

Dimension-Specific Shared Autonomy for Handling Disagreement in Telemanipulation

Michael Bowman  and Xiaoli Zhang , *Member, IEEE*

Abstract—One of the fundamental questions in shared control is how to allocate control power to the human and robot effectively. Conventional arbitration policies often define a uniform singular scalar for all 6 DOFs to blend human input and robot assistance. However, this singular scalar over-dominates some dimensions of the inputs and provides insufficient assistance in other dimensions. Thus, current shared control can support simple telemanipulation tasks such as pushing, pressing, and simple positional control but is limited in tasks with more DOFs like rotational motion. A dimension-specific arbitration policy is developed to customize the control arbitration along each DOF to fill the gap. It looks at whether the robotic assistance is too timid or aggressive along each DOF and determines the arbitration magnitude according to disagreement levels of control allocation and the user's willingness to accept assistance. The user's willingness is estimated from a feedback psychology model. The method has higher similarity and ratio of agreement between the human and robot (lower over-dominance) over existing methods and, simultaneously, improves the task performance. This arbitration strategy is expected to increase the adoption of teleoperation for object manipulation.

Index Terms—Telerobotics and teleoperation, motion control, human-centered robotics.

I. INTRODUCTION

TELEOPERATING a robot for a manipulation task is often difficult and complex due to indirect visualization, indirect manipulation with the robot, and physical discrepancies between a human and robot hand [1], [2]. Shared control injects autonomy into an operator's actions to overcome these issues. Shared autonomy is desirable for teleoperation applications, including industrial inspection and maintenance, disaster relief, or assistive living.

Object telemanipulation requires more than simple positional control, realizing the system requires rotational and finger control. These motion components provide a natural split in handling complex tasks. Operators handle these components in

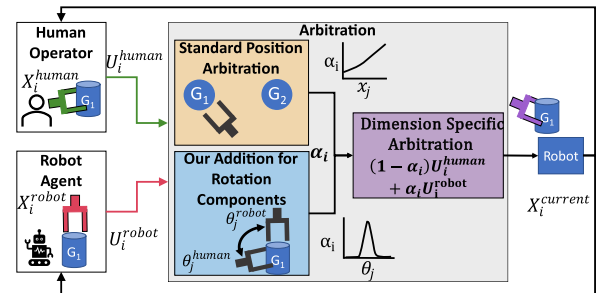


Fig. 1. The overview of the proposed telemanipulation control paradigm. Our contributions include partitioning dissimilar dimensions, devising independent arbitration policies to handle control allocation disagreement, and applying the appropriate dimension-specific control authority.

their preferred approach. For example, the operator may obtain a position near an object, then begin to handle the rotations and end with finger control. Here an iterative process occurs by the operator to refine each step until they achieve the goal. Certain combinations of refinement will be more intuitive to different operators, like either handling the depth before lining up the robot in an x-y plane or vice versa. Naturally, operators likely only adjust a few components while holding others nearly constant (i.e., holding a position but adjusting the pose angle or finger posture) to succeed in a grasp [6].

The autonomous actions may differ from the operator actions along all or a portion of these motion components. For instance, given the same high-level target pose, the autonomous agent and operator may have differing motion strategies to achieve it (Fig. 1: human operator and robot agent). Disagreement occurs due to misperception or preference of the operator. Neither agent is incorrect in their actions to accomplish the task, nor is the robot producing suboptimal action plans. This disagreement is fundamentally inherent in shared control, and precautions are necessary to prevent over-dominance from the robotic agent. Further, the robot and human may only partially agree, where they may agree in the positional component and differ in the rotational component. Fig. 2 shows this partial agreement.

Determining the appropriate level of autonomy in teleoperation is the key research topic to realize telemanipulation scenarios fully. The most common state-of-the-art strategy in shared control is to use a linear blending strategy with a universal arbitration term, α . The control authority is a singular scalar value that influences every dimension of the control rather than individual components. Specifically, it only includes positional components, not rotational ones, to dictate arbitration along a

Manuscript received 29 August 2022; accepted 2 January 2023. Date of publication 24 January 2023; date of current version 1 February 2023. This letter was recommended for publication by Associate Editor M. Selvaggio and Editor J. Ryu upon evaluation of the reviewers' comments. This work was supported by U.S. NSF under Grants 1652454 and 2114464. (Corresponding author: Xiaoli Zhang.)

This work involved human subjects or animals in its research. Approval of all ethical and experimental procedures and protocols was granted by Colorado School of Mines' Human Subject Committee.

The authors are with the Colorado School of Mines, Intelligent Robotics and Systems Lab, Golden, CO 80401 USA (e-mail: mibowman@mines.edu; xlzhang@mines.edu).

This letter has supplementary downloadable material available at <https://doi.org/10.1109/LRA.2023.3239313>, provided by the authors.

Digital Object Identifier 10.1109/LRA.2023.3239313

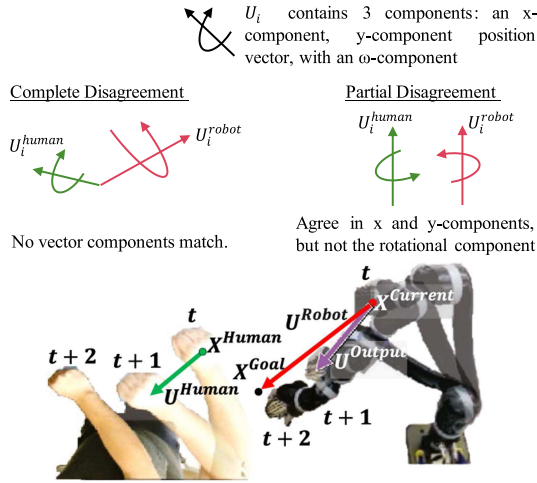


Fig. 2. Agreed positional components influence and dominate arbitration models, leaving rotational components more robot-dominated than intended. An example of an operator's and robot's actions for 3 timesteps with partial agreement in the vision-based telemanipulation scenario.

trajectory. This strategy cannot handle partial agreement that arises in telemanipulation scenarios, especially towards the end of the trajectory where conventionally, the robot assumes most control ($\alpha = 1$). Consider a case where the robot is near the positional goal, and the operator attempts to refine the position before adjusting the rotation. With the conventional strategies, the robot will dominate not only the positional components but the rotational aspect too as α approaches 1. This is despite the operator not yet focusing on the rotational components to grasp an object successfully and restricts their ability to alter them. The over-dominance restricts the operator in how they want to complete the task (i.e., grasp the object from the side rather than the top like in Fig. 1). This leads to an operator's attempt to refine actions but sees the control as too aggressive thus make drastic changes [3]. Over-adjustments and strong reactionary responses are due to the low control authority. However, the operator's initial intuitive refinement is to be subtle. This is especially common for telemanipulation applications [4], [5].

This work aims to provide autonomy to aspects the robot and user agree on and reduce assistance in areas where they disagree. Otherwise, the arbitration leads to operator resistance. Toward this aim, a dimension-specific control authority to focus on individual components to succeed is developed. The contributions are as follows:

- 1) Break down dimensional aspects with respective motion components, and quantify dimension-specific disagreement and the human's willingness for assistance.
- 2) Develop a dimension-specific arbitration strategy for non-dominating assistance along different motion components.
- 3) Validate and evaluate the dimension-specific arbitration strategy in telemanipulation tasks.

II. RELATED WORK

Current telemanipulation focuses on improving hand motion mapping capabilities and lacks shared autonomy strategies. These mapping strategies use model-based approaches such as

DexPilot [7] or use synergies [8], [9]. These methods consider orientation and finger mapping. However, deployment still relies on an operator to fully control a system that is difficult to maneuver and slower than normal actions an operator would take [7]. Despite improving the physical discrepancy issues, the operator struggles to overcome the disembodiment problem, where they face indirect perception and manipulation of the robot environment and lack dynamic interaction with objects.

Introducing shared control to circumvent the above issues has empowered the autonomous agent with its knowledge of the environment. However, shared control with linear blending has only been used on simplified approaching problems [10]. The arbitration is a single scalar based on the robot's confidence in a human's intended goal [11], [12], [13] or empirically defined to increase from 0 to 1, where the higher autonomy levels are near the goal [6]. This standard strategy leads to over-dominance of a robot's actions near goal states; an operator may disagree with the output actions. Operators have sometimes reported a desire to retain control despite lower task performance [14], [15]. Often seen as inferencing issues to achieve a better arbitration, research lacks questioning the structural issues for linear blending techniques and handling disagreement.

The work in [3] provides insight on reframing a multi-inference sub-policy to quantify disagreement between operators and robotics assistance plans for 2D approaching tasks. Other strategies quantify disagreement between operator and robot actions and then apply minimum assistance to ensure safe control [12], [16], [17]. The disagreement quantification works for both discrete and continuous domains [18]. Alternatively, handling disagreement can be viewed as mutual adaptation or adapt to adaptation [19], [20], [21]. One aim of mutual adaptation is to have an online adjustment of control authority based on the operator's action. However, all these strategies obfuscate the action's individual components by analyzing the entire vector and have rarely been used in telemanipulation fields.

III. METHODS

A. Disagreement of Control Allocation

Both agents have pose states denoted as $X \in \mathbb{R}^6$. Where X is split between 3D positional and three independent rotational dimensions (roll, pitch, yaw), thus $X \in \{x_j, \theta_j\}$, $j = 1, 2, 3$. Both the robot and operator apply actions $U \in \{\dot{x}_j, \dot{\theta}_j\}$, $j = 1, 2, 3$. We assume the goal state, X^{goal} , is given and known by the robot and the operator, and the optimal robot policy U^{robot} is set as to go towards X^{goal} as seen in Fig. 2. The operator is free to choose actions, U^{human} , other than those towards the X^{goal} to achieve the task, thus going to their own state, X^{human} (i.e., the current human hand pose as seen in Fig. 2). The standard linear blending arbitration is augmented to a dimension-specific one for both the U^{human} and U^{robot} , as in (1), where i is each dimension of the control vector ($i = 1, 2, \dots, 6$). The resultant action, U_i^{output} , applies to the robot dynamics with the current robot state, $X^{current}$, to obtain the next state, X^{next} in (2).

$$U_i^{output} = U_i^{robot} \alpha_i + U_i^{human} (1 - \alpha_i) \quad (1)$$

$$X^{next} = f(X^{current}, U^{output}) \quad (2)$$

α_i can be summarized as the L1 distance (or L2 if preferred) towards X^{goal} , from an initial robot state, X^{start} , and $X^{current}$, to determine the control authority. Eq. (3) shows an L1 distance using $X^{current}$ in the standard error function (erf).

$$\alpha_i = 1 - \text{erf} \frac{\|X^{goal} - X^{current}\|_1}{\|X^{goal} - X^{start}\|_1} \quad (3)$$

For the baseline case (denoted as the Standard (S) controller), all α_i are the same value as (3). The α_i is only from the robot's perspective as $X^{current}$ solely dictates it to promote as much assistance as possible ($\alpha_i = 1$). However, we must account for misalignment, as previously discussed. So, we alter α_i through quantifying disagreement of control allocation. The disagreement between both the current human (X^{human}) and robot ($X^{current}$) is the difference in expected control allocation. The expected control allocation is where the current robot expects to provide α help based on $X^{current}$ and the operator expects to gain α assistance based on X^{human} . There is a mismatch between $X^{current}$ and X^{human} , thus the expected α is different, due to misalignment. So, (4) makes two changes to (3). The first is to update it as a dimensional difference. The second is to represent each agent's current state position (x_j^{human} or $x_j^{current}$) as a probability of expected control allocation (as if each agent were to apply their position to (3) and produce an α_i), where A represents either the human (H) or the robot (R) agent and x_j^{agent} is either x_j^{human} or $x_j^{current}$.

$$P(A_i) = 1 - \text{erf} \frac{\|x_j^{goal} - x_j^{agent}\|_1}{\|x_j^{goal} - x_j^{start}\|_1} \forall i = \{1, 2, 3\}, j = \{1, 2, 3\} \quad (4)$$

Eq. (4) is inappropriate for quantifying rotation dimensions to account for circular variables. Therefore, a more natural, appropriate probability distribution for rotational control allocation is the Von Mises distribution [3] in (5). Functions with similar characteristics that can handle circular variables can also be applied. For instance, our rotations are independent, leading to a univariate probability function in (5).

$$P(A_i) = \frac{e^{\kappa \cos(\theta_j^{agent} - \theta_j^{goal})}}{2\pi I_0(\kappa)} \forall i = \{4, 5, 6\}, j = \{1, 2, 3\} \quad (5)$$

κ is an empirically defined variable, and $I_0(\kappa)$ is the modified Bessel function, and θ_j^{agent} refers to rotational dimensions of X . Eqs. (4) and (5) are shown in Fig. 3 as solid black lines. In Fig. 3, the human and robot have differing expectations of control allocation along a singular dimension. $P(H_i)$ and $P(R_i)$ are the bounds of disagreement for control authority seen by the operator and the actual robot pose. In conventional systems, assistance is either *timid* (the robot's provided assistance is behind the human's expected assistance in Fig. 3's top row) or *aggressive* (the robot's provided assistance is ahead of the human's expected assistance in Fig. 3's bottom row). With our system, we aim to determine whether assistance is truly timid or aggressive (ultimately determining the appropriate level of assistance to provide) by including the operator's willingness

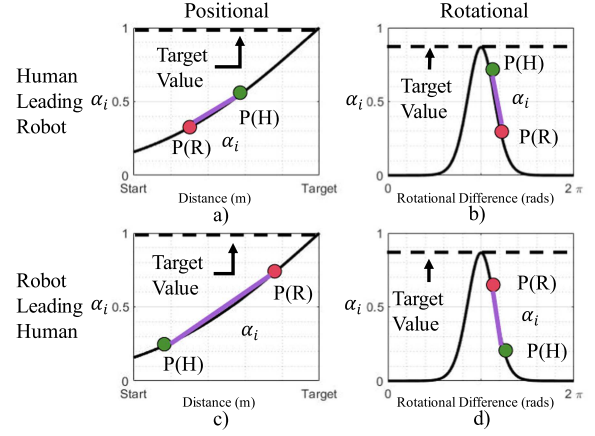


Fig. 3. The expected control authority by the human (H), and robot (R) using (4) (left column) and 5 (right column) for specific dimensions. The target value corresponds to the goal state known by both agents. Conventional systems take $P(R)$ as the true α . Top row: human leads robot, meaning lower assistance, α , in conventional systems ($P(R) < P(H)$). Bottom row: robot leads human, meaning higher assistance, α , in conventional systems ($P(R) > P(H)$). The purple regions are potential α from our strategy in (8). Our approach aims to find a better α than conventional means by evaluating the operator's willingness, B_i , which contextualizes the U_i^{human} based on their current progress in the task. A high B_i will lead α_i being closer to the conventional strategy ($P(R)$), while a low B_i will lead to α_i to be near $P(H)$.

to accept assistance (which contextualizes U_i^{agent} in terms of current progress toward the goal). For instance, if the operator has a low willingness to accept assistance for conventionally aggressive assistance (bottom row of Fig. 3), then the α will be set closer to $P(H)$. Alternatively, if they have a high willingness to accept assistance, then the conventionally aggressive assistance is not true, the operator is willing to gain aggressive assistance, and α can be set closer to $P(R)$.

B. User's Willingness to Accept Assistance

The operator's willingness to accept assistance stems from their desire to change the current control authority in the dynamic process. Determining the change to this authority level can temper too aggressive and boost too timid assistance. The operator's current desire and preference towards the robot's assistance is identified by looking at both agents' actions along each dimension. A probabilistic approach, inspired by psychology feedback models [22], [23], [24], is developed to determine willingness to accept assistance. The psychology model discusses a proposed feedback mechanism for how people change their behaviors in a task based on current progress, previous experiences, and the current rate of progress. Further, a Beta prime distribution uses these components to determine when a person will likely pick up the pace or back off. Lastly, the model discusses how the pace changes for the individual dimensions; as one dimension slows, another takes precedence. The feedback model was conceptual and discusses the pace in the abstract; we adapt them to a mathematical framework and apply them to the shared autonomy. Compared with the existing model-free data-driven approach that defines adaptability through a Mixed Observable Markov Decision Process [20], the primary benefit of using the mathematical model is to provide the robot with the operator's

willingness to accept the control authority in a human-like manner. The proposed model can be adapted in two ways. The first is to move it into the motion domain rather than a conceptual abstraction of rate changes to complete tasks. The second is to adapt the model to give a predictive or feedforward effect rather than a reactive calculation (i.e., computing acceleration from measured U_i^{human}). Together, both adaptations account for the expected way an operator would move toward the goal given their current progress (correctly transforming or normalizing the U_i^{agent} into the model's concept of speed change). Due to the predictive nature of the model and the transformation, this is referred to as the “expected speed change.” To make these adaptations to the psychology model, we identified four components: 1) current actions, 2) previous actions, 3) current progress in the task, and 4) aggregated previous experiences. The first and second components are instantaneous reactionary decisions, while the third term looks at a longer planning window for the current trial. The last term influences the operator's expectations in completing the trial based on previous experiences. The last two terms are what the operator relies on to know if they are slow or fast in this trial compared to previous trials. Eq. (6) combines these terms to quantify the agent's expected speed change based on their current (first term) and last actions (second term), current progress (third term), and aggregated previous experiences (last term). Where N is the total number of previous attempts in doing the task, and n is the index for a specific attempt. Eq. (6) aims to quantify/normalize whether the operator will increase or decrease their velocity based on their current progress compared to previous experiences in the task.

$$\begin{aligned} \psi_i^{agent} = & U_i^{agent} - \frac{X_{i,current} - X_{i,previous}}{\|X_{i,current} - X_{i,previous}\|_2} \\ & - \frac{X_{i,current} - X_{i,start}}{\|X_{i,current} - X_{i,start}\|_2} \\ & + \frac{1}{N} \sum_{n=0}^N \frac{X_{i,n}^{terminal} - X_{i,n}^{start}}{\|X_{i,n}^{terminal} - X_{i,n}^{start}\|_2} \end{aligned} \quad (6)$$

The variable *agent* refers to the human agent (U_i^{human}) or the robot agent (U_i^{robot}) actions. Each rate (ψ_i^{human} and ψ_i^{robot}) is then used in a Beta prime distribution with hyperparameters b , p , and q . The hyperparameters impact the shape and scale and are empirically chosen based on human speed in manipulation tasks and the robot's velocity limits. The heuristic described in [23] inspires the Beta prime distribution. The distribution allows for normalization and quantification compared to the heuristic in [23] and provides context toward what the operator is attempting to achieve with their speed change. ψ_i^{robot} is necessary as it is the reference to compare the sampling for ψ_i^{human} (which is used to see if it is faster or slower). The robot policy, U_i^{robot} , determines ψ_i^{robot} before the arbitration takes place which is discussed in the next section. ψ_i^{robot} is necessary as it influences the shape of the curve and is the mean value for the probability density function to evaluate the ψ_i^{human} (which is calculated after measuring U_i^{human}). The ψ_i^{human} evaluation produces a single deterministic value (which we call willingness, B_i) for the specific time point. However, this could be extended to a

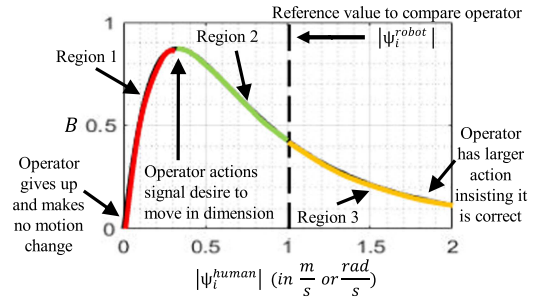


Fig. 4. The operator's willingness to accept assistance for a single dimension. The dashed line represents the reference which is the expected speed change of the robot, ψ_i^{robot} , from the optimal robot policy (U_i^{robot}). Three regions exist. The first (in red) shows when no expected speed change ($\psi_i^{human} = 0$) occurs, which implies the operator is unwilling to change and accept assistance. As the ψ_i^{human} increases, it implies they are willing to change and accept assistance in this dimension. The second region (in green) shows that as ψ_i^{human} gets closer to ψ_i^{robot} , corrections to ψ_i^{human} are less warranted. The third region (in yellow) shows when the ψ_i^{human} is faster than ψ_i^{robot} , they think their expected speed change is more desired as it is above the ψ_i^{robot} , and that the robot should be following suit. Note that there are two cases B_i can be 0, either when $\psi_i^{human} = 0$, or when ψ_i^{human} is very fast.

stochastic sampling approach by either 1) taking multiple time point samples or 2) utilizing uncertainty in the measurement of U_i^{human} . It should be noted that the rates are dynamic, in that, at each time point ψ_i^{robot} and ψ_i^{human} will change to new values. The willingness, B_i , to accept assistance is between 0 (fully unwilling) and 1 (fully willing).

$$B_i = \beta(|\psi_i^{human}|, |\psi_i^{robot}|, b, p, q) \quad (7)$$

We obtain a different B_i curve for each dimension. Fig. 4 shows (7) and its three distinct regions. The left part of the first region is when the operator has little or no ψ_i^{human} compared to ψ_i^{robot} . When the ψ_i^{human} is near 0, two aspects are needed: 1) constant U_i^{human} so the first and second terms of (6) can cancel out, and 2) $X_{i,current}$ configuration that leads to the third and fourth terms being either 0 (the third term is 0 when $X_{i,current} = X_{i,start}$ and the fourth term is 0 with no previous trials attempted) or cancel out ($X_{i,current} = X_{i,goal}$). ψ_i^{human} near 0 results in an unwillingness to accept assistance ($B_i = 0$). One example of when $\psi_i^{human} = 0$ occurs is near the goal—resulting in the third and last terms of (6) canceling out—with the operator no longer making adjustments to achieve the task, leading to the first and second terms approaching 0. Note that this is one example, but $\psi_i^{human} = 0$ could occur anywhere along the trajectory as long as the four terms cancel (e.g., when the operator moves in the y-direction in Fig. 5(b), the x-dimension ψ_i^{human} is approximately 0 leading to the purple distribution closely resembling the green one due to B_i being approximately 0). The other critical point is when ψ_i^{human} is slightly lower than ψ_i^{robot} ($B_i = 1$). In this scenario, the operator's expected speed change is not far off from the ψ_i^{robot} and the expectation is the operator desires to achieve ψ_i^{robot} .

The second region shows ψ_i^{human} nearing ψ_i^{robot} , where the operator's B_i is reduced and changes are less warranted. In this region, the psychology model of [23] states that the operator decreases their B_i as ψ_i^{human} nears the ψ_i^{robot} because they

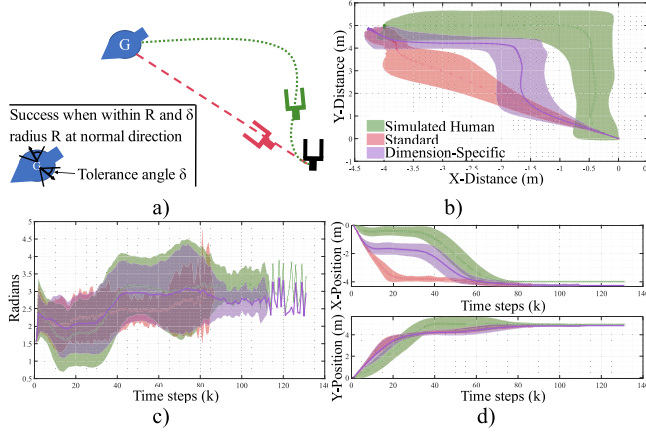


Fig. 5. (a) The simulation environment. The stochastic human has a policy in green and the robot policy is in red. The robot policy follows the shortest distance and smallest angle to the goal. The S controller follows (3) to set the α_i for all i . The starting pose is always identical (black). (b) The distributions of all simulated trajectories. (c) Is the heading across time, (d) Is the x and y directions across time. The plots in (d) also shows the trends of $P(H)$ in green, and $P(R)$ as red (S) and purple (DS). There is a large difference in the human's and S controller's x-position causing disagreement in control allocation of $P(H)$ and $P(R)$. There is less disagreement between the human and DS.

lose their sense of urgency to achieve it and their expected speed change is in an acceptable bound. The third region occurs when ψ_i^{human} is much faster than ψ_i^{robot} , implying the operator will have a strong desire to move in this direction. In this scenario, the operator does not care as much about the help because it moves slower. Note that in the third region, when ψ_i^{human} becomes extremely fast, B_i will go to 0.

C. Handling Disagreement in a Dimension-Specific Setting

The dimension-specific arbitration uses the disagreement between control authority allocation and the user's willingness. $P(H_i)$ and $P(R_i)$ act as bounds, where a better arbitration exists between them in (8). The new α_i tempers too aggressive control while boosting too timid control.

$$\alpha_i = (P(R_i) - P(H_i)) B_i + P(H_i) \quad (8)$$

After α_i from (8) is calculated, the output action from (1) is calculated. Each dimension will assume a different level of α_i . An example of possible α_i is shown in Fig. 3(a)–(d) in purple. For instance, consider if the current situation resembles Fig. 3(c), where an operator is behind the robot. If ψ_i^{human} is near 0 and thus $B_i = 0$ (region 1 of Fig. 4), then $\alpha_i = P(H_i)$. This gives the operator more control and tempering too aggressive assistance. If ψ_i^{human} begins to increase where $B_i = 1$ (near the max of Fig. 4), then at this time point, the control authority goes towards the robot at $\alpha_i = P(R_i)$. Providing more active assistance in tune with the operator. If ψ_i^{human} continues much faster at the next time point where a transition from Fig. 3(c) to (a) occurs, the B_i goes toward 0 (region 3 of Fig. 4) which pulls the robot toward the human's expected control authority as $\alpha_i = P(H_i)$ and boosts too timid assistance ($P(H_i)$ is dynamic, and changes, so, they are not the same $P(H_i)$). This does not mean the robot provides zero assistance but uses the operator's expected control authority allocation. In this transition phase, α_i first increases

from $P(H_i)$ to $P(R_i)$ in Fig. 3(c) then keeps increasing from $P(R_i)$ to $P(H_i)$ in Fig. 3(a). $P(H_i)$, $P(R_i)$ and B_i change along each dimension; thus, each α_i will be different. The linear model for (8) should not impact the operator's perception of rotational assistance as they could hold similar α_i but with two different poses. Rather the human-robot policy misalignment would have a higher influence (i.e., the robot wants $\pi - \theta_j^{robot}$, and the operator has the desired $\pi + \theta_j^{human}$).

IV. EXPERIMENTAL RESULTS

A. Simulation Setup

For easier visualization of the results and qualitative analysis of each α_i and the outcome of their respective actions, a simplified telemanipulation scenario with three DOF (two translational and one rotational) shown in Fig. 5(a) is used. Success occurs when the robotic hand is radius, R , away from the goal position and within a radian tolerance, δ , from the normal direction to the goal position. The tolerances resemble real scenarios where a point goal is not enough. The simulated human trajectories sample a probability distribution that attempts to grasp the object's handle. The robot's target pose has noise at the target position to simulate when a robot is not precise enough to achieve the desired pose. 1000 simulated trajectories were used to compare the baseline S controller (mentioned in III.A) and the Dimension-Specific (DS) controller. Fig. 5(b) shows all three trajectory distributions. The robot policy for the DS and S controllers is to follow the shortest distance and smallest angle to the goal. The DS controller starts with a B_i of 1 (fully willing) because no previous human actions indicate otherwise.

B. Simulation Results

Qualitative analysis comes in two forms. The first is the appearance of output robot trajectories for the same human trajectory. Our approach aims to align more with the principal axis the operator is trying to achieve. Fig. 5 shows this is evident compared to the singular α that over-dominates the trajectory (i.e., although the operator is moving primarily in the y-direction, the S controller forces motion in the x-direction). Fig. 5(d) also demonstrates the trends of $P(H)$ in green and $P(R)$ as red (S) and purple (DS). There is a large difference in the human's and S controller's x-position, causing disagreement in the control allocation of $P(H)$ and $P(R)$. There is less disagreement when considering humans and DS. The second qualitative test is to compare α_i . The α_i in Fig. 6 corresponds to the distributions in Fig. 5. The expected operator actions proceed in the y-direction before the x-direction. The DS controller should provide more control authority in y before the x-direction. In the top plot of Fig. 6, the α_i is relatively low. As the operator is not moving in this dimension, giving more control to the operator in this dimension can prevent the robot from over-extending its motion in an undesired path. Whereas, in the middle plot, the operator moves toward the component goal making the α_i go higher, and reduces the authority the operator experiences, allowing the robot to provide more assistance in this direction. Further, the S controller calculates (3) before developing an output action with

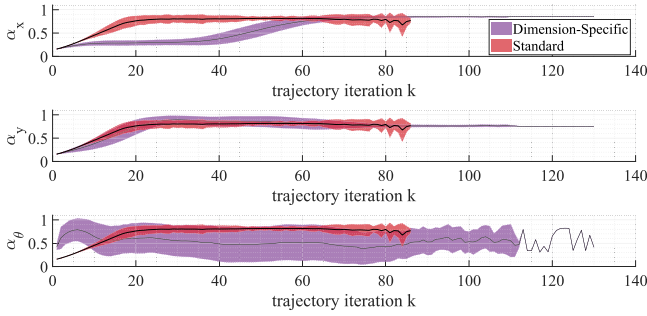


Fig. 6. α_i distributions for the trajectories shown in Fig. 5.

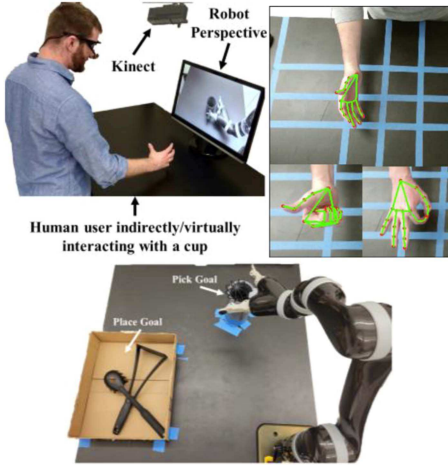


Fig. 7. Operators see the robot perspective on the screen while they move their hand freely over the table. The Xbox Kinect captures hand posture to extract the palm pose and the amount each finger is open or closed. The information is then sent to the robot to execute the motion.

(1), which is why the same alpha is along all components. The DS controller value for α_θ has two trends. The first trend is at the beginning of the trajectory; the human's current and target heading are close to one another. The human angle does not have significant deviations, so the robot is afforded more control with α_θ . The second trend occurs when the human heading begins to deviate, signaling to the robot that the human wants to assume more control to achieve a desired heading, driving α_θ lower overall. This becomes more apparent in a real-world scenario.

C. Telemanipulation Setup

Fig. 7 describes the telemanipulation experimental setup. An Xbox Kinect takes an RGB-D image, and Mediapipe [25] extracts features used as inputs for the robot control. The operator directly controls the hand pose (position and orientation) and the amount each finger opens and closes, eight variables in total. The operator aims to grasp kitchen utensils from a holder and place them in a bin. A failure occurs if an object does not reach the bin or lands on the table. The operator is free to move how they wish to accomplish the task. The robot is given pick and place goal poses with induced noise.

The institution's Institutional Review Board (IRB) approved a set of experiments. Before participating in the study, written consent was obtained, acknowledging they understood the robot

TABLE I
TELEOPERATION RESULTS FOR THEIR RESPECTIVE MEANS AND CONFIDENCE INTERVALS

Mode	Success Rate	Trial Time (s)	Cosine Distance	Ratio of Agreement
DS	0.446 [0.346,0.547]	152.4 [125.0,185.8]	0.076 [0.065,0.089]	0.981 [0.980,0.982]
S	0.326 [0.234,0.425]	223.4 [176.2,283.3]	0.400 [0.385,0.417]	0.895 [0.893,0.897]

DS means Dimension-Specific, S means Standard. Bold means this performs better in this metric by statistical significance.

setup, the purpose of the study, and the potential risks involved (arm soreness). Three trials were conducted for each arbitration strategy for each participant. A trial consisted of three runs; each run had a different object. A randomized order of the control modes was used to reduce learning effects. 10 participants volunteered (a total of 90 runs for each arbitration strategy) in the pilot study. The volunteer breakdown included 25.5 ± 4.15 years of age, with 4 women and 6 men involved.

D. Telemanipulation Results

Task performance. The quantitative analysis is broken down by success rate, completion time, cosine similarity, and the ratio of agreement. Due to the smaller number of trials compared to the simulation results, the metrics have a different analysis. The success rate uses a Laplace Estimate over the Maximum Likelihood Estimate to reflect the true success rate. Likewise, an adjusted-Wald 95% confidence interval creates the bounds of success. An N-1 Chi-Squared test determines statistical significance. All comparisons are in Table I. The DS controller achieves a higher success rate than the S strategy, but no statistical significance is found ($p = 0.092$). The success rate is low for both strategies due to a few factors that are not related to operator error as our scenario was designed to reenact real world challenges: 1) induced noise on the 2 goal locations, 2) induced time delay of the teleoperation scenario, 3) imperfect information as there was a single camera view so operators must overcome the depth perception. The induced noise of the goal state provides a more realistic scenario where imperfect perception causes the robot to have uncertainty about the goal. The induced time delay provides a realistic scenario. Therefore, both uncertainties reduce the success rate of the task.

Time data is notorious for not being positively skewed from a normal distribution [26]; thus, a log transform of the data must be used. The geometric mean and 95% confidence interval is used. A two-sample t-test is used to determine the statistical significance. The time data was the total trial data (the summed three runs); it includes both success and failures to determine if the control strategies limit an operator in multi-step tasks. The DS controller outperforms the S with a statistical significance of $p = 0.0141$. The task performance favors the DS controller with a better success rate and trial time than the S controller.

Assistance quality of the robot control. Objective measures for determining user agreement widen the gap further by exploring over-dominance. Two measures used to identify this are the

cosine distance and ratio of agreement. The former is defined:

$$\sum_{actions} 1 - \frac{U^{human} \cdot U^{output}}{\|U^{human}\| \|U^{output}\|}$$

The aim is to determine the alignment of the control vectors for the robot and the operator. A score of 0 is considered optimal. The ratio of agreement is defined:

$$\phi = \frac{\sum ag}{number\ of\ actions}, 0 \leq \phi < 1$$

$$ag = \begin{cases} 1, & \text{if } U^{human} \cdot U^{output} > 0 \\ 0, & \text{if } U^{human} \cdot U^{output} \leq 0 \end{cases}$$

This metric determines the number of output actions, U^{output} , which aligns with the human operator. The aim is to normalize the number of actions taken and determine a better ratio of agreeable actions. A score of 1 is considered optimal. The geometric mean and 95% confidence intervals for both metrics were generated and placed in Table I. The two-sample t-test is used to determine the statistical significance. The cosine similarity and ratio of agreement demonstrate that the DS controller follows human actions more than the S controller. The DS controller outperforms the S controller for both measures with $p < 0.001$. This strongly indicates that the DS controller assists in the direction the operator is moving while further showing the over-dominance of the robot actions occurring in the S controller. The DS controller shows greater alignment with the operator and improved task performance with the objective measures. The conclusion is that the DS controller benefits the operators in their preferred dimensions.

To further assess whether the assistance benefits the operator, qualitative analysis can be done by comparing the operator and robot trajectories. Our video shows the over-dominance issues regarding rotation differences. The robot output poses differ significantly for two similar human trajectories. The difference has major repercussions on the overall perception and adaptation of the operator as they must contend with a robot that dominates control near a goal location. The over-dominance causes some operators to hesitate as it is not similar to their hand posture; the operator moves to regain control. Another approach is to compare the trajectory differences of the operators and corresponding outputs, $X^{current}$, shown in Fig. 8. A Jensen-Shannon divergence measure compares the distributions which are bounded between 0 and 1. Two trends are apparent when analyzing the trajectory distributions. The first is that the S controller forces the robot behavior to act with a smaller deviation compared to the DS controller. This is evident by how Fig. 8(d) shows a rather straight-line trajectory between the 2 goal locations (pick and place goal), whereas Fig. 8(c) shows a more diverse suite of trajectories to aid the operator towards the goals. The divergence between the DS and the S controllers' outputs is 0.1896. The second discernable pattern is how much the robot conforms to the operators. The S controller (Fig. 8(b) to (d)) has a divergence of 0.3017 conforms less to the human than the DS controller (Fig. 8(a) to (c)) has a divergence of 0.2905. The higher divergence between the S controller means

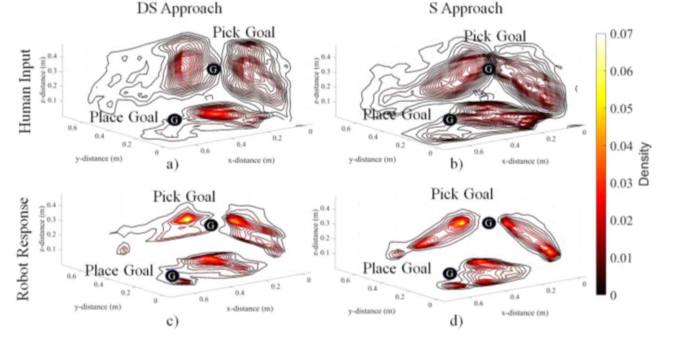


Fig. 8. Trajectory distributions in $\mathbb{R}^3$. (a) Human input for DS, (b) Human input for S, (c) Robot response for DS, (d) Robot response for S. A Jensen-Shannon divergence metric compares the distributions. (a) to (b) divergence is 0.1293, (c) to (d) is 0.1896, (a) to (c) is 0.2905, and (b) to (d) is 0.3017.

it conforms less than the DS controller. The more concentrated regions in Fig. 8(b) compared to Fig. 8(a) suggest that operators were fighting the robot near the pick goal location, which leads to the credence that the S controller over-dominates the control and does not allow the human to compensate. All the divergence values are relatively low and similar to one another as there are only 2 goal locations to move between (meaning there are limited strategies for reaching both goals).

V. DISCUSSION

A strategy to enable dimension-specific shared control has been developed. By breaking down the problem into subproblems, each dimension has an individual arbitration curve that allows the operator to adapt to the task. Further, we provide a safeguard to prevent over-dominance. This is evident by the improved similarity and agreement for the DS controller over the conventional approach. Although, it is difficult to say with certainty that the proposed strategy improves task performance. The low success rates for both strategies prove that this environment and task were challenging for operators, which may have also contributed to slower reported times. In ideal cases, the S controller should be faster as the robot has higher levels of authority; however, in the real-world, our approach yields faster times. This is likely due to the operator handling uncertainties of the robot and reducing unnecessary adjustment times. The proposed strategy could benefit remote factory and facility maintenance, telenursing, and assistive living tasks. For example, the assistive living tasks where operators may have varying degrees of desired assistance [15], and our approach can provide an assistance level more in line with the operator.

Our strategy aims to give appropriate assistance to improve performance in challenging tasks. The proposed rotational arbitration is necessary for providing an operator with seamless assistance and may benefit users with more limited mobility as the agent can understand when to assume more or less control of the system. The rotational arbitration's limitations are two-fold: 1) the hyperparameter tuning of (5) and (7), and 2) a temporal filter window on both the α_i and B_i to smooth out spikes that could occur in control authority. These are

relatively straightforward limitations and should be tuned based on the interface, setup, and tasks the shared control system needs to handle. Future work directly stemming from this paper should aim to improve the success rate by 1) devising new arbitration curves to handle specific control dimensions (translation vs. rotation) based on properties discussed in this work and 2) investigating alternative strategies for adjusting control allocation on the fly (i.e., alternative models to B_i) in a dimension-specific context to handle the misalignment issues.

Conventional arbitration may hold advantages in a few scenarios: 1) in very controlled settings with low uncertainty, 2) in lower DOF interfaces (i.e., a controller with 2 or 3 DOFs rather than the 6 DOFs), 3) in tasks that do not require a heavy influence on rotation such as tracing tasks on a fixed plane. First, the autonomous agent should yield an optimal action plan where an operator needs limited interaction. Towards the second, it may be beneficial for low-DOF interfaces to use a synergy-style approach [9] where more assistance is warranted. Towards the third, simpler tasks where automated subroutines shine may allow the autonomy to offload all user control. Regardless, our strategy provides an avenue for meaningful shared control in environments where an operator desires more control over telemanipulation systems.

REFERENCES

- [1] Y. Rybarczyk, E. Colle, and P. Hoppenot, "Contribution of neuroscience to the teleoperation of the rehabilitation robot," in *Proc. IEEE Int. Conf. Syst., Man Cybern.*, 2002, vol. 4, pp. 6, doi: [10.1109/ICSMC.2002.1173273](https://doi.org/10.1109/ICSMC.2002.1173273).
- [2] N. Healey, "Speculation on the neuropsychology of teleoperation: Implications for presence research and minimally invasive surgery," *Presence*, vol. 17, no. 2, pp. 199–211, 2008.
- [3] Y. Oh, M. Toussaint, and J. Mainprice, "Learning to arbitrate human and robot control using disagreement between sub-policies," in *Proc. IEEE/RSJ Int. Conf. Intell. Robots Syst.*, 2021, pp. 5305–5311.
- [4] M. Bowman, S. Li, and X. Zhang, "Intent-uncertainty-aware grasp planning for robust robot assistance in telemanipulation," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2019, pp. 409–415.
- [5] D. Song, C. H. Ek, K. Huebner, and D. Kragic, "Task based robot grasp planning using probabilistic inference," *IEEE Trans. Robot.*, vol. 31, no. 3, pp. 546–561, Jun. 2015.
- [6] A. Dragan and S. S. Srinivasa, "A policy blending formalism for shared control," *Int. J. Robot. Res.*, vol. 32, 2013, pp. 790–805.
- [7] A. Handa et al., "DexPilot: Vision-based teleoperation of dexterous robotic hand-arm system," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2020, pp. 9164–9170.
- [8] G. Gioioso, G. Salvietti, M. Malvezzi, and D. Prattichizzo, "Mapping synergies from human to robotic hands with dissimilar kinematics: An approach in the object domain," *IEEE Trans. Robot.*, vol. 29, no. 4, pp. 825–837, Aug. 2013.
- [9] D. P. Losey et al., "Learning latent actions to control assistive robots," *Auton. Robots*, vol. 46, pp. 115–147, 2022.
- [10] H. Wang and X. P. Liu, "Adaptive shared control for a novel mobile assistive robot," *IEEE/ASME Trans. Mechatron.*, vol. 19, no. 6, pp. 1725–1736, Dec. 2014.
- [11] D. Gopinath, S. Jain, and B. D. Argall, "Human-in-the-loop optimization of shared autonomy in assistive robotics," *IEEE Robot. Automat. Lett.*, vol. 2, no. 1, pp. 247–254, Jan. 2017.
- [12] A. A. Allaban, V. Dimitrov, and T. Padir, "A blended human-robot shared control framework to handle drift and latency," in *Proc. IEEE Int. Symp. Saf., Secur., Rescue Robot.*, 2019, pp. 81–87.
- [13] C. Schultz, S. Gaurav, M. Monfort, L. Zhang, and B. D. Ziebart, "Goal-predictive robotic teleoperation from noisy sensors," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2017, pp. 5377–5383.
- [14] S. Javdani, H. Admoni, S. Pellegrinelli, S. S. Srinivasa, and J. A. Bagnell, "Shared autonomy via hindsight optimization for teleoperation and teaming," *Int. J. Robot. Res.*, vol. 37, 2018, pp. 717–742.
- [15] D.-J. Kim et al., "How autonomy impacts performance and satisfaction: Results from a study with spinal cord injured subjects using an assistive robot," *IEEE Trans. Syst., Man, Cybern.*, vol. 42, no. 1, pp. 2–14, Jan. 2012.
- [16] A. Broad, T. Murphey, and B. Argall, "Operation and imitation under safety-aware shared control," in *Proc. Int. Workshop Algorithmic Found. Robot.*, 2020, pp. 905–920.
- [17] A. Broad, T. Murphey, and B. Argall, "Highly parallelized data-driven MPC for minimal intervention shared control," *Proc. Robot. Sci. Syst.*, 2019, doi: [10.15607/RSS.2019.XV.008](https://doi.org/10.15607/RSS.2019.XV.008).
- [18] Y. Oh, S.-W. Wu, M. Toussaint, and J. Mainprice, "Natural gradient shared control," in *Proc. IEEE Int. Conf. Robot Hum. Interactive Commun.*, 2020, pp. 1223–1229.
- [19] S. Jain and B. Argall, "Recursive Bayesian human intent recognition in shared-control robotics," in *Proc. IEEE/RSJ Int. Conf. Intell. Robots Syst.*, 2018, pp. 3905–3912.
- [20] S. Nikolaidis, Y. X. Zhu, D. Hsu, and S. Srinivasa, "Human-robot mutual adaptation in shared autonomy," in *Proc. 12th ACM/IEEE Int. Conf. Hum.-Robot Interact.*, 2017, pp. 294–302.
- [21] R. Balachandran et al., "Adaptive authority allocation in shared control of robots using Bayesian filters," in *Proc. IEEE Int. Conf. Robot. Automat.*, 2022, pp. 11298–11304.
- [22] C. S. Carver, "Negative affects deriving from the behavioral approach system," *Emotion*, vol. 4, no. 1, pp. 3–22, 2004.
- [23] C. S. Carver, "Control processes, priority management, and affective dynamics," *Emotion Rev.*, vol. 7, no. 4, pp. 301–307, 2015.
- [24] C. S. Carver, "Control theory, goal attainment, and psychopathology," *Psychol. Inquiry*, vol. 29, no. 3, pp. 139–144, 2018.
- [25] C. Lugaresi et al., "Mediapipe: A framework for perceiving and processing reality," in *Proc. 3rd workshop Comput. Vis. AR/VR IEEE Comput. Vis. Pattern Recognit.*, 2019.
- [26] J. Sauro et al., "Average task times in usability tests: What to report?," in *Proc. SIGCHI Conf. Hum. Factors Comput. Syst.*, 2010, pp. 2347–2350.