

# Swarm UIs: Impact of Assistance on Users' Sense of Agency

OPHÉLIE JOBERT, Univ. Grenoble Alpes, CNRS, Grenoble INP, LIG, France

AMINA KORGHLOU, YELLI COULIBALY, and THIBAUT LEONE, Univ. Grenoble Alpes, CNRS, Grenoble INP, LIG, France

ALIX GOGUEY, Univ. Grenoble Alpes, CNRS, Grenoble INP, LIG, France

BRUNO BERBERIAN, ONERA, DTIS/ICNA, France

JULIEN BOURGEOIS, Université Marie et Louis Pasteur, institut FEMTO-ST, CNRS, France

CELINE COUTRIX, Univ. Grenoble Alpes, CNRS, Grenoble INP, LIG, France

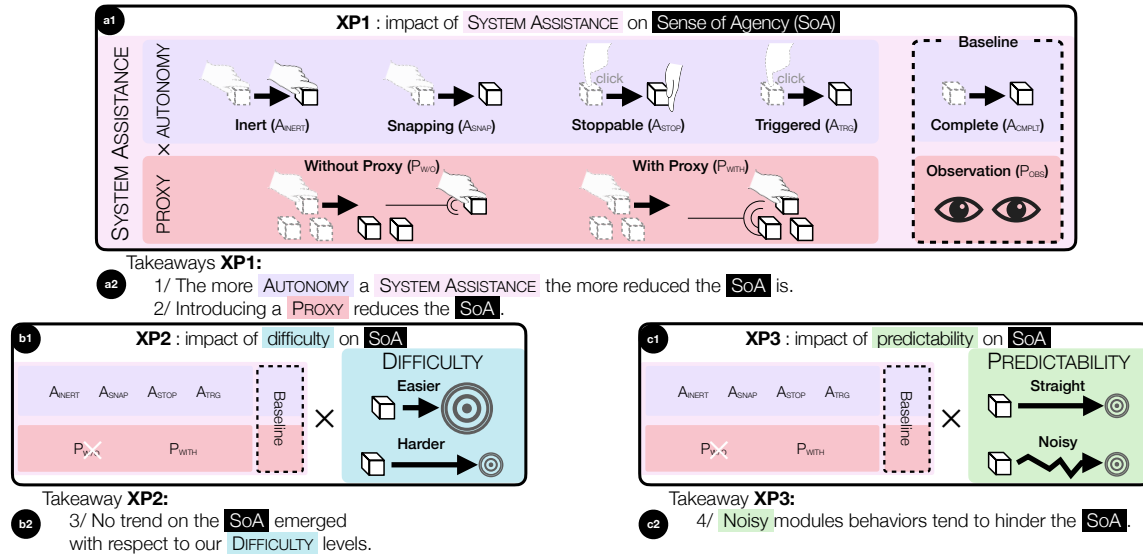


Fig. 1. This paper explores the relationships between the Sense of Agency (SoA) and System Assistance (XP1), Task Performance (XP2), and System Predictability (XP3) using a drag-and-drop task in which users moved Toio<sup>TM</sup> robots onto targets. **(a1)** In XP1, we varied the system assistance through the combination of 4 AUTONOMY ( $A_{INERT}$ ,  $A_{SNAP}$ ,  $A_{STOP}$ , and  $A_{TRG}$ ) levels  $\times$  2 PROXY ( $P_{W/O}$  and  $P_{WITH}$ ) levels resulting in 8 TYPES OF ASSISTANCE, to which we added a BASELINE (with its own  $A_{CMPLT}$  autonomy and  $P_{OBS}$  proxy levels). **(b1)** In XP2, the task performance varies with the targets' size and distance (task DIFFICULTY). We dropped the  $P_{W/O}$  level and tested the remaining 5 TYPES OF ASSISTANCE (including the BASELINE) on 2 task DIFFICULTY (*Easier* and *Harder*) levels. **(c1)** In XP3, the system predictability varies with two types of trajectories (system PREDICTABILITY). We tested the same 5 TYPES OF ASSISTANCE (including the BASELINE), as in XP2, on 2 system PREDICTABILITY (*Straight* and *Noisy*) levels. **(a2)**, **(b2)**, and **(c2)** summarize the takeaways of XP1, XP2 and XP3.

CCS Concepts: • **Human-centered computing**  $\rightarrow$  **Human computer interaction (HCI)**.

Additional Key Words and Phrases: Swarm UI, Sense of Agency, Autonomy, Assistance, Proxy

Authors' Contact Information: Ophélie Jobert, Univ. Grenoble Alpes, CNRS, Grenoble INP, LIG, Grenoble, France, joberto@univ-grenoble-alpes.fr; Amina Korghlou, Yelli Coulibaly; Thibaut Leone, Univ. Grenoble Alpes, CNRS, Grenoble INP, LIG, Grenoble, France, korghloua@univ-grenoble-alpes.fr; Alix Goguey, Univ. Grenoble Alpes, CNRS, Grenoble INP, LIG, Grenoble, France, gogueya@univ-grenoble-alpes.fr; Bruno Berberian, ONERA, DTIS/ICNA, Salon de Provence, France, bruno.berberian@onera.fr; Julien Bourgeois, Université Marie et Louis Pasteur, institut FEMTO-ST, CNRS, Montbéliard, France, julien.bourgeois@univ-fcomte.fr; Celine Coutrix, Univ. Grenoble Alpes, CNRS, Grenoble INP, LIG, Grenoble, France, celine.coutrix@cnrs.fr.

## Abstract

Swarm UIs provide assistance to support users in their tasks and are increasingly explored in HCI. This paper studies the extent to which this assistance impacts users' sense of agency. A reduced sense of agency can lead to non-use of the interface or a diminishing sense of responsibility regarding the consequences of users' actions. We conduct three experiments studying the impact of three factors on the sense of agency: the level of assistance, the task difficulty, and the predictability of modules. Our nine assistance levels vary in system autonomy and module coordination (proxy vs. no proxy). We find that higher assistance reduces users' sense of agency, and this effect is not impacted by task difficulty. Predictability only impacts the least assistive interaction techniques. Our results will foster users' acceptance, responsibility, and use of swarm UIs.

## 1 Introduction

Swarm UIs are interfaces made up of several modules that can move independently, without needing to be physically connected to one another. They operate collectively and respond to user commands as a group [67]. These interfaces enable new ways for users to enter commands to the system [20, 57, 67, 80, 87, 108], or for the system to provide feedback to users [58, 59, 67, 108]. The actuation allows swarm UIs to provide assistance to users. For instance, actuated modules can behave autonomously, or follow the movement of another module, as illustrated in recent HCI research [67].

We currently do not know the effect of such assistance on users' sense of agency (SoA) (Figure 1a<sub>1</sub>) –that is, the feeling of controlling an action and its effects in the environment [49]. For example, too much assistance could lead to a loss of control for users. The problem we address in this paper is the lack of balance between the assistance provided by Swarm UIs and users' SoA.

Additionally, the SoA may be influenced by two other factors: (1) users' task performance [54, 77, 115, 122] and (2) the predictability of the response of the swarm [23, 24, 101, 102, 123]. First, on the one hand, users' task performance can decrease with increasing task complexity [40], which in turn may reduce their SoA [77]. On the other hand, assistance can make up for increased task complexity, e.g., trajectory correction, and could therefore improve users' SoA [54, 122]. We currently do not know the combined effect on users' sense of agency of assistance and performance of swarm UIs (Figure 1b<sub>1</sub>). Second, a mismatch between users' predictions and the actual system response can decrease the SoA [24, 101, 102] –such as when turbulence happens in the system movement [102]. Assistance could worsen the predictability of swarm UIs, and further reduce users' SoA. We currently do not know the combined effect on users' SoA of assistance and predictability (Figure 1c<sub>1</sub>). This paper addresses the following research question: *To what extent does the assistance provided by swarm UIs lower the users' sense of agency and how task performance and system predictability modulate this effect?*

Answering this question is crucial, as a reduced SoA can lower users' perceived responsibility for the outcomes of their actions [22]. This is especially critical for high-stakes applications –such as surgery [104] or search-and-rescue coordination [52]– where users must retain a sense of control to make informed decisions taking into account any potential consequences. Moreover, user's ability to maintain control over a system impacts its acceptability [100].

A consensual grand challenge in shape-changing UIs is to be able to isolate experimental factors and to take into account concepts from Human-Robot Interaction [1]. To achieve this goal, we defined nine assistance types based on examples found in the literature [57, 67, 80, 87, 108]: eight assistance levels and one baseline shown in Figure 1. We structured these types along two factors: (1) the degree of modules' autonomy at different stages of the task, and (2) whether or not the swarm's movement is controlled through a single module (i.e., a module serving as a proxy for users

to control the satellite modules). We experimentally measure the impact on the SoA of these assistance types, as well as the impact of task difficulty –or users' performance–, and predictability of modules' trajectory (Figure 1).

To do so, we conducted three experiments (total N=105, each N=35) in which participants were asked to increase the size of a square made of a swarm of Toios<sup>TM</sup> [111] to a target size. The first experiment studies the impact on the SoA of the assistance type across our nine levels, from modules with no autonomy to fully autonomous modules. The second experiment studies the combined impact on the SoA of the assistance type and the difficulty of the task –through the position and the size of the target [40]. The third experiment studies the combined impact on the SoA of the assistance type and the predictability of the modules' trajectory.

This article presents the following contributions: First, we provide the first empirical evidence of a decrease in SoA with increasing assistance in the context of swarm UIs. Second, we are the first to shed light on three novel types of SoA that are specific to swarm UIs, when using a proxy: (a) a sense of agency for the proxy module, (b) one for the other modules in the swarm (satellites), and (c) one for the swarm as a whole. Third, we are also the first to show that the task difficulty has no impact on the SoA for swarm UIs, and that the predictability of the modules' trajectory impacts the SoA of the whole swarm, but not the ones of the proxy and satellite modules. This paper provides advice for swarm UI designers to introduce assistance while preserving users' SoA, therefore promoting users' responsibility and interface adoption.

## 2 Related Work

We review studies at the intersection between system assistance, swarm UIs, and sense of agency.

### 2.1 System assistance

Many applications introduced assistance to help users perform their tasks, such as driving [34, 73], text/image/video editing [46], or target acquisition (e.g., [29, 36]). In this section, we focus on assistance provided with swarm UIs. Assistance is the degree of autonomy left to a system to help the user perform a task. Sheridan et al. [99] suggest ten assistance levels based on the system automation, from no assistance (i.e., without automation, level 1) to complete assistance (i.e., full automation, level 10). Levels 2 to 4 define the balance between agents (human vs. system) in making decisions. Levels 5 to 7 define how the system performs the action. For example, Walker et al. [119] applied two of these assistance levels to swarm interfaces controlled by a graphical interface. However, all these levels have been established for teleoperated systems, as opposed to direct interaction with swarms that we study in this paper.

To identify levels of assistance relevant for direct manipulation of swarm UIs, we searched for examples in prior work. In addition to a systematic literature review [90], we additionally studied more recent work [56, 91], and work conducted in broader contexts (e.g., drones [20, 52] or larger robots [80]). From these papers, we kept the ones presenting usage scenarios [20, 44, 59, 67, 107, 108] and interaction techniques [29, 52, 56, 58, 59, 80, 91, 122], including those suggested by users [57, 87]. In this set of papers, we found that swarm UIs can:

- (1) Be completely inert [44, 57]. In this case, the user must move each module manually, for instance, to create a secret code [44].
- (2) Be completely autonomous, triggered either by users [57, 58, 87, 107, 108] or by the system [20, 59, 108]. For instance, when reading a book, pressing on a word can make the swarm UI take the form of the word being read [107, 108].

- (3) Assist the users by letting the system do most of the task, but allowing the users to stop the automation whenever they want [52, 57]. For instance, in a study on drones for victim search, automation was found beneficial, but users must retain the ability to take action at any moment [52].
- (4) Assist the user in achieving their goal, for example by improving final accuracy [122], correcting trajectories [122], or speeding up task completion by finalizing the objective [29].
- (5) Reduce the user's physical or cognitive effort. In this case, the swarm UI may be perceived as a coordinated group, where the movements of one or several modules guide the behavior of the swarm. The user directly manipulates only a limited number of modules in order to interact with the entire swarm or with subgroups [44, 57, 67, 80, 86, 91]. For example, the user can move a single module freely, while the swarm imitates the motion of this *proxy* module [57, 67, 80, 91].

We base the levels of our experimental variable on these assistance levels.

## 2.2 Swarm UIs

The actuation of swarm UIs enabled researchers to explore interaction techniques for their control [57, 80, 87], users' perception of their speed, their motion types, and haptic feedback [35, 58, 59], as well as potential use cases [20, 67, 108]. Prior studies focus on identifying effective interaction techniques but overlook users' feelings, such as emotion or sense of control. However, how people feel about a UI can strongly impact its adoption [81] and use [63, 71].

Prior work studied users' emotional perception of swarm UIs as a function of the modules' speed, fluidity, synchronization, or movement force [37, 58, 59, 64]. Kim and Follmer [59] investigated perception during social contact through remote touch via a swarm UI. Participants were instructed to control the swarm by moving their dominant fingers on a touchscreen to create a haptic pattern for each social interaction. After each trial, they were asked to rate the ease of control. Here, ease of control was therefore evaluated for touchscreen control, which is indirect and external to the swarm interface. Yet, eliminating the need for a third-party interface would enable numerous additional use cases through direct interaction with the robotic modules, as in [44, 57]. While these prior work studied the affective and subjective perception of swarm UIs, we found no prior work on the sense of agency.

## 2.3 Sense of agency (SoA)

The sense of agency is defined as the perception of control over one's own actions and, through these, over events in the external world [49]. It is considered a key factor in the adoption [81] and use of interfaces [71]. Moreover, SoA is a key factor in assigning causal responsibility and can have a drastic impact on users' engagement in their tasks [12]. For these reasons, an interface must allow users to maintain a sense of internal control [100]. In other words, users must have the feeling that their actions are at the origin of the system's responses.

**2.3.1 SoA and HCI.** HCI recently studied sense of agency [13–15, 71, 75, 82], for instance, to compare interaction techniques such as on-skin vs. physical buttons [16, 17, 29], or to evaluate the benefits of vibrotactile feedback [95]. These studies measured the SoA either implicitly [13–17, 29, 71, 75, 82] or explicitly [13–16, 71, 82, 95]. On the one hand, among the implicit measures, the *intentional binding* constitutes one of the main markers of SoA. *Intentional binding* leverages the fact that voluntary actions and their sensory effects are perceived as being temporally closer to each other, compared to the effects of involuntary or reflex actions [50]. Although questioned in the literature [61, 105, 110], this time compression remains one of the most robust markers of the SoA.

On the other hand, one can explicitly measure the SoA by asking users to quantify via a Likert scale the extent to which they felt at the origin/in control of the events [18, 27, 28]. Although it may be sensitive to the expectations of participants, explicit measurements are a more direct measure of control experience and have the advantage of being simpler to implement than the *intentional binding*. In our study, we use the explicit measurement.

In this paper, we study the impact on the SoA of the assistance type, the user's performance, and the predictability of the system. We now review prior work on the SoA exploring the predictability of the system, the performance when performing the task, and the assistance the system provides users to perform tasks, in particular when they cooperate.

**2.3.2 SoA, Predictability and Performance.** The sense of agency is a subjective experience that seems to depend on multiple cues [109]. Among these cues, the level of performance achieved, the effort put in, and the predictability of our actions seem to play a major role. Previous work showed a reduced SoA when the system behaves in unexpected ways [23, 24, 101, 102, 123]. However, no prior work studied the combined impact on users' SoA of the predictability (e.g., trajectory errors or bugs) and the level of assistance.

Conversely, prior work showed that performance can modulate the SoA [54, 77, 115, 122]. For example, van der Wel et al. [115] demonstrated that performance was positively associated with SoA: the higher the performance, the stronger the participants' SoA. The authors suggested that improved task performance likely leads to greater alignment between participants' performance expectations and their actual outcomes, thereby reinforcing SoA. This interpretation is consistent with the idea that SoA increases when the expected and actual consequences of one's actions coincide. However, the system assistance, which could improve performance, could also strengthen the users' SoA.

Moreover, prior work examined the relationship between SoA and task effort. While prior work showed that greater physical effort seems to lead to an increased SoA [32, 78], others showed that greater cognitive effort seems to lead to a decreased SoA [10, 11, 53]. Increasing the number of possibilities offered could decrease the subjective effort and improve the SoA [10].

Our experiments will explore the respective impact of robot predictability and performance on SoA while interacting with swarm UIs.

**2.3.3 SoA and Assistance.** Beyond the impact of predictability and performance, our main interest concerns the impact of the level and type of assistance on users' sense of agency. Prior work suggests that high levels of assistance may reduce the SoA [29, 71, 114, 126, 127]. Coyle et al. [29] investigated the effect of system assistance on the SoA in a target acquisition task. Their system provides assistance through a force drawing the cursor towards the target. Their results suggest that the SoA decreases as the assistance increases. Moreover, Berberian et al. [15] investigated the effect of system assistance in an aircraft supervision task. The assistance improved gradually from full control by the operator to full control by the system. Their results suggest a strong and gradual decrease in SoA when assistance improves.

Prior work also examined the effects of assistance, mental workload, and time delay on *intentional binding* when triggering a sound via a button [126, 127]. Results indicate a decrease in the SoA when assistance increases, regardless of the mental workload. Wen et al. [122] showed that, when task performance is taken into account, the SoA tends to increase with better performance, regardless of the assistance level. This suggests that assistance can impact the SoA either directly or indirectly by reinforcing performance. However, Ueda et al. [114] showed that this effect holds only up to a breaking point: when system assistance becomes excessive, the SoA decreases even if performance remains high. This implies that an optimal balance must be found between automation, performance, and SoA to foster effective between the system and the user [12, 13, 114].

It is important to note that these studies were conducted in the context of graphical interfaces. In the case of swarm UIs, the introduction of direct physical manipulation of modules may challenge existing findings on the sense of agency. As shown above, physical involvement in the task may increase users' SoA. Moreover, these studies considered the system as a tool rather than as a cooperative partner. The next section presents how the SoA is impacted by such a cooperation between users and a swarm.

**2.3.4 SoA and Cooperation.** Cooperation happens when “two or more individuals work together toward the attainment of a mutual goal or complementary goals” (APA [6]). Prior work studied the impact of cooperation on the sense of agency, in the case of cooperation between humans [82, 105] and between humans and machines [27, 28, 47, 83, 93, 96].

In social situations, a particular type of SoA emerges whenever the consequences of an individual's actions are linked to those of a partner [103]. The concept of *We-Agency* [82, 83] suggests that individual SoA can dissolve into a joint sense of agency when actions are performed cooperatively. However, this dissolution has been criticized due to insufficient empirical support [66]. When pursuing a common goal, studies show that joint SoA does not replace individual SoA, but supplements it through the development of (1) a *shared* sense of agency [66] and (2) a *vicarious* sense of agency for the partner's actions [82, 105].

Obhi and Hall [83] investigated which of these was at play when cooperating on a task with a computer [83]. Participants performed a joint action either with a human or with a computer. The authors measured the SoA indirectly and asked participants to estimate who was responsible for the action. The SoA for the co-agent's actions appeared when the partner was human, but not when it was a computer. Similarly, Sahai et al. [97] studied participants' SoA in a *Simon Social* task [?] under three conditions: individually, in cooperation with a human, or in cooperation with a computer. The results showed that participants had a vicarious SoA when cooperating with humans, but not when cooperating with computers.

Another hypothesis is the distribution of SoA among partners. Prior work showed that the individual SoA tends to decrease during joint action, being redistributed to other human partner [18, 72, 92, 128] or to robotic partner [27]. Such distributed and joint SoA can only emerge in synchronous tasks, i.e., when the behavior of other agents aligns with one's expectations [69, 92].

However, on the contrary to the systems considered in the above studies, swarm UIs can physically move in space to perform an action, creating a sensory effect similar to human actions [93].

Prior work investigated the SoA when cooperating with robots [27, 28, 47, 93, 96]. In most cases, prior work shows a negative effect of the robotic partner on the SoA compared to the absence of a partner [27, 28, 47, 96]. However, Roselli et al. [93] compared the SoA when the robotic partner physically performs the task, or when the same system was physically present but performed the task without physical motion. The results showed a *vicarious* sense of agency for the physical action, while it decreased when the task was performed without physical motion. Similarly, Sahai et al. [96] compared the SoA for three types of partners: (1) a servomotor, (2) a humanoid robot, and (3) a human. Results showed a significant difference between the three types of partners (human > humanoid robot > servomotor). This suggests that the type of autonomous system involved in the cooperation task has an effect on the SoA.

However, in these studies, partners did not interact with each other when cooperating. On the contrary, in existing swarm UIs, users can directly interact with the modules to achieve their goal (e.g., [57, 67, 108]). Moreover, as opposed to swarm UIs, prior studies focused on a single robot. To the best of our knowledge, no prior work investigated the impact on SoA of the assistance provided by a swarm UI during a cooperative task yet.

We now present three experiments, sharing similar procedures, apparatus, and experimental design (Figures 1a<sub>1</sub>, 1b<sub>1</sub> and 1c<sub>1</sub>).

Across these experiments, we explore the impact of the assistance on users' sense of agency when interacting with a swarm UI (Experiment 1), and how this is further impacted by performance (Experiment 2) and predictability (Experiment 3).

### 3 Experiment 1: Effect of the TYPE OF ASSISTANCE on SoA

#### 3.1 Method

The aim of this experiment is to evaluate the impact of the assistance provided by a swarm of robots on the users' sense of agency. To do this, participants performed a drag-and-drop task to resize a circle made of robots. They dragged robots initially arranged in a circle at the center of a platform, and dropped them into a circular target (6 cm in diameter) positioned 15.5 cm away (center to center) in front of each robot (Figures 2 and 3).

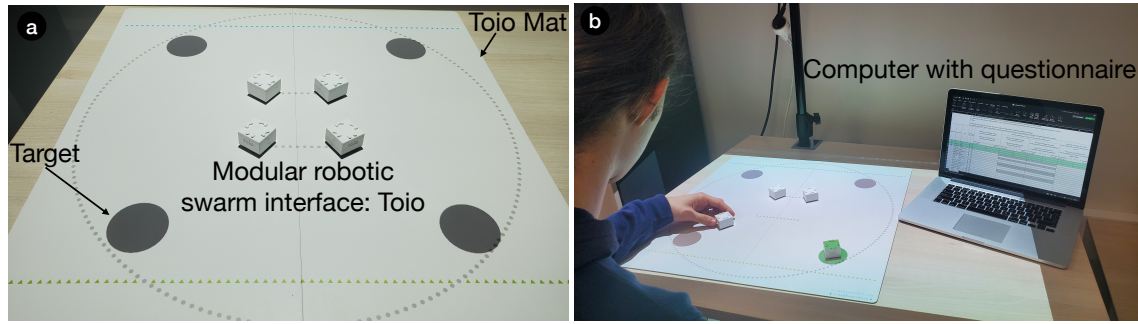


Fig. 2. Experimental setup: (a) Swarm UI consisting of 4 Toios modules, in a square starting position, in front of their gray target vertex, (b) Participant dragging a module, before they answer our sense of agency questionnaire on the right.

**3.1.1 Participants.** We recruited 35 participants by e-mail (18 self-identified as women, 17 as men,  $M = 29.26$  y.o.,  $SD = 7.57$ ). Each participant received a €15 gift voucher. Eighteen work in IT and 17 in other fields.

**3.1.2 Apparatus.** We used four rectangular Toio™ robots [111] (3 cm × 3 cm, height of 2 cm). We considered using more robots, but four was the limit to keep the experiment within an acceptable length (90 min on average). The Toios were controlled by Bluetooth via the *Toio SDK* for Unity [79]. We used Unity 2021.3.0f1 on a MacBook Pro M1 Max. To enable future replication, we provide the code of our experimental software in the supplementary material. The robots and their mat were placed on a table in front of the participants (Figure 2). To project the targets onto the mat, we placed a video projector above the mat.

**3.1.3 Experimental Design.** We used a within-subjects design, with two independent variables and one main dependent.

*Independent Variables: AUTONOMY and PROXY.* Our first experiment involves two independent variables controlling the assistance provided by the Toios: AUTONOMY and PROXY. We call a TYPE OF ASSISTANCE a combination of the AUTONOMY × PROXY factors (Figure 1a<sub>1</sub>).

We also added the “complete” TYPE OF ASSISTANCE ( $A_{\text{CMPLT}} \times P_{\text{OBS}}$ ) as our BASELINE: in this condition, the participants do not interact with the robots while they move towards their targets. In the results, we consider it as a distinct TYPE OF ASSISTANCE level, as well as a distinct AUTONOMY level and a distinct PROXY level.

*AUTONOMY.* The participants drag-and-drop the robots into their targets in five different ways (lines of Figure 3):

**INERT ( $A_{\text{INERT}}$ )** The robots are completely inert (i.e., no autonomy). Participants drag each robot from its initial position and drop it on their target. Participants control the activation, the end of the movement, and its trajectory.

**SNAPPING ( $A_{\text{SNAP}}$ )** The robots are semi-autonomous. Participants slide each robot from its initial position to the vicinity of its target. When the robot is sufficiently close to its target (i.e., within 9.5 cm of the target center), it moves autonomously to complete its movement (i.e., enters the target on its own). Participants control the activation and most of the trajectory.

**STOPPABLE ( $A_{\text{STOP}}$ )** The robots are semi-autonomous. Participants initiate the movement of each robot by clicking on it. The robot then moves autonomously towards its target. Participants then stop the robot by holding it still for 150 ms, when it has reached its target. Participants control the activation and end of the movement.

**TRIGGERED ( $A_{\text{TRG}}$ )** The robots are highly autonomous. Participants initiate the movement of each robot by clicking on it. The robot moves autonomously to its target. Participants control the activation, but not the end of movement or its trajectory.

**COMPLETE ( $A_{\text{CMPLT}}$ )** This level is the BASELINE of the AUTONOMY factor. The robots are completely autonomous. Participants do not interact with the robots. All at once, the robots move autonomously towards their targets. Participants have no control over the activation, trajectory, or end of movement.

*PROXY.* The robots imitate or not the movements of a leader, called *PROXY* (columns of Figure 3) :

**WITHOUT PROXY ( $P_{\text{w/o}}$ )** Participants interact with each robot to complete the task. All robots behave as described above, according to their level of AUTONOMY.

**WITH PROXY ( $P_{\text{WITH}}$ )** Participants interact *only with the proxy robot* to drag and drop *all of them* into their targets. Only the proxy robot behaves as described above in the AUTONOMY levels. The other robots, called satellites [44], mimic the proxy robot’s movements to reach their target. Participants define the proxy robot by clicking on it at the beginning of the task.

**OBSERVATION ( $P_{\text{OBS}}$ )** This level is the BASELINE of the PROXY factor. Participants do not interact with any robots. Participants only observe the robots. All at once, the robots move autonomously to their target.

*Dependent Variable.* Our dependent variable is participants’ SENSE OF AGENCY (SoA) on the modules’ movement<sup>1</sup>. We measure the overall SoA for each TYPE OF ASSISTANCE with the question “How much control did you have over the robots’ movements?” using an 8-point Likert scale [27, 28]. We also measure the SoA on the proxy and satellites for the TYPES OF ASSISTANCE  $P_{\text{WITH}}$  by the questions “How much control did you have over the Proxy movement?” and “How much control did you have over Other robots’ movements (satellites)?” using the same scale.

We choose a protocol with an explicit measurement of the sense of agency. As discussed in the related work (section 2.3.1), it is quick to set up, unlike implicit techniques, such as the *intentional binding*, and is the most direct measure of sense of agency [33].

<sup>1</sup>For more condensed writing, we omit the phrase “on modules’ movement” in the rest of the paper.



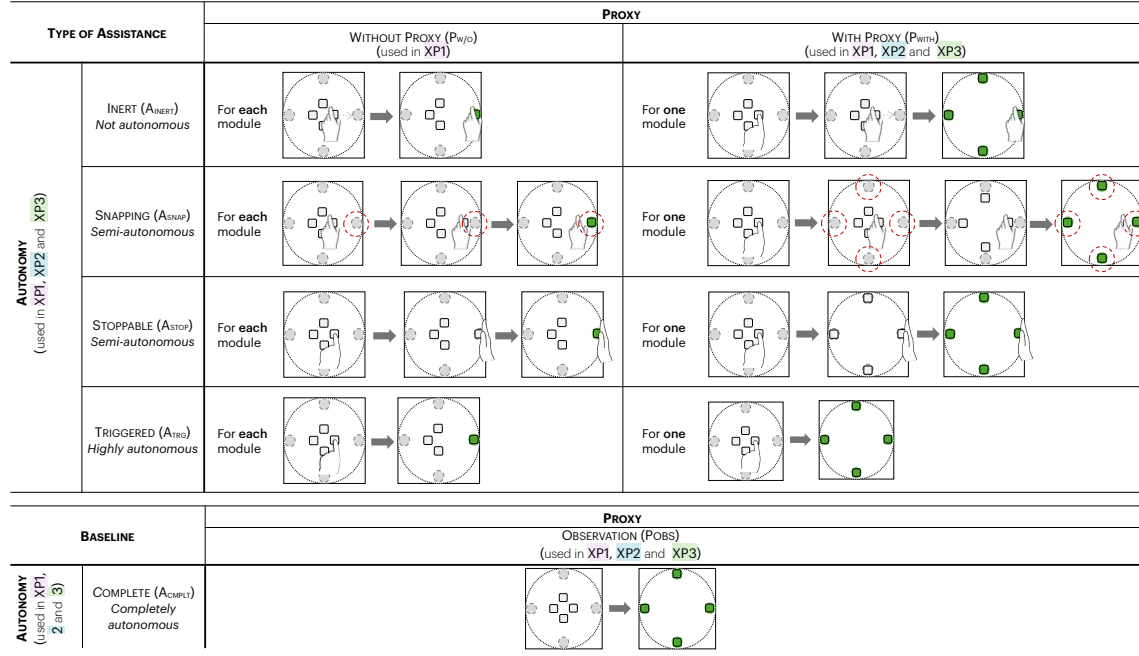


Fig. 3. Each TYPE OF ASSISTANCE (AUTONOMY  $\times$  PROXY) illustrated for the drag-and-drop task. In line, each level of AUTONOMY is reported and in column each level of PROXY. The representation of the BASELINE ( $A_{cmlt} \times P_{obs}$ ) is in the last table. The XP with the number specified the experiment where the assistance type is used.

At the end of the experiment, we additionally asked participants to rank the nine TYPES OF ASSISTANCE according to: (1) their preference, (2) their sense of control (i.e., sense of agency), (3) perceived performance (speed and accuracy), and (4) ease of use: “Rank the 9 methods according to your preferences (resp. your sense of control, perceived performance (speed and accuracy), ease of use). Explain your choices aloud.” The scale ranged from (1) favorite / total control / very good / very easy, to (9) least favorite / no control / no performance / very difficult. Their spoken explanations were recorded.

We additionally asked a question concerning their feelings about the success of the task to check that the trial went well: “How did you feel about the success of the task?”

**3.1.4 Hypotheses: effect of TYPES OF ASSISTANCE on SoA.** For the first experiment, we made the following hypotheses:

**Hyp1.1** The higher the level of AUTONOMY (i.e., complete autonomy), the lower the SoA.

**Hyp1.1bis** For semi-autonomous levels (i.e.,  $A_{snap}$  and  $A_{stop}$ ), having control over a large part of the journey rather than just the end of the movement reinforces the SoA.

**Hyp1.2** SoA  $P_{with}$  is lower than  $P_{w/o}$ .

**Hyp1.3** The SoA is higher for the module used as a proxy than for satellite modules.

**3.1.5 Procedure.** The experiment is divided into two phases: (1) a training phase, and (2) an experimental phase. At the start of the session, participants read and signed an informed consent form. For the duration of the experiment, participants sat at a table on which we placed the Toios mat (Figure 2).

The task performed by the participants is to position four Toios –without lifting them– in their respective targets as quickly and accurately as possible. This *drag-and-drop* task is illustrated in Figure 3. At the start of each trial, the Toios position themselves in a circle in the center of the mat. A target is projected in front of each robot (i.e., a circular gray area 6 cm in diameter, 15.5 cm apart). The test begins when the computer emits a sound. Each target turns green when the corresponding robot is entirely within the circular target area. If the robot leaves the target, it turns gray again. Once all targets have turned green, the computer emits a sound to indicate validation and the end of the trial. For the  $P_{WITH}$  condition, when the proxy robot is selected –by clicking on it– it emits a sound and its LED lights up for 500 ms.

The purpose of the training phase is to familiarize participants with the task and with each TYPE OF ASSISTANCE. Participants train on each TYPE OF ASSISTANCE in the same order for each participant: (1)  $A_{INERT}$ , (2)  $A_{SNAP}$ , (3)  $A_{STOP}$ , and (4)  $A_{TRG}$ . All levels of autonomy are presented first in the  $P_{W/O}$  condition, then all in the  $P_{WITH}$  condition. The BASELINE ( $A_{CMPLT} \times P_{OBS}$ ) was presented at the end. For each TYPE OF ASSISTANCE, we orally remind the participants how to interact with the Toios. Participants then performed the task at least once for each TYPE OF ASSISTANCE. If necessary, they repeat the trial until they fully understand the TYPE OF ASSISTANCE and feel comfortable with it.

The experimental phase is divided into 12 blocks. In each block, participants experience each TYPE OF ASSISTANCE. Their order is randomized within each block. At the start of each trial, the experimenter orally reminds a short reminder sentence to instruct participants how to interact with the Toios, according to the current TYPE OF ASSISTANCE.

At the end of each task, participants were asked to answer the questionnaire described in section 3.1.3.

This first experiment lasted on average 90 min (min 75 min, max 110 min).

### 3.2 Data analysis

To test our hypotheses (section 3.1.4), we aggregate the SoA scores by participant and by TYPE OF ASSISTANCE using the mean. As the aggregated data, after Greenhouse-Geisser correction, verified normality and sphericity (see Supplementary Materials), we report a two-factor repeated measures ANOVA ( $AUTONOMY \times PROXY$ ) to determine the significance of our two main factors and their interaction. We then run Tukey post hoc tests for our two factors and their interaction, to determine the significance of differences between each condition. For ANOVA ( $AUTONOMY \times PROXY$ ) and post-hoc tests, for each participant, we subtract their baseline SoA score. This gives us the gain in sense of agency compared with a fully autonomous system. For all our significance tests, we use the threshold commonly used in the community of 0.05 for our *p-value*.

We illustrate our results using the estimation technique based on means and 95% confidence intervals (CIs) –reported between [square brackets] in the text and by error bars in the figures. We also use pairwise differences: (1) between our factor levels and baselines (i.e.,  $A_{CMPLT}$  and/or  $P_{OBS}$ ), and (2) for TYPES OF ASSISTANCE WITH PROXY, between the proxy score and the satellite score. These techniques are recommended by the scientific community [7, 38], and in particular by the Group on Transparent Statistics [112].

To avoid unintentional false positive inflation of our results [41], we registered this experiment before we collected any data [3].

### 3.3 Results Experiment 1: Effect of the TYPE OF ASSISTANCE on SoA

In this section, we present the results for the effect of each  $AUTONOMY$ ,  $PROXY$ , and  $AUTONOMY \times PROXY$  (i.e., TYPE OF ASSISTANCE) on SoA, and the effect of the modules' role for  $AUTONOMY P_{WITH}$  on SoA (i.e., proxy module vs. satellite modules).

We report the complementary results for the perceived error rate, and rankings based on preference, sense of control, perceived performance, and ease of use in Appendix F.

**3.3.1 Effect of AUTONOMY and PROXY on SENSE OF AGENCY (Hyp1.1, Hyp1.1<sub>bis</sub> and Hyp1.2).** Figures 4a, 5a, and 6a respectively show the SoA ratings according to the AUTONOMY and PROXY conditions, and their interaction (AUTONOMY  $\times$  PROXY).

We report their means and CIs in Appendix A. The two-factor repeated measures ANOVA with the Greenhouse-Geisser correction shows:

- A large significant main effect of the AUTONOMY factor (**Hyp1.1 supported**,  $F(1.88,63.84)=74.5$ ,  $p < .001$ ,  $\eta_p^2 = .254$ , means and CIs in Appendix A.1),
- A small but significant main effect of the PROXY factor (**Hyp1.2 supported**,  $F(1,34)=28.5$ ,  $p < .001$ ,  $\eta_p^2 = .036$ , means and CIs in Appendix A.2),
- A small but significant interaction effect (AUTONOMY  $\times$  PROXY) ( $F(2.30,78.14)=16.9$ ,  $p < .001$ ,  $\eta_p^2 = .010$ , means and CIs in Appendix A.3).

Post hoc tests show that the levels of AUTONOMY are significantly different from each other (post hoc tests reported Appendix B.1). Figure 4a shows a drop in SoA as AUTONOMY increases. Each level shows a higher SoA than the BASELINE (pairwise differences greater than 0 in Figure 4b). **These findings support Hyp1.1 and Hyp1.1<sub>bis</sub>.**

Figure 5a shows that the SoA is greater  $P_{w/o}$  than  $P_{with}$ , and greater  $P_{with}$  than  $P_{obs}$  (i.e., the BASELINE). Each level shows a higher SoA than the BASELINE (pairwise differences greater than 0 in Figure 5b). **These findings support Hyp1.2.**

On the one hand, the SoA is significantly higher than the BASELINE as soon as the user interacts with at least one module, regardless of the type of AUTONOMY. On the other hand, controlling only one module (rather than all of them) has a significant negative effect on the SoA.

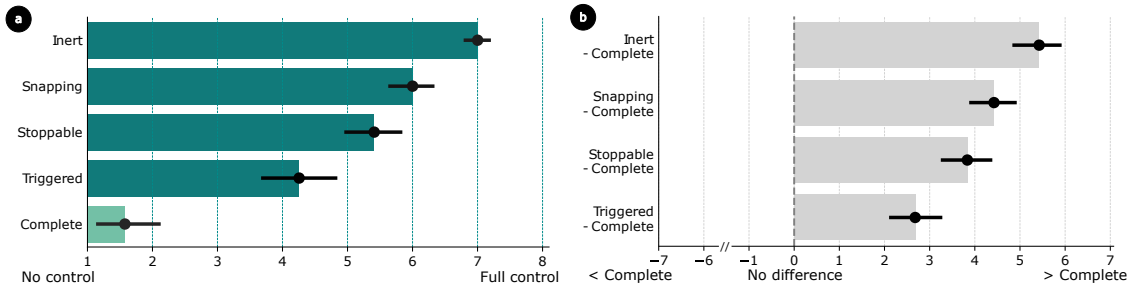


Fig. 4. SENSE OF AGENCY (SoA) depending on AUTONOMY. Error bars represent *bootstrapped* 95% confidence intervals. (a) Average SoA scores depending on AUTONOMY. (b) Pairwise differences in SoA depending on AUTONOMY. The difference is with the baseline:  $A_{CMPLT}$ .

For each AUTONOMY, the differences between  $P_{w/o}$  and  $P_{with}$  are significant (Appendix B.2). Figure 6a illustrates the interaction between AUTONOMY and PROXY, and shows that the SoA increases between levels of AUTONOMY with a small gap between  $P_{w/o}$  and  $P_{with}$  for each AUTONOMY.

These results suggest that even if a general trend emerges, the verification of Hyp1.2 seems mainly due to the difference between  $P_{w/o}$  and  $P_{with}$  for  $A_{INERT}$ .

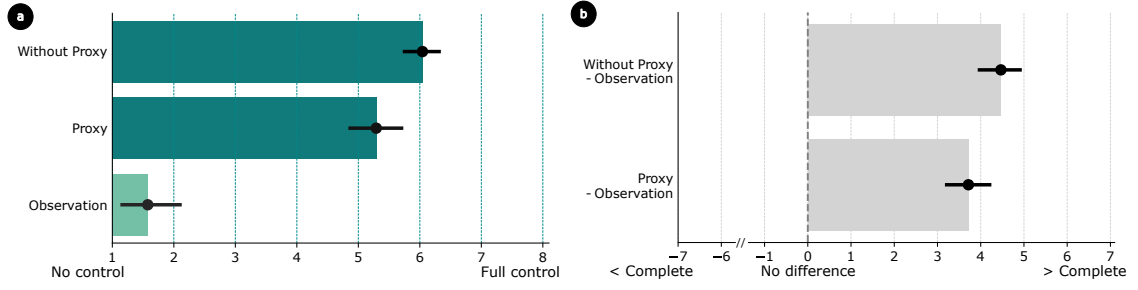


Fig. 5. SENSE OF AGENCY (SoA) depending on PROXY. Error bars represent *bootstrapped* 95% confidence intervals. For more condensed writing,  $P_{W/O}$  is named “Without Proxy” and  $P_{WITH}$  is named “Proxy”.

(a) Average SoA scores depending on PROXY. (b) Pairwise differences in SoA depending on PROXY. The difference is with the baseline:  $P_{OBS}$ .

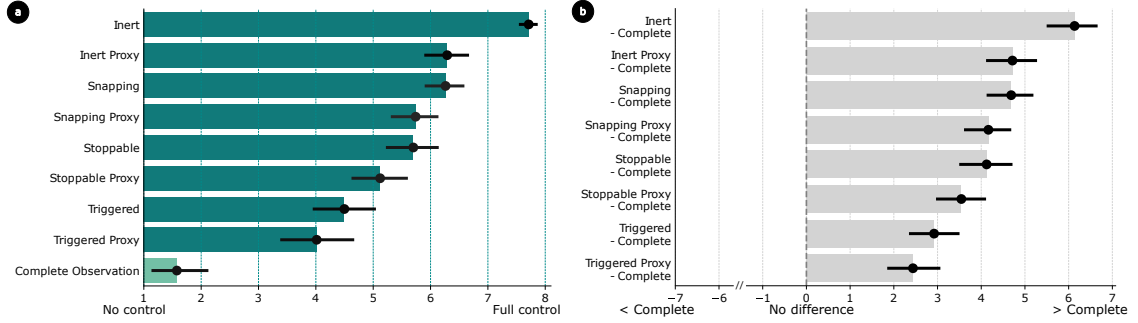


Fig. 6. SENSE OF AGENCY (SoA) depending on TYPE OF ASSISTANCE. Error bars represent *bootstrapped* 95% confidence intervals. For more condensed writing, TYPES OF ASSISTANCE  $P_{W/O}$  are named only by the name of the AUTONOMY and TYPES OF ASSISTANCE  $P_{WITH}$  are named as follows: “name of the AUTONOMY Proxy”.

(a) Average SoA scores depending on TYPE OF ASSISTANCE. (b) Pairwise differences of the SoA according to the TYPE OF ASSISTANCE. The difference is made with the BASELINE:  $A_{CMPLT} \times P_{OBS}$ .

**3.3.2 Effect of the MODULE’S ROLE on the SENSE OF AGENCY (Hyp1.3).** Figure 7a shows that the SoA increases as the AUTONOMY decreases. The gap between the SoA of the *proxy* module (in gray on Figure 7a), the *satellite* modules (in orange on Figure 7a), and the *swarm* (in green on Figure 7a) decreases with AUTONOMY. Moreover, for each AUTONOMY level, the SoA of the *proxy* is higher than the *satellites* and the *swarm*; and the SoA of the *swarm* is higher than the *satellites*. We conducted a two-factor repeated measures ANOVA (AUTONOMY  $\times$  MODULE’S ROLE) with Greenhouse-Geisser correction. The levels of MODULE’S ROLE were: *proxy*, *satellites*, and *swarm*. The results show:

- A significant main effect of the AUTONOMY factor ( $F(1.70,57.96)=61$ ,  $p < .001$ ,  $\eta_p^2 = .174$ ),
- A significant main effect of the MODULE’S ROLE factor ( $F(1.17,39.76)=46.6$ ,  $p < .001$ ,  $\eta_p^2 = .183$ ),
- A small but significant interaction effect (AUTONOMY  $\times$  MODULE’S ROLE) ( $F(2.46,83.79)=23.2$ ,  $p < .001$ ,  $\eta_p^2 = .019$ , means and CIs in Appendix A.4).

The post-hoc tests show that the pairwise differences between the *proxy*, the *satellites*, and the *swarm* are significant (Appendix B.3). The *proxy* has the highest SoA ( $M = 5.98$ , 95% CI [5.69, 6.26]), the *swarm* shows a slightly lower SoA ( $M = 5.29$ , 95% CI [5.01, 5.56]), and the *satellites* have the lowest SoA ( $M = 3.88$ , 95% CI [3.54, 4.23]). Moreover, for each

AUTONOMY, the differences between the *proxy*, the *satellites*, and the *swarm* are significant, except between the *satellites* and the *swarm* for  $A_{TRG}$  (Appendix B.4). **These findings support Hyp1.3.**

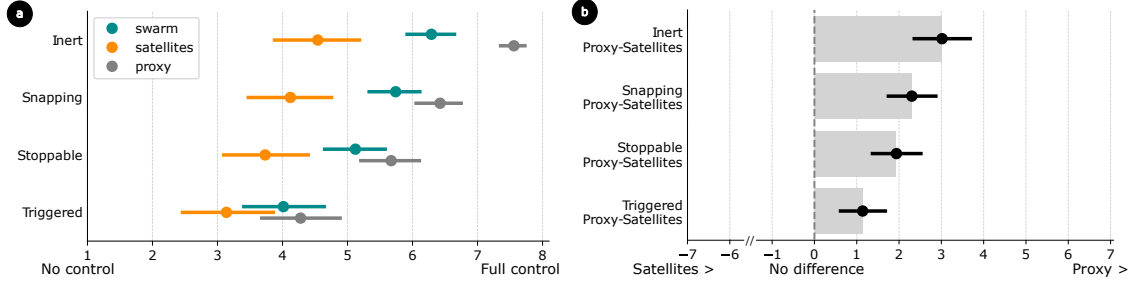


Fig. 7. SENSE OF AGENCY (SoA) depending on AUTONOMY and the role of modules in movement coordination (proxy module vs. satellite modules). Error bars represent *bootstrapped* 95% confidence intervals. For more condensed writing,  $P_{WITH}$  is named “Proxy”.

(a) Average SoA scores depending on AUTONOMY and the role of modules in movement coordination (proxy module vs. satellite modules). (b) Pairwise differences in SoA according to the role of modules in movement coordination (proxy vs. satellite modules). The difference is between the SoA score for the module that was the proxy module and the SoA score for the modules controlled through the proxy module (satellites).

In this experiment, we mainly demonstrated the impact of AUTONOMY on the SoA (Hyp 1.1) for swarm UIs. We also showed that SoA  $P_{WITH}$  is lower than  $P_{W/O}$  (Hyp 1.2). The SoA is higher for the module used as a proxy than for satellite modules (Hyp 1.3).

As TYPE OF ASSISTANCE can impact the performance, and, in turn, performance can impact the SoA, the next experiment studies the combined effect of TYPE OF ASSISTANCE and performance on the SoA.

## 4 Experiment 2: Combined effect of task DIFFICULTY and TYPE OF ASSISTANCE on SoA

### 4.1 Method, Apparatus and Data Analysis

The aim of this experiment is to evaluate the combined impact of the task DIFFICULTY and the TYPE OF ASSISTANCE provided by the UI on the sense of agency. According to Fitts' law, increasing the targets' distance and decreasing their size will increase participants' movement time, and therefore increase the task difficulty. Based on related work [54, 77, 115, 122], we expect a decrease in SoA when the difficulty of the task increases. To verify this, we used the same experimental task as in Experiment 1 (Figures 2 and 3).

The apparatus was the same as in Experiment 1 (Section 3.1.2).

We used the same approach for data analysis as in the first experiment (Section 3.2). To avoid unintentional false positive inflation of our results [41], we registered this experiment before we collected any data [5].

**4.1.1 Participants.** We recruited 35 participants by e-mail (12 self-identified as women, 23 as men,  $M = 23$  y.o.,  $SD = 4.5$ ). Each participant received a €15 gift voucher. Twenty-two work in IT and 13 in other fields. None of these participants participated in the first experiment.

**4.1.2 Experimental Design.** We used a within-subjects design, with two independent variables: AUTONOMY and DIFFICULTY.

As Figure 1b<sub>1</sub> illustrates, AUTONOMY refers to the same levels as in Experiment 1 (Section 3.1.3). However, we did not test the PROXY variable in this 2<sup>nd</sup> experiment and set it to the P<sub>WITH</sub> and baseline AUTONOMY levels. There were therefore five AUTONOMY levels.

In the *Harder* DIFFICULTY, targets have the same size and were at the same position as in the first experiment: Size = 6 cm, Distance = 15.5 cm (Fitts' ID = 3.3)<sup>2</sup>. In the *Easier* DIFFICULTY, targets were bigger and closer to Toios: Size = 12 cm, Distance = 9.5 cm (Fitts' ID = 1.2).

We had the same dependent variable as in Experiment 1 (Section 3.1.3).

**4.1.3 Hypotheses: effect of the DIFFICULTY on the SoA.** For the second experiment, we make the following hypotheses:

**Hyp2.1** SoA for *Harder* DIFFICULTY is weaker than *Easier* DIFFICULTY.

**Hyp2.2** The higher the level of AUTONOMY (i.e., complete autonomy), the less the SoA is impacted by DIFFICULTY.

**Hyp2.3** SoA is more impacted by DIFFICULTY for the module used as a proxy than for satellite modules.

**4.1.4 Procedure.** The experiment proceeded exactly as Experiment 1 (Section 3.1.5) except for the following elements: First participants trained on each TYPE OF ASSISTANCE in the same order for each participant and on *Harder* DIFFICULTY only. The order of the autonomy levels is as follows: (1) A<sub>INERT</sub>, (2) A<sub>SNAP</sub>, (3) A<sub>STOP</sub>, (4) A<sub>TRG</sub>, and (5) A<sub>CMPLT</sub>. Second, the experimental phase was divided into 12 blocks. In each block, participants experience each TYPE OF ASSISTANCE × DIFFICULTY combination. The order of the combinations is randomized within each block.

This experiment lasted an average of 60 min (min 45 min, max 90 min).

## 4.2 Results: Effect of the DIFFICULTY on SoA

**4.2.1 Effect of DIFFICULTY, AUTONOMY and their interaction on SENSE OF AGENCY (Hyp2.1 and Hyp2.2).** Figure 8a shows that the SoA increases as the AUTONOMY decreases (consistently with Experiment 1). However, the gap between *Easier* and *Harder* DIFFICULTY remains stable for each AUTONOMY. Figure 8b confirms these results, showing that all AUTONOMIES have a SoA similar for *Easier* and *Harder* DIFFICULTY (CIs of the pairwise differences cross 0). The two-factor repeated measures ANOVA shows:

- A significant effect of the AUTONOMY factor ( $F(2.59,88.05)=125.719$ ,  $p < .001$ ,  $\eta_p^2 = .621$ ),
- An inconclusive effect of the DIFFICULTY factor (**Hyp2.1 unsupported**,  $F(1,34)=0.129$ ,  $p=.722$ ,  $\eta_p^2 < .001$ , *Easier*:  $M = 4.65$ , 95% CI [4.35, 5.03], *Harder*:  $M = 4.66$ , 95% CI [4.37, 5.03]),
- An inconclusive interaction effect (AUTONOMY × DIFFICULTY) (**Hyp2.2 unsupported**,  $F(2.93,99.56)=1.195$ ,  $p=.315$ ,  $\eta_p^2 < .001$ , means and CIs in Appendix C.1).

**4.2.2 Effect of the DIFFICULTY on the SENSE OF AGENCY depending on MODULE'S ROLE (Hyp2.3).** Figure 8a shows, for each MODULE'S ROLE (proxy vs. satellites), a significant decrease in SoA when the AUTONOMY increases, but no effect of the DIFFICULTY. Figure 8b shows that the SoA experienced on the PROXY module does not seem to be more impacted by DIFFICULTY than that of the satellite modules.

We conducted a three-factor repeated measures ANOVA (AUTONOMY × DIFFICULTY × MODULE'S ROLE) with Greenhouse-Geisser correction. The levels of MODULE'S ROLE were: *proxy*, *satellites*, and *swarm*. The results show:

- A significant main effect of the AUTONOMY factor ( $F(1.89,64.17)=82.2345$ ,  $p < .001$ ,  $\eta_p^2 = .174$ ),

<sup>2</sup>We used the Shannon formulation of Fitts' IDs [74]:  $ID = \log_2(1 + \frac{AMPLITUDE}{TOLERANCE})$ . As TOLERANCE, we used the difference between the diameter of the target and the diagonal of the Toio footprint (i.e.,  $3\sqrt{2}$  cm) as per usual for drag-and-drop Fitts' tasks [43].

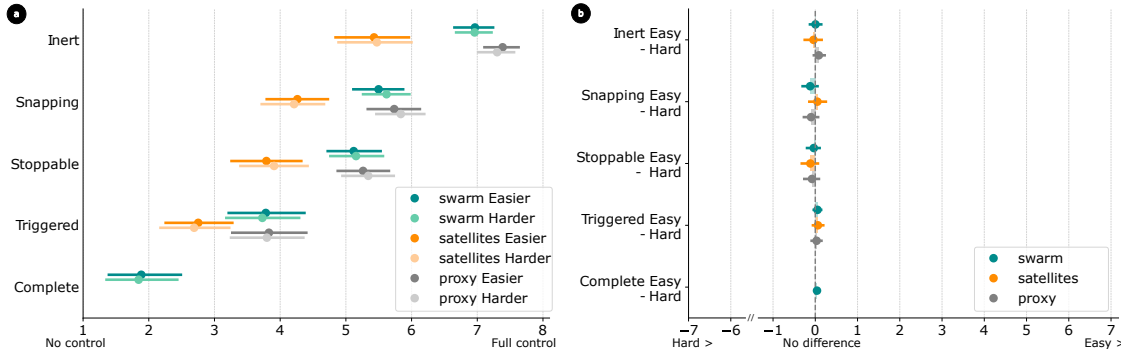


Fig. 8. SENSE OF AGENCY (SoA) separated by AUTONOMY  $\times$  DIFFICULTY for the swarm, the satellite modules and the proxy module. Error bars represent *bootstrapped* 95% confidence intervals.

(a) Average SoA scores depending on AUTONOMY  $\times$  DIFFICULTY for the swarm, the satellite modules and the proxy module. (b) Pairwise differences in SoA for the swarm, the satellite modules and the proxy module, according to TYPE OF ASSISTANCE. The difference is between the SoA score for the *Easier* DIFFICULTY and the *Harder* DIFFICULTY for each TYPE OF ASSISTANCE and for the swarm, the satellite modules and the proxy module.

- An inconclusive main effect of the DIFFICULTY factor ( $F(1,34)=0.2004$ ,  $p=.657$ ,  $\eta_p^2 < .001$ ),
- A significant main effect of the MODULE'S ROLE factor ( $F(1.21,40.99)=33.5745$ ,  $p < .001$ ,  $\eta_p^2 = .183$ ),
- A small but significant interaction effect of AUTONOMY  $\times$  MODULE'S ROLE ( $F(2.81,95.59)=7.6005$ ,  $p < .001$ ,  $\eta_p^2 = .004$ , means and CIs in Appendix C.3)
- Inconclusive interactions effects of: AUTONOMY  $\times$  DIFFICULTY ( $F(2.62,89.15)=0.7775$ ,  $p=.494$ ,  $\eta_p^2 < .001$ ), DIFFICULTY  $\times$  MODULE'S ROLE (**Hyp2.3 unsupported**,  $F(1.56,53.08)=0.0599$ ,  $p=.903$ ,  $\eta_p^2 < .001$ , means and CIs in Appendix C.2), and AUTONOMY  $\times$  DIFFICULTY  $\times$  MODULE'S ROLE ( $F(4.22,143.34)=2.1565$ ,  $p=.073$ ,  $\eta_p^2 < .001$ , means and CIs in Appendix C.4).

**These findings do not support Hyp2.3.** This experiment does not allow us to conclude that there is a main effect of the DIFFICULTY on the SoA (Hyp 2.1) for swarm UIs. The results are also inconclusive about the impact of DIFFICULTY depending on AUTONOMY (Hyp 2.2) and MODULE'S ROLE (Hyp 2.3).

As SoA can be impacted by issues in the system response, the next experiment studies the combined effect of the TYPE OF ASSISTANCE and predictability of the system's response on the SoA.

## 5 Experiment 3: Combined effect of PREDICTABILITY and TYPE OF ASSISTANCE on SoA

### 5.1 Method, Apparatus and Data Analysis

Prior work indicates that the consistency between users' expectations and the UI behavior increases the sense of agency [23, 24, 101, 102, 123]. For this reason, we expect the sense of agency to decrease as the UI predictability decreases. The aim of this experiment is to evaluate the combined impact of the predictability of the system and the AUTONOMY on the sense of agency. In particular, we would like to learn how this predictability impacts proxies and satellites. To do so, we used the same apparatus (Section 3.1.2) and experimental task (Figures 2 and 3) as in Experiment 1.

We used the same approach for data analysis as in the first experiment (Section 3.2). To avoid unintentional false positive inflation of our results [41], we registered this experiment before we collected any data [4].

**5.1.1 Participants.** We recruited 35 participants by e-mail (12 self-identified as women, 23 as men,  $M = 27.66$  y.o.,  $SD = 9.02$ ). Each participant received a €15 gift voucher. 17 work in IT and 18 in other fields. These participants did not participate in the first and second experiments.

**5.1.2 Experimental Design.** We used a within-subjects design, with two independent variables (Figure 1c<sub>1</sub>):

The AUTONOMY refers to the same level as in the second experiment (Section 4.1.2).

The PREDICTABILITY<sup>3</sup> refers to the noise in the modules' trajectory. In the *Straight* condition, the modules' trajectory is not affected by vibrations and is very predictable. The Toios act as in both previous experiments. In the *Noisy* condition, the trajectory is affected by Toio vibrations. Each Toio moves randomly around its position and is little predictable.

Our dependent variables are the same as in the first two experiments.

**5.1.3 Hypotheses: effect of the PREDICTABILITY on the SoA.** We make the following hypotheses:

**Hyp3.1** SoA for *Noisy* PREDICTABILITY is weaker than *Straight* PREDICTABILITY.

**Hyp3.2** The higher the level of AUTONOMY, the less the SoA is impacted by PREDICTABILITY.

**Hyp3.3** SoA is more impacted by PREDICTABILITY for the module used as a proxy than for satellite modules.

**5.1.4 Procedure.** The experiment proceeded as the first experiment (Section 3.1.5) except for the following: First, participants were trained on each AUTONOMY in the same order for each participant and on *Straight* PREDICTABILITY only. The order of the autonomy levels is as follows: (1)  $A_{INERT}$ , (2)  $A_{SNAP}$ , (3)  $A_{STOP}$ , (4)  $A_{TRG}$ , and (5)  $A_{CMPLT}$ . In this experiment, participants never interacted without a proxy, as this condition was not presented.

The experimental phase was divided into 12 blocks. In each block, participants experience each AUTONOMY  $\times$  PREDICTABILITY combination. The order of the combinations is randomized within each block.

This experiment lasted 60 min on average (min 45 min, max 90 min).

## 5.2 Results: Effect of the PREDICTABILITY on SoA

**5.2.1 Effect of PREDICTABILITY, AUTONOMY and their interaction on SENSE OF AGENCY (Hyp3.1 and Hyp3.2).** Figure 9a shows that the SoA increases as the AUTONOMY decreases (consistently with experiment 1). It also shows that the SoA with a *Straight* PREDICTABILITY is a little higher than the SoA with a *Noisy* PREDICTABILITY. Figure 9b shows that all levels of AUTONOMY exhibit a difference in the SoA between *Straight* and *Noisy* PREDICTABILITY (differences greater than 0), except for  $A_{CMPLT}$  (crosses 0). The two-factor repeated measures ANOVA shows:

- A significant effect of the AUTONOMY factor ( $F(1.96,66.50)=48.86$ ,  $p < .001$ ,  $\eta_p^2 = .262$ ),
- A small but significant effect of the PREDICTABILITY factor (**Hyp3.1 supported**,  $F(1,34)=10.35$ ,  $p=.003$ ,  $\eta_p^2 = .007$ , *Straight*:  $M = 4.79$ , 95% CI [4.25, 5.41], *Noisy*:  $M = 4.44$ , 95% CI [3.98, 4.99]).
- A small but significant interaction effect (AUTONOMY  $\times$  PREDICTABILITY) (**Hyp3.2 supported**,  $F(2.84,96.51)=7.61$ ,  $p < .001$ ,  $\eta_p^2 = .003$ , means and CIs in Appendix D.1).

Post hoc tests show significant differences between PREDICTABILITY for  $A_{INERT}$ ,  $A_{SNAP}$ , and  $A_{STOP}$ . However, post hoc tests are inconclusive for  $A_{TRG}$  and  $A_{CMPLT}$  levels (Hyp3.2, Appendix E.1). This suggests that when the AUTONOMY is too high the PREDICTABILITY does not seem to impact the SoA anymore.

<sup>3</sup>Some models distinguish between three classes of intention – distal, proximal, and motor – according to their level in the hierarchy of action control [84]. PREDICTABILITY refers to motor PREDICTABILITY: the discrepancy between the participant's intention and the sensory consequences observed on the modules (more details in Section 6.4).



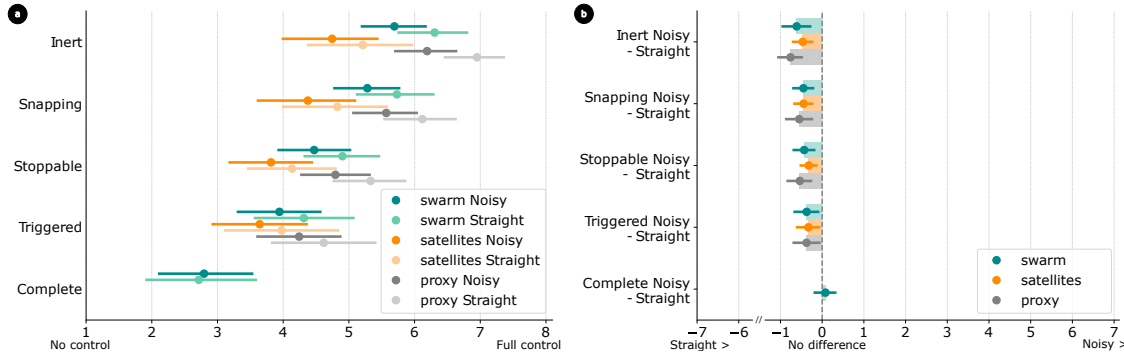


Fig. 9. SENSE OF AGENCY (SoA) depending on AUTONOMY × PREDICTABILITY for the swarm, the satellite modules and the proxy module. Error bars represent *bootstrapped* 95% confidence intervals.

(a) Average SoA scores depending on AUTONOMY × PREDICTABILITY for the swarm, the satellite modules and the proxy module. (b) Pairwise differences in SoA for the swarm, the satellite modules and the proxy module, according to TYPE OF ASSISTANCE. The difference is between the SoA score for the *Noisy* PREDICTABILITY and the *Straight* PREDICTABILITY for each TYPE OF ASSISTANCE and for the swarm, the satellite modules and the proxy module.

**5.2.2 Effect of the PREDICTABILITY on the SENSE OF AGENCY depending on MODULE'S ROLE (Hyp3.3).** Figure 9a shows a significant decrease in SoA as a function of AUTONOMY for each MODULE'S ROLE (proxy vs. satellites), but little impact of the PREDICTABILITY for each AUTONOMY. Figure 9b shows that the proxy and satellites are similarly impacted by PREDICTABILITY.

We conducted a three-factor repeated measures ANOVA (AUTONOMY × PREDICTABILITY × MODULE'S ROLE) with Greenhouse-Geisser correction. The levels of MODULE'S ROLE were: *proxy*, *satellites*, and *swarm*. The results show:

- A significant main effect of the AUTONOMY factor ( $F(1.67,56.80)=29.62$ ,  $p < .001$ ,  $\eta_p^2 = .100$ ),
- A small but significant main effect of the PREDICTABILITY factor ( $F(1,34)=19.14$ ,  $p < .001$ ,  $\eta_p^2 = .013$ ),
- A significant main effect of the MODULE'S ROLE factor ( $F(1.43,48.60)=23.84$ ,  $p < .001$ ,  $\eta_p^2 = .050$ ),
- A small but significant interaction effect of AUTONOMY × MODULE'S ROLE ( $F(2.09,71.14)=7.40$ ,  $p=.001$ ,  $\eta_p^2 = .005$ , means and CIs in Appendix D.2)
- A small but significant interaction effect of PREDICTABILITY × MODULE'S ROLE (**Hyp3.3 supported**  $F(1.69,57.62)=4.03$ ,  $p=.029$ ,  $\eta_p^2 < .001$ )
- Inconclusive interaction effect of: AUTONOMY × PREDICTABILITY ( $F(2.43,82.73)=1.79$ ,  $p=.165$ ,  $\eta_p^2 < .001$ ), and AUTONOMY × PREDICTABILITY × MODULE'S ROLE ( $F(4.55,154.83)=1.19$ ,  $p=.316$ ,  $\eta_p^2 < .001$ , means and CIs in Appendix D.3).

Post-hoc tests (Appendix E.2) show that the pairwise differences between *proxy Straight* ( $M = 5.75$ , 95% CI = [5.42, 6.08]) and *proxy Noisy* ( $M = 5.20$ , 95% CI = [4.91, 5.49]), between *swarm Straight* ( $M = 5.31$ , 95% CI = [4.98, 5.64]) and *swarm Noisy* ( $M = 4.85$ , 95% CI = [4.56, 5.14]), and between *satellites Straight* ( $M = 4.54$ , 95% CI = [4.13, 4.94]) and *satellites Noisy* ( $M = 4.14$ , 95% CI = [3.78, 4.51]) are significant. The gap in SoA between the *Straight* and *Noisy* is largest for the *proxy*, slightly smaller for the *swarm*, and smallest for the *satellites*. **These findings support Hyp3.3.** Moreover, the results of the *swarm* are aligned with the results confirming Hyp3.1.

In this experiment, we mainly demonstrated the main effect of PREDICTABILITY on the SoA (Hyp 3.1) for swarm UIs. We also showed that the level of AUTONOMY impacts the effect of PREDICTABILITY on the SoA (Hyp 3.2). The impact of

the PREDICTABILITY is higher for the PROXY module than for satellite modules (Hyp 3.3), but this effect is not impacted by AUTONOMY.

**Summary:** Users' SoA with swarm UIs decreases with increasing assistance. We introduce three novel types of SoA that are specific to swarm UIs, when using a proxy: (a) A sense of agency for the proxy module, (b) One for the other modules in the swarm (satellites), and (c) One for the swarm as a whole. Our task difficulty has an inconclusive impact on users' SoA with swarm UIs. The predictability of the modules' trajectory impacts the SoA of the proxy (a), the satellite modules (b) and the whole swarm (c). However, predictability does not impact SoA differently across autonomy levels, whether for the swarm, the proxy, or the satellites.

## 6 Discussion, Limitations and Future Work

### 6.1 Effect of assistance type on SoA

This paper aims to investigate the impact of the type of assistance provided by a swarm UI on the users' sense of agency (SoA), and how the performance or the predictability can weaken or strengthen this impact. To this end, we conducted three experiments (each  $N=35$ , total  $N=105$ ) focusing respectively on: (1) the impact of the assistance on the SoA, (2) the impact of the performance on the SoA for each assistance, and (3) the impact of the system predictability on the SoA for each assistance. We define a type of assistance as consisting of two dimensions: (1) the degree of control available to the user over different phases of the movement (AUTONOMY), and (2) whether or not the swarm's movements (satellites) are coordinated through a single module (proxy). For the second and the third experiments, we focused solely on the AUTONOMY with proxy.

**6.1.1 General impact.** Our results reinforce prior knowledge that AUTONOMY impacts SoA [12, 13, 29, 71, 114, 126, 127] and extend it to swarm UIs (Hyp1.1 supported). In swarm UIs (this paper) as in other UIs (prior work), excessive automation appears to substantially reduce users' SoA. Importantly, we found that, for swarm UIs, simply pressing the module to trigger its movement already restores some of the users' SoA. Our paper shows that semi-autonomous levels of swarm UIs impact SoA differently (Hyp1<sub>bis</sub> supported). Control over part of the trajectory of modules ( $A_{SNAP}$ ), as opposed to the final phase of their movement ( $A_{STOP}$ ), does significantly improve SoA. This paper therefore demonstrates that to reinforce users' SoA, and subsequently, users' responsibility and system adoption, the swarm UI designers should propose autonomy at the end of the task. For example, one could propose autonomy to improve the precision or conclude a trajectory following users' decision.

Compared with prior literature, this paper introduced the evaluation of the impact of proxies on SoA. We demonstrate that manipulating modules through a proxy has an effect on SoA (Hyp1.2 supported). Manipulating only a proxy module tends to reduce SoA compared to manipulating all modules directly. However, for equivalent levels of AUTONOMY, performing the task with a proxy seems to have a slight impact on SoA, except when the system is inert. Overall, we show that the effect of AUTONOMY is stronger than that of PROXY. The use of a proxy provides a specific type of AUTONOMY to the satellite modules shifting from full user control (without proxy) to partial loss of control (with proxy).

Our results do not allow us to conclude on the impact of task performance on SoA, regardless of the assistance type (Hyp2.1 and 2.2 unsupported). We also demonstrate that the system predictability has a slight impact on the SoA for each assistance type (Hyp3.1 and 3.2 supported). Yet, this effect decreases when AUTONOMY increases. For example, Kim et al.'s participants proposed two different techniques for the swarm to follow a specific trajectory [57]: either

drawing the path with one module, or drawing the path with the finger on the table (stronger autonomy). In the first case, users' SoA will be strong. On the contrary, if they draw the path with the finger (stronger autonomy), users SoA will be weak. Our results help designers choose between these design alternatives depending on their application (game vs. search and rescue). However, we show that if modules do not behave as expected, the SoA will be more impacted when they draw the path with one module than with the finger. If the system behaves unpredictably in a search and rescue coordination center scenario, users SoA who directly control one module on the city map will be more impacted than those who are using the second, finger technique. Subsequently, users controlling a proxy may tend to shift responsibility away from themselves.

**6.1.2 Impact of modules' role.** Prior work in HCI proposed the use of a proxy to interact with a swarm UI [44, 57, 67, 91]. Studying the impact of this type of interaction on the SoA is therefore timely to avoid any decrease in users' sense of responsibility or in the adoption of swarm UIs.

When using proxies, we demonstrated that users do not experience the same SoA with all modules (Hyp1.3 supported). The sense of agency is higher for directly controlled modules (proxy) than for indirectly controlled modules (satellites). We also show that autonomy exerts a stronger effect on the proxy module than on its satellites. For example, imagine users trying to arrange a swarm UI in a circle as with Zooids [67]. To do so, Le Goc et al. proposed to directly control only 2 modules [67]. In this case, users feel a stronger control over these two modules, and a weak control over the other modules controlled indirectly through the two proxies. Imagine users moving the proxies with some assistance, like the system proposing a rough circle size and position for users to adjust it at the end. In this case, users will feel much weaker control for the proxies, but not for the satellites. In a search and rescue coordination center scenario, users who directly control a single module to guide a swarm will experience a strong SoA over the proxy module, but a weaker SoA over the satellites. If automation is introduced by suggesting optimal paths to locate victims, it will have little impact on the SoA over the satellites –since it was already low. However, autonomy will affect the SoA over the proxy module. If all modules end up taking the wrong path, in particular the proxy, users may be more likely to shift responsibility away from themselves, attributing the failure to incorrect system recommendations.

Interestingly, we found in this paper that the sense of agency experienced over the whole swarm's movements tends to align with the sense of agency experienced over the proxy. Yet, as autonomy decreases, the sense of agency associated with the swarm diverges from that of the proxy. With low autonomy, the sense of agency associated with the swarm lies at an intermediate level –between the sense of agency reported for the proxy and its satellites. Here, users who place the swarm in a circle will feel as though they control the whole swarm almost as much as the two proxies. However, if none of the proxies have any autonomy, users will feel much more control over the proxies than over the whole swarm. In the rescue coordination center scenario, if the precise placement of all modules is critical to save lives, designers should rather provide as little autonomy as possible to the proxies so that users feel more responsible for the placement of the whole swarm. The SoA for the proxy is therefore no more affected by performance or predictability than that for satellites (Hyp 2.3 and 3.3 unsupported).

## 6.2 What type of SoA is felt when interacting with a proxy?

We can make several hypotheses about the sense of agency (SoA) felt when interacting with a proxy. First, we could consider that users perceive the swarm as hierarchical. Users would then be the *leader* in accomplishing the task through the proxy module and could consider the other modules as *followers* of their action. Prior work shows that a human *leader* develops vicarious SoA for the actions of those humans under their command [21, 26, 65], even if this is less

than their individual SoA [21, 65]. In the case of swarm UIs, this analogy would explain why users may feel SoA for the movements of the satellite modules, even if it is less than the SoA for the proxy module. Furthermore, Le Bars et al. [65] show that during a joint action, there is the creation of a shared SoA, which in our case corresponds to the SoA felt for the swarm. This SoA is expected to be lower than the individual SoA. The level of motor involvement in the task has a stronger influence on individual SoA than on shared SoA [65]. In our case, this would explain why the SoA felt for the proxy is higher and more impacted by AUTONOMY than the SoA felt for the swarm.

Based on this analogy with joint action, we proposed three types of SoA involved when interacting with a swarm UI through one of its modules acting as a proxy:

**Leader SoA** felt for the proxy module, impacted by the level of autonomy ( $\approx$  individual SoA),

**Follower SoA** felt for the satellite modules, little impacted by the level of autonomy and with a fairly low level ( $\approx$  vicarious SoA),

**Shared SoA** felt for the whole swarm, impacted by the level of autonomy.

Swarm UI designers now know that they need to take these three different SoA into account for their design. Designers should prepare for the possibility that the system does not behave as users expect. Users are more likely to attribute an error in the satellites to a system malfunction rather than to their own input –regardless of the system autonomy. Conversely, preparing for the possibility that proxies might not behave as intended (e.g., drifting too far from the target position), designers now know that users are more inclined to feel responsible for this error –unless the system exhibits a high level of autonomy, in which case responsibility is again attributed to the system. This same mechanism of perceived responsibility for proxies applies to the swarm as a whole. For the system to be adopted, designers should avoid errors that can be attributed to the system, and therefore consider lower autonomy or introducing a proxy.

Second, we could consider that users would have a SoA distributed among the modules, due to the synchronous aspect of our task (i.e., the concordance of the movements of the satellite modules with those of the proxy). The individual SoA is expected to decrease during a joint action synchronized with another human agent [18, 72, 92, 128] or robotic agent [27], as if the SoA was distributed among the agents. In our case, users would perceive the satellite modules as having their own SoA, and the SoA would be distributed among proxy and satellite modules.

However, the creation of distributed SoA and diffusion of responsibility only happen in the case of synchronous tasks –as in our first and second experiments– and not in the case of asynchronous tasks [69, 92]. Furthermore, variability or discrepancies between our actions and those of our partners would impact joint SoA (shared SoA and vicarious SoA) rather than SoA for our actions (individual SoA) [65]. Our third experiment (predictability) partially supports this analogy. Indeed, when we look at the effect size in Figure 9, the effect of predictability on the SoA seems weaker for the satellites ( $\approx$  vicarious) and the swarm ( $\approx$  shared) than for the proxy ( $\approx$  individual). Moreover, when comparing the level SoA between the first and third experiments (between around 3 and 8 in Figure 7 vs. between around 4 and 7 in Figure 9), experiencing unpredictability seems to impact users' proxy (8 to 7) and satellites' SoA (3 to 4). This suggests that an asynchronous, unpredictable task tends to cancel the difference between the three types of SoA in swarm UIs (Leader, Follower, and Shared SoA). In the case of a search and rescue scenario, imagine that robots make a decision (autonomy), failing to take users' input into account (unpredictability). Users would have less SoA for swarm, proxy, and satellites. They might therefore be more likely to attribute the outcomes to the system and disengage from its use. On the contrary, if robots take users' input into account and become predictable, designers should be aware that users' responsibility for the actions of the proxy and satellites will increase, but depart from each other.

### 6.3 DIFFICULTY and SoA for swarm UIs

To test the impact of performance on SoA, we applied a Fitts' law like experiment protocol [40], expecting target size and distance to affect users' performance. We verified the validity of our approach: the DIFFICULTY had an impact on the actual performance. Figure 10b shows that participants were faster in the easy condition compared to the hard condition, regardless of the autonomy. However, Figure 10a shows that the sense of agency is not correlated with this actual performance, as easy and hard difficulty lead to the same SoA. A trend seems to emerge from Figure 10b: the actual performance improves with the autonomy, while simultaneously in Figure 10a, the sense of agency decreases. Designers are now aware that they need to balance performance and SoA when designing autonomy for their Swarm UIs.

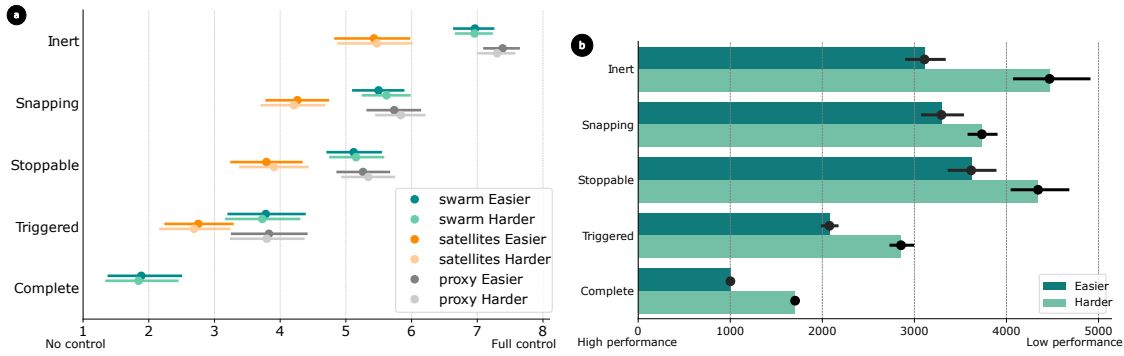


Fig. 10. SENSE OF AGENCY (SoA) and Performance of the task depending on AUTONOMY and DIFFICULTY. Error bars represent *bootstrapped* 95% confidence intervals. (a) Average SoA scores depending on AUTONOMY  $\times$  DIFFICULTY for the swarm, the satellite modules, and the proxy module. (b) Average movement time (ms) to perform the task, depending on AUTONOMY and DIFFICULTY.

To keep the length of the second experiment reasonable, we decided to use only two pairs of size and distance. However, even if the gap between our two Fitts' IDs (*Harder*: 3.3, and *Easier*: 1.2) was large enough to impact, as expected, the measured performance, there is a possibility that this gap was simply not large enough to see an impact on the SoA. One can hypothesize that SoA could be somehow robust to low to mild difficulty but starts being degraded with harder levels of difficulty. Our results with these Fitts' IDs are already actionable by the HCI community, as swarm UIs can be within arm's reach, as in our experiment: e.g., [67, 108]. Future work could replicate this experiment on a larger surface, to further vary these pairs of size and distance. Moreover, Yamada et al. [124] demonstrated, for a pointing task, that both target distance and target size can affect performance when a drone swarm is operated remotely using a handheld controller. However, they also reported that Fitts' Law does not adequately model this interaction. We may therefore ask whether, during direct control performed at a distance (for instance, in mid-air), the resulting performance would resemble that observed during direct control with physical contact or during indirect supervisory control of the swarm. This could help to further measure the impact of the difficulty on SoA in a wider variety of applications. For instance, in the search and rescue scenario, directly manipulating the swarm could be difficult because of the large area to cover. One could test larger distances to inform the design of these applications.

Future work can also test different approaches to vary the difficulty of the task, e.g., asking users to create shapes of varying complexity with the swarm UI.

#### 6.4 PREDICTABILITY and SoA for swarm UIs

In our experiment, PREDICTABILITY is manipulated by introducing motor noise, i.e. noise in the modules' trajectory. This PREDICTABILITY reflects the congruence between users' intentions and the outcomes. The motor noise causes a discrepancy between the users' intention and the sensory consequences (i.e., the movement of the modules). Numerous studies have demonstrated the role of congruence between our intentions and the sensory consequences of our actions in the development of the sense of agency (for a review, see [25]). In our situation, we find a fairly weak, albeit present, effect of PREDICTABILITY on the sense of agency. Several reasons can be put forward. First, the motor noise introduced remains fairly low. It is possible that noise of greater amplitude would have a greater disruptive effect on the development of SoA. Similarly, the introduction of temporal incongruence (i.e., a lag) could also increase the effect of this incongruence on SoA. One factor that we believe may be more important concerns the level of intention. Some models distinguish between three classes of intention according to their level in the hierarchy of action control: motor (i.e., *the modules' movements are what I expected*), distal (i.e., *the modules reach the target regardless of their movement*) and proximal (i.e., *the action unfolds as planned, there is no perceived latency between my actions and those of the modules*) [84]. In our case, the incongruence introduced concerns only the motor level of intentions. For example, for the distal level (i.e., at the overall goal level), the modules perform the intended action (i.e., reaching the target). We can therefore imagine a more significant effect of the incongruence if the modules did not reach the chosen target. It is also likely that the weight of each level of intention depends on the role of the modules. In the case of the proxy, the participant performs direct control over the module, so that the congruence between their motor intention and the module's action is certainly paramount. In the case of the satellites, the participant performs remote control over the modules, so that the importance may be placed more on the congruence between their distal intention (the goal to be achieved) and the behavior of the modules. Future work could also replicate this experiment with modules that would not necessarily reach the intended target or stop before the target. These experiments could help assess how distal intentions impact the SoA. This would help HCI designers focus first and foremost on the dimensions that have the most impact on SoA.

#### 6.5 Sense of agency beyond swarm UIs

Assistance will redefine how we perceive and interact with computers. Assistive systems will no longer be seen merely as tools for completing tasks, but as partners with whom we collaborate to achieve a common goal [85]. Among these, swarm robotics gained attention and are now actively studied. Beyond swarm UIs already available to the general public such as Sony Toios –a gaming console for children– many research prototypes would directly benefit from our results, e.g., swarm UIs –2D or 3D– to interact with data visualization [57–59, 67, 88, 108], for architecture [20], for remote telepresence [45], to actuate everyday objects [125], to create animations [55], and to learn physics [70]. Beyond swarms, robotics systems with a collective behavior –even when modules are attached to each other, such as Kilobot [94] or Catoms [89]– can also benefit from our results.

Although we studied assistance with a swarm UI, our results could be applied to other AI systems. Among these assistive systems, some share common characteristics with our experimental setup. For instance, human-robot multi-agent systems [30] could draw on our results. This can be particularly relevant when modeling agent behaviors. In such models, humans should not be expected to react in the same way depending on the type of assistance granted to the robots with which they interact. The model should therefore account for these expected variations in human behavior. For example, if the goal is for humans to retain a sense of responsibility while still benefiting from robots' autonomy,

using a proxy could be a way to provide them with a form of control. It may help restore their sense of responsibility by increasing their sense of agency. Conversely, if responsibility is not a major issue at a given stage of the modeling while autonomy is crucial (e.g., to save time), the model should anticipate these changes of state and their consequences on human behavior.

Although one of the main principles of human–AI interaction is to maintain global control in the hands of users [2], other assistive systems share too few characteristics with our experimental setup to allow us to hypothesize how our results would transfer, such as generative AI application for writing or image/video editing [46]. Moreover, in assistive systems in general, autonomy is known to impact factors other than the sense of agency. For example, autonomy can impact users' engagement: with an autonomous car, the automation of the system can lead to a disengagement of users in the task, named the “out-of-the-loop” phenomenon [19, 76, 98]. Future work could study how assistance articulates with disengagement and sense of agency, in the case of swarm and non-swarm assistive systems.

## 6.6 Limitation of explicit self-reports of sense of agency

Two types of measures are typically used to measure sense of agency, referring to two distinct dimensions of this phenomenal experience: explicit and implicit. Explicit measure consists of asking participants about the degree of control they feel they have over their actions and their consequences. These explicit reports provide access to what is referred to as the Judgement of Agency (JoA). The implicit measure is based on perceptual changes induced by the performance of a voluntary action (in particular, sensory attenuation or intentional binding). These implicit measures make it possible to capture what is known as the feeling of agency (FoA). We chose to use an explicit measurement of SoA for three reasons. First, while implicit measures, such as intentional binding, have been shown to be sensitive to factors such as intentionality and free choice, recent work has questioned their specificity as a measure of agency [60, 62, 106]. In particular, intentional binding effects have been observed even in the absence of voluntary action, suggesting that it may reflect broader mechanisms such as temporal prediction or multisensory integration.

Secondly, several studies seem to indicate a close link between the judgment of agency and the acceptability of technological systems [68, 100], but also between the judgment of agency and the willingness to engage in interaction with technology [39].

Finally, explicit measures of agency are much less complex to implement than implicit measures, involving a shorter experiment and, consequently, a lower risk of participant fatigue. An implicit measure of SoA, like intentional binding, takes longer to log, and in our case would result in a 50% or more increase in the experiment time. Given that our experiments lasted already from 45 to 110 minutes, we chose to maintain an acceptable experimental length for participants by using an explicit measure of SoA.

However, while explicit self-reports are the more direct measure of SoA, they are known to be influenced by cognitive biases and social desirability. For instance, they can lead participants to overestimate their capacity for action and to attribute events unrelated to their own actions [48, 113, 121]. Moreover, the risk of demand characteristics of reported explicit judgments, such as the Likert scales used in our experiments, is commonly discussed [120]. Indeed, participants may be tempted to answer explicit questions in ways they believe are aligned with the experimenter's expectations.

In the end, depending on the experimenter's focus, both explicit and implicit measurements have their strengths and limitations. It would be interesting for future work to replicate our experiments with an implicit measure.

## 6.7 Subjective perception of types of assistance and sense of agency

Beyond the sense of agency, other factors can influence the adoption and use of an interface [71]. These include perceived ease of use and performance [31, 63, 116–118].

**6.7.1 Balance between sense of agency and system autonomy.** Our final questionnaire shows, for the three experiments, negative correlations between, on the one hand, the sense of agency of the swarm and perceived performance, and, on the other hand, perceived ease of use. Perceived performance and ease of use are positively correlated for the three experiments (Appendix F). Subjective results lead us to believe that when an action involves several (potentially autonomous) modules, usability is not determined solely by the experience of controlling each element, but also by the achievement of a common goal, among others. For example, on the one hand, the use of a proxy module and the observation situation minimizes control over the other modules. On the other hand, they increase perceived performance and ease of use. Designers should also take into account such subjective feelings about the swarm UI.

However, striving at all costs to promote complete automation or the use of an intermediate module would be a mistake –at the expense of the sense of agency. In particular, the sense of agency is now considered an important determinant of the responsibility we feel for our actions, and can thus drastically impact human behavior and task engagement [42, 51]. For example, Caspar et al. [22] showed that cadets felt less sense of agency than officers, and also felt less responsible for the consequences of their actions. Drawing on this analogy, excessive automation of swarm UIs, leading to a decrease in the sense of agency (as our experiments demonstrate) could therefore lead to a decrease in the sense of responsibility for the consequences of our actions –a necessary condition for acting ethically and morally [8, 9]. It is therefore important for designers to make design choices based on the criticality of their application (Figure 11). For example, a swarm UI dedicated to the search for victims in rubble-strewn areas requires high performance and therefore high autonomy. But who would be responsible if a landslide occurred because the system chose one route over another without consulting a user? From the user’s perspective, at what point would their sense of agency be too weak to feel responsible for the consequences of the swarm’s actions, even if their supervision was required? On the contrary, in non-critical situations such as games, for example, it can be ok for a designer to vary the performance of their system without users necessarily feeling a strong sense of agency.

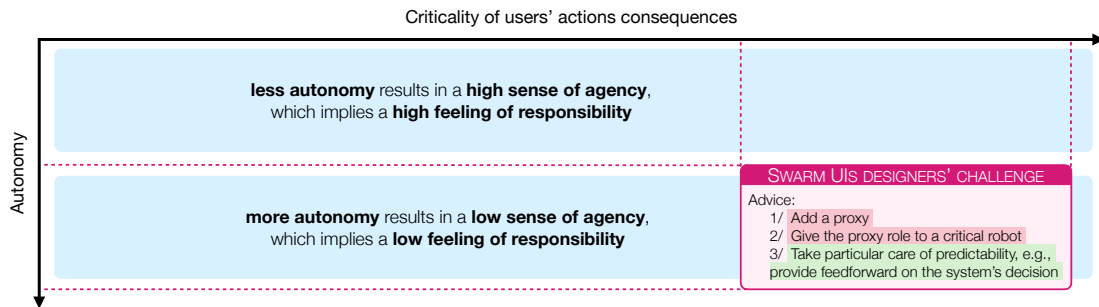


Fig. 11. Guidelines to design autonomous Swarm UIs with a high user’ SoA for critical applications. In blue: main takeaways depending on the autonomy level (low vs. high). In pink: guidelines to design autonomous Swarm UIs for critical applications while maintaining a high user’ SoA. Pieces of advice 1 and 2 (in red) are informed by XP1, advice 3 (in green) is informed by XP3.

**6.7.2  $A_{STOP}$   $P_{w/o}$ :** *Least preferred, perceived as least efficient and least easy to use.* The type of assistance  $A_{STOP}$   $P_{w/o}$  was inspired by existing examples of swarm UIs, such as stop Zooids during their movements [57]. In all our three



experiments, this type of assistance was rated as the least preferred, and perceived as the least efficient and easy to use in the final questionnaire. This could be explained by the implementation of this condition in our experiments: our system requires users to maintain a module still for 150 ms for the system to actually stop their movement. However, we noticed participants removing their hands too early. Fifty-nine participants (out of the 105 total participants) reported that this way of stopping the robots made the interaction:

- Uncomfortable (e.g., P1.4 “I don’t like the fact that it bumps into my hand and doesn’t stop.”, P2.11 “The  $A_{STOP}$  is very uncomfortable, because we have the feeling to stop something in the movement [...], so we have to force against the robot’s motor.”, P3.21 “Honestly, I didn’t find it very pleasant to use. It was more about having to stop it”),
- Slow (e.g., P1.19: “When you put your hand down, it takes a few seconds to stop, and that’s a bit frustrating because you know you have to move quickly, especially when you have to do all four [...] that’s the one where, from a speed perspective, you think, well, it’s less efficient.”, P3.12 “Most of the time it just wouldn’t pick it up right, so we’d lose a bit of time.”), and
- Prone to error (e.g., P1.30 “I even made a mistake because you try to say, ‘OK, that’s good, it’s going to stop,’ you take your hand away, but in fact it doesn’t stop.”, P2.6 “Sometimes it’s not stable enough, and it just goes.”).

Before completely discarding this type of assistance, future work should improve the stopping interaction, such as clicking on the robot.

## 7 Conclusion and General Future Work

Swarm UIs offer new ways for users to interact with systems. Through their actuation, they assist users. However, too much actuation can impact users’ sense of agency (i.e., feeling of control), and, subsequently, the acceptability of swarm UIs, and lead to a loss of users’ responsibility in the actions with the swarm UI. In this article, we studied the impact of the type of assistance on the users’ sense of agency regarding the movement of modules during a joint action with a swarm UI (human-robot cooperative task). We also studied how the task performance and the system predictability modulate this impact. To do this, we extracted nine types of assistance from the literature that can be applied to swarm UIs during a cooperative situation with users. These types of assistance are based on two factors: the autonomy of the system (i.e., the type of control the user can have over different phases of the modules’ movement) and the coordination of movements between modules (i.e., interaction with a single module as a proxy vs. interaction with each module). We assessed subjective SENSE OF AGENCY (SoA) using direct questions for these TYPES OF ASSISTANCE. Our results show a main effect of the level of autonomy, a main effect of the use of a proxy on the SoA, and an interaction effect between them. Our results also show that when using a proxy, the SoA on the movement of the proxy will be greater than that on the swarm, and the SoA on both –SoA on proxy and SoA on swarm– will be greater than that on the satellite modules. The task complexity does not modulate these effects. The system predictability increases the swarm, proxy, and satellite modules SoA. Figure 11 summarizes the outcomes of our paper.

This study opens up several avenues for further research to deepen our understanding of the relationship between sense of agency and the assistance provided by swarm UIs.

In our study, we decided to focus on direct and subjective measurement of SoA. Future work could replicate our study by focusing on indirect and objective measurement of SoA in order to compare the results obtained between the two measurements. In addition, to lay the foundations for the relationship between SoA and the type of assistance provided by a swarm UI, we decided to use a simple and common task in HCI (drag and drop) and to strengthen internal

validity through laboratory testing. On the one hand, future work could try to replicate our results on a more complex task requiring satellite modules to make decisions on their own to optimize task success. On the other hand, future work could try to replicate our results in a more realistic situation, e.g., on a search and rescue application.

## Acknowledgments

This work was supported by the French National Research Agency (ANR) project MolecUI (ANR-21-CE33-0009).

## References

- [1] Jason Alexander, Anne Roudaut, Jürgen Steimle, Kasper Hornbæk, Miguel Bruns Alonso, Sean Follmer, and Timothy Merritt. 2018. Grand challenges in shape-changing interface research. In *Proceedings of the 2018 CHI conference on human factors in computing systems*. 1–14.
- [2] Saleema Amershi, Dan Weld, Mihaela Vorvoreanu, Adam Fourney, Besmira Nushi, Penny Collisson, Jina Suh, Shamsi Iqbal, Paul N Bennett, Kori Inkpen, et al. 2019. Guidelines for human-AI interaction. In *Proceedings of the 2019 chi conference on human factors in computing systems*. 1–13.
- [3] Anonymous. 2025. The Effect of System's Assistance Type on Users' Sense of Agency with modular UIs. [https://osf.io/sw356/?view\\_only=77172b00278b4dd092dcdf67549db1ce](https://osf.io/sw356/?view_only=77172b00278b4dd092dcdf67549db1ce) Last retrieved: 2025-02-27.
- [4] Anonymous. 2025. The Effect of System's Assistance Type on Users' Sense of Agency with modular UIs Depending on System Issue (or trajectory). [https://osf.io/9u4gc/?view\\_only=30a96f70650c478da852cdf543de09fc](https://osf.io/9u4gc/?view_only=30a96f70650c478da852cdf543de09fc) Last retrieved: 2025-08-17.
- [5] Anonymous. 2025. The Effect of System's Assistance Type on Users' Sense of Agency with modular UIs Depending on Task Difficulty. [https://osf.io/7xvpb/?view\\_only=a42aed7c49a64c14a3c876c2ab7d57aa](https://osf.io/7xvpb/?view_only=a42aed7c49a64c14a3c876c2ab7d57aa) Last retrieved: 2025-08-17.
- [6] APA dictionary - cooperation definition 2018. <https://dictionary.apa.org/cooperation>.
- [7] American Psychological Association et al. 2010. *Publication manual of the American psychological association*. American Psychological Association.
- [8] Albert Bandura. 1999. Moral disengagement in the perpetration of inhumanities. *Personality and social psychology review* 3, 3 (1999), 193–209.
- [9] Albert Bandura, Claudio Barbaranelli, Gian Vittorio Caprara, and Concetta Pastorelli. 1996. Mechanisms of moral disengagement in the exercise of moral agency. *Journal of personality and social psychology* 71, 2 (1996), 364.
- [10] Zeynep Barlas, William E Hockley, and Sukhvinder S Obhi. 2017. The effects of freedom of choice in action selection on perceived mental effort and the sense of agency. *Acta psychologica* 180 (2017), 122–129.
- [11] Zeynep Barlas and Sukhvinder S Obhi. 2013. Freedom, choice, and the sense of agency. *Frontiers in human neuroscience* 7 (2013), 514.
- [12] Bruno Berberian. 2019. Man-Machine teaming: a problem of Agency. *IFAC-PapersOnLine* 51, 34 (2019), 118–123.
- [13] Bruno Berberian, Patrick Le Blaye, Nicolas Maille, and Jean-Christophe Sarrazin. 2012. Sense of Control in Supervision Tasks of Automated Systems. *Aerospace Lab* 4 (2012), p–1.
- [14] Bruno Berberian, Patrick Le Blaye, Christian Schulte, Nawfel Kinani, and Pern Ren Sim. 2013. Data transmission latency and sense of control. In *Engineering Psychology and Cognitive Ergonomics. Understanding Human Cognition: 10th International Conference, EPCE 2013, Held as Part of HCI International 2013, Las Vegas, NV, USA, July 21-26, 2013, Proceedings, Part I* 10. Springer, 3–12.
- [15] Bruno Berberian, Jean-Christophe Sarrazin, Patrick Le Blaye, and Patrick Haggard. 2012. Automation technology and sense of control: a window on human agency. *PloS one* 7, 3 (2012), e34075.
- [16] Joanna Bergström, Jarrod Knibbe, Henning Pohl, and Kasper Hornbæk. 2022. Sense of agency and user experience: Is there a link? *ACM Transactions on Computer-Human Interaction (TOCHI)* 29, 4 (2022), 1–22.
- [17] Joanna Bergstrom-Lehtovirta, David Coyle, Jarrod Knibbe, and Kasper Hornbæk. 2018. I really did that: Sense of agency with touchpad, keyboard, and on-skin interaction. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*. 1–8.
- [18] Frederike Beyer, Nura Sidarus, Sofia Bonicalzi, and Patrick Haggard. 2017. Beyond self-serving bias: diffusion of responsibility reduces sense of agency and outcome monitoring. *Social cognitive and affective neuroscience* 12, 1 (2017), 138–145.
- [19] Francesco N Biondi, Monika Lohani, Rachel Hopman, Sydney Mills, Joel M Cooper, and David L Strayer. 2018. 80 MPH and out-of-the-loop: Effects of real-world semi-automated driving on driver workload and arousal. In *Proceedings of the human factors and ergonomics society annual meeting*, Vol. 62. Sage Publications Sage CA: Los Angeles, CA, 1878–1882.
- [20] Sean Braley, Calvin Rubens, Timothy Merritt, and Roel Vertegaal. 2018. GridDrones: A Self-Levitating Physical Voxel Lattice for Interactive 3D Surface Deformations.. In *UIST*, Vol. 18. D200.
- [21] Emilie A Caspar, Axel Cleeremans, and Patrick Haggard. 2018. Only giving orders? An experimental study of the sense of agency when giving or receiving commands. *PloS one* 13, 9 (2018), e0204027.
- [22] Emilie A Caspar, Salvatore Lo Bue, Pedro A Magalhães De Saldanha da Gama, Patrick Haggard, and Axel Cleeremans. 2020. The effect of military training on the sense of agency and outcome processing. *Nature communications* 11, 1 (2020), 4366.
- [23] Valerian Chambon, Elisa Filevich, and Patrick Haggard. 2014. What is the human sense of agency, and is it Metacognitive? In *The cognitive neuroscience of metacognition*. Springer, 321–342.
- [24] Valerian Chambon and Patrick Haggard. 2012. Sense of control depends on fluency of action selection, not motor performance. *Cognition* 125, 3 (2012), 441–451.

- [25] Valérian Chabon, Nura Sidarus, and Patrick Haggard. 2014. From action intentions to action effects: how does the sense of agency come about? *Frontiers in human neuroscience* 8 (2014), 320.
- [26] Thierry Chaminade and Jean Decety. 2002. Leader or follower? Involvement of the inferior parietal lobule in agency. *Neuroreport* 13, 15 (2002), 1975–1978.
- [27] Francesca Ciardo, Frederike Beyer, Davide De Tommaso, and Agnieszka Wykowska. 2020. Attribution of intentional agency towards robots reduces one's own sense of agency. *Cognition* 194 (2020), 104109.
- [28] Francesca Ciardo, Davide De Tommaso, Frederike Beyer, and Agnieszka Wykowska. 2018. Reduced sense of agency in human-robot interaction. In *Social Robotics: 10th International Conference, ICSR 2018, Qingdao, China, November 28-30, 2018, Proceedings 10*. Springer, 441–450.
- [29] David Coyle, James Moore, Per Ola Kristensson, Paul Fletcher, and Alan Blackwell. 2012. I did that! Measuring users' experience of agency in their own actions. In *Proceedings of the SIGCHI conference on human factors in computing systems*. 2025–2034.
- [30] Abhinav Dahiya, Alexander M Aroyo, Kerstin Dautenhahn, and Stephen L Smith. 2023. A survey of multi-agent human-robot interaction systems. *Robotics and Autonomous Systems* 161 (2023), 104335.
- [31] Fred D Davis, Richard P Bagozzi, and Paul R Warshaw. 1989. User acceptance of computer technology: A comparison of two theoretical models. *Management science* 35, 8 (1989), 982–1003.
- [32] Jelle Demanet, Paul S Muhle-Karbe, Margaret T Lynn, Iris Blotenberg, and Marcel Brass. 2013. Power to the will: How exerting physical effort boosts the sense of agency. *Cognition* 129, 3 (2013), 574–578.
- [33] John A Dewey and Günther Knoblich. 2014. Do implicit and explicit measures of the sense of agency measure the same thing? *PLoS one* 9, 10 (2014), e110118.
- [34] Xuan Di and Rongye Shi. 2021. A survey on autonomous vehicle control in the era of mixed-autonomy: From physics-based to AI-guided driving policy learning. *Transportation research part C: emerging technologies* 125 (2021), 103008.
- [35] Griffin Dietz, Jane L E, Peter Washington, Lawrence H Kim, and Sean Follmer. 2017. Human perception of swarm robot motion. In *Proceedings of the 2017 CHI conference extended abstracts on human factors in computing systems*. 2520–2527.
- [36] Morgan Dixon, James Fogarty, and Jacob Wobbrock. 2012. A general-purpose target-aware pointing enhancement using pixel-level analysis of graphical interfaces. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems (Austin, Texas, USA) (CHI '12)*. Association for Computing Machinery, New York, NY, USA, 3167–3176. doi:10.1145/2207676.2208734
- [37] Marco Dorigo, Dario Floreano, Luca Maria Gambardella, Francesco Mondada, Stefano Nolfi, Tarek Baaboura, Mauro Birattari, Michael Bonani, Manuele Brambilla, Arne Brutschy, et al. 2013. Swarmanoid: a novel concept for the study of heterogeneous robotic swarms. *IEEE Robotics & Automation Magazine* 20, 4 (2013), 60–71.
- [38] Pierre Dragicevic. 2016. Fair statistical communication in HCI. *Modern statistical methods for HCI* (2016), 291–330.
- [39] Baruch Eitam, Patrick M Kennedy, and E Tory Higgins. 2013. Motivation from control. *Experimental brain research* 229, 3 (2013), 475–484.
- [40] Paul M Fitts. 1954. The information capacity of the human motor system in controlling the amplitude of movement. *Journal of experimental psychology* 47, 6 (1954), 381.
- [41] Wolfgang Forstmeier, Eric-Jan Wagenmakers, and Timothy H. Parker. 2017. Detecting and avoiding likely false-positive findings – a practical guide. *Biological Reviews* 92 (November 2017), 1941–1968. Issue 4. doi:10.1111/brv.12315
- [42] Chris D Frith. 2014. Action, agency and responsibility. *Neuropsychologia* 55 (2014), 137–142.
- [43] Alix Goguet, Mathieu Nancel, Géry Casiez, and Daniel Vogel. 2016. The Performance and Preference of Different Fingers and Chords for Pointing, Dragging, and Object Transformation. In *Proceedings of the 2016 CHI Conference on Human Factors in Computing Systems (San Jose, California, USA) (CHI '16)*. Association for Computing Machinery, New York, NY, USA, 4250–4261. doi:10.1145/2858036.2858194
- [44] Alix Goguet, Cameron Steer, Andrés Lucero, Laurence Nigay, Deepak Ranjan Sahoo, Céline Coutrix, Anne Roudaut, Sriram Subramanian, Yutaka Tokuda, Timothy Neate, et al. 2019. PickCells: A physically reconfigurable cell-composed touchscreen. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*. 1–14.
- [45] Antonio Gomes, Calvin Rubens, Sean Braley, and Roel Vertegaal. 2016. Bitdrones: Towards using 3d nanocopter displays as interactive self-levitating programmable matter. In *Proceedings of the 2016 CHI Conference on Human Factors in Computing Systems*. 770–780.
- [46] Roberto Gozalo-Brizuela and Eduardo C Garrido-Merchán. 2023. A survey of Generative AI Applications. *arXiv preprint arXiv:2306.02781* (2023).
- [47] Ouriel Grynszpan, Aisha Sahai, Nasmeh Hamidi, Elisabeth Pacherie, Bruno Berberian, Lucas Roche, and Ludovic Saint-Bauzel. 2019. The sense of agency in human-human vs human-robot joint action. *Consciousness and cognition* 75 (2019), 102820.
- [48] Patrick Haggard. 2017. Sense of agency in the human brain. *Nature Reviews Neuroscience* 18, 4 (2017), 196–207.
- [49] Patrick Haggard and Valerian Chabon. 2012. Sense of agency. *Current biology* 22, 10 (2012), R390–R392.
- [50] Patrick Haggard, Sam Clark, and Jeri Kalogeras. 2002. Voluntary action and conscious awareness. *Nature neuroscience* 5, 4 (2002), 382–385.
- [51] Patrick Haggard and Manos Tsakiris. 2009. The experience of agency: Feelings, judgments, and responsibility. *Current Directions in Psychological Science* 18, 4 (2009), 242–246.
- [52] Maria-Theresa Oanh Hoang, Niels Van Berkel, Mikael B Skov, and Timothy R Merritt. 2023. Challenges and requirements in multi-drone interfaces. In *Extended Abstracts of the 2023 CHI Conference on Human Factors in Computing Systems*. 1–9.
- [53] Emma E Howard, S Gareth Edwards, and Andrew P Bayliss. 2016. Physical and mental effort disrupts the implicit sense of agency. *Cognition* 157 (2016), 114–125.

- [54] Kazuya Inoue, Yuji Takeda, and Motohiro Kimura. 2017. Sense of agency in continuous action: Assistance-induced performance improvement is self-attributed even with knowledge of assistance. *Consciousness and cognition* 48 (2017), 246–252.
- [55] Matthew Jeung, Anup Sathya, Wanli Qian, Steven Arellano, Luke Jimenez, and Ken Nakagaki. 2025. Shape n'Swarm: Hands-on, Shape-aware Generative Authoring with Swarm UI and LLMs. In *Proceedings of the 38th Annual ACM Symposium on User Interface Software and Technology*. 1–16.
- [56] Lawrence H Kim, Veronika Domova, Yuqi Yao, and Parsa Rajabi. 2023. SwarmFidget: Exploring Programmable Actuated Fidgeting with Swarm Robots. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology*. 1–15.
- [57] Lawrence H Kim, Daniel S Drew, Veronika Domova, and Sean Follmer. 2020. User-defined swarm robot control. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*. 1–13.
- [58] Lawrence H Kim and Sean Follmer. 2017. Ubiswarm: Ubiquitous robotic interfaces and investigation of abstract motion as a display. *Proceedings of the ACM on Interactive, Mobile, Wearable and Ubiquitous Technologies* 1, 3 (2017), 1–20.
- [59] Lawrence H Kim and Sean Follmer. 2019. Swarmhaptics: Haptic display with swarm robots. In *Proceedings of the 2019 CHI conference on human factors in computing systems*. 1–13.
- [60] Wladimir Kirsch, Wilfried Kunde, and Oliver Herbort. 2019. Intentional binding is unrelated to action intention. *Journal of Experimental Psychology: Human Perception and Performance* 45, 3 (2019), 378.
- [61] Gaiqing Kong, Cheryne Aberkane, Clément Desoche, Alessandro Farnè, and Marine Vernet. 2023. No evidence in favour of the existence of 'intentional' binding. *bioRxiv* (2023), 2023–02.
- [62] Gaiqing Kong, Cheryne Aberkane, Clément Desoche, Alessandro Farnè, and Marine Vernet. 2024. No evidence in favor of the existence of "intentional" binding. *Journal of experimental psychology: human perception and performance* 50, 6 (2024), 626.
- [63] Songpol Kulviwat, Gordon C Bruner II, Anand Kumar, Suzanne A Nasco, and Terry Clark. 2007. Toward a unified theory of consumer acceptance technology. *Psychology & Marketing* 24, 12 (2007), 1059–1084.
- [64] Vali Lalioti, Ken Nakagaki, Ramarko Bhattacharya, and Yasuaki Kakehi. 2024. [e] Motion: Designing Expressive Movement in Robots and Actuated Tangible User Interfaces. In *Proceedings of the Eighteenth International Conference on Tangible, Embedded, and Embodied Interaction*. 1–3.
- [65] Solène Le Bars, Alexandre Devaux, Tena Nevidal, Valerian Chambon, and Elisabeth Pachier. 2020. Agents' pivotality and reward fairness modulate sense of agency in cooperative joint action. *Cognition* 195 (2020), 104117.
- [66] Alexis Le Besnerais, James W Moore, Bruno Berberian, and Ouriel Grynszpan. 2024. Sense of agency in joint action: a critical review of we-agency. *Frontiers in psychology* 15 (2024), 1331084.
- [67] Mathieu Le Goc, Lawrence H Kim, Ali Parsaei, Jean-Daniel Fekete, Pierre Dragicevic, and Sean Follmer. 2016. Zooids: Building blocks for swarm user interfaces. In *Proceedings of the 29th annual symposium on user interface software and technology*. 97–109.
- [68] Kevin Le Goff, Arnaud Rey, Patrick Haggard, Olivier Oullier, and Bruno Berberian. 2018. Agency modulates interactions with automation technologies. *Ergonomics* 61, 9 (2018), 1282–1297.
- [69] Marin Le Guillou, Laurent Prévot, and Bruno Berberian. 2023. Bringing together ergonomic concepts and cognitive mechanisms for human–AI agents cooperation. *International Journal of Human–Computer Interaction* 39, 9 (2023), 1827–1840.
- [70] Jiatong Li, Ryo Suzuki, and Ken Nakagaki. 2023. Physica: Interactive Tangible Physics Simulation based on Tabletop Mobile Robots Towards Explorable Physics Education. In *Proceedings of the 2023 ACM Designing Interactive Systems Conference*. 1485–1499.
- [71] Hannah Limerick, David Coyle, and James W Moore. 2014. The experience of agency in human-computer interactions: a review. *Frontiers in human neuroscience* 8 (2014), 643.
- [72] Janeen D Loehr. 2022. The sense of agency in joint action: An integrative review. *Psychonomic Bulletin & Review* 29, 4 (2022), 1089–1117.
- [73] Yifang Ma, Zhenyu Wang, Hong Yang, and Lin Yang. 2020. Artificial intelligence applications in the development of autonomous vehicles: A survey. *IEEE/CAA Journal of Automatica Sinica* 7, 2 (2020), 315–329.
- [74] I Scott MacKenzie. 1989. A note on the information-theoretic basis for Fitts' law. *Journal of motor behavior* 21, 3 (1989), 323–330.
- [75] John E McEneaney. 2013. Agency effects in human–computer interaction. *International Journal of Human–Computer Interaction* 29, 12 (2013), 798–813.
- [76] Natasha Merat, Bobbie Seppelt, Tyron Louw, Johan Engström, John D Lee, Emma Johansson, Charles A Green, Satoshi Katagaki, Chris Monk, Makoto Itoh, et al. 2019. The "Out-of-the-Loop" concept in automated driving: proposed definition, measures and implications. *Cognition, Technology & Work* 21, 1 (2019), 87–98.
- [77] Janet Metcalfe and Matthew Jason Greene. 2007. Metacognition of agency. *Journal of Experimental Psychology: General* 136, 2 (2007), 184.
- [78] Rin Minohara, Wen Wen, Shunsuke Hamasaki, Takaki Maeda, Motochiro Kato, Hiroshi Yamakawa, Atsushi Yamashita, and Hajime Asama. 2016. Strength of intentional effort enhances the sense of agency. *Frontiers in psychology* 7 (2016), 1165.
- [79] Morikatron. 2024. DATAtab: Online Statistics Calculator. [https://morikatron.github.io/toio-sdk-for-unity/docs\\_EN/](https://morikatron.github.io/toio-sdk-for-unity/docs_EN/) Last retrieved: 2024-08-27.
- [80] Sasanka Nagavalli, Meghan Chandarana, Katia Sycara, and Michael Lewis. 2017. Multi-operator gesture control of robotic swarms using wearable devices. In *Proceedings of the Tenth International Conference on Advances in Computer-Human Interactions*. IARIA.
- [81] Don Norman. 2004. *Emotional design: Why we love (or hate) everyday things*. Basic books.
- [82] Sukhvinder S Obhi and Preston Hall. 2011. Sense of agency and intentional binding in joint action. *Experimental brain research* 211 (2011), 655–662.
- [83] Sukhvinder S Obhi and Preston Hall. 2011. Sense of agency in joint action: Influence of human and computer co-actors. *Experimental brain research* 211 (2011), 663–670.

- [84] Elisabeth Pacherie. 2008. The phenomenology of action: A conceptual framework. *Cognition* 107, 1 (2008), 179–217.
- [85] Marine Pagliari, Valérian Chambon, and Bruno Berberian. 2022. What is new with Artificial Intelligence? Human-agent interactions through the lens of social agency. *Frontiers in Psychology* 13 (2022), 954444.
- [86] Brian Pendleton and Michael Goodrich. 2013. Scalable human interaction with robotic swarms. In *AIAA Infotech@ Aerospace (I@ a) Conference*. 4731.
- [87] Ekaterina Peshkova and Martin Hitz. 2017. Exploring user-defined gestures to control a group of four UAVs. In *2017 26th IEEE International Symposium on Robot and Human Interactive Communication (RO-MAN)*. IEEE, 169–174.
- [88] Alan Pham, Yuxiao Li, Miyu Fukuoka, and Ken Nakagaki. 2025. Buoyancé: Reeling Helium-Inflated Balloons with Mobile Robots on the Ground for Mid-Air Tangible Display, Interaction, and Assembly. In *Proceedings of the 38th Annual ACM Symposium on User Interface Software and Technology*. 1–19.
- [89] Benoit Piranda and Julien Bourgeois. 2018. Designing a quasi-spherical module for a huge modular robot to create programmable matter. *Autonomous Robots* 42, 8 (2018), 1619–1633.
- [90] Laura Prusko, Céline Coutrix, Yann Laurillau, Benoit Piranda, and Julien Bourgeois. 2021. Molecular hci: Structuring the cross-disciplinary space of modular shape-changing user interfaces. *Proceedings of the ACM on Human-Computer Interaction* 5, EICS (2021), 1–33.
- [91] Samira Pulatova and Lawrence H Kim. 2024. Co-Designing Programmable Fidgeting Experience with Swarm Robots for Adults with ADHD. In *Proceedings of the 26th International ACM SIGACCESS Conference on Computers and Accessibility*. 1–14.
- [92] Paul Reddish, Eddie MW Tong, Jonathan Jong, and Harvey Whitehouse. 2020. Interpersonal synchrony affects performers' sense of agency. *Self and Identity* 19, 4 (2020), 389–411.
- [93] Cecilia Roselli, Francesca Ciardo, and Agnieszka Wykowska. 2019. Robots improve judgments on self-generated actions: an Intentional Binding Study. In *Social Robotics: 11th International Conference, ICSR 2019, Madrid, Spain, November 26–29, 2019, Proceedings 11*. Springer, 88–97.
- [94] Michael Rubenstein, Christian Ahler, Nick Hoff, Adrian Cabrera, and Radhika Nagpal. 2014. Kilobot: A low cost robot with scalable operations designed for collective behaviors. *Robotics and Autonomous Systems* 62, 7 (2014), 966–975.
- [95] Nihar Sabnis, Ata Otaran, Dennis Wittchen, Johanna K. Didion, Jürgen Steimle, and Paul Strohmeier. 2025. Foot Pedal Control: The Role of Vibrotactile Feedback in Performance and Perceived Control. In *Proceedings of the Nineteenth International Conference on Tangible, Embedded, and Embodied Interaction*. 1–15.
- [96] Aisha Sahai, Emilie Caspar, Albert De Beir, Ouriel Grynszpan, Elisabeth Pacherie, and Bruno Berberian. 2023. Modulations of one's sense of agency during human-machine interactions: a behavioural study using a full humanoid robot. *Quarterly journal of experimental psychology* 76, 3 (2023), 606–620.
- [97] Aisha Sahai, Andrea Desantis, Ouriel Grynszpan, Elisabeth Pacherie, and Bruno Berberian. 2019. Action co-representation and the sense of agency during a joint Simon task: Comparing human and machine co-agents. *Consciousness and cognition* 67 (2019), 44–55.
- [98] Damien Schnebelen, Camilo Charron, and Franck Mars. 2020. Estimating the out-of-the-loop phenomenon from visual strategies during highly automated driving. *Accident Analysis & Prevention* 148 (2020), 105776.
- [99] Thomas B Sheridan, William L Verplank, and TL Brooks. 1978. Human/computer control of undersea teleoperators. In *NASA. Ames Res. Center The 14th Ann. Conf. on Manual Control*.
- [100] Ben Shneiderman and Catherine Plaisant. 2004. *Designing the user interface: strategies for effective human-computer interaction*. Pearson Addison-Wesley éd.
- [101] Nura Sidorus and Patrick Haggard. 2016. Difficult action decisions reduce the sense of agency: A study using the Eriksen flanker task. *Acta psychologica* 166 (2016), 1–11.
- [102] Nura Sidorus, Matti Vuorre, Janet Metcalfe, and Patrick Haggard. 2017. Investigating the prospective sense of agency: Effects of processing fluency, stimulus ambiguity, and response conflict. *Frontiers in Psychology* 8 (2017), 545.
- [103] Crystal A Silver, Benjamin W Tatler, Ramakrishna Chakravarthi, and Bert Timmermans. 2021. Social agency as a continuum. *Psychonomic Bulletin & Review* 28, 2 (2021), 434–453.
- [104] Metin Sitti, Hakan Ceylan, Wenqi Hu, Joshua Giltinan, Mehmet Turan, Sehyuk Yim, and Eric Diller. 2015. Biomedical applications of untethered mobile milli/microrobots. *Proc. IEEE* 103, 2 (2015), 205–224.
- [105] Lars Strother, Kristin A House, and Sukhvinder S Obhi. 2010. Subjective agency and awareness of shared actions. *Consciousness and cognition* 19, 1 (2010), 12–20.
- [106] Keisuke Suzuki, Peter Lush, Anil K Seth, and Warrick Roseboom. 2019. Intentional binding without intentional action. *Psychological science* 30, 6 (2019), 842–853.
- [107] Ryo Suzuki, Junichi Yamaoka, Daniel Leithinger, Tom Yeh, Mark D Gross, Yoshihiro Kawahara, and Yasuaki Kakehi. 2018. Dynablock: Dynamic 3d printing for instant and reconstructable shape formation. In *Proceedings of the 31st annual ACM symposium on user interface software and technology*. 99–111.
- [108] Ryo Suzuki, Clement Zheng, Yasuaki Kakehi, Tom Yeh, Ellen Yi-Luen Do, Mark D Gross, and Daniel Leithinger. 2019. Shapebots: Shape-changing swarm robots. In *Proceedings of the 32nd annual ACM symposium on user interface software and technology*. 493–505.
- [109] Matthias Synofzik, Gottfried Vosgerau, and Axel Lindner. 2009. Me or not me—an optimal integration of agency cues? *Consciousness and cognition* 18, 4 (2009), 1065–1068.

- [110] Takumi Tanaka, Takuya Matsumoto, Shintaro Hayashi, Shiro Takagi, and Hideaki Kawabata. 2019. What makes action and outcome temporally close to each other: A systematic review and meta-analysis of temporal binding. *Timing & Time Perception* 7, 3 (2019), 189–218.
- [111] Toio (official website) 2018. <https://www.sony.com/en/SonyInfo/design/stories/toio/>.
- [112] Transparent Statistics in Human–Computer Interaction [n. d.]. <https://transparentstatistics.org>.
- [113] Manos Tsakiris, Patrick Haggard, Nicolas Franck, Nelly Mainy, and Angela Sirigu. 2005. A specific role for efferent information in self-recognition. *Cognition* 96, 3 (2005), 215–231.
- [114] Sayako Ueda, Ryoichi Nakashima, and Takatsune Kumada. 2021. Influence of levels of automation on the sense of agency during continuous action. *Scientific reports* 11, 1 (2021), 2436.
- [115] Robrecht PRD van der Wel, Natalie Sebanz, and Guenther Knoblich. 2012. The sense of agency during skill learning in individuals and dyads. *Consciousness and cognition* 21, 3 (2012), 1267–1279.
- [116] Viswanath Venkatesh and Fred D Davis. 2000. A theoretical extension of the technology acceptance model: Four longitudinal field studies. *Management science* 46, 2 (2000), 186–204.
- [117] Viswanath Venkatesh, Michael G Morris, Gordon B Davis, and Fred D Davis. 2003. User acceptance of information technology: Toward a unified view. *MIS quarterly* (2003), 425–478.
- [118] Viswanath Venkatesh, James YL Thong, and Xin Xu. 2012. Consumer acceptance and use of information technology: extending the unified theory of acceptance and use of technology. *MIS quarterly* (2012), 157–178.
- [119] Phillip Walker, Steven Nunnally, Michael Lewis, Nilanjan Chakraborty, and Katia Sycara. 2013. Levels of automation for human influence of robot swarms. In *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, Vol. 57. SAGE Publications Sage CA: Los Angeles, CA, 429–433.
- [120] Stephen J Weber and Thomas D Cook. 1972. Subject effects in laboratory research: an examination of subject roles, demand characteristics, and valid inference. *Psychological bulletin* 77, 4 (1972), 273.
- [121] Daniel M Wegner and Thalia Wheatley. 1999. Apparent mental causation: Sources of the experience of will. *American psychologist* 54, 7 (1999), 480.
- [122] Wen Wen, Atsushi Yamashita, and Hajime Asama. 2015. The sense of agency during continuous action: performance is more important than action-feedback association. *PloS one* 10, 4 (2015), e0125226.
- [123] Dorit Wenke, Stephen M Fleming, and Patrick Haggard. 2010. Subliminal priming of actions influences sense of control over effects of action. *Cognition* 115, 1 (2010), 26–38.
- [124] Kaito Yamada, Hiroki Usuba, and Homei Miyahita. 2019. Modeling Drone Pointing Movement with Fitts’ Law. In *Extended Abstracts of the 2019 CHI Conference on Human Factors in Computing Systems*. 1–6.
- [125] Lilith Yu, Chenfeng Gao, David Wu, and Ken Nakagaki. 2023. AeroRigUI: Actuated TUIs for spatial interaction using rigging swarm robots on ceilings in everyday space. In *Proceedings of the 2023 CHI conference on human factors in computing systems*. 1–18.
- [126] Debora Zanatto, Mark Chattington, and Jan Noyes. 2021. Human-machine sense of agency. *International Journal of Human-Computer Studies* 156 (2021), 102716.
- [127] Debora Zanatto, Mark Chattington, and Jan Noyes. 2021. Sense of agency in human-machine interaction. In *Advances in Neuroergonomics and Cognitive Engineering: Proceedings of the AHFE 2021 Virtual Conferences on Neuroergonomics and Cognitive Engineering, Industrial Cognitive Ergonomics and Engineering Psychology, and Cognitive Computing and Internet of Things, July 25-29, 2021, USA*. Springer, 353–360.
- [128] Laura Zapparoli, Eraldo Paulesu, Marika Mariano, Alessia Ravani, and Lucia M Sacheli. 2022. The sense of agency in joint actions: A theory-driven meta-analysis. *Cortex* 148 (2022), 99–120.

## A Tables of means and confidence intervals for the first experiment

### A.1 Means and 95% *bootstrapped* confidence intervals for AUTONOMY (Hyp1.1 and 1.1<sub>bis</sub>)

| TYPE OF ASSISTANCE |                                                     | SoA Mean [CI]     |
|--------------------|-----------------------------------------------------|-------------------|
| AUTONOMY           | $A_{\text{INERT}}$ ( <i>Not autonomous</i> )        | 7 [6.81, 7.18]    |
|                    | $A_{\text{SNAP}}$ ( <i>Semi-autonomous</i> )        | 6 [5.60, 6.29]    |
|                    | $A_{\text{STOP}}$ ( <i>Semi-autonomous</i> )        | 5.41 [4.92, 5.79] |
|                    | $A_{\text{TRG}}$ ( <i>Highly autonomous</i> )       | 4.26 [3.72, 4.85] |
| BASELINE           |                                                     | SoA Mean [CI]     |
| AUTONOMY           | $A_{\text{CMPLT}}$ ( <i>Completely autonomous</i> ) | 1.58 [1.22, 2.23] |

### A.2 Means and 95% *bootstrapped* confidence intervals for PROXY (Hyp1.2)

| TYPE OF ASSISTANCE | SoA Mean [CI]     |                   |
|--------------------|-------------------|-------------------|
|                    | PROXY             |                   |
|                    | $P_{\text{w/o}}$  | $P_{\text{WITH}}$ |
| SoA Mean [CI]      | 6.04 [5.72, 6.30] | 5.29 [4.85, 5.70] |
| BASELINE           | PROXY             |                   |
|                    | $P_{\text{OBS}}$  |                   |
|                    | 1.58 [1.22, 2.23] |                   |

### A.3 Means and 95% *bootstrapped* confidence intervals for TYPE OF ASSISTANCE

|                    |                                                    | SoA Mean<br>[CI]     |                      |
|--------------------|----------------------------------------------------|----------------------|----------------------|
| TYPE OF ASSISTANCE |                                                    | PROXY                |                      |
|                    |                                                    | P <sub>W/O</sub>     | P <sub>WITH</sub>    |
| AUTONOMY           | A <sub>INERT</sub><br><i>Not autonomous</i>        | 7.71<br>[7.54, 7.83] | 6.29<br>[5.89, 6.63] |
|                    | A <sub>SNAP</sub><br><i>Semi-autonomous</i>        | 6.26<br>[5.87, 6.55] | 5.74<br>[5.29, 6.09] |
|                    | A <sub>STOP</sub><br><i>Semi-autonomous</i>        | 5.70<br>[5.18, 6.09] | 5.12<br>[4.62, 5.58] |
|                    | A <sub>TRG</sub><br><i>Highly autonomous</i>       | 4.50<br>[3.97, 5.04] | 4.01<br>[3.44, 4.68] |
| BASELINE           |                                                    | PROXY                |                      |
|                    |                                                    | P <sub>OBS</sub>     |                      |
| AUTONOMY           | A <sub>CMPLT</sub><br><i>Completely autonomous</i> | 1.58<br>[1.22, 2.23] |                      |

### A.4 Means and 95% *bootstrapped* confidence intervals for TYPE OF ASSISTANCE × MODULE'S ROLE

|                                  |                    | SoA Mean<br>[CI]     |                      |                      |
|----------------------------------|--------------------|----------------------|----------------------|----------------------|
| TYPE OF ASSISTANCE               |                    | MODULE'S ROLE        |                      |                      |
|                                  |                    | PROXY                | SATELLITES           | SWARM                |
| AUTONOMY<br>(P <sub>WITH</sub> ) | A <sub>INERT</sub> | 7.56<br>[7.30, 7.71] | 4.54<br>[3.87, 5.17] | 6.29<br>[5.89, 6.63] |
|                                  | A <sub>SNAP</sub>  | 6.42<br>[5.97, 6.71] | 4.12<br>[3.47, 4.76] | 5.74<br>[5.29, 6.09] |
|                                  | A <sub>STOP</sub>  | 5.67<br>[5.15, 6.08] | 3.73<br>[3.11, 4.40] | 5.12<br>[4.62, 5.58] |
|                                  | A <sub>TRG</sub>   | 4.28<br>[3.70, 4.91] | 3.14<br>[2.49, 3.90] | 4.01<br>[3.44, 4.68] |

## B Post hoc tests for the first experiment

### B.1 Tukey's post hoc tests for AUTONOMY

Table 1 shows Tukey's post hoc tests for AUTONOMY in the first experiment.



Table 1. Tukey's post hoc tests for AUTONOMY in the first experiment. In red, the insignificant p-values and in orange the borderline p-values (between 0.04 and 0.06).

| Comparison  |            | t          | $p_{Tukey}$ |
|-------------|------------|------------|-------------|
| AUTONOMY 1  | AUTONOMY 2 | (ddl = 34) |             |
| $A_{INERT}$ | $A_{SNAP}$ | 8.45       | <.001       |
|             | $A_{STOP}$ | 9.00       | <.001       |
|             | $A_{TRG}$  | 10.91      | <.001       |
| $A_{SNAP}$  | $A_{STOP}$ | 3.57       | .006        |
|             | $A_{TRG}$  | 7.89       | <.001       |
| $A_{STOP}$  | $A_{TRG}$  | 6.99       | <.001       |

## B.2 Tukey's post hoc tests for TYPES OF ASSISTANCE for first experiment

Table 2 shows Tukey's post hoc tests for TYPES OF ASSISTANCE in the first experiment.

Table 2. Tukey’s post hoc tests for TYPES OF ASSISTANCE in the first experiment. In red, the insignificant p-values and in orange the borderline p-values (between 0.04 and 0.06). For a more condensed notation, TYPES OF ASSISTANCE WITHOUT PROXY are named only by the name of the AUTONOMY, and TYPES OF ASSISTANCE WITH PROXY are named as follows: “name of the AUTONOMY P<sub>WITH</sub>”.

| Comparison                           |                                      | t<br>(ddl = 34) | ptukey |
|--------------------------------------|--------------------------------------|-----------------|--------|
| TYPE OF ASSISTANCE 1                 | TYPE OF ASSISTANCE 2                 |                 |        |
| A <sub>INERT</sub>                   | A <sub>INERT</sub> P <sub>WITH</sub> | 6.520           | <.001  |
|                                      | A <sub>SNAP</sub>                    | 8.249           | <.001  |
|                                      | A <sub>SNAP</sub> P <sub>WITH</sub>  | 8.512           | <.001  |
|                                      | A <sub>STOP</sub>                    | 8.994           | <.001  |
|                                      | A <sub>STOP</sub> P <sub>WITH</sub>  | 9.756           | <.001  |
|                                      | A <sub>TRG</sub>                     | 10.802          | <.001  |
|                                      | A <sub>TRG</sub> P <sub>WITH</sub>   | 10.420          | <.001  |
| A <sub>INERT</sub> P <sub>WITH</sub> | A <sub>SNAP</sub>                    | 0.163           | 1.000  |
|                                      | A <sub>SNAP</sub> P <sub>WITH</sub>  | 5.946           | <.001  |
|                                      | A <sub>STOP</sub>                    | 2.519           | 0.222  |
|                                      | A <sub>STOP</sub> P <sub>WITH</sub>  | 7.340           | <.001  |
|                                      | A <sub>TRG</sub>                     | 8.349           | <.001  |
|                                      | A <sub>TRG</sub> P <sub>WITH</sub>   | 9.385           | <.001  |
| A <sub>SNAP</sub>                    | A <sub>SNAP</sub> P <sub>WITH</sub>  | 3.658           | .017   |
|                                      | A <sub>STOP</sub>                    | 2.956           | .092   |
|                                      | A <sub>STOP</sub> P <sub>WITH</sub>  | 5.339           | <.001  |
|                                      | A <sub>TRG</sub>                     | 7.455           | <.001  |
|                                      | A <sub>TRG</sub> P <sub>WITH</sub>   | 7.534           | <.001  |
| A <sub>SNAP</sub> P <sub>WITH</sub>  | A <sub>STOP</sub>                    | 0.179           | 1.000  |
|                                      | A <sub>STOP</sub> P <sub>WITH</sub>  | 3.904           | .009   |
|                                      | A <sub>TRG</sub>                     | 6.284           | <.001  |
|                                      | A <sub>TRG</sub> P <sub>WITH</sub>   | 7.678           | <.001  |
| A <sub>STOP</sub>                    | A <sub>STOP</sub> P <sub>WITH</sub>  | 3.432           | .030   |
|                                      | A <sub>TRG</sub>                     | 6.124           | <.001  |
|                                      | A <sub>TRG</sub> P <sub>WITH</sub>   | 6.064           | <.001  |
| A <sub>STOP</sub> P <sub>WITH</sub>  | A <sub>TRG</sub>                     | 5.342           | <.001  |
|                                      | A <sub>TRG</sub> P <sub>WITH</sub>   | 6.746           | <.001  |
| A <sub>TRG</sub>                     | A <sub>TRG</sub> P <sub>WITH</sub>   | 3.528           | .024   |

### B.3 Tukey’s post hoc tests for MODULE’s ROLE for first experiment

Table 3 shows Tukey’s post hoc tests for MODULE’s ROLE in the first experiment.

Table 3. Tukey's post hoc tests for MODULE'S ROLE in the first experiment. In red, the insignificant p-values and in orange the borderline p-values (between 0.04 and 0.06).

| Comparison      |                   | t          | $p_{tukey}$ |
|-----------------|-------------------|------------|-------------|
| MODULE'S ROLE 1 | MODULE'S ROLE 2   | (ddl = 34) |             |
| <i>Swarm</i>    | <i>Proxy</i>      | -7.29      | <.001       |
|                 | <i>Satellites</i> | 5.73       | <.001       |
| <i>Proxy</i>    | <i>Satellites</i> | 7.52       | <.001       |

#### B.4 Tukey's post hoc tests for AUTONOMY $\times$ MODULE'S ROLE for first experiment

Table 4 shows Tukey's post hoc tests for TYPES OF ASSISTANCE  $\times$  MODULE'S ROLE in the first experiment.

Table 4. Tukey's post hoc tests for AUTONOMY  $\times$  MODULE'S ROLE in the first experiment. In bold, the comparison of MODULE'S ROLE 1 with other MODULE'S ROLE for AUTONOMY 1. In red, the insignificant p-values and in orange the borderline p-values (between 0.04 and 0.06). The autonomy correspond to TYPES OF ASSISTANCE WITH PROXY ( $P_{WITH}$ ).

| Comparison        |                          |                   |                          | t<br>(ddl = 34) | ptukey          |
|-------------------|--------------------------|-------------------|--------------------------|-----------------|-----------------|
| MODULE'S ROLE 1   | AUTONOMY 1               | MODULE'S ROLE 2   | AUTONOMY 2               |                 |                 |
| <b>Swarm</b>      | <b>A<sub>INERT</sub></b> | <i>Swarm</i>      | A <sub>SNAP</sub>        | 5.946           | <.001           |
|                   |                          | <i>Swarm</i>      | A <sub>STOP</sub>        | 7.340           | <.001           |
|                   |                          | <i>Swarm</i>      | A <sub>TRG</sub>         | 9.385           | <.001           |
|                   |                          | <b>Proxy</b>      | <b>A<sub>INERT</sub></b> | <b>-6.396</b>   | <b>&lt;.001</b> |
|                   |                          | <b>Satellites</b> | <b>A<sub>INERT</sub></b> | <b>6.995</b>    | <b>&lt;.001</b> |
|                   | <b>A<sub>SNAP</sub></b>  | <i>Swarm</i>      | A <sub>STOP</sub>        | 3.904           | 0.018           |
|                   |                          | <i>Swarm</i>      | A <sub>TRG</sub>         | 7.678           | <.001           |
|                   |                          | <b>Proxy</b>      | <b>A<sub>SNAP</sub></b>  | <b>-6.112</b>   | <b>&lt;.001</b> |
|                   |                          | <b>Satellites</b> | <b>A<sub>SNAP</sub></b>  | <b>6.333</b>    | <b>&lt;.001</b> |
|                   | <b>A<sub>STOP</sub></b>  | <i>Swarm</i>      | A <sub>TRG</sub>         | 6.746           | <.001           |
|                   |                          | <b>Proxy</b>      | <b>A<sub>STOP</sub></b>  | <b>-5.843</b>   | <b>&lt;.001</b> |
|                   |                          | <b>Satellites</b> | <b>A<sub>STOP</sub></b>  | <b>5.336</b>    | <b>&lt;.001</b> |
|                   | <b>A<sub>TRG</sub></b>   | <b>Proxy</b>      | <b>A<sub>TRG</sub></b>   | <b>-4.570</b>   | <b>0.003</b>    |
|                   |                          | <b>Satellites</b> | <b>A<sub>TRG</sub></b>   | <b>3.240</b>    | <b>0.092</b>    |
| <b>Proxy</b>      | <b>A<sub>INERT</sub></b> | <i>Proxy</i>      | A <sub>SNAP</sub>        | 6.493           | <.001           |
|                   |                          | <i>Proxy</i>      | A <sub>STOP</sub>        | 7.522           | <.001           |
|                   |                          | <i>Proxy</i>      | A <sub>TRG</sub>         | 10.024          | <.001           |
|                   |                          | <b>Satellites</b> | <b>A<sub>INERT</sub></b> | <b>8.670</b>    | <b>&lt;.001</b> |
|                   | <b>A<sub>SNAP</sub></b>  | <i>Proxy</i>      | A <sub>STOP</sub>        | 3.786           | 0.025           |
|                   |                          | <i>Proxy</i>      | A <sub>TRG</sub>         | 7.987           | <.001           |
|                   |                          | <b>Satellites</b> | <b>A<sub>SNAP</sub></b>  | <b>7.720</b>    | <b>&lt;.001</b> |
|                   | <b>A<sub>STOP</sub></b>  | <i>Proxy</i>      | A <sub>TRG</sub>         | 7.191           | <.001           |
|                   |                          | <b>Satellites</b> | <b>A<sub>STOP</sub></b>  | <b>6.420</b>    | <b>&lt;.001</b> |
|                   | <b>A<sub>TRG</sub></b>   | <b>Satellites</b> | <b>A<sub>TRG</sub></b>   | <b>4.132</b>    | <b>0.010</b>    |
| <b>Satellites</b> | <b>A<sub>INERT</sub></b> | <i>Satellites</i> | A <sub>SNAP</sub>        | 4.773           | 0.002           |
|                   |                          | <i>Satellites</i> | A <sub>STOP</sub>        | 5.195           | <.001           |
|                   |                          | <i>Satellites</i> | A <sub>TRG</sub>         | 5.595           | <.001           |
|                   | <b>A<sub>SNAP</sub></b>  | <i>Satellites</i> | A <sub>STOP</sub>        | 2.862           | <b>0.200</b>    |
|                   |                          | <i>Satellites</i> | A <sub>TRG</sub>         | 4.546           | 0.003           |
|                   | <b>A<sub>STOP</sub></b>  | <i>Satellites</i> | A <sub>TRG</sub>         | 3.812           | 0.023           |

### C Tables of means and confidence intervals for the second experiment

#### C.1 Means and 95% *bootstrapped* confidence intervals TYPE OF ASSISTANCE $\times$ DIFFICULTY

| TYPE OF ASSISTANCE               |                                                    | SoA Mean<br>[CI]     |                      |
|----------------------------------|----------------------------------------------------|----------------------|----------------------|
|                                  |                                                    | DIFFICULTY           |                      |
|                                  |                                                    | Easier               | Harder               |
| AUTONOMY<br>(P <sub>WITH</sub> ) | A <sub>INERT</sub><br><i>Not autonomous</i>        | 6.96<br>[6.67, 7.25] | 6.96<br>[6.68, 7.22] |
|                                  | A <sub>SNAP</sub><br><i>Semi-autonomous</i>        | 5.50<br>[5.11, 5.87] | 5.62<br>[5.26, 5.97] |
|                                  | A <sub>STOP</sub><br><i>Semi-autonomous</i>        | 5.12<br>[4.73, 5.54] | 5.15<br>[4.77, 5.57] |
|                                  | A <sub>TRG</sub><br><i>Highly autonomous</i>       | 3.78<br>[3.22, 4.38] | 3.73<br>[3.19, 4.30] |
| BASELINE                         |                                                    |                      |                      |
| AUTONOMY<br>(P <sub>OBS</sub> )  | A <sub>CMPLT</sub><br><i>Completely autonomous</i> | 1.89<br>[1.40, 2.48] | 1.85<br>[1.36, 2.43] |

#### C.2 Means and 95% *bootstrapped* confidence intervals for DIFFICULTY $\times$ MODULE'S ROLE (Hyp 2.3)

|            |        | SoA Mean<br>[CI]     |                      |                      |
|------------|--------|----------------------|----------------------|----------------------|
|            |        | MODULE'S ROLE        |                      |                      |
|            |        | PROXY                | SATELLITES           | SWARM                |
| DIFFICULTY | Easier | 5.55<br>[5.26, 5.85] | 4.06<br>[3.76, 4.36] | 4.65<br>[5.06, 5.63] |
|            | Harder | 5.57<br>[5.28, 5.86] | 4.07<br>[3.77, 4.38] | 4.66<br>[5.09, 5.65] |

**C.3 Means and 95% *bootstrapped* confidence intervals for AUTONOMY  $\times$  MODULE'S ROLE**

| TYPE OF ASSISTANCE                                 |                    | SoA Mean<br>[CI]     |                      |                      |
|----------------------------------------------------|--------------------|----------------------|----------------------|----------------------|
|                                                    |                    | MODULE'S ROLE        |                      |                      |
|                                                    |                    | PROXY                | SATELLITES           | SWARM                |
| <b>AUTONOMY</b><br><br>( <b>P<sub>WITH</sub></b> ) | $A_{\text{INERT}}$ | 7.34<br>[7.15, 7.53] | 5.45<br>[5.05, 5.83] | 6.96<br>[6.75, 7.16] |
|                                                    | $A_{\text{SNAP}}$  | 5.79<br>[5.51, 6.06] | 4.24<br>[3.90, 4.56] | 5.56<br>[5.29, 5.82] |
|                                                    | $A_{\text{STOP}}$  | 5.30<br>[5.03, 5.58] | 3.85<br>[3.49, 4.22] | 5.14<br>[4.86, 5.42] |
|                                                    | $A_{\text{TRG}}$   | 3.81<br>[3.43, 4.21] | 2.72<br>[2.49, 3.90] | 3.75<br>[3.36, 4.16] |

**C.4 Means and 95% *bootstrapped* confidence intervals for AUTONOMY  $\times$  DIFFICULTY  $\times$  MODULE'S ROLE**

| TYPE OF ASSISTANCE                                 |                    |               | SoA Mean<br>[CI]     |                      |                      |
|----------------------------------------------------|--------------------|---------------|----------------------|----------------------|----------------------|
|                                                    |                    |               | MODULE'S ROLE        |                      |                      |
|                                                    |                    |               | PROXY                | SATELLITES           | SWARM                |
| <b>AUTONOMY</b><br><br>( <b>P<sub>WITH</sub></b> ) | $A_{\text{INERT}}$ | <i>Easier</i> | 7.39<br>[7.11, 7.63] | 5.43<br>[4.85, 5.96] | 6.96<br>[6.67, 7.25] |
|                                                    |                    | <i>Harder</i> | 7.30<br>[7.10, 7.56] | 5.47<br>[4.91, 6]    | 6.96<br>[6.68, 7.22] |
|                                                    | $A_{\text{SNAP}}$  | <i>Easier</i> | 5.74<br>[5.33, 6.13] | 4.26<br>[3.81, 4.71] | 5.50<br>[5.11, 5.87] |
|                                                    |                    | <i>Harder</i> | 5.84<br>[5.46, 6.20] | 4.21<br>[3.73, 4.67] | 5.62<br>[5.26, 5.97] |
|                                                    | $A_{\text{STOP}}$  | <i>Easier</i> | 5.26<br>[4.89, 5.66] | 3.79<br>[3.26, 4.31] | 5.12<br>[4.73, 5.54] |
|                                                    |                    | <i>Harder</i> | 5.34<br>[4.96, 5.73] | 3.91<br>[3.40, 4.41] | 5.15<br>[4.77, 5.57] |
|                                                    | $A_{\text{TRG}}$   | <i>Easier</i> | 3.83<br>[3.30, 4.40] | 2.75<br>[2.27, 3.28] | 3.78<br>[3.22, 4.38] |
|                                                    |                    | <i>Harder</i> | 3.80<br>[3.28, 4.36] | 2.69<br>[2.20, 3.23] | 3.73<br>[3.19, 4.30] |

**D Tables of means and confidence intervals for the third experiment****D.1 Means and 95% *bootstrapped* confidence intervals TYPE OF ASSISTANCE  $\times$  PREDICTABILITY**

| TYPE OF ASSISTANCE                                 |                                                    | SoA Mean<br>[CI]     |                      |
|----------------------------------------------------|----------------------------------------------------|----------------------|----------------------|
|                                                    |                                                    | PREDICTABILITY       |                      |
|                                                    |                                                    | <i>Straight</i>      | <i>Noisy</i>         |
| <b>AUTONOMY</b><br><br>( <b>P<sub>WITH</sub></b> ) | $A_{\text{INERT}}$<br><i>Not autonomous</i>        | 6.30<br>[5.77, 6.79] | 5.69<br>[5.22, 6.16] |
|                                                    | $A_{\text{SNAP}}$<br><i>Semi-autonomous</i>        | 5.73<br>[5.13, 6.29] | 5.28<br>[4.79, 5.78] |
|                                                    | $A_{\text{STOP}}$<br><i>Semi-autonomous</i>        | 4.90<br>[4.32, 5.47] | 4.71<br>[3.91, 5.03] |
|                                                    | $A_{\text{TRG}}$<br><i>Highly autonomous</i>       | 4.31<br>[3.57, 5.08] | 3.94<br>[3.31, 4.59] |
| <b>BASELINE</b>                                    |                                                    |                      |                      |
| <b>AUTONOMY</b><br>( <b>P<sub>OBS</sub></b> )      | $A_{\text{CMPLT}}$<br><i>Completely autonomous</i> | 2.72<br>[1.92, 3.61] | 2.80<br>[2.11, 3.55] |

**D.2 Means and 95% *bootstrapped* confidence intervals for AUTONOMY  $\times$  MODULE'S ROLE**

| TYPE OF ASSISTANCE                                 |                    | SoA Mean<br>[CI]     |                      |                      |
|----------------------------------------------------|--------------------|----------------------|----------------------|----------------------|
|                                                    |                    | MODULE'S ROLE        |                      |                      |
|                                                    |                    | PROXY                | SATELLITES           | SWARM                |
| <b>AUTONOMY</b><br><br>( <b>P<sub>WITH</sub></b> ) | $A_{\text{INERT}}$ | 6.57<br>[6.22, 6.89] | 4.98<br>[4.44, 5.51] | 6<br>[5.64, 6.35]    |
|                                                    | $A_{\text{SNAP}}$  | 5.84<br>[5.47, 6.20] | 4.60<br>[4.06, 5.15] | 5.51<br>[5.12, 5.89] |
|                                                    | $A_{\text{STOP}}$  | 5.06<br>[4.68, 5.45] | 3.98<br>[3.52, 4.44] | 4.69<br>[4.29, 5.09] |
|                                                    | $A_{\text{TRG}}$   | 4.43<br>[3.93, 4.93] | 3.81<br>[3.25, 4.38] | 4.13<br>[3.64, 4.62] |

### D.3 Means and 95% *bootstrapped* confidence intervals for AUTONOMY $\times$ PREDICTABILITY $\times$ MODULE'S ROLE

|                                |             | SoA Mean<br>[CI] |                      |                      |                      |
|--------------------------------|-------------|------------------|----------------------|----------------------|----------------------|
| TYPE OF ASSISTANCE             |             | PREDICTABILITY   | MODULE'S ROLE        |                      |                      |
|                                |             |                  | PROXY                | SATELLITES           | SWARM                |
| AUTONOMY<br><br>( $P_{WITH}$ ) | $A_{INERT}$ | <i>Straight</i>  | 6.95<br>[6.46, 7.35] | 4.75<br>[4.39, 5.97] | 6.30<br>[5.77, 6.79] |
|                                |             | <i>Noisy</i>     | 6.19<br>[5.71, 6.64] | 5.21<br>[4, 5.46]    | 5.69<br>[5.22, 6.16] |
|                                | $A_{SNAP}$  | <i>Straight</i>  | 6.12<br>[5.56, 6.62] | 4.82<br>[4.01, 5.60] | 5.73<br>[5.13, 6.29] |
|                                |             | <i>Noisy</i>     | 5.57<br>[5.06, 6.04] | 4.38<br>[3.62, 5.12] | 5.28<br>[4.79, 5.78] |
|                                | $A_{STOP}$  | <i>Straight</i>  | 5.33<br>[4.76, 5.86] | 4.14<br>[3.46, 4.81] | 4.90<br>[4.32, 5.47] |
|                                |             | <i>Noisy</i>     | 4.80<br>[4.27, 5.31] | 3.81<br>[3.18, 4.45] | 4.47<br>[3.91, 5.03] |
|                                | $A_{TRG}$   | <i>Straight</i>  | 4.62<br>[3.83, 5.39] | 3.97<br>[3.12, 4.85] | 4.31<br>[3.57, 5.08] |
|                                |             | <i>Noisy</i>     | 4.24<br>[3.61, 4.88] | 3.64<br>[2.93, 4.37] | 3.94<br>[3.31, 4.59] |

## E Post hoc tests for the third experiment

### E.1 Tukey's post hoc tests for AUTONOMY $\times$ PREDICTABILITY for third experiment

Table 5 shows Tukey's post hoc tests for AUTONOMY  $\times$  PREDICTABILITY in the third experiment. For more condensed writing, we only report the comparison of *Straight Noisy* PREDICTABILITY for each AUTONOMY.

Table 5. Tukey's post hoc tests for AUTONOMY and PREDICTABILITY in the third experiment. In red, the insignificant p-values and in orange the borderline p-values (between 0.04 and 0.06).

| Tukey test<br>for <i>Straight</i> and <i>Noisy</i> comparison |                 |             |
|---------------------------------------------------------------|-----------------|-------------|
| AUTONOMY                                                      | t<br>(ddl = 34) | $p_{tukey}$ |
| $A_{INERT}$                                                   | -3.606          | .029        |
| $A_{SNAP}$                                                    | -3.635          | .027        |
| $A_{STOP}$                                                    | -3.406          | .047        |
| $A_{TRG}$                                                     | -2.572          | .269        |
| $A_{CMPLT}$                                                   | 0.621           | 1           |



## E.2 Tukey's post hoc tests for PREDICTABILITY $\times$ MODULE'S ROLE for third experiment

Table 6 shows Tukey's post hoc tests for PREDICTABILITY  $\times$  MODULE'S ROLE in the third experiment.

Table 6. Tukey's post hoc tests for PREDICTABILITY  $\times$  MODULE'S ROLE in the first experiment. In bold, the comparison of MODULE'S ROLE 1 with the same MODULE'S ROLE for PREDICTABILITY 1 and PREDICTABILITY 2 that are different. In orange the borderline p-values (between 0.04 and 0.06).

| Comparison      |                  |                 |                  | t<br>(ddl = 34) | ptukey       |
|-----------------|------------------|-----------------|------------------|-----------------|--------------|
| MODULE'S ROLE 1 | PREDICTABILITY 1 | MODULE'S ROLE 2 | PREDICTABILITY 2 |                 |              |
| <b>Swarm</b>    | <b>Noisy</b>     | <b>Swarm</b>    | <b>Straight</b>  | <b>-4.034</b>   | <b>0.004</b> |
|                 |                  | Proxy           | Noisy            | -3.080          | 0.043        |
|                 |                  | Sat             | Noisy            | 3.903           | 0.005        |
|                 | Straight         | Proxy           | Straight         | -4.208          | 0.002        |
|                 |                  | Sat             | Straight         | 4.458           | 0.001        |
| <b>Proxy</b>    | <b>Noisy</b>     | <b>Proxy</b>    | <b>Straight</b>  | <b>-4.373</b>   | <b>0.001</b> |
|                 |                  | Sat             | Noisy            | 5.117           | <.001        |
|                 | Straight         | Sat             | Straight         | 5.866           | <.001        |
| <b>Sat</b>      | <b>Noisy</b>     | <b>Sat</b>      | <b>Straight</b>  | <b>-4.220</b>   | <b>0.002</b> |

## F Ranking of TYPES OF ASSISTANCE and correlations between SoA, preference, perceived control, perceived performance, ease of use and perceived success for the three experiment

### F.1 Ranking and correlations for the first experiment

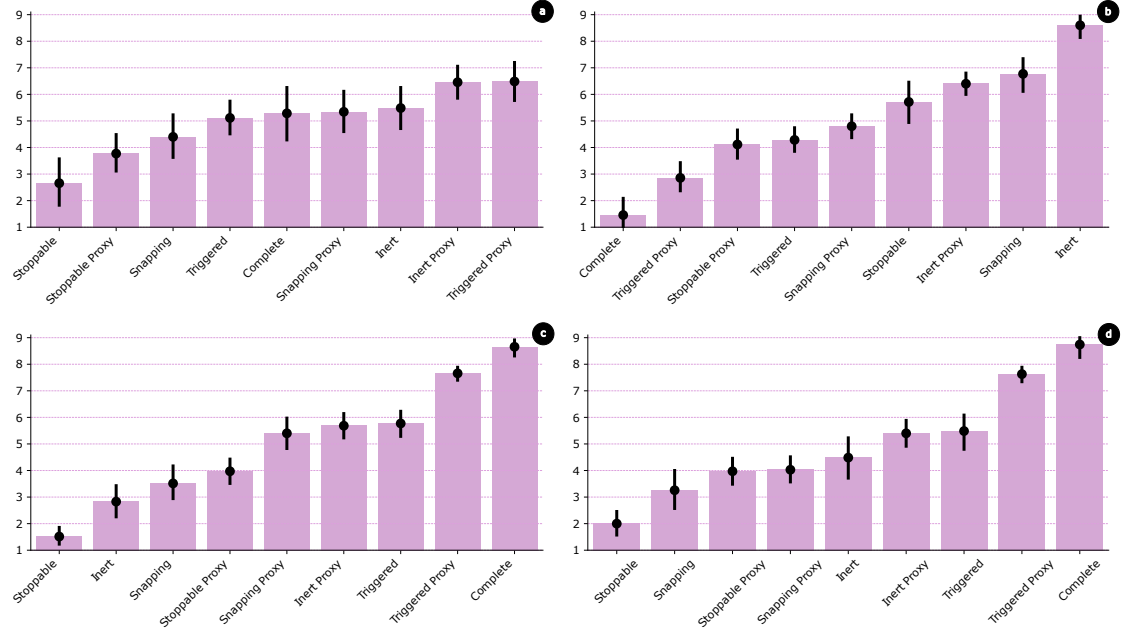


Fig. 12. Rankings of (a) preference, (b) perceived control, (c) perceived performance (speed × accuracy), and (d) perceived ease of use for the first experiment. Error bars represent bootstrapped 95% confidence intervals. For more concise writing, TYPES OF ASSISTANCE  $P_{w/o}$  are named only by the AUTONOMY name, and TYPES OF ASSISTANCE  $P_{with}$  are named as “AUTONOMY name Proxy”.

|                       | SoA       | Preference | Perceived control | Perceived performance | Perceived ease of use | Perceived success |
|-----------------------|-----------|------------|-------------------|-----------------------|-----------------------|-------------------|
| SoA                   | 1         |            |                   |                       |                       |                   |
| Preference            | 0.077     | 1          |                   |                       |                       |                   |
| Perceived control     | 0.560***  | 0.092*     | 1                 |                       |                       |                   |
| Perceived performance | -0.313*** | 0.311***   | -0.383***         | 1                     |                       |                   |
| Perceived ease of use | -0.252*** | 0.318***   | -0.299***         | 0.620***              | 1                     |                   |
| Perceived success     | -0.215*** | 0.165***   | -0.276***         | 0.394***              | 0.326***              | 1                 |

Table 7. Kendall's Tau (correlation) between the score of SoA, preference, sense of control, perceived performance, perceived ease of use and perceived task success for the first experiment (\*  $p < .05$ , \*\*  $p < .01$ , \*\*\*  $p < .001$ ).

## F.2 Ranking and correlations for the second experiment

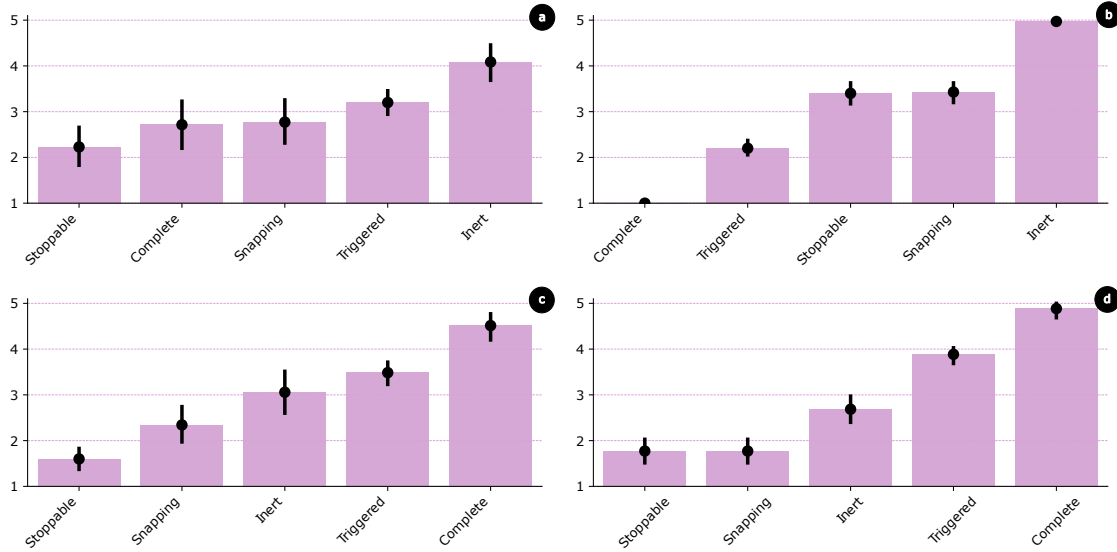


Fig. 13. Rankings of (a) preference, (b) perceived control, (c) perceived performance (speed  $\times$  accuracy), and (d) perceived ease of use for the second experiment. Error bars represent bootstrapped 95% confidence intervals. For more concise writing, TYPES OF ASSISTANCE  $P_{w/o}$  are named only by the AUTONOMY name, and TYPES OF ASSISTANCE  $P_{with}$  are named as “AUTONOMY name Proxy”.

| SoA                   | Preference | Perceived control | Perceived performance | Perceived ease of use | Perceived success |   |
|-----------------------|------------|-------------------|-----------------------|-----------------------|-------------------|---|
| SoA                   | 1          |                   |                       |                       |                   |   |
| Preference            | 0.200**    | 1                 |                       |                       |                   |   |
| Perceived control     | 0.628***   | 0.256***          | 1                     |                       |                   |   |
| Perceived performance | -0.230***  | 0.182***          | -0.334***             | 1                     |                   |   |
| Perceived ease of use | -0.357***  | 0.159**           | -0.467***             | 0.594***              | 1                 |   |
| Perceived success     | -0.019     | 0.135*            | -0.034                | 0.166**               | 0.153**           | 1 |

Table 8. Kendall's Tau (correlation) between the score of SoA, preference, sense of control, perceived performance, perceived ease of use and perceived task success for the second experiment (\*  $p < .05$ , \*\*  $p < .01$ , \*\*\*  $p < .001$ ).

### F.3 Ranking and correlations for the third experiment

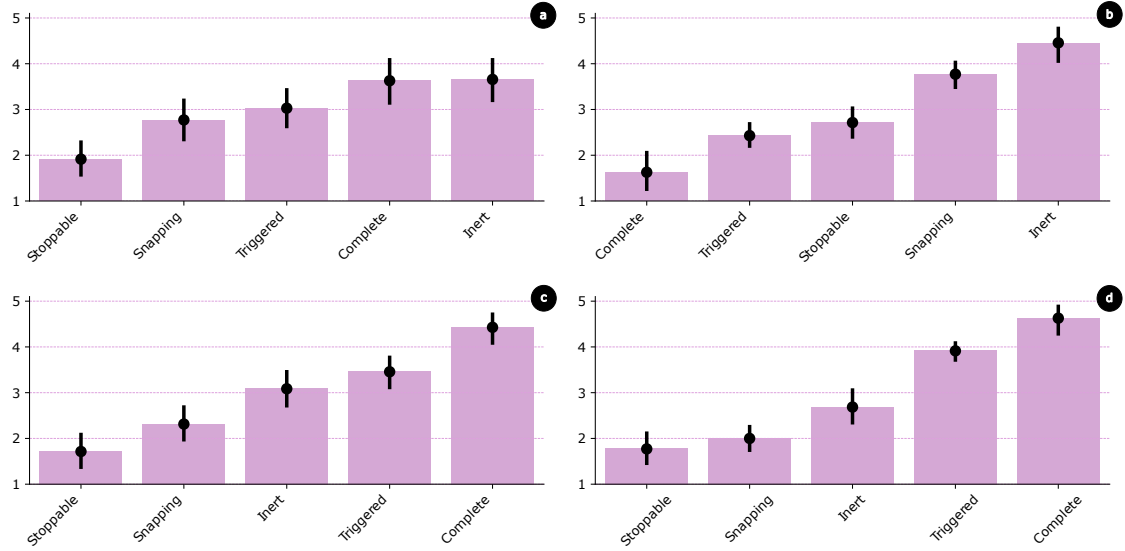


Fig. 14. Rankings of (a) preference, (b) perceived control, (c) perceived performance (speed  $\times$  accuracy), and (d) perceived ease of use for the third experiment. Error bars represent bootstrapped 95% For more concise writing, TYPES OF ASSISTANCE  $P_{w/o}$  are named only by the AUTONOMY name, and TYPES OF ASSISTANCE  $P_{with}$  are named as “AUTONOMY name Proxy”.

| SoA                   | Preference | Perceived control | Perceived performance | Perceived ease of use | Perceived success |
|-----------------------|------------|-------------------|-----------------------|-----------------------|-------------------|
| SoA                   | 1          |                   |                       |                       |                   |
| Preference            | 0.034      | 1                 |                       |                       |                   |
| Perceived control     | 0.346***   | 0.196**           | 1                     |                       |                   |
| Perceived performance | -0.127*    | 0.443***          | -0.164**              | 1                     |                   |
| Perceived ease of use | -0.245***  | 0.387***          | -0.250***             | 0.569***              | 1                 |
| Perceived success     | 0.192***   | 0.155**           | 0.085                 | 0.154**               | 0.111* 1          |

Table 9. Kendall’s Tau (correlation) between the score of SoA, preference, sense of control, perceived performance, perceived ease of use and perceived task success for the third experiment (\*  $p < .05$ , \*\*  $p < .01$ , \*\*\*  $p < .001$ ).