

Effect of Robotic Modular Interface Assistance Type on Sense of Agency

Effet du type d'assistance d'une interface modulaire robotisée sur le sens d'agentivité

Effet du type d'assistance d'une interface modulaire robotisée sur le sens d'agentivité

Ophélie Jobert
Université Grenoble Alpes
CNRS
Grenoble, France
ophelie.jobert@univ-grenoble-alpes.fr

Thibaut Leone
Université Grenoble Alpes
Grenoble, France
leone.thibaut@yahoo.com

Alix Goguey
LIG, Grenoble Computer Science
Laboratory
Université Grenoble Alpes
Grenoble, France
alix.goguey@univ-grenoble-alpes.fr

Bruno Berberian
ONERA
Salon-de-Provence, France
bruno.berberian@onera.fr

Julien Bourgeois
Univ. Bourgogne Franche-Comté,
FEMTO-ST Institute, CNRS
Montbéliard, France
julien.bourgeois@femto-st.fr

Céline Coutrix
Université Grenoble Alpes
CNRS
Grenoble, France
celine.coutrix@imag.fr

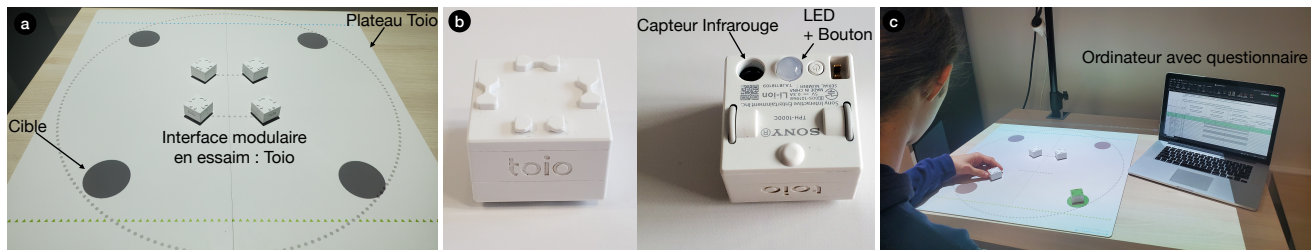


Figure 1: Présentation du matériel et de la configuration : (a) Essaim composé de 4 modules Toios™ dans leur position de départ avec, devant chacun, sa cible à atteindre, (b) Toio™ vu de dessus et de dessous avec ses capteurs, (c) Participant-e faisant glisser un module, avec, à droite, le questionnaire d'expérience à remplir à la fin de chaque essai.

Abstract

Robotic modular interfaces, increasingly studied in Human-Computer Interaction, offer assistance to support users in their tasks. However, this assistance can harm the sense of agency (i.e., feeling of control), leading to non-use of the interface or a diminishing sense of responsibility regarding the consequences of users' actions. The impact of robotic modular interface assistance, during a cooperative task, on the user's sense of agency has not yet been studied. In this article, we propose to remedy this through the use of swarm robotic interfaces. We focus on nine levels of assistance, varying system autonomy and module coordination. Our study

shows that: (1) the higher the autonomy, the lower the sense of agency, (2) coordination seems to have an impact on the sense of agency, and (3) three types of sense of agency emerge depending on the coordination of the modules.

Résumé

Les interfaces modulaires robotisées, de plus en plus étudiées en Interaction Humain-Machine, proposent de l'assistance pour soutenir les utilisatrices dans leurs tâches. Cependant, cette assistance peut nuire au sens d'agentivité (c.-à-d. sentiment de contrôle), amenant à la non-utilisation de l'interface ou à la déresponsabilisation de l'utilisatrice face aux conséquences de ses actes. L'impact de l'assistance des interfaces modulaires robotisées, lors d'une tâche de coopération, sur le sens d'agentivité de l'utilisatrice n'a pas encore été étudiée. Dans cet article, nous proposons d'y remédier via l'utilisation d'interfaces modulaires en essaim. Nous nous intéressons à neuf niveaux d'assistance, variant autonomie du système et coordination des modules. Notre étude montre que : (1) plus



This work is licensed under a Creative Commons Attribution 4.0 International License.
IHM '25, Toulouse, France

© 2025 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-1522-8/25/11
<https://doi.org/10.1145/3765712.3765725>

l'autonomie est élevée plus le sens d'agentivité sera faible, (2) la coordination semble avoir un impact sur le sens d'agentivité et, (3) trois types de sens d'agentivité émergent en fonction de la coordination des modules.

CCS Concepts

• **Human-centered computing** → **Human computer interaction (HCI)**.

Keywords

Modular interfaces, Sense of Agency, Autonomy, Assistance, Proxy

Mots clés

Interfaces modulaires, Sens Agentivité, Autonomie, Assistance, Module intermédiaire

ACM Reference Format:

Ophélie Jobert, Thibaut Leone, Alix Goguet, Bruno Berberian, Julien Bourgeois, and Céline Coutrix. 2025. Effect of Robotic Modular Interface Assistance Type on Sense of Agency: Effet du type d'assistance d'une interface modulaire robotisée sur le sens d'agentivité. In *The 36th Conference on l'Interaction Humain-Machine (IHM '25), November 03–07, 2025, Toulouse, France*. ACM, New York, NY, USA, 18 pages. <https://doi.org/10.1145/3765712.3765725>

1 Introduction

Les interfaces modulaires à changement de forme sont des interfaces faites de plusieurs modules pouvant se réarranger spatialement, en autonomie ou non [54]. Ces interfaces rendent donc possibles de nouvelles façons d'interagir tant en entrée qu'en sortie. Les travaux portant sur ces interfaces se sont pour l'instant intéressés aux caractéristiques physiques, telles que la taille ou la forme [55, 56], aux techniques d'interaction les utilisant en entrée et/ou en sortie [33–35, 48, 53], ou encore aux scénarios d'application de ces interfaces [14, 26, 41, 68]. Dans cet article, nous nous intéressons plus particulièrement aux interfaces modulaires robotisées.

Un des avantages des interfaces modulaires robotisées est d'octroyer une certaine autogestion aux modules les composant pour assister l'utilisatrice dans l'accomplissement de tâches plus ou moins complexes. Cette autogestion peut passer par exemple par l'autonomie dont disposent les modules à certains moments de la tâche et/ou par la coordination ou non des mouvements du groupe sur le mouvement d'un seul module. Par exemple, l'utilisatrice des Zooids déplace un seul module pendant que les autres modules suivent la même trajectoire [41]. Cependant, une trop grande autogestion dans un système peut engendrer une perte de contrôle des utilisatrices. Comme le soulignent Shneiderman et Plaisant [64], l'acceptabilité d'un système est très largement dépendante de la capacité des utilisateurs à exercer du contrôle sur ce système. En outre, le sentiment de contrôle est également un facteur important dans l'intention comportemental et dans la réalisation du comportement lui-même [1], notamment dans le cas des interfaces [73]. De plus, se sentir moins en contrôle peut diminuer le sentiment de responsabilité des conséquences de ses actes [16]. Or, dans le cas d'utilisation d'un système critique (p. ex. lors d'une chirurgie ou de recherches de victimes), il est important que l'utilisatrice puisse garder un sentiment de contrôle sur la situation

afin de se sentir responsable de ses actions. L'autogestion des interfaces modulaires robotisées doit donc pouvoir assister l'utilisatrice dans sa tâche sans pour autant détériorer son expérience de contrôle. Nous proposons donc d'explorer le développement du sens d'agentivité¹ – c.-à-d. le sentiment de contrôler une action et ses effets dans l'environnement [28] – en fonction du type d'assistance apportée par les modules lors de situations de coopération avec l'utilisatrice. En particulier, nous appelons *type d'assistance* cette proportion d'autogestion de l'interface. Notre question de recherche vise donc à connaître **comment le type d'assistance en situation de coopération impact le sens d'agentivité**. Les interfaces modulaires robotisées étant pour la plupart à un stade précoce de développement, nous nous concentrons sur un type particulier : les interfaces en essaim. Ces interfaces sont caractérisées par des modules pouvant se déplacer par eux mêmes sans besoin d'être reliés aux autres modules de l'interface, qui se déplacent collectivement et réagissent aux commandes de l'utilisatrices [41].

Dans cet article, nous explorons l'impact du type d'assistance de l'essaim sur le développement du sens d'agentivité. Pour cela, nous avons réalisé une expérience où les participant-es amenaient les modules d'un essaim de Toios™ [70] dans leurs cibles (Figures 1). Nous faisons varier le type d'assistance selon différents niveaux : de modules sans aucune autonomie à complètement autonomes. Nous avons mesuré le sentiment de contrôle ressenti par les participant-es sur les mouvements des modules pour chaque type d'assistance.

Le défi majeur de cette expérience réside dans l'identification et l'extraction des types d'assistances dans le cas des interfaces modulaires robotisées. Nous avons identifié neuf types d'assistance (huit niveaux de type d'assistance et une baseline) en nous basant sur la littérature [33, 41, 48, 53, 68]. Ces niveaux résultent de deux facteurs : (1) l'autonomie dont disposent les modules à certains moments de la tâche (c.-à-d. le type de contrôle que peut avoir l'utilisatrice sur différentes phases du mouvement des modules), et (2) la coordination ou non des mouvements de l'essaim sur le mouvement d'un seul module (c.-à-d. l'utilisation d'un module comme intermédiaire ou non). Nos résultats suggèrent un effet de l'autonomie des modules, et un effet de l'utilisation d'un module intermédiaire sur le sens d'agentivité, sans permettre de conclure sur l'interaction entre type d'autonomie et utilisation d'un intermédiaire sur le sens d'agentivité.

Cet article présente les contributions suivantes :

- (1) L'identification et la définition de neuf types d'assistance (huit types d'assistance et une baseline) qu'une interface modulaire robotisée peut apporter à l'utilisatrice,
- (2) La mesure de l'impact de ces neuf types d'assistance sur le sens d'agentivité,
- (3) Pour les types d'assistances avec un module intermédiaire, trois types de sens d'agentivité semblent entrer en jeu : (a) un sens d'agentivité pour le module intermédiaire, (b) un pour les autres modules de l'essaim, et (c) un pour l'essaim entier.

2 Travaux Précédents

Nous positionnons notre travail à l'intersection entre les travaux sur les interfaces modulaires, notamment robotisées, et ceux sur le

¹traduction du concept anglais de *Sense of Agency* (SoA)

sens d'agentivité, notamment avec des agents externes à l'individu lors de la réalisation d'une tâche.

2.1 Les interfaces modulaires (robotisées)

Les interfaces à changement de forme, dont les interfaces modulaires, ont la capacité de changer de forme physique, en entrée et/ou en sortie du système, à l'initiative de l'utilisatrice ou du système si elles sont robotisées. Les interfaces modulaires se caractérisent par le fait qu'elles sont constituées de plusieurs modules pouvant se réarranger dans l'espace [54]. Les travaux antérieurs en IHM sur les interfaces modulaires se sont concentrés sur l'identification de scénarios d'utilisation [13, 14, 26, 41, 43, 68], sur les caractéristiques physiques des modules [55, 56], ou encore sur les techniques d'interaction en entrée et en sortie [22, 33–35, 48, 53].

Dans cet article, nous nous intéressons aux interfaces modulaires *robotisées*. La miniaturisation des interfaces modulaires robotisées est un défi technologique majeur [2]. Pourtant, il en existe un type assez avancé pour permettre des tests utilisateurs : les interfaces en essaim qui sont caractérisées par des modules pouvant se déplacer par eux-mêmes sans besoin d'être reliés aux autres, qui se déplacent collectivement, et réagissent aux commandes de l'utilisatrice [41]. La robotisation du changement de forme de ces interfaces a permis aux chercheuses et chercheurs d'explorer des techniques d'interaction pour les contrôler [33, 48, 53], la perception par les utilisatrices de leur vitesse, leur type de déplacement et de retour haptique [22, 34, 35], et des cas d'utilisation [13, 43, 68].

Les études présentées ci-dessus se concentrent sur l'identification d'interactions efficaces, mais laissent de côté leur perception. Or, la perception (p. ex. émotionnelle ou de contrôle) suscitée par une interface peut avoir un impact sur son adoption et son utilisation [37, 49]. Des études se sont intéressées à la perception émotionnelle qu'ont les utilisatrices de ces interfaces en fonction de la vitesse, la fluidité, la synchronisation, ou la force de mouvement des modules [23, 34, 35, 38]. Kim et Follmer [35] se sont intéressés à la perception lors de contacts sociaux par le toucher à distance via ces interfaces. Les participant·es devaient contrôler l'essaim par des mouvements des doigts de leur main dominante sur un écran tactile afin de créer un motif haptique pour chaque contact social. Après chaque essai, ils devaient évaluer la facilité de contrôle. Ici, la facilité du contrôle était donc évaluée pour un contrôle par écran tactile, donc indirect et externe à l'interface en essaim. Or, pouvoir s'affranchir d'une interface tierce permettrait de nombreux cas d'utilisation supplémentaires, grâce à une interaction directe avec les modules robotisés. Même si ces travaux ont étudié la perception affective et subjective des interfaces en essaim, nous n'avons trouvé aucune étude sur le sens d'agentivité, bien qu'il contribue également à l'adoption et l'utilisation de la technologie [44]. De plus, l'un des avantages des interfaces modulaires robotisées est de pouvoir s'appuyer sur l'autonomie totale ou partielle de l'interface (c.-à-d. différents niveaux d'assistance) pour aider l'utilisatrice dans sa tâche. Cette autonomie peut remettre en cause le sens d'agentivité des utilisatrices. Dans ce contexte, il nous semble critique d'explorer le lien entre le sentiment de contrôle de l'utilisatrice et le niveau d'assistance des interfaces modulaires robotisées.

2.2 Le sens d'agentivité

Le sens d'agentivité se définit comme la perception de contrôle de ses propres actions et, au travers de celles-ci, des événements du monde extérieur [28]. Il est considéré comme un facteur clé dans l'adoption et l'utilisation des interfaces. En effet, une interface doit permettre de maintenir chez l'utilisatrice un sentiment de contrôle interne [64]. C'est-à-dire que l'utilisatrice doit avoir le sentiment que ses actions sont à l'origine des réponses du système. En outre, le sens d'agentivité constitue un élément clé pour l'attribution de la responsabilité causale et peut impacter de manière drastique l'engagement des utilisatrices dans leur tâche [6].

Le sens d'agentivité peut être mesuré de façon implicite ou explicite. Parmi les mesures implicites, le liage intentionnel (*Intentional Binding*) constitue l'un des principaux marqueurs du sens d'agentivité. Le liage intentionnel se traduit par le fait que les actions volontaires et leurs effets sensoriels sont perçus comme étant temporellement plus proches l'un de l'autre, par rapport aux effets d'actions involontaires ou des réflexes [29]. Bien que parfois remise en cause dans la littérature [36, 66, 69], cette compression temporelle reste l'un des marqueurs du sens d'agentivité les plus robustes.

La mesure explicite du sens d'agentivité se fait en demandant à l'utilisatrice de quantifier via une échelle de Likert à quel point elle s'est sentie à l'origine (ou en contrôle) des événements, p. ex. comme dans [12, 18, 19]. Dans notre étude, nous utilisons la mesure explicite. Bien que pouvant être sensible aux attentes des participant·es, les jugements d'agentivité constituent une mesure plus directe de l'expérience de contrôle et ont l'avantage d'être plus simples à mettre en place que le liage intentionnel.

Le concept du sens d'agentivité a été appliqué au domaine de l'IHM [7–9, 44, 46, 50], p. ex. pour comparer l'effet de nouvelles techniques d'interaction (p. ex. bouton sur la peau vs. bouton physique ou pavé tactile) [10, 11, 20], ou pour évaluer le bénéfice d'un retour vibrotactile [60].

Lors de l'utilisation d'un système autonome dans l'accomplissement d'une tâche, deux notions peuvent avoir un effet sur le sens d'agentivité : (1) l'assistance que peut fournir le système pour effectuer les tâches, et (2) la coopération entre l'utilisatrice et le système. Nous définissons l'**assistance** comme le degré d'autonomie laissé à un système pour aider l'utilisatrice à effectuer sa tâche. Cette assistance peut avoir plusieurs niveaux allant d'une assistance complète (c.-à-d. le système est complètement autonome et l'utilisatrice n'intervient pas ou seulement pour activer le système au début de la tâche) à aucune assistance (c.-à-d. le système est complètement inerte et l'utilisatrice doit faire la tâche du début à la fin) [42, 63]. Nous définissons la **coopération** comme le fait que deux agentes (dans notre cas, une agente humaine et un système autonome) s'allient pour parvenir au même but. La coopération intervient lors d'une tâche où les deux agentes sont actives dans la tâche.

2.3 Les niveaux d'assistance d'une interface

Sheridan et al. [63] proposent dix niveaux d'assistance basés sur l'automatisation des systèmes, allant d'aucune assistance (c.-à-d. sans automatisation –niveau 1) à une assistance complète (c.-à-d. automatisation totale –niveau 10). Les niveaux 2 à 4 définissent

l'équilibre entre les agents (humaine vs. système) pour prendre la décision. Les niveaux 5 à 7 définissent comment le système exécute l'action. Tous ces niveaux ont été établis pour des systèmes téléopérés. Par exemple, Walker et al. [77] ont appliqué deux de ces niveaux d'assistance à des interfaces en essaim contrôlées à travers une interface graphique. Au contraire, nous étudions ici la manipulation directe des modules, comme dans Pickcells [26] par exemple.

Afin d'identifier des niveaux d'assistance pertinents pour les interfaces en essaim manipulées directement, nous sommes partis de la revue systématique de la littérature de Prusko et al. [54] pour extraire nos 9 catégories de niveaux d'assistance. Nous avons ensuite solidifié cette catégorisation avec des travaux de la littérature plus récents ou issu d'un cadre plus large que l'étude Prusko et al. [54] (p. ex. drones ou robots plus grands). Nous avons plus particulièrement examiné les travaux proposant des scénarios d'utilisation et des types d'interaction [14, 20, 26, 31, 32, 34, 35, 41, 48, 57, 67, 68, 78], y compris suggérées par les utilisatrices [33, 53]. Il en ressort qu'une interface en essaim peut :

- (1) Être complètement inerte [26, 33]. L'utilisatrice doit alors bouger elle-même chaque module, p. ex. pour faire son code secret [26].
- (2) Être complètement autonome, soit suite à une activation de l'utilisatrice [33, 34, 53, 67, 68], soit suite à une information reçue du système [14, 35, 68]. P. ex., lors de la lecture d'un livre, appuyer sur un mot permettrait à l'interface modulaire de prendre la forme du mot lu [67, 68].
- (3) Assister l'utilisatrice en laissant le système faire une grande partie de la tâche, mais en laissant à l'utilisatrice la possibilité de stopper l'automatisation quand elle le souhaite [31, 33]. P. ex., lors d'une étude sur des drones pour la recherche de victimes, il a été noté que l'automatisation du système peut être bénéfique mais que les utilisatrices doivent pouvoir intervenir à tout moment [31].
- (4) Assister l'utilisatrice dans l'atteinte de son objectif, par exemple, en améliorant la précision finale [78], en corrigeant sa trajectoire [78] ou encore en l'aidant à atteindre l'objectif plus rapidement en finissant la tâche [20].
- (5) Assister l'utilisatrice en lui permettant de réduire ses efforts physiques ou cognitifs. Pour cela, l'interface modulaire peut être perçue comme un groupe se coordonnant sur les mouvements d'un ou plusieurs modules pour se déplacer. L'utilisatrice interagit directement avec un nombre limité de modules afin d'interagir avec l'essaim entier ou avec des sous-groupes [26, 33, 41, 48, 52, 57]. Par exemple, la personne peut ne déplacer qu'un module comme elle le souhaite et l'essaim bouge en imitant le module dit *intermédiaire* [33, 41, 48, 57].

2.4 Sens d'agentivité et Assistance

Les études antérieures semblent montrer qu'un niveau d'assistance élevé peut diminuer le sens d'agentivité [20, 44, 72, 79, 80]. Coyle et al. [20] se sont intéressés à l'effet de l'assistance sur le sens d'agentivité dans le cas d'une tâche de pointage et de sélection aussi rapide et précise que possible. Ici, l'assistance était modulée via une force attirant le curseur vers la cible. Les résultats semblent

montrer une baisse du sens d'agentivité quand l'assistance augmente. D'autres études se sont intéressées à l'effet de l'assistance, de la charge mentale et du délai sur le liage intentionnel dans le cadre d'une tâche de déclenchement d'un son via un bouton [79, 80]. Les résultats semblent montrer un effet de l'assistance sur le sens d'agentivité, et ce, peu importe le niveau de charge mentale.

Wen et al. [78] ont montré que le sens d'agentivité était plus important quand la performance était meilleure, quel que soit le niveau d'assistance. Cela montre que le niveau d'assistance peut impacter le sens d'agentivité directement ou indirectement en renforçant la performance. Cependant, Ueda et al. [72] montrent cela jusqu'à un point de rupture où l'assistance du système, devenant trop importante, entraîne une baisse du sens d'agentivité même si la performance est supérieure. Cela sous-entend donc qu'il faudrait trouver un équilibre entre automatisation, performance et sens d'agentivité pour optimiser la coopération entre le système et l'utilisatrice [6, 7, 72]. Il est important de noter que ces études ont été menées sur des interfaces graphiques. Dans le cas des interfaces modulaires robotisées, l'introduction de la manipulation directe des modules physiques pourrait remettre en cause les résultats existants sur le sens d'agentivité. En outre, ces études se sont intéressées au système comme un outil et non comme un partenaire coopérant pour l'atteinte d'un but commun.

2.5 Sens d'agentivité et Coopération

L'impact de la coopération sur le sens d'agentivité a été étudié dans le cas d'interactions humaine-humaine [50, 66] et humaine-machine [18, 19, 27, 51, 59, 61]. Lors de situation sociale, un sens d'agentivité particulier peut se développer, dans toutes les situations où les conséquences des actions de la personne sont liés à une partenaire [65]. Le concept de *We-Agency* [50, 51] signifie que le sens d'agentivité individuel se dissout dans un sens d'agentivité commun lors de la réalisation d'action en coopération. Cependant, le concept de *We-Agency* a été remis en question par manque de preuves [40]. Le Besnerais et al. [40] proposent de délaisser le terme de *We-Agency* au profit de *sens d'agentivité joint*. Ce sens d'agentivité joint se manifeste lors d'une action jointe (c.-à-d. fait de poursuivre un but commun), par exemple, par le développement d'une expérience d'agentivité pour les actions de la partenaire (sens d'agentivité vicariant), ou par le développement d'une expérience d'agentivité commune en plus de celle individuelle. Par exemple, des études utilisant le liage intentionnel ont montré qu'une personne a un grand sens d'agentivité quand elle effectue une tâche de coopération, même si l'action est effectuée par une autre personne [50, 66].

Partant de ce constat, Obhi et Hall [51] se sont demandés lequel de ces deux phénomènes était à l'œuvre lors d'une tâche de coopération avec un système informatique [51]. Les participant-es devaient effectuer une action jointe avec une co-agente humaine ou informatique. Les auteurs ont mesuré le sens d'agentivité indirectement et ont demandé aux participant-es d'estimer qui était à l'origine de l'action. Le sens d'agentivité pour l'action de la co-agente semblait exister quand celle-ci était humaine mais pas lorsque c'était l'ordinateur. De façon similaire, Sahāi et al. [62] ont étudié le sens d'agentivité des participant-es lors d'une tâche de *Simon Social*²

²https://en.wikipedia.org/wiki/Simon_effect

pour trois conditions : seul (condition individuelle), en coopération avec un partenaire humain (condition humaine-humaine), ou en coopération avec partenaire artificiel (condition humaine-ordinateur). Les résultats ont montré qu'un sens d'agentivité vicariant semblait exister dans la condition humaine-humaine mais pas dans la condition humaine-ordinateur. Le sens d'agentivité lors de tâches coopératives peut donc être impacté par la nature –humaine ou artificielle– du partenaire.

Néanmoins, dans ces études, les co-agents sont programmés pour effectuer une action. Cette programmation de l'action la rend non intentionnelle pour l'agente humaine, ce qui sera le cas pour tout type de système informatique. Cependant, les interfaces en essaim peuvent bouger physiquement dans l'espace pour exécuter une action, créant un effet sensoriel proche des actions humaines [59].

Des études se sont intéressées au sens d'agentivité lors de la présence d'un co-agent robotique [18, 19, 27, 59, 61]. Dans la plupart des cas, les études semblent montrer un effet négatif du co-agent robotique sur le sens d'agentivité comparé à l'absence de co-agent [18, 19, 27, 61]. Cependant, Roselli et al. [59] ont comparé le sens d'agentivité quand un co-agent robotique faisait physiquement la tâche ou quand celui-ci était physiquement présent mais effectuait la tâche par bluetooth. Les résultats ont montré un sens d'agentivité *vicariant* dans le cas de l'action physique et une baisse du sens d'agentivité quand la tâche était effectuée par bluetooth. De façon similaire, Sahaï et al. [61] ont comparé le sens d'agentivité pour trois types de co-agents : (1) servomoteur, (2) robot humanoïde, et (3) humain. Leur résultats ont montré qu'il semble exister une différence significative entre les trois types de co-agents (humaine > robot humanoïde > servomoteur). Cela laisse donc penser que le type de système autonome impliqué dans la tâche de coopération pourrait avoir un effet plus ou moins important sur le sens d'agentivité.

Cependant, ces études se sont concentrées sur la notion de coopération quand le système robotique et la personne effectuaient la tâche sans communiquer. Leurs actions se répercutaient sur le résultat, mais à aucun moment la personne ne pouvait interagir avec le robot pour effectuer la tâche. Dans le cas des interfaces modulaires robotisées, les travaux en IHM permettent d'interagir avec l'interface dans l'atteinte d'un objectif par la personne (p. ex. [33, 41, 68]). De plus, les études présentées ci-dessus se sont intéressées à des systèmes robotiques utilisant un unique robot, au contraire des interfaces en essaim. À notre connaissance, aucune étude ne s'intéresse à l'impact de TYPES D'ASSISTANCE d'un partenaire artificiel lors d'une tâche jointe (c.-à-d. de coopération) sur le développement du sens d'agentivité de la partenaire humaine. Or, l'essor des interfaces en essaim montre qu'il est temps de se poser cette question. Dans cet article, nous nous intéressons pour la première fois à l'effet de l'automatisation d'un système sur le sens d'agentivité de l'utilisatrice, quand ce système est un partenaire dans l'atteinte d'un but.

3 Méthode

Le but de cette expérience est d'évaluer l'impact du TYPE D'ASSISTANCE d'un essaim de robots, sur le sentiment de contrôle des utilisatrices (c.-à-d. le sens d'agentivité). Pour cela, nos participant-es ont effectué une tâche de glisser-déposer. Cette tâche

consistait à faire glisser quatre robots rectangulaires (3 cm × 3 cm, haut de 2 cm), initialement placés en cercle au centre d'un plateau (Figure 1), pour les déposer dans une cible circulaire (Figure 2) de 6cm de diamètre positionnée face à chaque robot à 15,5 cm de distance (centre à centre).

3.1 Participant-es

Nous avons recruté 35 participant-es par e-mail (18 s'identifiant comme femme, 17 comme homme, $M = 29.26$, $SD = 7.57$). Dix-huit travaillent dans l'informatique, 17 dans d'autres domaines. Leur participation était rétribuée 15 €.

3.2 Matériel

Nous avons contrôlé par Bluetooth les robots Toio™ de Sony [70] via le *Toio™ SDK* pour Unity [47] 2021.3.0f1 sur un MacBook Pro M1 Max. Nous fournissons le code de notre expérience dans les fichiers additionnels. Les robots, et le plateau permettant leur déplacement, ont été placés sur une table (Figure 1), au dessus de laquelle nous avons placé un vidéo-projecteur pour projeter les cibles.

3.3 Plan Expérimental

Nous utilisons un plan expérimental intra-sujets, avec deux variables indépendantes et une dépendante principale.

3.3.1 Variables Indépendantes.

Notre expérience comporte deux variables indépendantes contrôlant la façon d'interagir avec les Toios™ : AUTONOMIE et INTERMÉDIAIRE. Nous appelons un TYPE D'ASSISTANCE une combinaison AUTONOMIE × INTERMÉDIAIRE (Figure 2).

Nous ajoutons également un TYPE D'ASSISTANCE considéré comme notre BASELINE. Cette BASELINE représente un niveau d'assistance complet : la personne n'interagit pas avec les robots. Dans les résultats, nous la considérerons donc aussi bien comme un niveau de TYPE D'ASSISTANCE distinct, que comme un niveau d'AUTONOMIE distinct et un niveau INTERMÉDIAIRE distinct.

AUTONOMIE. Les participant-es doivent glisser-déposer les robots dans leurs cibles, de cinq manières différentes (lignes de la Figure 2) :

INERTE Les robots sont complètement inertes (c.-à-d. il n'y a aucune autonomie de la part des robots). Les participant-es font glisser chaque robot de sa position initiale jusqu'à sa cible. Les participant-es contrôlent le déclenchement, la fin du mouvement et le trajet.

AIMANTÉE Les robots sont semi-autonomes. Les participant-es font glisser chaque robot de sa position initiale jusqu'à proximité de sa cible. Lorsque le robot est suffisamment proche de sa cible (c.-à-d. à moins de 9,5 cm du centre de leur cible), il se déplace de manière autonome pour terminer son mouvement (c.-à-d. il va de lui-même dans la cible). Les participant-es contrôlent le déclenchement et la majorité du trajet.

ARRÊTABLE Les robots sont semi-autonomes. Les participant-es initient le mouvement de chaque robot en cliquant dessus. Puis, de façon autonome le robot va vers sa cible. Les participant-es doivent alors arrêter le robot en l'immobilisant pendant 150 ms, quand elles estiment que celui-ci est dans

sa cible. Les participant-es contrôlent le déclenchement et la fin du mouvement.

DÉCLENCHÉE Les robots sont fortement autonomes. Les participant-es initient le mouvement de chaque robot en cliquant dessus. Le robot va de façon autonome dans sa cible. Les participant-es contrôlent le déclenchement, mais ne contrôlent ni le trajet, ni la fin du mouvement.

COMPLÈTE Ce niveau est la BASELINE du facteur AUTONOMIE. Les robots sont complètement autonomes. Les participant-es n'interagissent pas avec les robots. Tous en même temps, les robots vont de façon autonome dans leur cible. Les participant-es ne contrôlent ni le déclenchement, ni le trajet, ni la fin du mouvement.

INTERMÉDIAIRE. Les robots imitent ou non les mouvement d'un leader, nommé *INTERMÉDIAIRE* (colonnes de la Figure 2) :

SANS INTERMÉDIAIRE Les participant-es interagissent avec chaque robot pour accomplir la tâche. Tous les robots se comportent comme décrit ci-dessus selon le niveau d'AUTONOMIE.

AVEC INTERMÉDIAIRE Les participant-es interagissent *uniquement avec le robot intermédiaire* afin de tous les glisser-déposer dans leurs cibles. Seul le robot *intermédiaire* se comporte comme décrit ci-dessus dans les niveaux d'AUTONOMIE. Les autres robots, appelés satellites [26], imitent les mouvements de l'*intermédiaire* pour atteindre leur cible. Les participant-es définissent le robot *intermédiaire* en cliquant dessus au début de la tâche.

OBSERVATION Ce niveau est la BASELINE du facteur INTERMÉDIAIRE. Les participant-es n'interagissent pas avec les robots. Les participant-es doivent seulement observer les robots. Tous en même temps, les robots vont de façon autonome dans leur cible.

3.3.2 Variable Dépendante.

Notre variable dépendante est le SENS D'AGENTIVITÉ (SoA) sur le mouvement des modules³. Nous mesurons le SoA global pour chaque TYPE D'ASSISTANCE par la question "Quel était votre degré de contrôle sur les mouvements des robots ?" utilisant une échelle de Likert en 8 points [18, 19]. Nous mesurons également le SoA sur l'intermédiaire et les satellites pour les TYPES D'ASSISTANCE AVEC INTERMÉDIAIRE par les questions "Quel était votre degré de contrôle sur le mouvement de l'intermédiaire ?" et "Quel était votre degré de contrôle sur les mouvements des autres robots (satellites) ?" en utilisant la même échelle. Nous choisissons un protocole avec une mesure directe du sens d'agentivité. Par sa subjectivité cette mesure peut donc avoir des limitations, mais elle permet d'être rapide à mettre en place contrairement aux techniques indirectes, comme le liage intentionnel, et est la mesure la plus directe du sens d'agentivité.

En plus de notre variable dépendante principale, nous avons demandé aux participant-es, à la fin de l'expérience, de classer les neuf TYPES D'ASSISTANCE en fonction de : (1) leur préférence, (2) leur sentiment de contrôle (c.-à-d. sens d'agentivité), (3) la performance (rapidité et précision) perçue, et (4) la facilité d'utilisation.

³Pour une écriture plus condensée, nous omettons la mention "sur le mouvement des modules" dans la suite du papier.

L'échelle allait de (1) la préférée / contrôle total / très performante / très facile, à (9) la moins aimée / pas de contrôle / pas performante / très difficile. Leurs explications à voix haute étaient enregistrées par dictaphone.

Une question concernant le ressenti quant à la réussite de la tâche était posée pour vérifier le bon déroulement de l'essai : "Quel est votre sentiment quant à la réussite de la tâche ?".

3.4 Question de Recherche et Hypothèses

Notre question de recherche est la suivante :

RQ Dans quelle mesure chaque TYPE D'ASSISTANCE a-t-il un impact sur le SENS D'AGENTIVITÉ ?

Nous émettons les hypothèses suivantes :

Hyp1 Plus le niveau d'AUTONOMIE est élevé (c.-à-d. complètement autonome) plus le SoA est faible.

Hyp1_{bis} Pour les niveaux semi-autonomes (c.-à-d. AUTONOMIE AIMANTÉE et AUTONOMIE ARRÊTABLE), avoir le contrôle sur une grande partie du trajet plutôt qu'uniquement sur la fin du trajet renforce le SoA.

Hyp2 Le SoA AVEC INTERMÉDIAIRE est plus faible que SANS INTERMÉDIAIRE.

Hyp3 Le SoA est plus élevé pour le module utilisé comme INTERMÉDIAIRE que pour les modules satellites.

3.5 Procédure

Les participant-es sont d'abord invité-es à lire et signer le formulaire de consentement. Pendant toute la durée de l'expérience, les participant-es sont assis-es à une table sur laquelle est placé le plateau des ToiosTM (Figure 1).

La tâche qu'effectuent les participant-es consiste à positionner quatre ToiosTM –sans les soulever– dans leurs cibles respectives le plus rapidement et précisément possible. Cette tâche de *glisser-déposer* est illustrée dans la figure 2. Au début de chaque essai, les ToiosTM se positionnent en cercle au centre du plateau. Face à chaque robot est projeté une cible (c.-à-d. une zone grise circulaire de 6 cm de diamètre, distante de 15,5 cm). L'essai commence lorsqu'un son est émis par l'ordinateur. Chaque cible devient verte quand le robot affilié se trouve entièrement dans la zone circulaire. Si le robot ressort de sa cible, cette dernière redevient grise. Une fois toutes les cibles passées au vert, un son est émis par l'ordinateur pour indiquer la validation et la fin de l'essai. Pour la condition AVEC INTERMÉDIAIRE, lorsque le ToioTM intermédiaire est sélectionné – par un clic sur celui-ci– ce dernier émet un son et sa LED s'allume pendant 500 ms. L'expérience se déroule en deux phases : une phase d'entraînement, puis une phase expérimentale.

Pendant la phase d'entraînement, les participant-es effectuent la tâche avec tous les TYPES D'ASSISTANCE en suivant le même ordre pour chaque participant-e. Tous les niveaux d'AUTONOMIE sont d'abord présentés SANS INTERMÉDIAIRE, puis tous AVEC INTERMÉDIAIRE. L'ordre des niveaux d'AUTONOMIE est le suivant : (1) INERTE, (2) AIMANTÉE, (3) ARRÊTABLE, (4) DÉCLENCHÉE, et (5) COMPLÈTE. Pour chaque TYPE D'ASSISTANCE, une explication de la façon d'interagir est fournie à voix haute aux participant-es. Puis, la tâche est effectuée au moins une fois pour chaque TYPE D'ASSISTANCE. Si nécessaire, l'essai est répété jusqu'à ce que les participant-es comprennent parfaitement le TYPE D'ASSISTANCE.

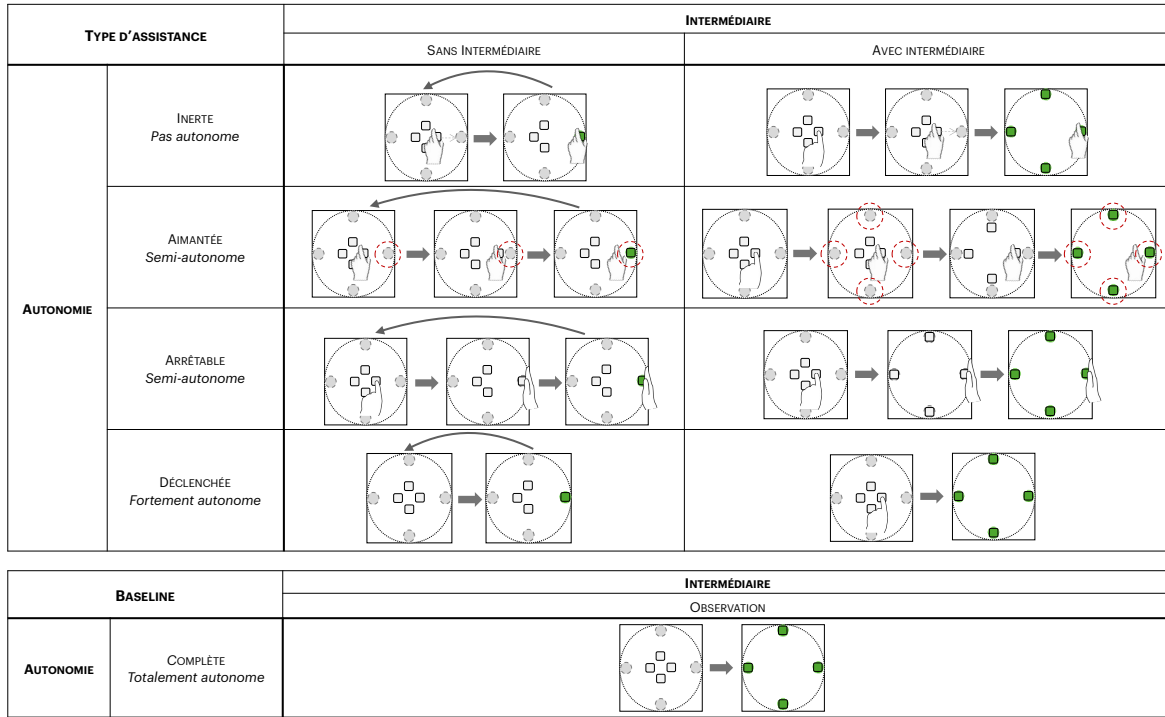


Figure 2: Les variables indépendantes AUTONOMIE et INTERMÉDIAIRE illustrées pour la tâche de glisser-déposer.

La phase expérimentale est divisée en 12 blocs. Dans chaque bloc, les participant-es expérimentent chaque TYPE D'ASSISTANCE. L'ordre des TYPES D'ASSISTANCE est aléatoire au sein de chaque bloc. Au début de chaque essai, l'expérimentatrice prononce une courte phrase de rappel pour indiquer aux participant-es comment interagir avec les Toios™, en fonction du TYPE D'ASSISTANCE en cours. À la fin de chaque tâche, les participant-es répondent au questionnaire décrit en partie 3.3.2. L'expérience durait en moyenne 1h30 (min 1h15, max 1h50).

4 Analyse des données

Pour répondre à notre question de recherche (section 3.4), nous agrégeons les scores du SoA par participants et par TYPE D'ASSISTANCE en utilisant la moyenne. Nous effectuons une ANOVA à deux facteurs (AUTONOMIE $\times$ INTERMÉDIAIRE), afin de connaître la significativité de nos deux facteurs principaux et de leur interaction. Puis, nous effectuons des tests post hoc de Tukey pour nos deux facteurs et leur interaction, afin de connaître la significativité des différences entre chaque condition. Nous prenons le seuil communément utilisé dans la communauté de 0.05 pour notre valeur p . Pour les tests de significativité, nous retranchons les scores du SoA des baselines, aux scores des différents niveaux de nos facteurs. Cela nous permet de connaître le gain du sens d'agentivité par rapport à un système sans interaction.

Nous appuyons nos résultats en utilisant la technique d'estimation basée sur les moyennes et les intervalles de confiance (ICs) *bootstrappés* à 95% –reportés entre crochets “[]” dans le texte et par les barres d'erreurs dans les figures. Nous présentons

également les différences par paires : (1) entre les niveaux de nos facteurs et les baselines (c.-à-d. AUTONOMIE COMPLÈTE et/ou INTERMÉDIAIRE OBSERVATION), et (2) pour les TYPES D'ASSISTANCE AVEC INTERMÉDIAIRE, entre le score de l'intermédiaire et le score des satellites. Ces techniques sont recommandées par la communauté [3, 24], et notamment par le groupe de travail international sur les statistiques transparentes en IHM [71].

5 Résultats

Nous nous intéressons tout d'abord à l'effet de nos deux variables AUTONOMIE et INTERMÉDIAIRE sur le sens d'agentivité globale (i.e., le sentiment de contrôle sur le mouvement des robots). L'ANOVA à deux facteurs (AUTONOMIE $\times$ INTERMÉDIAIRE) montre :

- Un grand effet principal significatif du facteur AUTONOMIE sur le SoA ($F(3,272)=32.85$, $p < .001$, $\eta_p^2 = .266$),
- Un petit effet principal significatif du facteur INTERMÉDIAIRE sur le SoA ($F(1,272)=14.08$, $p < .001$, $\eta_p^2 = .049$),
- L'absence d'effet significatif lorsqu'on considère l'interaction entre ces deux facteurs ($F(3,272)=1.26$, $p=.290$, $\eta_p^2 = .014$).

Les sections suivantes présentent les différences entre chaque condition. L'annexe A présente les valeurs des moyennes et ICs. L'annexe D présente les résultats du taux d'erreurs perçues, et des classements en fonction de la préférence, du sentiment de contrôle, de la performance perçue et de la facilité d'utilisation.

5.1 Effet du facteur AUTONOMIE sur le SENS d'AGENTIVITÉ (hypothèses 1 et 1_{bis})

Les tests post hoc (annexe B.1) montrent que les niveaux d'AUTONOMIE sont significativement différents deux à deux dans la majorité des cas et tous significativement différents de la baseline (Hyp1). La figure 3a illustre nos résultats en montrant une baisse du SoA lorsque l'AUTONOMIE augmente. Seuls les niveaux semi-autonomes (AIMANTÉE et ARRÊTABLE) ne sont pas significativement différents entre eux (Hyp1_{bis}) (chevauchement des ICs en figure 3a).

La figure 3b montre que tous les niveaux du facteur AUTONOMIE ont un SoA supérieur à la BASELINE (différences supérieures à 0). Il semble donc que le SoA est significativement supérieur à la BASELINE dès que l'utilisatrice doit interagir avec les modules, et cela peu importe le type d'AUTONOMIE. L'Hyp1 est donc vérifiée, contrairement à l'Hyp1_{bis}.

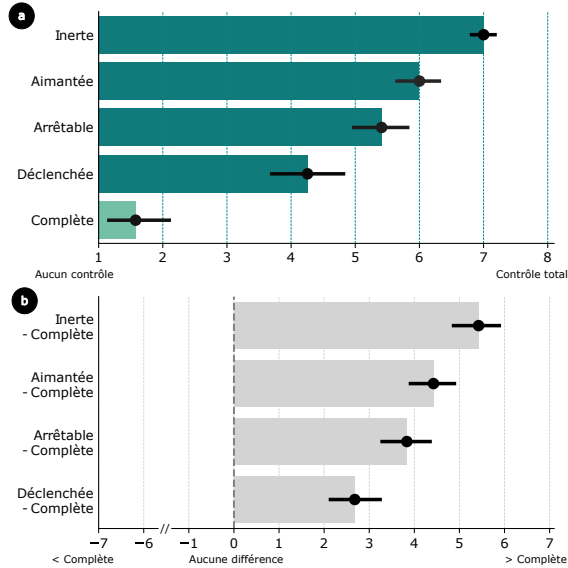


Figure 3: SENS d'AGENTIVITÉ (SoA) en fonction de l'AUTONOMIE. Les barres d'erreurs représentent les intervalles de confiance *bootstrapés* à 95%.

(a) Moyennes des scores du SoA en fonction de l'AUTONOMIE. (b) Différences par paire du SoA en fonction de l'AUTONOMIE. La différence se fait avec la BASELINE : AUTONOMIE COMPLÈTE.

5.2 Effet du facteur INTERMÉDIAIRE sur le SENS d'AGENTIVITÉ (hypothèse 2)

Les tests post hoc montrent que les conditions SANS INTERMÉDIAIRE et AVEC INTERMÉDIAIRE sont significativement différents ($t(272)=3.75, p<.001$) et que tous les niveaux sont significativement différents de la BASELINE. Nous retrouvons ce résultat sur la figure 4 où les ICs ont un écart significatif. La figure 4a illustre nos résultats en montrant que le SoA est plus important SANS INTERMÉDIAIRE qu'AVEC INTERMÉDIAIRE et plus important AVEC INTERMÉDIAIRE que pour la BASELINE OBSERVATION. La figure 4b montre que SANS

INTERMÉDIAIRE et AVEC INTERMÉDIAIRE ont un SoA supérieur à la BASELINE OBSERVATION (différences supérieures à 0). Dès que l'utilisatrice interagit avec au moins un module, le SoA semble donc significativement supérieur à la BASELINE OBSERVATION. Le fait de contrôler un seul module (plutôt que tous) semble donc avoir un effet négatif significatif sur le sens d'agentivité. L'Hyp2 semble donc se vérifier.

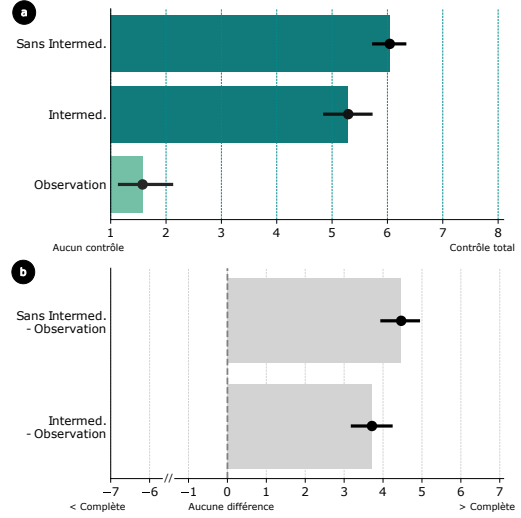


Figure 4: SENS d'AGENTIVITÉ (SoA) en fonction de l'INTERMÉDIAIRE. Les barres d'erreurs représentent les intervalles de confiance *bootstrapés* à 95%. (a) Moyennes des scores du SoA en fonction de l'INTERMÉDIAIRE. (b) Différences par paire du SoA en fonction de l'INTERMÉDIAIRE. La différence se fait avec la BASELINE : INTERMÉDIAIRE OBSERVATION.

5.3 Effet d'interaction entre les facteurs AUTONOMIE et INTERMÉDIAIRE sur le SENS d'AGENTIVITÉ

La figure 5a suggère que le SoA augmente entre les niveaux d'AUTONOMIE mais l'écart entre SANS INTERMÉDIAIRE et AVEC INTERMÉDIAIRE reste stable pour la majorité des AUTONOMIES, sauf pour le niveau INERTE. Pour vérifier cette observation, nous avons aussi effectué des tests de Tukey entre les TYPES d'ASSISTANCE (Annexe B.2). Les tests post hoc montrent également que tous les niveaux sont significativement différents à la BASELINE. La figure 5b illustre ce résultat en montrant que tous les TYPES d'ASSISTANCE ont un SoA supérieur à la BASELINE COMPLÈTE × OBSERVATION (différences supérieures à 0). Le SoA semble donc significativement supérieur à la BASELINE COMPLÈTE × OBSERVATION dès que la personne interagit avec au moins un module, et cela, peu importe le type d'AUTONOMIE.

Ces résultats renforcent notre Hyp1, la non vérification de notre Hyp1_{bis}, et suggèrent que même si une tendance générale se dessine, la vérification de l'Hyp2 semble principalement due à la différence entre SANS INTERMÉDIAIRE et AVEC INTERMÉDIAIRE pour l'AUTONOMIE INERTE.

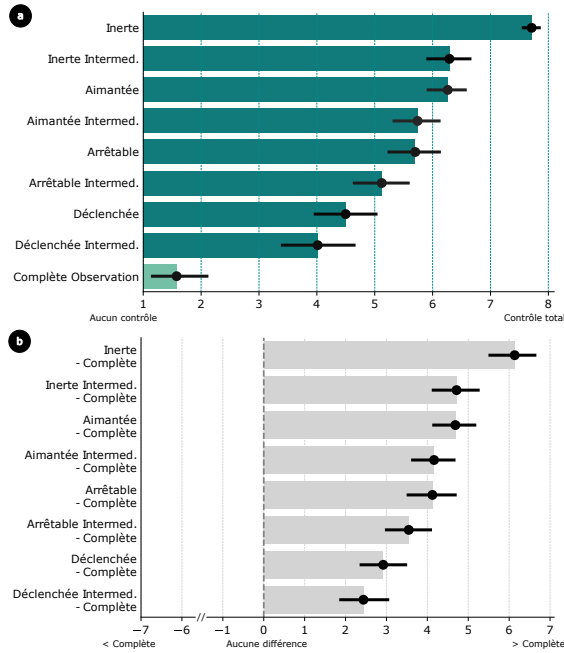


Figure 5: SENS D'AGENTIVITÉ (SoA) en fonction du TYPE D'ASSISTANCE. Les barres d'erreurs représentent les intervalles de confiance *bootstrappés* à 95%. Pour une écriture plus condensée, les TYPES D'ASSISTANCE SANS INTERMÉDIAIRE sont nommés seulement par le nom de l'AUTONOMIE. (a) Moyennes des scores du SoA en fonction du TYPE D'ASSISTANCE. (b) Différences par paire du SoA en fonction du TYPE D'ASSISTANCE. La différence se fait avec la BASELINE : AUTONOMIE COMPLÈTE $\times$ INTERMÉDIAIRE OBSERVATION.

5.4 Effet du rôle du module dans la coordination des mouvements sur le SENS D'AGENTIVITÉ (hypothèse 3)

Nous nous sommes ensuite focalisé sur la condition AVEC INTERMÉDIAIRE afin de comparer l'expérience de contrôle ressentie pour le robot intermédiaire (i.e., le robot contrôlé directement) par rapport au sentiment de contrôle ressenti pour les robots satellites (i.e., les robots contrôlés indirectement via le module intermédiaire). Le test de Wilcoxon montre une différence significative entre le module intermédiaire et les modules satellites ($W=9264$, $p < .001$), avec un SoA plus important pour l'intermédiaire (5.98 [5.68, 6.27]) que pour ses satellites (3.88 [3.54, 4.23]).

Afin de mieux comprendre, l'effet que chaque rôle (intermédiaire ou satellite) peut avoir sur le SoA, en fonction de l'AUTONOMIE, nous avons exploré plusieurs éléments :

- (1) les différences de SoA entre chaque niveau d'AUTONOMIE pour l'intermédiaire (gris sur la figure 6a),
- (2) les différences de SoA entre chaque niveau d'AUTONOMIE pour les satellites (orange sur la figure 6a),
- (3) les différences de SoA entre chaque niveau d'AUTONOMIE pour l'essaim complet (voir section 5.3, vert sur la figure 6a), et

- (4) les différences de SoA entre l'intermédiaire, les satellites et l'essaim complet pour chaque niveau d'AUTONOMIE (figure 6b).

Pour les 1^{er} et 2^e éléments, nous avons mené des tests de Kruskal-Wallis (Annexe C.1). Les tests montrent un effet significatif de l'AUTONOMIE sur le SoA du module intermédiaire ($\chi^2(3) = 73.7$, $p < .001$, $\eta_p^2 = .5305$), et un effet significatif de l'AUTONOMIE sur le SoA des modules satellites ($\chi^2(3) = 10.3$, $p = .016$, $\eta_p^2 = .0744$). Cependant, les comparaisons par paires montrent que les différences entre les AUTONOMIES sont non significatives pour le SoA des satellites, sauf entre INERTE et DÉCLENCHÉE (Annexe C.1 Tableau 3b). À l'inverse, toutes les comparaisons sont significatives pour le SoA de l'intermédiaire, sauf entre AIMANTÉE et ARRÊTABLE (Annexe C.1 Tableau 3a). La figure 6a illustre nos résultats en montrant une baisse importante du SoA sur les modules intermédiaires quand l'AUTONOMIE augmente, et une légère baisse du SoA sur les modules satellites quand l'AUTONOMIE augmente.

Pour le 4^e élément, nous avons effectué des tests de Friedman entre les scores de SoA de l'intermédiaire, des satellites et de l'essaim pour chaque AUTONOMIE. Les tests montrent des différences significatives entre les scores de SoA, et les comparaisons par paires montrent des différences significatives entre chaque SoA deux à deux pour chaque AUTONOMIE :

- INERTE, $\chi^2(2) = 57.9$, $p < .001$ (comparaisons par paires en annexe C.2)
- AIMANTÉE, $\chi^2(2) = 45.6$, $p < .001$ (comparaisons par paires en annexe C.3)
- ARRÊTABLE, $\chi^2(2) = 38.9$, $p < .001$ (comparaisons par paires en annexe C.4)
- DÉCLENCHÉE, $\chi^2(2) = 38.1$, $p < .001$ (comparaisons par paires en annexe C.5)

La figure 6a illustre nos résultats en montrant que, pour toutes les AUTONOMIES, le score de SoA des modules satellites est inférieure à celui du module intermédiaire et de l'essaim et que le score de SoA de l'essaim est inférieur à celui de l'intermédiaire. La figure 6b montre que cette différence est d'autant plus importante quand le niveau d'AUTONOMIE est faible. L'Hyp3 est donc vérifiée.

6 Discussion

6.1 Effet du type d'assistance sur le SoA

Le but de notre étude est de déterminer comment le type d'assistance d'une interface modulaire robotisée peut impacter le sens d'agentivité. Pour cela nous avons mené une étude sur les interfaces en essaim. Nous considérons qu'un type d'assistance est constitué : (1) du type de contrôle que peut avoir l'utilisatrice sur différentes phases du mouvement (AUTONOMIE), et (2) de la coordination ou non des mouvements de l'essaim (satellites) sur un seul module (INTERMÉDIAIRE).

L'AUTONOMIE a un impact sur le SoA (Hyp1 vérifiée), ce qui est cohérent avec la littérature [6, 7, 20, 44, 72, 79, 80]. Une trop grande automatisation impacterait donc fortement le sens d'agentivité de l'utilisatrice. Cependant, nous ne pouvons pas conclure sur quel niveau semi-autonome permettrait d'avoir le plus de sens d'agentivité (Hyp1_{bis}). Le contrôle sur une partie du trajet, plutôt

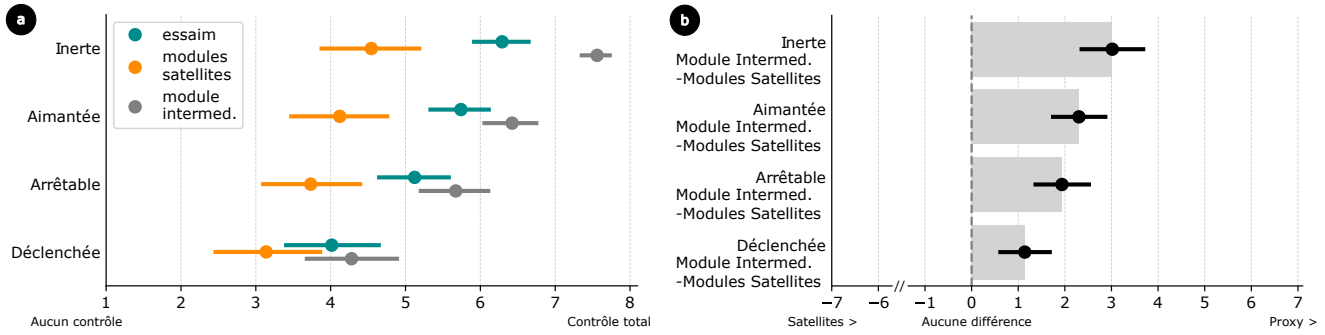


Figure 6: SENS D'AGENTIVITÉ (SoA) en fonction de l'AUTONOMIE et du rôle des modules (intermédiaire vs. satellites). Les barres d'erreurs représentent les intervalles de confiance *bootstrap* à 95%. (a) Moyennes du SoA en fonction de l'AUTONOMIE et du rôle des modules (intermédiaire vs. satellites). (b) Différences par paire du SoA en fonction du rôle des modules (intermédiaire vs. satellites). La différence se fait entre le score de SoA pour l'intermédiaire et le score de SoA pour les satellites.

que sur la fin du mouvement, ne semble pas renforcer significativement le sens d'agentivité. Il est important de noter que le simple fait d'appuyer sur le module pour déclencher son mouvement redonne du sens d'agentivité à l'utilisatrice.

L'INTERMÉDIAIRE a un impact sur le SoA (Hyp2 vérifiée). Ne manipuler qu'un module intermédiaire semble diminuer le sens d'agentivité (vs. manipuler tous les modules). Cependant, pour la même autonomie, le fait d'effectuer la tâche avec ou sans intermédiaire ne semble pas impacter le sens d'agentivité, sauf lorsque le système est inerte. L'impact de l'AUTONOMIE est donc plus fort que celui de l'INTERMÉDIAIRE. L'intermédiaire donne un type particulier d'autonomie aux modules satellites, passant d'un contrôle total de l'utilisatrice (SANS INTERMÉDIAIRE) à une perte partielle de contrôle (AVEC INTERMÉDIAIRE). Nous discutons maintenant l'impact du rôle des modules sur le sens d'agentivité. Le rôle des modules a un impact sur le SoA (Hyp3 vérifiée). Le sens d'agentivité pour les modules contrôlés directement (intermédiaires) est plus élevé que pour les modules contrôlés indirectement (satellites). Cependant, ce résultat est à pondérer en fonction de l'autonomie. En effet, l'autonomie a un impact plus important sur le module intermédiaire que sur ses satellites. Il est également intéressant de noter que le sens d'agentivité ressenti sur les mouvements de l'essaim se rapproche du niveau du sens d'agentivité ressenti sur les mouvements du module intermédiaire. Cependant, plus l'autonomie diminue, plus le sens d'agentivité de l'essaim s'éloigne du sens d'agentivité de l'intermédiaire. Il atteint alors un niveau entre le sens d'agentivité ressenti pour le module intermédiaire et le sens d'agentivité ressenti pour ses satellites.

6.2 Quel type de SoA est ressenti lors de l'interaction avec un intermédiaire ?

Nous pouvons émettre plusieurs hypothèses quant au sens d'agentivité ressenti lors de l'interaction avec un intermédiaire.

Premièrement, nous pourrions considérer que l'utilisatrice percevrait l'essaim comme hiérarchisé. L'utilisatrice serait alors *leader* dans l'accomplissement de la tâche à travers le module intermédiaire et considérerait les autres modules comme les *suiveurs*

de son action. Il a été montré que la personne *leader* développe un sens d'agentivité vicariant pour les actions des personnes sous son commandement [15, 17, 39], même si celui-ci est moindre que son propre sens d'agentivité [15, 39]. Dans notre cas, cela expliquerait pourquoi les utilisatrices peuvent ressentir du sens d'agentivité sur les mouvements des modules satellites même si celui-ci est moindre que le sens d'agentivité pour le module intermédiaire. De plus, Le Bars et al. [39] montre que lors d'une action jointe, il existerait la création d'un sens d'agentivité dit joint (c.-à-d. un sens d'agentivité partagé, dans notre cas correspond au sens d'agentivité perçu sur l'essaim). Ce sens d'agentivité serait moins élevé que le sens d'agentivité personnel. Le niveau d'implication motrice dans la tâche influe plus fortement sur le sens d'agentivité personnel que sur le sens d'agentivité joint [39]. Dans notre cas, cela expliquerait pourquoi le sens d'agentivité perçu sur le module *leader* serait plus élevé mais plus impacté par l'AUTONOMIE que le sens d'agentivité perçu sur l'essaim.

Ceci suggérerait que lors d'interactions avec une interface modulaire robotisée via l'un de ses modules servant d'intermédiaire, trois types de sens d'agentivité pourraient être en jeu, comme lors d'une action jointe avec une personne *leader* :

- Sens d'agentivité leader** ressenti pour le module intermédiaire, influencé par le niveau d'autonomie,
- Sens d'agentivité suiveur** ressenti pour les modules satellites, peu influencé par le niveau d'autonomie et avec un niveau assez faible (renvoie au sens d'agentivité vicariant),
- Sens d'agentivité joint** ressenti pour l'ensemble des modules, influencé par le niveau d'autonomie.

Deuxièmement, nous pourrions considérer que l'utilisatrice aurait un sens d'agentivité distribué entre les modules, par l'aspect synchrone de notre tâche (c.-à-d. la concordance des mouvements des modules satellites avec ceux de l'intermédiaire). Il a été montré dans la littérature que le sens d'agentivité personnel diminue lors d'une action jointe synchrone avec une autre agente humaine [12, 45, 58, 81] ou robotique [18], comme si le sens d'agentivité était distribué entre les agent-es. Dans notre cas, l'utilisatrice percevrait les modules satellites comme ayant leur propre sens d'agentivité, et le sens d'agentivité serait distribué entre les

modules. Cependant, la création d'un sens d'agentivité distribué et d'une diffusion de la responsabilité n'est possible que dans le cas de tâche synchrone et non dans le cas de tâches asynchrones [42, 58]. De plus, la variabilité ou l'écart entre nos actions et celles de nos partenaires impacterait plutôt le sens d'agentivité joint que le sens d'agentivité pour nos propres actions [42]. Nous pouvons donc nous demander à quel point un problème affectant les modules satellites (p. ex. du bruit dans la trajectoire) impacterait le sens d'agentivité joint et le sens d'agentivité ressenti sur les satellites.

6.3 Perception subjective des types d'assistance et sens d'agentivité

Au delà du sens d'agentivité, d'autres facteurs peuvent influencer l'adoption et l'utilisation d'une interface [44]. Parmi eux, la facilité d'utilisation et la performance perçues [21, 37, 74–76].

6.3.1 Compromis entre sens d'agentivité et automatisation des systèmes. Nos résultats montrent qu'il semble exister des corrélations négatives entre le sens d'agentivité et la performance perçue, d'un côté, et la facilité d'utilisation perçue, d'un autre côté, qui elles sont corrélées positivement (Annexe D). Cela nous amène à penser que lors d'une action faisant intervenir plusieurs agentes (potentiellement autonomes), l'utilisabilité n'est pas déterminée uniquement par l'expérience de contrôle de chaque élément mais également sur l'atteinte d'un but commun. Par exemple, dans notre cas, d'un côté l'utilisation d'un module intermédiaire et la situation d'observation minimisent le contrôle sur les autres modules. D'un autre côté, elles augmentent la performance et la facilité d'utilisation perçues.

Cependant, vouloir à tout prix favoriser l'automatisation complète ou l'utilisation d'un module intermédiaire, serait une erreur – au détriment du sens d'agentivité. En particulier, le sens d'agentivité est aujourd'hui considéré comme un déterminant important de la responsabilité ressentie vis-à-vis de nos propres actions, et peut ainsi moduler de manière drastique le comportement humain et l'engagement dans la tâche [25, 30]. Par exemple, Caspar et al. [16] ont montré que des élèves officiers ressentaient moins de sens d'agentivité que des officiers, et se sentaient également moins responsables des conséquences de leurs actes. Une trop grande automatisation des systèmes, entraînant donc une baisse du sens d'agentivité [6, 7, 72], pourrait par conséquent entraîner une baisse du sentiment de responsabilité des conséquences de nos actions – condition nécessaire pour agir de façon éthique et morale [4, 5]. Il est donc important pour les conceptrices de faire des choix en fonction de la criticité de leur système. Par exemple, imaginons un système en essaim servant dans la recherche de victimes en zones de décombres. Nous voudrions un système très performant et donc peut-être le plus autonome possible. Mais qui devient responsable si un effondrement de terrain se produit car le système a choisi une route plutôt qu'une autre sans l'avis d'une superviseuse humaine ? Et du côté de la superviseuse humaine, à quel moment son sens d'agentivité sera trop faible pour se sentir responsable des conséquences des actions de l'essaim, même si sa supervision était requise ? Au contraire, lors de situations non critiques telles que des situations de jeux par exemple, la conceptrice pourra faire varier la performance de son système sans forcément que les utilisatrices ressentent un sens d'agentivité important.

6.3.2 ARRÊTABLE SANS INTERMÉDIAIRE : moins préféré, performant et facile à utiliser ? Le type d'assistance ARRÊTABLE SANS INTERMÉDIAIRE semble être le moins apprécié et jugé comme le moins performant et le moins facile d'utilisation. Cela pourrait s'expliquer par l'implémentation de cette condition : un délai de 150 ms était nécessaire pour l'arrêt des robots. Les robots transmettaient leur position à pas de temps régulier, et si leur position ne changeait pas (c.-à-d., les participant·es bloquait le robot), une instruction leur était envoyée afin de les arrêter. Cependant, il arrivait que les participant·es enlèvent leur main trop tôt. 24 participant·es ont rapporté que cette façon d'arrêter les robots rendait l'interaction : inconfortable (p. ex. P4 "J'aime pas du tout le fait qu'il butte contre ma main, qu'il s'arrête pas."), lente (p. ex. P19 "quand on pose la main ça met quelques secondes avant de s'arrêter et du coup c'est un peu frustrant parce qu'on sait qu'il faut aller vite surtout quand on doit faire les 4 [...] c'est celle où d'un point de vue rapidité on se dit bah c'est moins efficace"), et sujette à l'erreur (p. ex. P30 "Je me suis même trompé·e parce que t'essaies de dire ok c'est bon il va s'arrêter t'enlèves la main en fait il s'arrête pas, [...] je trouve que c'est pas évident à utiliser, et puis comme tu veux aller vite tu veux pas laisser la main trop longtemps parce que le but c'est quand même de passer vite au robot d'après."). Une autre option à explorer sera de permettre l'arrêt de façon plus fiable, comme par exemple en cliquant sur le module.

7 Conclusion et Perspectives

Les interfaces modulaires robotisées à changement de forme offrent de nouvelles possibilités d'interaction et de coopération entre l'humaine et la machine, en tirant partie de leur autonomie. Cependant, une trop grande autonomie peut impacter le sens d'agentivité de l'utilisatrice. Cela pourrait avoir des effets sur l'acceptabilité de ces interfaces et entraîner une désresponsabilisation de l'utilisatrice. Dans cet article, nous avons étudié l'impact du type d'assistance sur le sens d'agentivité ressenti par l'utilisatrice sur le mouvement de modules lors d'une action jointe (tâche coopérative humaine-robot) avec une interface en essaim. Pour cela, nous avons extrait de la littérature neuf TYPES D'ASSISTANCE pouvant être appliqués aux interfaces modulaires robotisées lors d'une situation de coopération avec une agente humaine. Ces TYPES D'ASSISTANCE s'appuient sur deux facteurs : l'autonomie du système (c.-à-d. le type de contrôle que peut avoir l'utilisatrice sur différentes phases du mouvement des modules) et la coordination des mouvements entre modules (c.-à-d. interaction avec un seul module comme un intermédiaire vs. interaction avec chaque module). Nous avons évalué le SENS D'AGENTIVITÉ (SoA) subjectif par questions directes pour ces TYPES D'ASSISTANCE. Nos résultats montrent qu'il semble exister un effet principal du niveau de l'autonomie et un effet principal de l'utilisation d'un module intermédiaire sur le SoA sur le mouvement des robots, mais pas d'effet d'interaction entre les deux variables. Nos résultats montrent également que lors de l'utilisation d'un module intermédiaire, le SoA sur le mouvement du module intermédiaire sera plus important que celui sur l'essaim, et tout deux – SoA intermédiaire et SoA essaim – seront plus importants que celui sur les modules satellites.

Cette étude permet d'envisager plusieurs pistes de réflexion afin d'approfondir nos connaissances sur le lien entre sens d'agentivité

et assistance fournie par le système lors d'une tâche de coopération humaine-machine. Dans notre étude nous avons décidé de nous concentrer sur la mesure directe et subjective du SoA. Des travaux futurs pourraient reproduire notre étude en s'intéressant à la mesure indirect et objective du SoA afin de comparer les résultats obtenus entre les deux mesures.

De plus, nous avons décidé d'utiliser une tâche simple et courante en IHM (glisser-déposer) et de prioriser la validité interne par une étude en laboratoire afin de poser les fondations du lien entre SoA et type d'assistance d'une interface modulaire robotisée. D'un côté, de futurs travaux pourraient donc s'intéresser à la répliquabilité de nos résultats pour une tâche plus complexe demandant notamment aux modules satellites de prendre des décisions par eux-mêmes pour optimiser la réussite de la tâche. D'un autre côté, de futurs travaux pourraient s'intéresser à la répliquabilité de nos résultats dans des conditions réelles.

Nous avons maintenant une idée de la performance perçue de chacun de nos TYPES D'ASSISTANCE. Cependant, des travaux futurs pourraient s'intéresser à la performance objective, permettant ainsi d'identifier le bon compromis entre sens d'agentivité et performance réelle. Notre étude a également mis en lumière que, lors de l'interaction avec un module intermédiaire, les niveaux de sens d'agentivité pour chaque type de modules (intermédiaire, satellites, essaim complet) seraient différents. Nous émettons l'hypothèse que ces résultats pourraient être dus à la vision hiérarchisée de l'essaim par l'utilisatrice, avec la création d'un sens d'agentivité pour le module intermédiaire, un pour les modules satellites et un pour l'essaim (dit joint), et/ou à une distribution du sens d'agentivité entre les modules. Or, des études antérieures ont montré que la création d'un sens d'agentivité joint et la diffusion du sens d'agentivité sont possibles dans le cas de tâches synchrones, c'est-à-dire où les comportements des autres agentes correspondent à nos attendus [42, 58]. Des travaux futurs pourraient donc s'intéresser à l'impact d'un problème survenant dans le système (p. ex. du bruit dans la trajectoire des modules) sur le développement du sens d'agentivité pour chaque TYPE D'ASSISTANCE et notamment lors de l'utilisation d'un module intermédiaire.

Acknowledgements

Cette recherche a été financée par l'Agence Nationale de la Recherche (ANR-21-CE33-0009).

References

- [1] Icek Ajzen. 1991. The theory of planned behavior. *Organizational behavior and human decision processes* 50, 2 (1991), 179–211.
- [2] 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.
- [3] American Psychological Association et al. 2010. *Publication manual of the American psychological association*. American Psychological Association.
- [4] Albert Bandura. 1999. Moral disengagement in the perpetration of inhumanities. *Personality and social psychology review* 3, 3 (1999), 193–209.
- [5] 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.
- [6] Bruno Berberian. 2019. Man-Machine teaming: a problem of Agency. *IFAC-PapersOnLine* 51, 34 (2019), 118–123.
- [7] 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.
- [8] 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.
- [9] 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.
- [10] 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.
- [11] Joanna Bergström-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.
- [12] 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.
- [13] Ramarko Bhattacharya, Jonathan Lindstrom, Ahmad Taka, Martin Nisser, Stefanie Mueller, and Ken Nakagaki. 2024. FabRobotics: Fusing 3D Printing with Mobile Robots to Advance Fabrication, Robotics, and Interaction. In *Proceedings of the Eighteenth International Conference on Tangible, Embedded, and Embodied Interaction*. 1–13.
- [14] Sean Braley, Calvin Rubens, Timothy Merritt, and Roel Vertegaal. 2018. Grid-Drones: A Self-Levitating Physical Voxel Lattice for Interactive 3D Surface Deformations. In *UIST*, Vol. 18. D200.
- [15] 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.
- [16] 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.
- [17] Thierry Chaminade and Jean Decety. 2002. Leader or follower? Involvement of the inferior parietal lobule in agency. *Neuroreport* 13, 15 (2002), 1975–1978.
- [18] Francesca Ciaro, 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.
- [19] Francesca Ciaro, 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.
- [20] 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.
- [21] 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.
- [22] 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.
- [23] 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.
- [24] Pierre Dragicevic. 2016. Fair statistical communication in HCI. *Modern statistical methods for HCI* (2016), 291–330.
- [25] Chris D Frith. 2014. Action, agency and responsibility. *Neuropsychologia* 55 (2014), 137–142.
- [26] Alix Goguey, 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.
- [27] Ouriel Grynszpan, Aisha Sahaï, 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.
- [28] Patrick Haggard and Valerian Chambon. 2012. Sense of agency. *Current biology* 22, 10 (2012), R390–R392.
- [29] Patrick Haggard, Sam Clark, and Jeri Kalogeras. 2002. Voluntary action and conscious awareness. *Nature neuroscience* 5, 4 (2002), 382–385.
- [30] Patrick Haggard and Manos Tsakiris. 2009. The experience of agency: Feelings, judgments, and responsibility. *Current Directions in Psychological Science* 18, 4 (2009), 242–246.
- [31] 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.
- [32] Lawrence H Kim, Veronika Domova, Yuqi Yao, and Parsa Rajabi. 2023. Swarm-Fidget: Exploring Programmable Actuated Fidgeting with Swarm Robots. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology*. 1–15.
 - [33] 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.
 - [34] 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.
 - [35] 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.
 - [36] Gaïqing 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.
 - [37] 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.
 - [38] 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.
 - [39] 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.
 - [40] 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.
 - [41] Mathieu Le Goc, Lawrence H Kim, Ali Parsaei, Jean-Daniel Fekete, Pierre Dragicevic, and Sean Follmer. 2016. Zoids: Building blocks for swarm user interfaces. In *Proceedings of the 29th annual symposium on user interface software and technology*. 97–109.
 - [42] 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.
 - [43] 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.
 - [44] 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.
 - [45] Janeen D Loehr. 2022. The sense of agency in joint action: An integrative review. *Psychonomic Bulletin & Review* 29, 4 (2022), 1089–1117.
 - [46] John E McEneaney. 2013. Agency effects in human–computer interaction. *International Journal of Human–Computer Interaction* 29, 12 (2013), 798–813.
 - [47] Morikatron. 2024. DATAtab: Online Statistics Calculator. https://morikatron.github.io/toio-sdk-for-unity/docs_EN/ Last retrieved: 2024-08-27.
 - [48] 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.
 - [49] Don Norman. 2004. *Emotional design: Why we love (or hate) everyday things*. Basic books.
 - [50] Sukhvinder S Obhi and Preston Hall. 2011. Sense of agency and intentional binding in joint action. *Experimental brain research* 211 (2011), 655–662.
 - [51] 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.
 - [52] Brian Pendleton and Michael Goodrich. 2013. Scalable human interaction with robotic swarms. In *ALAA Infotech@ Aerospace (I@ a) Conference*. 4731.
 - [53] 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.
 - [54] Laura Prusko, Céline Coutrix, Yann Laurillau, Benoît 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.
 - [55] Laura Prusko, Hongri Gu, Julien Bourgeois, Yann Laurillau, and Céline Coutrix. 2023. Modular Tangible User Interfaces: Impact of Module Shape and Bonding Strength on Interaction. In *Proceedings of the Seventeenth International Conference on Tangible, Embedded, and Embodied Interaction*. 1–15.
 - [56] Laura Prusko, Yann Laurillau, Benoît Piranda, Julien Bourgeois, and Céline Coutrix. 2021. Impact of the Size of Modules on Target Acquisition and Pursuit for Future Modular Shape-changing Physical User Interfaces. In *Proceedings of the 2021 International Conference on Multimodal Interaction*. 297–307.
 - [57] 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.
 - [58] 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.
 - [59] 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.
 - [60] 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.
 - [61] Aisha Sahai, Emilie Caspar, Albert De Beir, Ouriel Grynszpan, Elisabeth Pachier, 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.
 - [62] Aisha Sahai, Andrea Desantis, Ouriel Grynszpan, Elisabeth Pachier, 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.
 - [63] 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*.
 - [64] Ben Shneiderman and Catherine Plaisant. 2004. *Designing the user interface: strategies for effective human-computer interaction*. Pearson Addison-Wesley éd.
 - [65] Crystal A Silver, Benjamin W Tatler, Ramakrishna Chakravarthy, and Bert Timmermans. 2021. Social agency as a continuum. *Psychonomic Bulletin & Review* 28, 2 (2021), 434–453.
 - [66] 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.
 - [67] 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.
 - [68] 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.
 - [69] 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.
 - [70] Toio (site officiel) 2018. <https://www.sony.com/en/SonyInfo/design/stories/toio/>.
 - [71] Transparent Statistics in Human–Computer Interaction [n.d.]. <https://transparentstatistics.org>.
 - [72] 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.
 - [73] Viswanath Venkatesh. 2000. Determinants of perceived ease of use: Integrating control, intrinsic motivation, and emotion into the technology acceptance model. *Information systems research* 11, 4 (2000), 342–365.
 - [74] 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.
 - [75] 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.
 - [76] 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.
 - [77] Phillip Walker, Steven Nummally, 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.
 - [78] 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.
 - [79] Debora Zanatto, Mark Chattington, and Jan Noyes. 2021. Human-machine sense of agency. *International Journal of Human-Computer Studies* 156 (2021), 102716.
 - [80] 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.

- [81] 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 Annexe : Tableaux des moyennes et intervalles de confiance

A.1 Hyp1 et 1_{bis} : Moyennes et intervalles de confiance *bootstraps* à 95% du SENS d’AGENTIVITÉ (SoA) pour chaque niveau d’AUTONOMIE

TYPE D’ASSISTANCE		Moyenne SoA [IC]
AUTONOMIE	INERTE (<i>Pas autonome</i>)	7 [6.81, 7.18]
	AIMANTÉE (<i>Semi-autonome</i>)	6 [5.60, 6.29]
	ARRÊTABLE (<i>Semi-autonome</i>)	5.41 [4.92, 5.79]
	DÉCLENCHÉE (<i>Fortement autonome</i>)	4.26 [3.72, 4.85]
BASELINE		Moyenne SoA [IC]
AUTONOMIE	COMPLÈTE (<i>Totalement autonome</i>)	1.58 [1.22, 2.23]

A.2 Hyp2 : Moyennes et intervalles de confiance *bootstraps* à 95% du SENS d’AGENTIVITÉ (SoA) pour chaque niveau de INTERMÉDIAIRE

TYPE D’ASSISTANCE	INTERMÉDIAIRE	
	SANS INTERMÉDIAIRE	AVEC INTERMÉDIAIRE
Moyenne SoA [IC]	6.04 [5.72, 6.30]	5.29 [4.85, 5.70]
BASELINE		INTERMÉDIAIRE
		OBSERVATION
Moyenne SoA [IC]	1.58 [1.22, 2.23]	

A.3 Hyp3 : Moyennes et intervalles de confiance *bootstraps* à 95% pour chaque niveau de TYPE D’ASSISTANCE

		Moyenne SoA [IC]	
TYPE D’ASSISTANCE		INTERMÉDIAIRE	
		SANS INTERMÉDIAIRE	AVEC INTERMÉDIAIRE
AUTONOMIE	INERTE	7.71	6.29
	<i>Pas autonome</i>	[7.54, 7.83]	[5.89, 6.63]
	AIMANTÉE	6.26	5.74
	<i>Semi-autonome</i>	[5.87, 6.55]	[5.29, 6.09]
	ARRÊTABLE	5.70	5.12
	<i>Semi-autonome</i>	[5.18, 6.09]	[4.62, 5.58]
	DÉCLENCHÉE	4.50	4.01
	<i>Fortement autonome</i>	[3.97, 5.04]	[3.44, 4.68]
BASELINE		INTERMÉDIAIRE	
		OBSERVATION	
AUTONOMIE	COMPLÈTE	1.58	
	<i>Totalement autonome</i>	[1.22, 2.23]	

B Annexe : Tests post hoc de l’ANOVA

B.1 Tests post hoc de Tukey pour le facteur AUTONOMIE.

Comparaison			
AUTONOMIE 1	AUTONOMIE 2	t (ddl = 272)	ptukey
INERTE	AIMANTÉE	3.53	.003
	ARRÊTABLE	5.62	<.001
	DÉCLENCHÉE	9.70	<.001
AIMANTÉE	ARRÊTABLE	2.09	.160
	DÉCLENCHÉE	6.16	<.001
ARRÊTABLE	DÉCLENCHÉE	4.08	<.001

Table 1: En rouge, les valeur *p* non significatives et en orange les valeur *p* tendancieuse (entre 0.04 et 0.06).

B.2 Test post hoc de Tukey pour les TYPES D’ASSISTANCE.

Comparaison			
TYPE D’ASSISTANCE 1	TYPE D’ASSISTANCE 2	t (ddl = 272)	ptukey
INERTE	INERTE INTERMÉDIAIRE	3.5504	0.011
	AIMANTÉE	3.6218	0.008
	AIMANTÉE INTERMÉDIAIRE	4.9242	<.001
	ARRÊTABLE	5.0253	<.001
	ARRÊTABLE INTERMÉDIAIRE	6.4705	<.001
	DÉCLENCHÉE	8.0286	<.001
INERTE INTERMÉDIAIRE	DÉCLENCHÉE INTERMÉDIAIRE	9.2359	<.001
	AIMANTÉE	0.0714	1.000
	AIMANTÉE INTERMÉDIAIRE	1.3738	0.868
	ARRÊTABLE	1.4749	0.820
	ARRÊTABLE INTERMÉDIAIRE	2.9200	0.073
	DÉCLENCHÉE	4.4782	<.001
AIMANTÉE	DÉCLENCHÉE INTERMÉDIAIRE	5.6855	<.001
	AIMANTÉE INTERMÉDIAIRE	1.3024	0.897
	ARRÊTABLE	1.4035	0.855
	ARRÊTABLE INTERMÉDIAIRE	2.8487	0.088
	DÉCLENCHÉE	4.4068	<.001
	DÉCLENCHÉE INTERMÉDIAIRE	5.6141	<.001
AIMANTÉE INTERMÉDIAIRE	ARRÊTABLE	0.1011	1.000
	ARRÊTABLE INTERMÉDIAIRE	1.5463	0.781
	DÉCLENCHÉE	3.1044	0.043
	DÉCLENCHÉE INTERMÉDIAIRE	4.3117	<.001
ARRÊTABLE	ARRÊTABLE INTERMÉDIAIRE	1.4452	0.835
	DÉCLENCHÉE	3.0033	0.058
	DÉCLENCHÉE INTERMÉDIAIRE	4.2106	<.001
ARRÊTABLE INTERMÉDIAIRE	DÉCLENCHÉE	1.5581	0.775
	DÉCLENCHÉE INTERMÉDIAIRE	2.7654	0.108
DÉCLENCHÉE	DÉCLENCHÉE INTERMÉDIAIRE	1.2073	0.929

Table 2: En rouge, les valeur *p* non significatives et en orange les valeur *p* tendancieuse (entre 0.04 et 0.06). Pour une écriture plus condensée, les TYPES D’ASSISTANCE SANS INTERMÉDIAIRE sont nommés seulement par le nom de l’AUTONOMIE et les TYPES D’ASSISTANCE AVEC INTERMÉDIAIRE sont nommés de la façon suivante “nom de l’AUTONOMIE INTERMÉDIAIRE”.

C Annexe : Tests post hoc de l’effet du rôle du module dans la coordination des mouvements sur le SENS d’AGENTIVITÉ (SoA)

C.1 Tests post hoc de Dwass–Steel–Critchlow–Fligner pour le facteur TYPE D’ASSISTANCE pour le module intermédiaire et les modules satellites

Comparaison		W	p
AUTONOMIE 1	AUTONOMIE 2		
INERTE	AIMANTÉE	-7.82	<.001
	ARRÊTABLE	-8.96	<.001
	DÉCLENCHÉE	-9.30	<.001
AIMANTÉE	ARRÊTABLE	-3.43	0.072
	DÉCLENCHÉE	-6.78	<.001
ARRÊTABLE	DÉCLENCHÉE	-4.64	0.006

(a) Résultats des tests post hoc de Dwass–Steel–Critchlow–Fligner pour le module intermédiaire en fonction des TYPES D’ASSISTANCE.

Comparaison		W	p
AUTONOMIE 1	AUTONOMIE 2		
INERTE	AIMANTÉE	-1.36	0.770
	ARRÊTABLE	-2.41	0.322
	DÉCLENCHÉE	-4.15	0.017
AIMANTÉE	ARRÊTABLE	-1.31	0.790
	DÉCLENCHÉE	-3.07	0.132
ARRÊTABLE	DÉCLENCHÉE	-2.29	0.370

(b) Résultats des tests post hoc de Dwass–Steel–Critchlow–Fligner pour les modules satellites en fonction des TYPES D’ASSISTANCE.

Table 3: En rouge, les valeur *p* non significatives et en orange les valeur *p* tendancieuse (entre 0.04 et 0.06).

C.2 Comparaisons par paires (Durbin-Conover) pour l’AUTONOMIE INERTE entre les scores de SoA du module intermédiaire, des modules satellites et de l’essaim

Comparaison		t	p
Type de module(s) 1	Type de module(s) 2		
Intermédiaire	Satellites	18.02	<.001
	Essaim	8.79	<.001
Satellites	Essaim	9.23	<.001

C.3 Comparaisons par paires (Durbin-Conover) pour l’AUTONOMIE AIMANTÉE entre les scores de SoA du module intermédiaire, des modules satellites et de l’essaim

Comparaison		t	p
Type de module(s) 1	Type de module(s) 2		
Intermédiaire	Satellites	11.29	<.001
	Essaim	5.49	<.001
Satellites	Essaim	5.80	<.001

C.4 Comparaisons par paires (Durbin-Conover) pour l’AUTONOMIE ARRÊTABLE entre les scores de SoA du module intermédiaire, des modules satellites et de l’essaim

Comparaison		t	p
Type de module(s) 1	Type de module(s) 2		
Intermédiaire	Satellites	9.17	<.001
	Essaim	3.76	<.001
Satellites	Essaim	5.41	<.001

C.5 Comparaisons par paires (Durbin-Conover) pour l’AUTONOMIE DÉCLENCHÉE entre les scores de SoA du module intermédiaire, des modules satellites et de l’essaim

Comparaison		t	p
Type de module(s) 1	Type de module(s) 2		
Intermédiaire	Satellites	8.98	<.001
	Essaim	3.95	<.001
Satellites	Essaim	5.03	<.001

D Annexe : Résultats complémentaires : taux d’erreurs perçues et classements des TYPES D’ASSISTANCE

D.1 Analyse des données

Nous reportons le taux d’erreurs perçues. Pour cela, nous considérons comme une erreur perçue les essais dont la note au succès perçu de la tâche est en-dessous de sa moyenne pour le-a participant-e. Pour analyser le classement des neuf TYPES D’ASSISTANCE en fonction de la préférence, du sentiment de contrôle, de la performance perçue et de la facilité d’utilisation, nous effectuons des tests de Kruskal-Wallis pour chaque question afin de vérifier la significativité des différences entre chaque TYPE D’ASSISTANCE. Nous couplons nos tests à l’estimation des moyennes avec les ICs *bootstrappés* à 95%. Nous utilisons les scores inversés –une interaction classée en première position (score initial de 1) aura un score inversé de 9. Cela nous permet d’avoir des interactions identifiables facilement où le score le plus grand correspond à l’interaction la plus plébiscitée par les participant-es. Nous regardons également le consensus entre les participant-es en utilisant le coefficient de concordance (Kendall’s W). Le résultat est considéré comme négligable en-dessous de 0.3, faible entre 0.3 et 0.5, modéré entre 0.5 et 0.7, élevé entre 0.7 et 0.9 et très élevé au-dessus de 0.9.

Pour tous nos tests de significativité, nous prenons le seuil communément utilisé dans la communauté de 0.05 pour notre valeur *p*.

D.2 Taux d’erreurs perçues

Les moyennes de pourcentage d’erreurs par TYPE D’ASSISTANCE (Tableau 4) montre que :

- (1) Les TYPES D’ASSISTANCE avec le plus d’erreurs perçues sont ARRÊTABLE et INERTE.
- (2) Les TYPES D’ASSISTANCE AVEC INTERMÉDIAIRE diminuent le nombre d’erreurs perçues par rapport leur équivalent SANS INTERMÉDIAIRE.

D.3 Classements des TYPES D’ASSISTANCE

Le test de Kruskal-Wallis montre des différences significatives pour la préférence, le sentiment de contrôle, la performance perçue et la facilité d’utilisation perçue entre les niveaux du TYPE D’ASSISTANCE :

TYPE D'ASSISTANCE	% erreurs
ARRÊTABLE	51.43
INERTE	44.76
AIMANTÉE	30.48
ARRÊTABLE INTERMÉDIAIRE	27.14
AIMANTÉE INTERMÉDIAIRE	17.38
DÉCLENCHÉE	14.76
INERTE INTERMÉDIAIRE	12.86
DÉCLENCHÉE INTERMÉDIAIRE	5.48
COMPLÈTE	3.10

Table 4: Taux d'erreurs perçues. Les scores sont calculés à partir des résultats du succès perçu à la tâche en considérant comme une erreur tout essai dont la note est en-dessous la moyenne du succès perçu du/de la participant-e. Pour une écriture plus condensée, les TYPES D'ASSISTANCE SANS INTERMÉDIAIRE sont nommés seulement par le nom de l'AUTONOMIE et les TYPES D'ASSISTANCE AVEC INTERMÉDIAIRE sont nommés de la façon suivante "nom de l'AUTONOMIE INTERMÉDIAIRE".

- *Préférence* : $\chi^2(8) = 63.5$, $p < .001$, $\eta_p^2 = .202$. La taille d'effet très faible nous laisse à penser que même si le test est significatif, l'effet du TYPE D'ASSISTANCE sur la préférence sera très faible, et possiblement expliqué par d'autres facteurs.
- *Sentiment de contrôle* : $\chi^2(8) = 193.9$, $p < .001$, $\eta_p^2 = .617$
- *Performance perçue* : $\chi^2(8) = 218.7$, $p < .001$, $\eta_p^2 = .697$
- *Facilité d'utilisation perçue* : $\chi^2(8) = 186.4$, $p < .001$, $\eta_p^2 = .594$

Les scores de concordances entre les participant-es sont les suivants :

- *Préférence* : 0.20 (pas d'accord entre les participant-es)
- *Sentiment de contrôle* : 0.62 (accord modéré entre les participant-es)
- *Performance perçue* : 0.70. (accord élevé entre les participant-es)
- *Facilité d'utilisation perçue* : 0.59 (accord modéré entre les participant-es)

La figure 7a, récapitulant les scores moyens pour chaque TYPE D'ASSISTANCE pour la *préférence*, illustre que :

- (1) Deux TYPES D'ASSISTANCE se détachent comme étant préférés : INERTE AVEC INTERMÉDIAIRE et DÉCLENCHÉE AVEC INTERMÉDIAIRE.
- (2) Deux TYPES D'ASSISTANCE se détachent comme étant les moins appréciés : ARRÊTABLE SANS INTERMÉDIAIRE et ARRÊTABLE AVEC INTERMÉDIAIRE.
- (3) Les TYPES D'ASSISTANCE AVEC INTERMÉDIAIRE sont plus appréciés que leur équivalent SANS INTERMÉDIAIRE.

Plus de détail sur la significativité des différences entre chaque condition est à retrouver dans les fichiers additionnels.

La figure 7b, récapitulant les scores moyens pour chaque TYPE D'ASSISTANCE pour le *sentiment de contrôle*, illustre que :

- (1) Les participant-es se sentent le plus en contrôle avec INERTE AVEC INTERMÉDIAIRE.

- (2) Les participant-es se sentent le moins en contrôle avec COMPLÈTE OBSERVATION.
- (3) Les TYPES D'ASSISTANCE AVEC INTERMÉDIAIRE sont évalués comme procurant un sentiment de contrôle moindre que leur équivalent SANS INTERMÉDIAIRE.

Plus de détail sur la significativité des différences entre chaque condition est à retrouver dans les fichiers additionnels.

La figure 7c, récapitulant les scores moyens pour chaque TYPE D'ASSISTANCE pour la *performance perçue*, montre que :

- (1) En majorité, plus le TYPE D'ASSISTANCE a un niveau élevé d'AUTONOMIE, plus la performance perçue est élevée.
- (2) Les participant-es se sentent moins performant-es avec ARRÊTABLE SANS INTERMÉDIAIRE.
- (3) Les TYPES D'ASSISTANCE AVEC INTERMÉDIAIRE sont évalués comme plus performants que leur équivalent SANS INTERMÉDIAIRE.

Plus de détail sur la significativité des différences entre chaque condition est à retrouver dans les fichiers additionnels.

La figure 7d, récapitulant les scores moyens pour chaque TYPE D'ASSISTANCE pour la *facilité d'utilisation perçue*, montre que :

- (1) DÉCLENCHÉE AVEC INTERMÉDIAIRE et COMPLÈTE OBSERVATION sont perçues comme les plus faciles à utiliser.
- (2) ARRÊTABLE SANS INTERMÉDIAIRE et AIMANTÉE SANS INTERMÉDIAIRE sont perçues comme les plus difficiles à utiliser.
- (3) Les TYPES D'ASSISTANCE AVEC INTERMÉDIAIRE sont perçus comme plus faciles à utiliser que leur équivalent SANS INTERMÉDIAIRE.

Plus de détail sur la significativité des différences entre chaque condition est à retrouver dans les fichiers additionnels.

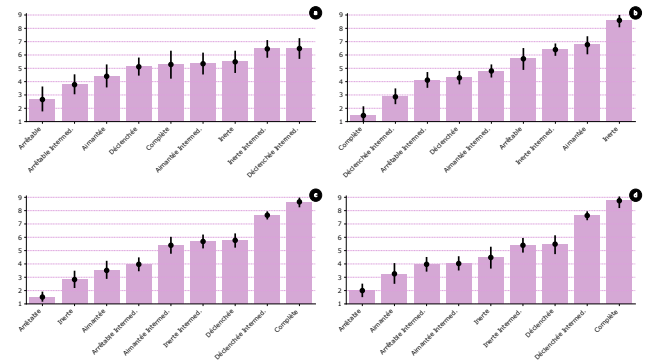


Figure 7: Classements de (a) la performance, (b) le contrôle perçu, (c) la performance perçue (rapidité × précision), et (d) la facilité d'utilisation perçue. Les barres d'erreurs représentent les intervalles de confiance bootstrappés à 95%. Pour une écriture plus condensée, les TYPES D'ASSISTANCE SANS INTERMÉDIAIRE sont nommés seulement par le nom de l'AUTONOMIE et les TYPES D'ASSISTANCE AVEC INTERMÉDIAIRE sont nommés de la façon suivante "nom de l'AUTONOMIE Intermed.".

	SoA	Préférence	Sentiment de contrôle	Performance perçue	Facilité d'util. perçue	Succès perçu
SoA	1					
Préférence	0.077	1				
Sentiment de contrôle	0.560***	0.092*	1			
Performance perçue	-0.313***	0.311***	-0.383***	1		
Facilité d'util. perçue	-0.252***	0.318***	-0.299***	0.620***	1	
Succès perçu	-0.215***	0.165***	-0.276***	0.394***	0.326***	1

Table 5: Tau de Kendall (corrélation) entre le score du SoA, de la préférence, du sentiment de contrôle, de la performance perçue, de la facilité d'utilisation perçue et du succès perçu à la tâche (* $p < .05$, ** $p < .01$, * $p < .001$).**

D.4 Corrélations entre SoA, préférence, contrôle perçu, performance perçue, facilité d'utilisation et succès perçu

Le tableau 5 reporte les corrélations deux à deux entre : SoA, préférence, contrôle perçu, performance perçue, facilité d'utilisation

et succès perçu. Nous constatons que toutes les corrélations sont significativement faibles à modérées, sauf pour la corrélation entre SoA et préférence à la corrélation n'est pas significative. Il semble donc qu'il n'y ait pas un facteur qui puisse être expliqué exclusivement par un autre.