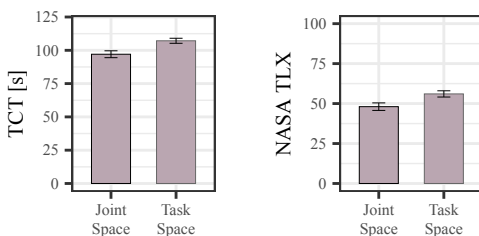


Fig. 9: Control Mode main effects.

interactions. In the application task, 95% of participants selected joint-space control.

Participants also commented on the separation of base and TDCR motion in joint-space control. While some appreciated this separation, others suggested that future systems should allow optional simultaneous control.

2) *View Mode*: Providing the additional top view significantly reduced TCT and workload and increased performance (Fig. 10). All participants selected the top view in the application task. Participants reported improved understanding of robot shape, position, and contact location (Table II), reducing the need for mental reconstruction.

Although the top view was generally described as helpful, limitations included challenges in depth perception and tip orientation, particularly in the application scenario.

3) *Force Display Mode*: No significant main effects of Force Display Mode were observed. Although participants described the feedback as helpful for collision awareness and motion understanding, duration, performance, and perceived workload did not improve significantly. In the application task, 76% selected the force feedback, suggesting that it may have benefits beyond the recorded measures. Visual feedback was generally preferred over vibrotactile feedback, which was sometimes perceived as stressful or disruptive.

4) *Interaction Effects*: A significant but small three-way interaction effect on the NASA TLX is shown in Fig. 7. Descriptively, force feedback reduced workload for task-space control when the top view was available. However, given the small effect size, the overall pattern of main effects remains dominant.

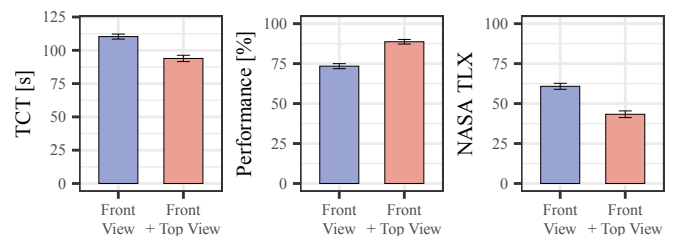


Fig. 10: View Mode main effects.

VI. DISCUSSION

Our results indicate that joint-space control enables faster interaction and lower perceived workload than task-space control. This contradicts previous findings by Majewicz *et al.* [6] who reported that task-space control was faster and more accurate, although some participants subjectively preferred joint-space control because they felt more in control. While this perceived sense of control aligns with our findings, we observed longer task completion times for task-space control. This discrepancy likely arises from differences in task constraints, highlighting the importance of evaluating control strategies under realistic operational conditions. Majewicz *et al.* [6] evaluated reachability in an unobstructed environment, whereas our study required navigation around obstacles. In the presence of geometric constraints and contact interactions, the effective mapping between operator input and tip motion becomes less predictable, increasing the value of explicit shape manipulation provided by joint-space control. Furthermore, the lower workload observed for joint-space control in our study, where the robot shape can be adjusted more directly, is consistent with findings by El-Hussieny *et al.* [5], who reported reduced workload when using a dedicated shape-control with a custom haptic input device. The distal camera segment remained under joint-space control in both conditions, since the positional task-space controller cannot command orientation; the task-space condition is therefore a hybrid approach. As this segment was identical across conditions, the comparison still isolates the paradigm: shape in joint space vs. position in task space.

Although the use of additional viewpoints to enhance spatial awareness has been proposed in prior work [4], [11], [18], [20], [28], their effectiveness has not been empirically validated. Our findings provide empirical evidence that explicit visualization of robot shape within the environment substantially improves teleoperation performance. Previous work suggests that immersive interfaces enable more efficient control and higher situational awareness [18]. In our study, the limitations of a desktop display became apparent, particularly in complex scenarios where the 2D top view was sometimes described as insufficient, and challenges in depth perception arose. These limitations could be addressed using augmented reality, as proposed in related work [4]. Such approaches may further enable more flexible viewpoints through interactive 3D visualization, for example, allowing users to observe the scene from the base robot's perspective, as suggested by several participants.

Despite participants predominantly choosing to have additional information about contact forces, this feedback did not significantly improve performance, duration, or workload. Although our study did not explicitly compare different feedback modalities, participants reported that visual feedback was easier to interpret than vibrotactile feedback, which aligns with prior findings in the broader HRI literature [23], [29]. Furthermore, previous work suggests that multimodal feedback may even reduce efficiency due to information overload [29], [30]. As participants generally described the feedback as helpful, future work should investigate improved feedback designs and

evaluate whether a simpler, unimodal approach could reduce cognitive load and potentially enhance performance.

The results for the more complex *application task* indicate that, even under challenging conditions, participants were able to visualize most target structures when using their preferred control and feedback modes. This indicates that the identified control and visualization strategies generalize beyond simplified abstract tasks and remain effective in a structurally complex inspection scenario. Simple input devices, such as a gamepad, can be effective when paired with appropriate control and feedback strategies. However, as discussed previously, the viewing modes require further investigation to determine how best to provide feedback on TDCR position and shape within the environment.

Translating the benefits observed in this study to real-world deployment requires accurate sensing and state estimation. Accordingly, digital twin or mapping frameworks must be supported by reliable shape estimation [31], localization [32], and force estimation [33] to provide such additional viewpoints and contact force feedback.

Limitations: This study focused on control efficiency and user preference rather than safety-related outcomes. Accordingly, force feedback was evaluated with respect to task performance, task completion time, workload, and subjective feedback. While safety is a critical consideration in contact-rich continuum robot applications, it was not explicitly assessed in the present study. Future work should therefore incorporate contact-related metrics, such as peak contact force, collision frequency, or collision duration, to evaluate the potential safety benefits of force feedback.

This study is further limited by the relatively small and homogeneous participant sample ($N = 21$), consisting primarily of individuals with an academic engineering background. While this population ensures a basic technical understanding of the system, it may not fully represent the target user group. Differences in domain expertise, prior experience, and operational strategies may influence interaction behavior and workload perception. Future work should therefore include larger and more diverse participant groups.

Implications: Based on the study results, we propose the following design implications for control and feedback approaches in human-robot interaction with continuum robots.

Employ joint-space control for efficient interaction. Our results suggest that joint-space control enables more efficient interaction than task-space control, with participants highlighting its predictability and intuitiveness. When using task-space control, steps should be taken to make the robot's motion and shape more comprehensible.

Provide users with additional information about position and shape. Findings show that providing an additional top view that conveys the TDCR's shape within the environment eases interaction and allows for more efficient control with lower workload, suggesting that such shape and position information is essential for effective control.

Provide subtle visual sensory substitution. Although giving contact force feedback did not significantly improve efficiency, participants predominantly chose to have it present and described it as helpful. Participant feedback indicates that

visual feedback is easier to interpret than vibrotactile feedback, and care should be taken to ensure that the display of force information is not overwhelming or disruptive.

VII. CONCLUSION

This study evaluated control and feedback strategies for continuum robot teleoperation under constrained, contact-aware, and partially observable conditions. The results show that explicitly conveying robot shape and position has the strongest impact on efficiency and workload, while joint-space control enables faster and less demanding interaction in constrained environments. Although contact force feedback did not significantly improve objective performance, it was widely preferred and perceived as supportive.

Together, these findings underscore that teleoperation performance in realistic tasks depends critically on reliable observability of the continuum robot configuration within the environment. Enabling such observability in real systems will require robust self-localization and shape estimation to support digital twin or mapping-based visualization frameworks for contact-rich continuum robot applications.

ACKNOWLEDGMENTS

We thank Professor Christian Hansen for his support in securing funding for this project and for initiating and supporting the collaboration between the involved institutions.

REFERENCES

- [1] J. Burgner-Kahrs, D. C. Rucker, and H. Choset, "Continuum robots for medical applications: A survey," *IEEE Trans. Rob.*, vol. 31, no. 6, pp. 1261–1280, 2015.
- [2] M. Russo, S. M. H. Sadati, X. Dong, A. Mohammad, I. D. Walker, C. Bergeles, K. Xu, and D. A. Axinte, "Continuum Robots: An Overview," *Adv. Intell. Syst.*, vol. 5, no. 5, p. 2200367, 2023.
- [3] M. Xie, C. Girerd, and T. K. Morimoto, "A 3-d haptic trackball interface for teleoperating continuum robots," in *IEEE RAS/EMBS Intl. Conf. Biomed. Rob. Biomech. (BioRob)*. IEEE, 2022, pp. 1–8.
- [4] Z. Lin, A. Gao, X. Ai, H. Gao, Y. Fu, and W. Chen, "ARei: Augmented-Reality-Assisted Touchless Teleoperated Robot for Endoluminal Intervention," *IEEE/ASME Trans. Mechatron.*, vol. 27, no. 5, pp. 3144–3154, 2021.
- [5] H. El-Hussieny, U. Mehmood, Z. Mehdi, S.-G. Jeong, M. Usman, E. W. Hawkes, A. M. Okamura, and J.-H. Ryu, "Development and evaluation of an intuitive flexible interface for teleoperating soft growing robots," in *IEEE/RSJ IROS*, 2018, pp. 4995–5002.
- [6] A. Majewicz and A. M. Okamura, "Cartesian and joint space teleoperation for nonholonomic steerable needles," in *World Haptics Conf. (WHC)*, 2013, pp. 395–400.
- [7] C. Fellmann, D. Kashi, and J. Burgner-Kahrs, "Evaluation of input devices for teleoperation of concentric tube continuum robots for surgical tasks," in *Proc. SPIE Med. Imaging*, vol. 9415, 2015, pp. 411–419.
- [8] C. G. Frazelle, A. D. Kapadia, and I. D. Walker, "A haptic continuum interface for the teleoperation of extensible continuum manipulators," *IEEE RA-L*, vol. 5, no. 2, pp. 1875–1882, 2020.
- [9] D. Black, S. Lilge, C. Fellmann, A. V. Reinschluessel, L. Kreuer, A. Nabavi, H. K. Hahn, R. Kikinis, and J. Burgner-Kahrs, "Auditory display for telerobotic transnasal surgery using a continuum robot," *J. Med. Robot. Res.*, vol. 04, no. 02, p. 1950004, 2019.
- [10] C. G. Frazelle, A. D. Kapadia, K. E. Fry, and I. D. Walker, "Teleoperation mappings from rigid link robots to their extensible continuum counterparts," in *IEEE ICRA*, 2016, pp. 4093–4100.
- [11] S. Bhattacharjee, S. Chattopadhyay, V. Rao, S. Seth, S. Mukherjee, A. Sengupta, and S. Bhaumik, "Kinematics and teleoperation of tendon driven continuum robot," *Procedia Comput. Sci.*, vol. 133, pp. 879–886, 2018.
- [12] B. Ouyang, Y. Liu, H.-Y. Tam, and D. Sun, "Design of an interactive control system for a multisection continuum robot," *IEEE/ASME Trans. Mechatron.*, vol. 23, no. 5, pp. 2379–2389, 2018.
- [13] Y. Yamamoto, S. Wakimoto, T. Kanda, and D. Yamaguchi, "A Soft Robot Arm with Flexible Sensors for Master-Slave Operation," *Eng. Proc.*, vol. 10, no. 1, p. 34, 2021.
- [14] P. D. S. H. Gunawardane, R. E. A. Pallewela, and N. T. Medagedara, "Tele-Operable Controlling System for Hand Gesture Controlled Soft Robot Actuator," in *IEEE-RAS RoboSoft*, 2019, pp. 14–18.
- [15] H.-S. Yoon and B.-J. Yi, "Design of a master device for controlling multi-modulated continuum robots," *Proc. Inst. Mech. Eng. C, J. Mech. Eng. Sci.*, vol. 231, no. 10, pp. 1921–1931, 2017.
- [16] Q. Guan, H. H. Cheng, B. Dai, and J. Hughes, "Control the soft robot arm with its physical twin," *arXiv preprint arXiv:2503.17227*, 2025.
- [17] A. D. Kapadia, I. D. Walker, and E. Tatlicioglu, "Teleoperation control of a redundant continuum manipulator using a non-redundant rigid-link master," in *IEEE/RSJ IROS*, 2012, pp. 3105–3110.
- [18] A. Martín-Barrio, J. J. Roldán, S. Terrile, J. Del Cerro, and A. Barrientos, "Application of immersive technologies and natural language to hyper-redundant robot teleoperation," *Virtual Reality*, vol. 24, no. 3, pp. 541–555, 2020.
- [19] W. Liu, Y. Duo, J. Liu, F. Yuan, L. Li, L. Li, G. Wang, B. Chen, S. Wang, H. Yang *et al.*, "Touchless interactive teaching of soft robots through flexible bimodal sensory interfaces," *Nature Comm.*, vol. 13, no. 1, p. 5030, 2022.
- [20] J. M. Bern, W. C. May, A. Osborn, F. Stella, S. Zargazadeh, and J. Hughes, "A soft robot inverse kinematics for virtual reality," in *IEEE ICRA*, 2024, pp. 14957–14963.
- [21] H. Ham, M. Park, T. Park, X. Gao, Y.-L. Park, and M. J. Park, "Teleoperation of soft robots with real-time fingertip haptic feedback using small batteries," *Adv. Mater. Technol.*, vol. 8, no. 15, p. 2300070, 2023.
- [22] M. C. Yip and D. B. Camarillo, "Model-less feedback control of continuum manipulators in constrained environments," *IEEE Trans. Rob.*, vol. 30, no. 4, pp. 880–889, 2014.
- [23] T. Howard and J. Szewczyk, "Assisting Control of Forces in Laparoscopy Using Tactile and Visual Sensory Substitution," in *New Trends in Med. and Serv. Rob.*, 2016, pp. 151–164.
- [24] C. Shentu, N. Baldassini, T. Zheng, P. Rao, and J. Burgner-Kahrs, "Do rigid-body simulators dream of soft robots? learning contact-rich manipulation for tendon-driven continuum robots," *arXiv:2606.22397*, 2026.
- [25] R. Grassmann, A. Senyk, and J. Burgner-Kahrs, "Clarke transform—a fundamental tool for continuum robotics," *arXiv preprint arXiv:2409.16501*, 2024.
- [26] R. R. Wilcox, "Chapter 8 - comparing multiple dependent groups," in *Introduction to Robust Estimation and Hypothesis Testing*, fifth edition ed., R. R. Wilcox, Ed. Academic Press, 2022, pp. 467–539.
- [27] J. Algina, H. J. Keselman, and R. D. Penfield, "An alternative to cohen's standardized mean difference effect size: A robust parameter and confidence interval in the two independent groups case," *Psychological Methods*, vol. 10, no. 3, pp. 317–328, 2005.
- [28] A. Martín-Barrio, J. J. Roldán-Gómez, I. Rodríguez, J. del Cerro, and A. Barrientos, "Design of a Hyper-Redundant Robot and Teleoperation Using Mixed Reality for Inspection Tasks," *Sensors*, vol. 20, no. 8, p. 2181, 2020.
- [29] T. Mielke, F. Heinrich, and C. Hansen, "SensARy Substitution: Augmented Reality Techniques to Enhance Force Perception in Touchless Robot Control," *IEEE Trans. Visual. Comput. Graphics*, vol. 31, no. 5, pp. 3235–3244, 2025.
- [30] Y. Jonetzko, J. Hartfill, N. Fiedler, F. Zhong, F. Steinicke, and J. Zhang, "Evaluating Visual and Auditory Substitution of Tactile Feedback During Mixed Reality Teleoperation," in *Cognitive Comput. Sys.*, 2023, pp. 331–345.
- [31] J. M. Ferguson, D. C. Rucker, and R. J. Webster, "Unified shape and external load state estimation for continuum robots," *IEEE Transactions on Robotics*, vol. 40, pp. 1813–1827, 2024.
- [32] S. Teetaert, G. Caroleo, M. Pontin, S. Lilge, J. Burgner-Kahrs, T. D. Barfoot, and P. Maiolino, "Continuum Robot Localization using Distributed Time-of-Flight Sensors," *arXiv:2602.07209*, Feb. 2026.
- [33] Z. Lin, J. Xia, Z. Xu, Y. Zou, C. Zhou, J. Chen, L. T.-L. Tong, S. Huang, H. Liu, W. Chen *et al.*, "Real-time whole-body contact estimation of continuum robots using multiplexed fibers for embodied actuation and sensing to quantify interactions," *Soft Robotics*, p. 21695172251388808, 2025.