



I am Robot, Your Health Adviser for Older Adults: Do You *Trust* My Advice?

Ioanna Giorgi^{1,3} · Aniello Minutolo² · Francesca Tiroto¹ · Oksana Hagen¹ · Massimo Esposito² · Mario Gianni¹ · Marco Palomino¹ · Giovanni L. Masala³

Accepted: 15 May 2023 / Published online: 8 June 2023
© The Author(s) 2023

Abstract

Artificial intelligence and robotic solutions are seeing rapid development for use across multiple occupations and sectors, including health and social care. As robots grow more prominent in our work and home environments, whether people would favour them in receiving useful advice becomes a pressing question. In the context of human–robot interaction (HRI), little is known about people’s advice-taking behaviour and trust in the advice of robots. To this aim, we conducted an experimental study with older adults to measure their trust and compliance with robot-based advice in health-related situations. In our experiment, older adults were instructed by a fictional human dispenser to ask a humanoid robot for advice on certain vitamins and over-the-counter supplements supplied by the dispenser. In the first experimented condition, the robot would give only *information-type advice*, i.e., neutral informative advice on the supplements given by the human. In the second condition, the robot would give *recommendation-type advice*, i.e., advice in favour of more supplements than those suggested initially by the human. We measured the trust of the participants in the type of robot-based advice, anticipating that they would be more trusting of information-type advice. Moreover, we measured the compliance with the advice, for participants who received robot-based recommendations, and a closer proxy of the actual use of robot health advisers in home environments or facilities in the foreseeable future. Our findings indicated that older adults continued to trust the robot regardless of the type of advice received, highlighting a type of protective role of robot-based recommendations on their trust. We also found that higher trust in the robot resulted in higher compliance with its advice. The results underpinned the likeliness of older adults welcoming a robot at their homes or health facilities.

Keywords Robot adviser · Robot-based advice · Advice taking · Human–robot interaction · HRI · Trust · Compliance · Senior care

1 Introduction

Social robots are increasingly seen as a viable support resource in the healthcare sector. These robots have been studied in multiple contexts such as assisting people with mobility and household tasks [1] or using them as companions [2]. Since the older population is growing faster than the

younger population [3], social robots have been extensively studied in their relationship with older adults to support the shortage of caregivers. These studies have mostly focused on dimensions (i.e., predictors) that most likely influence older adults’ acceptance of social robots [4–6]. In the present study, we focused on a specific type of acceptance, i.e., the acceptance of advice from social robots in the healthcare sector by primarily investigating the perceived trust in a robot giving advice. Advice-giving by robots may be an added value to encourage and maintain healthier behaviours in individuals, by helping them to proactively self-manage their health and well-being, thus reducing the necessity to depend on health-care professionals each time. By imparting unbiased, tailored advice to assess the specific needs of individuals, they can help people make informed decisions about their health, even when time constraints and/or socioeconomic status prevent

✉ Ioanna Giorgi
i.giorgi@kent.ac.uk

¹ School of Engineering, Computing and Mathematics, University of Plymouth, Plymouth, UK

² Institute for High Performance Computing and Networking (ICAR), National Research Council of Italy (CNR), Naples, Italy

³ School of Computing, University of Kent, Canterbury, UK

them from visiting a health facility. This form of robot–human teaming may help relieve the burden on caregivers, while still allowing professionals to access and understand the health patterns of the individuals they care for through robot mediation.

1.1 Theoretical Background

The degree to which people take advice from others is a common research problem in cognitive sciences. People’s inability to predict entirely the knowledge and the intentions of an adviser promote feelings of uncertainty toward the adviser and the advice received. This uncertainty seems to, however, reduce when people perceive high confidence [7] and high trust [8, 9] in the adviser. From a multi-disciplinary perspective, trust can be defined as “a psychological state comprising the intention to accept vulnerability based upon positive expectations of the intentions or behaviour of another” ([10], p. 395). In the advice give-and-take synergy, trust relies greatly, among others, on the perceived cognitive competence of the adviser in providing valuable and customised advice in a specific domain [11, 12]. Therefore, people are more willing to accept advice from a person that they consider competent [13, 14].

Moreover, compliance with the advice taken i.e., the likeliness to use the advice, also depends on the type of advice [15] and the type of task, such as context or task difficulty [16]. For example, people react with differential preferences to the “*recommendation for/against*” type of advice, i.e., advice in favour or against something, and the “*information*” type of advice, i.e., neutral advice providing information about something but not suggestion in favour or opposition of it. Information-type advice is often preferred [15]. In turn, people would be likelier to weigh an expert’s opinion more heavily if faced with a difficult task than when the task is simple and straightforward to them [16]. However, if the nature of the challenging task is more critical, i.e., heavier consequences of poor advice, more trust in the adviser is required to rely on them.

1.2 Adviser—Addressee in Human–Machine Interaction

When it comes to nonhuman advisers, there is still much to learn about the level of trust and compliance humans place in the advice they receive. Some studies have concluded that people exert more reliance on AI-algorithm-based advisers in analytical tasks given their perceived superiority in statistical deductions and ability to process greater information in real-time compared to humans [17]. However, whether humans will use a technology, in general or specifically to receive advice, depends on certain factors, such as the characteristics of the task conducted by the technology and its

fit, i.e., the appropriateness of that technology to solve or assist with solving the task [18, 19]. Moreover, as for the human–human relationship, the human–technology rapport requires considerable trust in advice-taking. In this regard, people’s behavioural tendencies and receptivity to agent-based advice remain less clear.

Trust in a machine is modulated by its perceived expertise [20, 21] and its approximation of human characteristics or anthropomorphism [22, 23]. Machines that embody human features receive more trust [23] and people are likelier to comply with their instructions [24] compared to mechanical physiology. Humanlike appearance also strengthens the dependence of people on agents [25] and trust resilience, or the ability to withstand breaches of trust [26]. However, when the other agent-related features beyond morphology, including behaviour, do not resemble human likenesses, trust in the machine is impaired due to inconsistency of expectations [27]. In addition, trust in the advice given by a machine can be impacted by the mismatch between the agent’s displayed morphology and attributed abilities with its actual conduct in a given task [24]. When an agent is perceived to have expert abilities in executing a task or supporting humans in their task goals, its advice is considered more credible, i.e., higher trust [28, 29]. Overall, the surveyed works suggest that non-human advisers may not necessarily be forthwith ill-favoured or distrusted, but trust in their advice is contingent on their anthropomorphism and perceived expertise in a given context.

1.3 ...and in Human–Robot Interaction (HRI)

Following the above, research insights from Human–Robot Interaction (HRI) further stress the importance of human resemblance in shaping people’s behaviour towards robots. Studies have found that the degree of anthropomorphism alters the way people interact with robots and can lead to more human-like behaviours than when interacting with a machine [30, 31]. For example, Hertz and Wiese found a linear change in advice-taking behaviour as anthropomorphism increased from a computer display to a robot adviser and a human adviser [32]. In their study, they investigated people’s willingness to receive advice from nonhuman agents, comparing their choice of and compliance with a human, a robot, and a computer agent in a social and analytical task. They reported that participants would prefer the human more often when they were given the option to choose the adviser before knowing the task (“agent-first condition”). Instead, when the task was known a priori (“task-first condition”), their choices were calibrated based on how “fit” they perceived the adviser to be for that task. For example, for a *social task*, they preferred the human adviser more often, whereas, for an *analytical task*, they would opt for a nonhuman agent, leaning more frequently towards the robot. Although, at first sight,

it appears that human advisers remain favoured in *social domains*, it has been demonstrated that people's acceptance and trust of robots in such domains is strengthened when the robot has a humanoid morphology, as well as when it exhibits appropriate social and psychological behaviour. For example, robots with enhanced social involvement can trigger higher reactance and compliance in persuasive attempts [33]. Additionally, people exert higher acceptance, trust and less physical proximity to a robot with appropriate communicative behaviour (verbal/non-verbal) [34].

Nevertheless, a valuable insight drawn from the above studies is that the type of task and the perceived adviser-task fit significantly influence people's advice-taking behaviours [32] and that a robot's dialogue initiative is only acceptable in some, but not all dialogue contents [34]. While robots may be considered superior and their advice may be utilised for complex statistical calculations, such as stock price predictions [35], their advice may be deemed inappropriate or met with scepticism in situations where human judgement is necessary, such as health-related contexts. For example, in an exploratory study with hospitalised children with Type 1 diabetes, the authors scrutinised the acceptability of a humanoid robot as an assistant in children's disease management and the advice or coaching delivered by the robot [36]. Their findings revealed that the acceptability of taking advice from a robot about blood glucose levels and patterns was considerable (above 90%), but not as high when it came to more critical aspects like calculating insulin doses.

However, the positive effect that persuasive social robots could have in supporting a healthy lifestyle must not be overlooked. For example, the advice given by a social robotic agent in a game designed by Ghazali et al. [37] led to making healthier versions of the beverage after each decision. Similarly, other authors showed that a social robot, which used persuasive strategies of empathy, would stimulate participants to drink significantly more water [38] or that robots with an adaptive linguistic style could be effective persuaders to encourage and foster a healthy diet [39]. Moreover, socially-enhanced robot advisers were found to help combat sedentary behaviours in office-based workplaces using interactive exercise experiences [40] and produce positive attitudes and compliance during meal-eating activities [41].

Trust in robot-based advice in health contexts is an important and growing real-world problem. The above studies emphasise the potential of social robots as such advisers. They also suggest that persuasion is more effective when the adviser is humanised, while also drawing importance on the context in which they are applied. Our study is motivated by the above conclusions, for example in our choice of robot and the designed experimental conditions but differs from the surveyed literature in several respects. First, their designed

experiments, even when these are health-related, are less critical (e.g., physical activity, drinking water, game) than our considered scenario of medication intake. The significant preoccupation with health maintenance in all societies has highlighted the importance of automation, thus, the impact that robots could play in the real world and even critical health activities should be studied in-depth. Second, many of the aforementioned studies have been carried out across the younger population (or a mix involving both younger and older adults), but much less in the older population, who may have substantially different expectations, beliefs about robots and experiences with them. Third, they have mainly investigated people's advice-taking of robot-based advice considering people's compliance or utilisation of the advice, using compliance as an indirect measure of trust toward the robot (see [42] for a recent review). However, compliance and trust remain two different, separated, constructs with trust being the predictor of compliance [43, 44]. Lastly, some of the studies explore how the attributes of a nonhuman adviser like humanlike morphology or social engagement impact advice-taking behaviours. Differently, our study concerns less the causal relation between anthropomorphism and advice-taking and instead it aims to understand how different types of advice (i.e., information-type and recommendation-type) provided by robots may influence people's trust in them where other conditions such as context-sensitivity can determine which advice is more acceptable for a population of older adults. Our research highlights potential outcomes and needs for further exploration of this topic, as well as the efficacy of robot-based advice in health contexts, to ensure that older adults are able to access reliable and trustworthy advice as healthcare demands increase.

2 The Present Study

2.1 Primary Aims and Research Questions

We aimed to explore older adults' *trust* and *compliance* with robot-based advice on non-prescription medicines (i.e., supplements, and vitamins). We did not make a direct comparison between human-based advice and robot-based advice. Instead, we measured trust and compliance with the type of advice given by the robot: *information-type* advice based on the instructions of a human expert and *recommendation-type* advice, where the robot's recommendation overreaches to the advice provided by the human. Our intent was to explore whether older adults would react with differential preferences to the two types of advice regarding their health. To this aim, we formulated the following research question:

RQ.1: Does the type of robot-based advice (information vs recommendation type) affect the trust of older adults in the robot on supplements?

The body of literature surveyed in Sect. 1 suggested that people's trust mainly depends on anthropomorphism and the perceived competence of the adviser. However, people's beliefs and preferences for robot's anthropomorphism are highly sensitive to the context, as context induces different expectations towards robots [45]. For example, in a previous pilot study, we found that a robot's anthropomorphism was less likely to affect the trust of older adults in the robot to help them with supplement intake [46]. Specifically, for this type of health-related scenario, older adults would attribute more significance to the reliability of the robot to perform the task, while the robot's social agency (e.g., social-psychological involvement) could not recover their trust when the robot committed an error. Given this motivation and considering that we explored the same context in this work, we focused mainly on the perceived competence of the robot as another desirable attribute for advisers but indirectly accounting for the effect of anthropomorphism, when specifically choosing a humanoid like NAO over other conversational machines. Our choice of adviser was also motivated by claims that an AI's *recommendations* are more effective when the AI is humanised [47], given that we also explored recommendation-type advice in our design.

Thus drawing onto the adviser's competence as our attribute, we intentionally aimed to influence how older adults' would perceive the competence and fit of the robot in the specific health-related task by involving human influence. We speculated that, in relatively sensitive tasks like advice on non-prescription medicine, the trust of older adults in the robot would be more resilient if the robot is labelled as an expert by a human specialist, which may result in higher compliance with its advice [28, 29]. Studies have shown that source expertise and credibility positively impact people's attitudes and increase their behavioural compliance [48, 49] (see [50] for a meta-analysis). Moreover, it has been shown that people also perceive non-human agents as such sources and that human-computer interaction is "unmediated and directly social" ([51], p. 687). In particular, human-robot interaction endows such characteristics that enable people to perceive the interaction with robots same as that between humans [52, 53]. In this sense, the social credibility of the robot as a source is similar to the credibility attributed to a human adviser, which, in turn, increases if the adviser is perceived as an official representative or expert in a given context [54].

To the best of our knowledge, this is the first attempt to explore the trust of older adults in a robot health adviser. We made the following hypothesis following literature insights

suggesting that information-type advice is often preferred [15]:

H1: Older adults who only received robot-based *information-type* advice on supplements would exert positive differences in trust compared to the pre-interaction baseline.

Moreover, we explored the strength of the following associations:

RQ. 2: To what degree is the older adults' trust in robot-based advice on supplements/vitamins associated with their willingness to use a robot in health facilities (such as pharmacies) in the future?

RQ. 3: To what degree is the older adults' trust in robot-based advice on supplements/vitamins associated with their willingness to use a robot at home in the future?

Finally, we aimed to gain a deeper insight into the compliance of older adults with the robot-recommended advice on supplements/vitamins. We consider compliance as a potential behavioural consequence of trust in the robot; hence, if they trust the robot's advice, they will be more receptive to its advice according to literature insights discussed earlier. Our following research hypothesis considers only participants who received *recommendation-type* advice from the robot and excludes those who only received *information-type* advice.

H2: The trust of older adults in the robot is likely accompanied by higher compliance with its advice, i.e., acceptance of the advice and low levels of negative feelings towards the advice received.

2.2 Exploratory (Secondary) Aims and Research Questions

A secondary aim of the present research study concerned the investigation of the *perceived trust in robots for prescription medicines*, including those for severe illnesses. We accounted that receiving recommendations from a robot only for supplements/vitamins could be perceived as having low-risk consequences in health contexts. Thus, we speculated that if the consequences of robot-based advice are negligible, such as in the case of non-prescription medicines, the trust of older adults in the robot might not change significantly before and after interacting with the robot, regardless of the type of advice given by the robot.

Maintaining the same experimental design, we have thus introduced another dependent variable concerning older people's trust in robot advice on prescription medicines. We decided to not simulate a scenario when prescription medicines were suggested by the robots for ethical reasons.

Thus, the results of our exploratory research questions should be read in light of participants experiencing only the scenario of robot recommendations on supplements/vitamins. Based on the anchoring effect [55], (i.e., where a decision is grounded on a more, or less related situation), we expected the close context of supplements/vitamins to prime people's response toward prescription medicine. The anchoring effect is a well-established psychological principle that states that people tend to rely on initial or anchor information as a reference point when making subsequent judgments for scenarios of similar nature [55]. Moreover, to the best of our knowledge, no empirical investigations on the trust of older adults in robots for prescription medicines were reported in previous work.

On these premises, we formulated the following exploratory research question: *Does the related experience of robot-based advice for supplements affect the trust of older adults in receiving robot-based advice for prescription medicines?*

To this aim, we explored RQ.1 to RQ3 considering medicines in lieu of supplements:

RQ.1': Does the type of robot-based advice (information vs recommendation type) affect the trust of older adults in the robot on prescription medicines?

RQ.2': To what degree is the older adults' trust in robot-based advice on prescription medicine associated with their willingness to use a robot in health facilities (such as pharmacies) in the future?

RQ.3': To what degree is the older adults' trust in robot-based advice on prescription medicine associated with their willingness to use a robot at home in the future?

3 Methodology

We conducted an interactive experiment between elder participants and a humanoid robot, which either informed or recommended non-prescription medicines (here vitamins, and over-the-counter supplements) to the older adult. We manipulated the type of robot advice as information-type advice, i.e., neutral without suggestions, and recommendation-type advice, i.e., robot-initiated suggestions in favour of an additional supplement not supplied by a fictional human dispenser at the start of the experiment (see Sect. 3.3.2).

3.1 Participants

The volunteers were recruited through purposive and snowball sampling with the help of the Plymouth Community Homes for a study on testing a robotic health supplement adviser. The general aim of the study was announced as “the

intention to understand their opinion on robotic use for quality ageing of older adults” and more detailed information on the experiment was provided when they agreed to participate in the study.

Excluding participants under 60 years old and those with severe cognitive impairments, we were able to recruit a total of $N = 30$ older adults for our study. The population involved 15 females and 15 males (range: 60–80 years; $\text{Mean}_{\text{age}} = 69.13$; $\text{SD}_{\text{age}} = 7.80$). Most participants held a university degree ($n = 18$), 5 had completed A-levels, and 7 had a primary/secondary education. 12 participants reported having had previous experience with robots, out of which 10 had been involved in our previous pilot study [46]. We used this measure to assign the participants evenly to the conditions while considering balancing the samples on the sex distribution. The final samples consisted of $n = 15$ for each of the designed experimental conditions (see Sect. 3.3).

3.2 Ethics Statement

Considering the sensitive nature of our pilot study, the ethical approval process enacted by the University of Plymouth required the use of the Plymouth Ethics Online System (PEOS). The experimental procedure was approved under the title ‘AGE IN Robot Home’ (project ID 3162) in November 2021 following the amendment of the documentation as per the recommendations of the Research Ethics and Integrity Committee. Ethical approval was granted for all the pilot studies relevant to the project throughout its entire duration. In turn, written informed consent was collected from all participants, including the publication rights of these case details.

3.3 Study Design

The experimental procedure involved a one-to-one interaction between an older adult participant and the NAO robot in a laboratory setup. The lab called Robot Home at the University of Plymouth is designed to resemble a real living room/home. NAO is an Aldebaran (former Softbank Robotics) anthropomorphic robot, extensively used for social engagement such as in healthcare and education applications [56, 57]. It is 22.6 inches tall, weighs 12.1 lb and has 25 degrees of freedom (DOF).

Our study design considered a health-related scenario, given the findings that older adults consent to tasks being carried out by robots only for specific contexts [58]. That said, trust in robot-based advice in a critical scenario, for example regarding health, might be substantially different from trust toward the robot's advice in a mild-severity interaction like gaming and entertainment.

3.3.1 Experimental Manipulation

The robot would give *information-type advice* (first condition) or *recommendation-type advice* (second condition) to the participant. The participants were blind to either condition. To allow for the difference in the robot-based type of advice, we used a cover story prior to the experiment. The participants were introduced to a human actor (different from the experimenter) who pretended to be a dispenser, i.e., an assistant to a licensed Pharmacist for the dispensing of pharmaceutical products to patients, but not a health professional in the UK health system. The intention to include a fictional dispenser was further motivated by studies which have shown that people are more likely to follow instructions given by an official representative in a scenario than an unofficial one [54]. We investigated two conditions (see Sect. 3.3.2 for details):

1. **Information-type advice condition:** the robot followed exactly the instructions of the dispenser, neutrally informing the participants only about the supplements provided by the human dispenser when they inquired about them, but not giving its own recommendations.
2. **Recommendation-type advice condition:** the robot followed the instructions of the dispenser, but at the end of the interaction it suggested an additional supplement to complement the health advice of the human, based on the participants' self-reported symptoms prior to the interaction.

3.3.2 Procedure

The participants were welcomed to the University by two researchers and the acting dispenser. We recorded their demographic data and asked them to read and sign the ethics consent form and a pre-interaction questionnaire. The participants were informed that their interaction with the robot would be video and voice-recorded at all stages for post-hoc analyses while guaranteeing their anonymity and data confidentiality.

First, our acting dispenser offered the participants a 'Well-being Questionnaire' and asked them to think about how they felt physically in the last seven days. The participants should select the category or categories of symptoms in the questionnaire if they experienced at least one of the symptoms in that category in the indicated period. For example, category 1 included symptoms like dry skin, blurry or decreased vision; category 2 irritability, confusion, lethargy, or fatigue; category 3 common cold or flu symptoms, etc. We also included an eighth category "None of the above" in case participants did not experience any symptoms or preferred not to answer. The participants were requested to indicate up to three categories of symptoms that they may have been experiencing. We assumed symptoms that are common among older adults

and can be tackled through beneficial nutrients like vitamins and minerals, found in a healthy diet and/or supplements. This information is publicly available in the NHS general health advice on Vitamins and Minerals at <https://www.nhs.uk/conditions/vitamins-and-minerals/>.

After, each individual participant was invited inside the Robot Home to familiarise themselves with the lab. They were greeted by the Softbank Robotics' humanoid robot Pepper, which described to the participant all the smart devices (e.g., the Google Home assistant, a robot Hoover, an automated plant system) and other robots present in the lab. This was done intentionally to relax the participants before the experiment and increase their awareness of the (robotic) technology to ideally avoid any feelings of inadequacy or anxiety when interacting with the robot. Given the similarity between the robot platforms of Pepper and NAO, the researcher gave helpful tips to the participant on how to interact with the robot, for example, to speak loudly and clearly, face the robot during the experiment, to allow it some time to respond to their questions or to repeat the questions if they believed the robot did not listen the first time. After the introductory demo with Pepper, the experimenter invited the participant to sit on the sofa in front of the NAO robot, as in Fig. 1.

Information-type advice condition: If the participants were assigned to the *information-type* advice condition, the dispenser would select *two* coloured boxes from a cabinet in the lab and place them on the table between the robot and the participant. Each coloured box represented a fictional over-the-counter supplement. The dispenser told the participant that they were recommending the health supplements based on the categories of symptoms that they had selected in the well-being questionnaire. In principle, each supplement was purposed for a particular category of symptoms, for example, vitamin A was recommended if participants selected symptoms of category 1, vitamin B Complex for symptoms in category 2, and so on. If the participants had selected category 8 "None of the above", the dispenser would give two default NHS-advised supplements for the general health of older adults, e.g., vitamin B12 and a (calcium + vitamin D) supplement. In turn, if participants had selected only one category, the dispenser would recommend a second supplement from category 8. The participants were intentionally not informed what the supplements were and for what purpose they were recommended. The dispenser told the participant that the robot was acting as a professional health adviser on these supplements, and they should seek its guidance on the use of the supplements, their benefits, the recommended dose, and diet. They also strictly instructed the participant to not swallow the supplements under no circumstance as they were used for simulation only.

Recommendation-type advice condition: If participants were enrolled in the *recommendation-type* advice condition,

the dispenser would follow the exact same routine, supplying only *two* coloured (supplement) boxes as in the first condition, asking the participant to interact with the robot for advice on the supplements. However, at the end of the interaction, i.e., when participants asked no more questions about the two supplements, the robot would ask the participants to retrieve another (third) supplement from the cabinet, indicating its colour. The robot informed the participant that it was recommending this additional supplement to them, based on the symptoms they had selected in the “Wellbeing Questionnaire”. The function of the robot-recommended supplements was true to the selected symptoms from the participants themselves, as these were recorded in a database accessible by the robot prior to the interaction. The robot offered to explain the supplement more in detail to the participant if they wished so. All participants in this group requested information on the robot-recommended supplement (manipulation check); nevertheless, their compliance with the robot’s recommendation was measured separately in the post-interaction questionnaire (see Sect. 5).

The experiment was started remotely by the experimenter in a separate room. The robot-participant interaction was not interfered with by any human presence (experimenter or dispenser) and the robot handled the session autonomously. The session was recorded using five GoPro 7 cameras and Sony audio recorders distributed around the lab at favourable angles. After the experiment, which was suitably interrupted naturally by the participant themselves if they wished to end the interaction, the participants were requested to fill in a post-interaction questionnaire. Finally, the researcher debriefed the participants on the overall experience and offered them the choice of withdrawing their data if they wished to do so.

4 Implementation of the Robot-Adviser

4.1 Robot Vision

To recognise the coloured boxes used in the experiment, the NAO robot was trained using YOLOv5 [59] on over 3 k labelled images and was further enhanced with colour recognition competence on 800 + shades. The recognised index identified and mapped uniquely onto one of the considered supplements (7 in total). During the experiment, the participant was instructed to show the fictional supplement box to the robot at least once when asking a question about it (Fig. 2). In turn, they were asked to show the supplement box they were inquiring about out of the available two on the table if they were switching between boxes. Once the robot had recognised the box and, hence, the type of supplement, it would remember this for the participant’s subsequent questions, until and unless the participant would change the box.



Fig. 1 Graphical representation of the Robot Home illustrating the experimental setup of the interaction between an older adult participant and the humanoid robot NAO (up). Image of the real laboratory (down)

The index of the recognised supplement was used alongside the spoken utterances of the participants to retrieve relevant information on the queried supplement using a Multimodal Conversational Agent (MCA), as explained in the next section.

4.2 Robot dialogue: Multimodal Conversational Agent (MCA)

As illustrated in Fig. 3, the robot invokes the *Multimodal Conversational Agent (MCA)* to interpret and maintain a conversation with the participants on the supplements. The MCA received two types of inputs: verbal (textual) input (transcripts of the spoken utterances produced by participants) and non-verbal input (the object recognition produced by the robot), which it contextualised within the current state of the conversation.

The MCA is characterised by a *retrieval-based* architecture [60]. It aims to find the responses of the highest relevance to the user inputs by exploring a certified knowledge base that contains candidate responses and by exploiting the dialogue context. Using only a certified and closed knowledge base ensures that no erroneous information about the entities of interest is generated, given that the answers are not extracted from the web or from similar uncertified sources, but inside of the intended domain. In turn, this prevents grammatical or insignificant errors in the produced answers.

Components: The *Multimodal Conversation Manager (MCM)* assumes the principal role in the MCA architecture

Fig. 2 Demonstration of the real experimental setup of the humanoid robot NAO (left) recognising the coloured fictional supplement boxes (right), by means of the interfaced YOLOv5 object recognition module

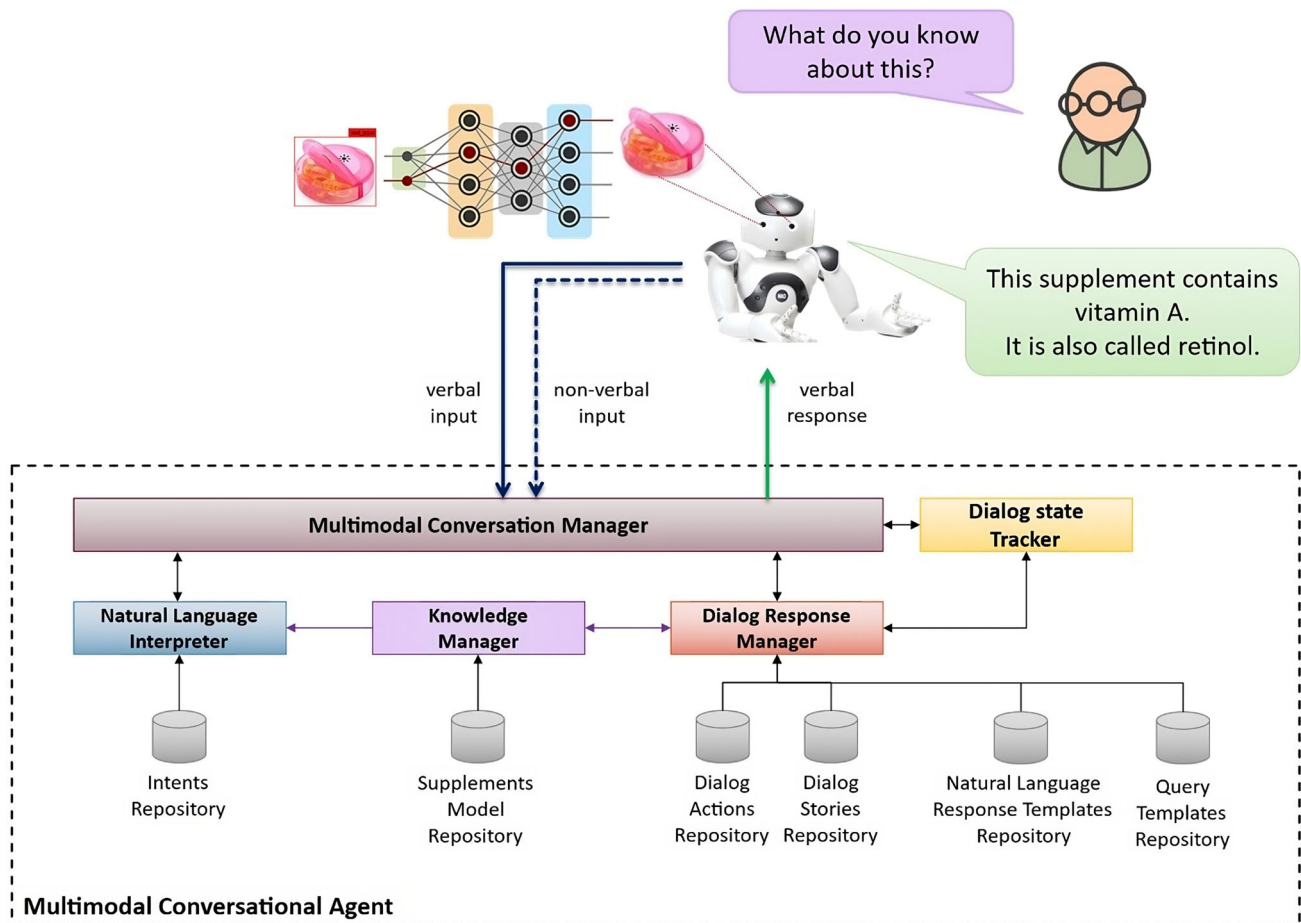
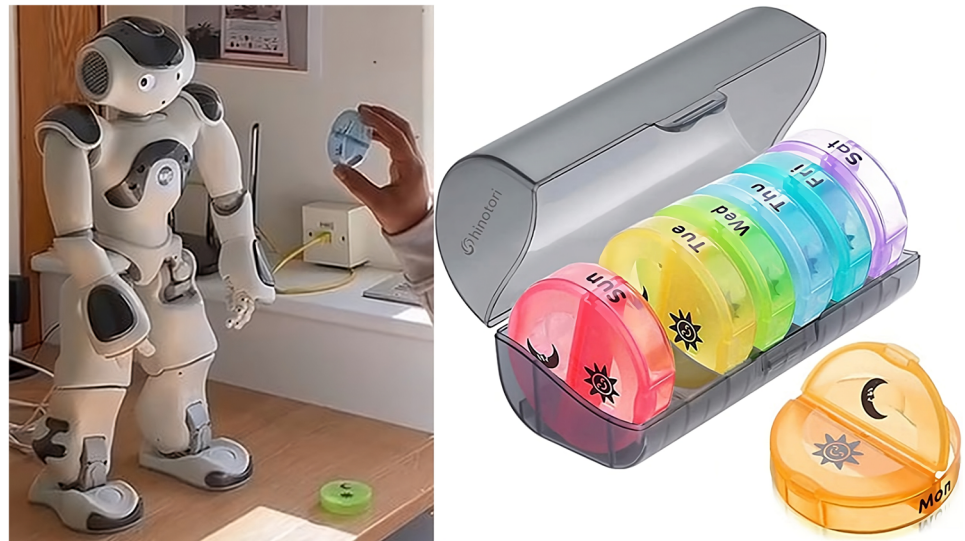


Fig. 3 The main components of the proposed autonomous robotic system coupled with the conversational agent and YOLO vision module

since it is responsible for orchestrating the information flow within the agent. In detail, when the user inputs (both verbal and non-verbal) are submitted to the agent, the MCM performs a set of conversational tasks by invoking (Fig. 3):

1. the *Natural Language Interpreter* (NLI) to interpret and understand the user verbal input to estimate the associated intent and extract any entity present in it;

2. the *Dialog State Tracker* (DST) to update the conversation state with the information about the detected user intent, the entities eventually mentioned by the user, the entity eventually showed to the robot by the user, and the information pertaining to their validation;
3. the *Dialog Response Manager* (DRM) to choose the most appropriate response based on the current conversation state and expected sequences of dialog steps, and next arrange the selected response in a natural language form that the user can easily understand.

Dataset and Training: The dataset used in this study has been elicited from our research team and domain experts, using the public information from NHS guidelines on supplements and vitamins. In the MCA model, this dataset is included under the *Supplements Model Repository*, which is handled by the *Domain Knowledge Handler* (DKH) illustrated in Fig. 3. This knowledge base has been formalised and encoded into a domain Knowledge Graph (KG), stored as RDF (Resource Description Framework, <https://www.w3.org/RDF/>) triples, and queried through the SPARQL language, recommended by the World Wide Web Consortium on RDF data sources.

The natural language inputs produced during the human–robot interaction are interpreted by the NLI component to estimate the associated intent and extract any entities eventually present (i.e., supplements, vitamins, or other domain entities of interest). To this aim, the NLI component has been designed to jointly perform both tasks, widely referred to as intent classification and slot filling, through a BERT-based [61] deep learning model architecture. This model has been trained with a data set composed of a list of textual inputs and questions, and the corresponding intent assigned by a human to each of them. Further details on this model can be found in [62].

For our study aims, three main types of intents have been foreseen:

- "dialog command" intents (such as *greetings*, *thanks*, *repeat_request*) model user commands for the robot to manage the current dialog state, whose response is pre-defined and unrelated to domain knowledge. To illustrate, the verbal input "please repeat" can be used to ask the robot to repeat the last answer/question generated. In contrast, the user input "forget it" allows aborting the current information request, and so on;
- "information request" intents (such as *description_request*, *functions_request*) model user questions addressed to the robot on at least one of the domain entities, for example, "What do you know about retinol?", whose proper answers are extracted from the results of specific queries over the knowledge base;

- "information completion" intents model user answers to questions posed by the robot for asking missing information required to satisfy a previous user request. For instance, in response to an incomplete question formulated by the user, such as "what are its benefits?" without showing the target or if object recognition fails, the robot will respond by asking the user the missing entity of interest, such as, for instance, "which supplement are you referring to?". The subsequent user input containing the missing information can be explicit, for instance, the entity "red box", or implicit, for example, "this one," and simultaneously show the robot the supplement of interest. Both these user verbal inputs are classified as "information completion" intents.

Each textual sample has been annotated with a unique intent label. In contrast, each of its tokens is tagged according to the IOB (Inside Outside Beginning) format depending on its membership to an entity of interest, included in the ontological model or not. In detail, each token is tagged with "B" when it is the beginning of a known entity, with "I" if it is inside a known entity, or with "O" if it is outside of any known entities. "B" and "I" tags are followed by a suffix representing the class/type of entity identified. A textual sample from the NLI dataset is reported in Table 1.

Once the user's verbal input has been interpreted, the MCM updates the dialogue state with the information about (1) the detected intent of the user (e.g., name of the supplement, benefits, dose, or diet), (2) the entities mentioned by the user if present (e.g., supplement, vitamin), and (3) the nonverbal input that is submitted eventually (i.e., the entity showed to the robot by the user). Next, the MCM invokes the DRM to firstly determine the proper dialogue response based on the status of the conversation and the expected dialogue flow, and, next, to arrange the selected response in a natural language form that the user can easily understand. To this aim, first, the DRM selects the most appropriate procedure (typically known as action) to be executed to fulfil a user's request. The possible actions are stored in the *Dialog Actions Repository*, which maintains their label names and the associated codes for the execution. A total of 10 actions have been implemented, with one action built for each defined intent. The DRM determines the most appropriate action through a supervised action selection model relying on an existing transformer-based architecture named TED (Transformer Embedding Dialogue) [63] to encode multi-turn dialogues and select the most appropriate action to perform based on the current dialogue flow.

Our model was trained with a data set composed of 20 dialog flows (typically known as story), annotated with action labels for each of their turns. The possible dialog flows are stored in the *Dialog Stories Repository*. For instance, in the following, a dialog flow is reported consisting of a

Table 1 A textual sample from the Natural Language Interpreter (NLI) dataset, annotated with its intent label (description_request), tokenised and tagged using IOB formatting

Text	Intent	Tokens	Tag
What is the orange supplement for?	description_request	What	O
		is	O
		the	O
		orange	B-Supplement
		supplement	I-Supplement
		for	O
		?	O

single intent-action exchange between the user and the system, which describes the case when the user submits an input text classified as a *dose_request* intent, i.e. “*Can you tell me the dose of retinol?*”, the most appropriate action *action_what_dose_request*, i.e., a query for the *Supplements Model Repository* sent through DKH, should be generated to arrange the correct answer:

<i>Can you tell me the dose of retinol?</i>	← the user text input
story_dose_request_001	← the name of the dialog story
<i>dose_request</i>	← the intent starting the dialog story
<i>action_what_dose_request</i>	← the desired action

When the most appropriate action has been selected, the DRM executes it to generate the associated dialog response. A dialog action may require or not the generation of a natural language response. When response-generating actions do not involve intents mentioned by the user in the conversation, e.g., greeting, error, thanksgiving messages, predefined message templates are used, which are stored in the *Natural Language Response Templates Repository*. These templates are composed of a list of different messages sharing the same information content. For example, to inform the user that the input received has not been understood, a natural language response is generated by randomly extracting one message from the list of predefined messages associated with the information content *Lack of understanding*, of which an excerpt is reported in Table 2:

In the case of response-generating actions that involve entities mentioned by the user (participant), the successful retrieval of these entities from the dialog state is a precondition to generating a proper response. This kind of actions is associated with a SPARQL (SPARQL Protocol and RDF Query Language) query template, which uses the entities retrieved from the dialog state for querying the *Supplements Model Repository* via DKM, and the result of this query is

Table 2 An excerpt of the predefined template to manage information content “Lack of understanding”

Information content	Message templates
	I don’t think I understand, can you rephrase it?
	I don’t understand, rephrase it
Lack of understanding	What? Repeat please
	Can you repeat please
	...

used to build the final response of the action. These query templates are stored in the *Query Template Repository*.

5 Measures

We collected the data from the pre- and post-interaction questionnaires. The self-reported measures were used for our post-hoc quantitative analysis.

5.1 Perceived Trust

We measured the general perceived trust of the participants in the pre-interaction questionnaire using two items as a baseline in a 7-point Likert agreement scale, where 1 indicated strong disagreement and 7 indicated strong agreement:

1. *Primary aim*: I would TRUST a robot if it gave me advice about health supplements/vitamins—**trust for supplements/vitamins**
2. *Secondary aim*: I would TRUST a robot if it gave me advice on my overall medication plan (including medication for severe illness)—**trust for prescription medicine**

These items were queried again in the post-interaction questionnaire to measure the participants’ perceived trust immediately after interacting with the robot. The *perceived trust* measure was used to address the research question RQ.1 and research hypothesis H.1 as our primary aim, and the anchoring effect on trust for our secondary aim (RQ. 1’).

Thus, item 2 (trust for prescription medicines) was only measured indirectly, given our initial assumption that differences in the participants' perceived trust for supplements before and after the experiment could potentially generalise to differences in the perceived trust for medicines.

5.2 Intention to Use Robots

The *intention to use robots* was measured as the willingness of older adults to use the robot in prospect in health facilities and/or at home. To this aim, we designed two items ex-Novo in the post-interaction questionnaire. The participants were asked to indicate their level of agreement on a 7-point Likert agreement scale:

1. *How likely are you to use a robot in your home in the next 12 months if it was available free of economic cost?*
2. *How likely are you to use a robot in health facilities (such as pharmacies) in the next 12 months if it was available?*

This measure was used to assess the associations between the level of trust in the robot with the proxy variable of actual use of robots at home or in facilities, which addresses the research questions RQ.2 and RQ. 3, and the related exploratory research questions.

5.3 Compliance

Participants' compliance with the robot-recommended supplement was measured for participants in the recommendation-type advice condition and was determined from the post-interaction questionnaire using the two following items:

1. *I would have taken the suggested supplement/vitamin if these were true*
2. *I felt uncomfortable, robots should NOT give any advice regarding supplements/vitamins*

To assess the level of agreement between trust specifically in the recommendation-type advice and compliance with this advice (i.e., actual acceptance of the advice; negative feelings after the advice), we included the following item regarding trust:

3. *I felt I could TRUST its advice*

The above items were measured using a 7-point Likert scale. This measure was used to address hypothesis H.2. We considered compliance as a potential behavioural consequence of trust in the robot.

5.4 Other Measures

Participants were additionally asked to indicate their level of agreement with the following item in the post-interaction questionnaire: *I found it easy to interact with the robot*. We included this measure to preliminary verify if the experimental conditions (information-type advice, recommendation-type advice) influenced the participants' perceived easiness of interacting with the robot, considering that this could result in potential differences in trust in the robot between groups, which must be accounted for in the post-hoc analysis of the results.

6 Data Analysis

The examination of the Shapiro–Wilk test and the descriptive statistic revealed that our data was not normally distributed.

For our repeated measures design, we applied Linear Mixed-Effects Regression Modelling (LMER) which includes random effects that can account for individual variation, and fixed effects, which explain variance across individuals. LMER is a more robust analytical approach compared to traditional linear modelling [64, 65] and is robust to violation of assumptions [66]. Moreover, we corroborated our results using non-parametric permutation mixed-ANOVAs [67]. This test does not rely on assumptions about the distribution of the data and performs better than traditional ANOVAs with small sample sizes [68].

We preliminary found that correlations between socio-demographics (*age, gender, education*) and the dependent variable (*perceived trust*) were weak and not significant. Consequently, we did not include these as covariates in the model to preserve statistical power. Moreover, our **other measure** of *perceived easiness of interacting with the robot* revealed that its effect on trust did not depend on the experimental group, i.e., participants of both groups perceived the robot as equally easy to interact with. Before our analysis, we verified that the experimental groups showed no differences in the overall level of trust in robots prior to the interaction, using non-parametric permutation ANOVAs.

All analyses were carried out within the R environment (4.2.1, 2020-06-23; R Core Team, 2020). The package *lme4* [69] was used for the LMER, the package *permuco* [68] for the permutation test, and the package *emmeans* [69] for pairwise comparisons. Correlations were explored using Spearman's Rank correlation coefficient (RQ. 2 and RQ. 3).

7 Results

7.1 Primary Results

Trust in the Robot-Based Advice for Supplements/Vitamins (RQ1 And H1)

For the Mixed-Effect Models, *time* and *group* were treated as fixed effects, while *factor ID* as random effects (i.e., random intercept of subject). The analyses concerned nested model comparisons against the null model (i.e., intercept-only model) based on chi-square difference tests (i) and the results of the regression model (ii).

First, we tested whether adding the fixed effect of *time* to the null model would significantly improve the model fit. The results of this analysis showed a non-significant improvement ($\chi^2(1) = 0.1307, p = 0.72$). The null model was then compared to a model that included the fixed effect of *group*. Neither was this model a better fit to the data than the null model ($\chi^2(1) = 0.2661, p = 0.61$). Next, the null model was compared to a model that included both main effects (*time* + *group*, $\chi^2(2) = 0.3968, p = 0.82$) and to a model additionally including the interaction term (*time*group*, $\chi^2(1) = 1.1187, p = 0.77$). Both models did not improve the model fit significantly.

Our research hypothesis H1 concerned the interaction effect between *group* and *time* on *trust in robots for supplements*. We expect that only participants in the information-based condition would trust the robot's advice for supplements more after interacting with the robot, compared to not interacting with the robot (baseline). Therefore, in this hypothesis we did not expect differences in time for the participants in the second group (recommendation-type advice) compared to the participants enrolled in the first group (information-type advice). That is, trust in the robot for supplements after the interaction would not change for participants in the recommendation-type condition.

Our results of the mixed design, including fixed and random effects, confirmed the following pattern: no difference between groups ($\beta = -0.47, CI [-1.56 \text{ to } 0.63], p = 0.396$), between times ($\beta = -0.60, CI [-2.39 \text{ to } 1.19], p = 0.505$), and no interaction ($\beta = 0.47, CI [-0.67 \text{ to } 1.60], p = 0.413$). We examined participants' ratings of the dependent variable of trust across Time 1 (pre-interaction) and Time 2 (post-interaction) and between groups (i.e., main effects) but not including the interaction term. Also in this model, the main effects of *group* ($\beta = -0.23, CI [-1.17 \text{ to } 0.70], p = 0.619$) and *time* ($\beta = 0.10, CI [-0.46 \text{ to } 0.66], p = 0.723$) remained non-significant. Data visualisation of these results in Fig. 4 confirmed that participants in the *first condition* (i.e., participants that interacted with a robot which strictly followed the human prescription of the supplements) did not differ in their level of *trust in robots for supplements* compared to participants in the *second condition* (i.e., participants that interacted

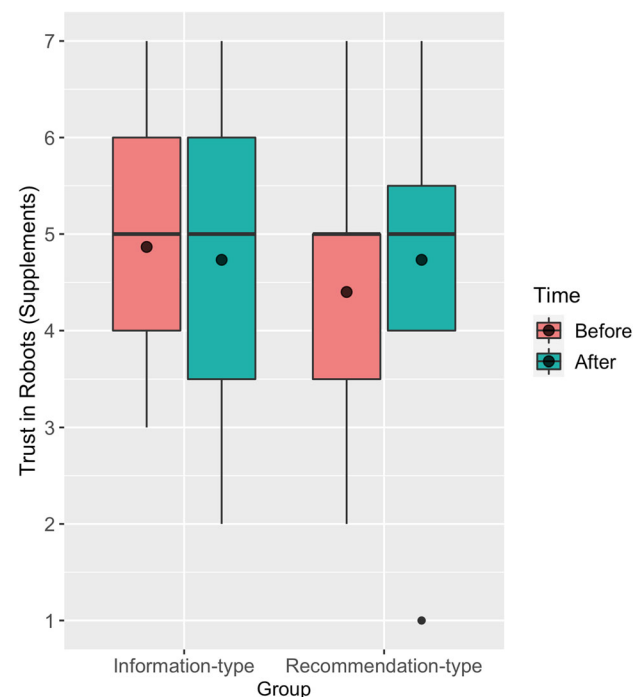


Fig. 4 Boxplot distributions of the interaction between group and time on trust in the robot for supplements. Each box represents data between the 25th and 75th percentiles. The boxplot shows the median (black horizontal lines inside each box), and limits of the interquartile range (lower and upper hinges of the boxes, respectively). The dots inside the boxes represent the mean values of each distribution

with a robot that made self-recommendations to complement the human's advice with additional supplements). Our power analysis revealed that the power for the effects considered was limited, with estimates falling below the commonly accepted 0.80 threshold. Nevertheless, the graphical inspection in Fig. 4 showed no tendency toward any meaningful difference between conditions, which might suggest that one should replicate the same null effects even at larger sample sizes.

Given that differences between *times* within *groups* did not emerge, our hypothesis H1 is partially supported: trust in the robot's advice for supplements remained somewhat high even when the robot *recommended* more supplements to the participants than the human dispenser supplied. However, no positive differences in trust (i.e., higher levels of trust) were observed when the robot simply informed participants about the supplements, in contrast to our hypothesis.

Compliance (H2)

We aimed to explore the compliance with the robot-recommended advice on supplements of the participants enrolled in the *second group*. Therefore, the following descriptive data concerns 15 participants only.

As illustrated in Fig. 5, the average level of agreement toward *trust for the advice received* ("I felt that I could

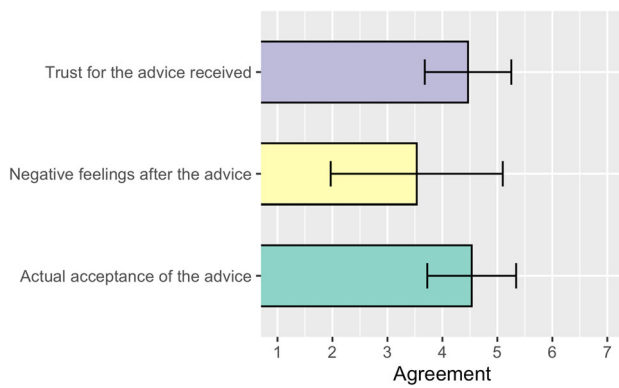


Fig. 5 Descriptive mean differences concerning participants' attitudes and feelings toward the robot of the group within the recommendation-type advice condition. Error bars represent 95% confidence intervals

TRUST its advice”) and the *actual acceptance of the advice* (“I would have taken the suggested supplement/vitamin if these were true”) could be interpreted as moderated, meaning beyond the central point of the scale. In contrast, participants indicated, on average, low levels of negative feelings as represented by the bar plot *negative feelings after the advice* (“I felt uncomfortable, robots should NOT give any advice regarding supplements/vitamins”), although with higher variation in response compared to the other measures. Although we did not measure a direct relationship between trust and compliance due to the relatively small sample size ($n = 15$), our results indicate that compliance with the robot-based advice follows a similar agreement pattern with their trust in the robot. More specifically, the recipient of the advice follows a positive trend with higher trust, with associated low levels of negative feelings about the advice. These results confirm our hypothesis H2: that trust is accompanied by higher levels of compliance with the advice, even when this advice arises from the robot’s self-recommendations. This is in line with other findings that people would comply more with the instructions of machines, which receive more of their trust [24].

7.2 Secondary Results

Trust in the Robot-Based Advice For Prescription Medicines (RQ1’)

This model included *trust in robots for prescription medicines* as the dependent variable. We applied the same data analysis procedure used for the *trust in robots for non-prescription medicine*.

When comparing the null model against the model that included the fixed effect *time*, we found a significant improvement in the model fit ($\chi^2(1) = 5.9264$, $p = 0.015$). Instead, the model that included the fixed effect of *group* was not a better fit to the data than the null model ($\chi^2(1) = 0.0675$, p

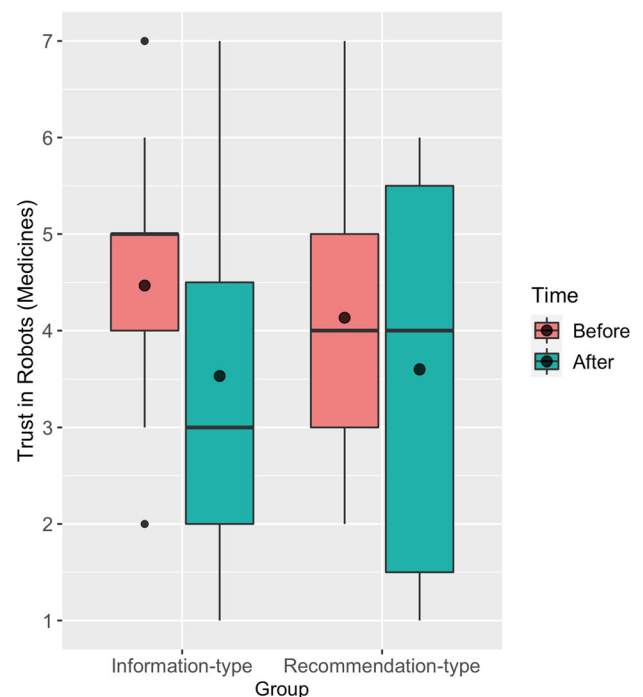


Fig. 6 Boxplot distributions of the interaction between group and time on trust in the robot for medicines. Each box represents data between the 25th and 75th percentiles. The boxplot shows: the median (black horizontal lines inside each box), limits of the interquartile range (lower and upper hinges of the boxes, respectively). The dots inside the boxes represent the mean values of each distribution

$= 0.80$). A significant improvement to the null model was found when comparing the model that included both main effects ($\text{time} + \text{group}$, $\chi^2(2) = 5.9939$, $p = 0.049$) but not in the model that included the interaction term ($\text{time} * \text{group}$, $\chi^2(2) = 6.4853$, $p = 0.090$). Consistently with the nested model comparisons, the model that included the interaction between the fixed effect *time* and *group*, accounting for the random effects, showed no difference between groups ($\beta = -0.33$, CI $[-1.55 \text{ to } 0.88]$, $p = 0.585$), between times ($\beta = -1.33$, CI $[-3.20 \text{ to } 0.53]$, $p = 0.157$), and no interaction effect ($\beta = 0.40$, CI $[-0.78 \text{ to } 1.58]$, $p = 0.499$). However, the model including only the main effects showed a significant effect of *time* on *trust in robot for medicines* ($\beta = -0.73$, CI $[-1.32 \text{ to } -0.15]$, $p = 0.015$) but the fixed effect *group* remained not significant ($\beta = -0.13$, CI $[-1.20 \text{ to } 0.93]$, $p = 0.803$), $R^2_{\text{marginal}} = 0.049$; $R^2_{\text{conditional}} = 0.560$.¹

These results were corroborated by non-parametric permutation mixed-ANOVAs, using 6000 Monte Carlo iterations, confirming thus the accuracy of our results. That is, the main effect of *time* on *trust in robots for medicines*

¹ Marginal R^2 represents the variance explained by the fixed effects; Conditional R^2 represents the variance explained by both the fixed and the random effects.

($F(6.33394)$) was significant when considering both the parametric p -value (0.018) and the resampled p -value (0.015) in the model that did not include the interaction term. Nevertheless, the interaction plot in Fig. 6 showed that the interaction effect might in fact have occurred.

It is possible that the decreased degrees of freedom due to the additional interaction term might have negatively impacted the statistical power to detect the interaction considering the small sample size. In fact, the power analysis showed that estimates fell below the 0.80 threshold also in this model. Therefore, given the exploratory nature of our pilot study, we proceeded with further analyses looking at specific contrasts: levels of *group* within each level of *time*. To this aim, we performed contrasts for the estimated marginal means (EMMs) extracted from the model that included the interaction term, using Dunnett-style comparisons (Fig. 7).

Results of this post-hoc test showed that participant in the first group (*information-type advice*) were less likely to trust the robot for medication prescription after interacting with the robot ($M = 3.53$, $SE = 0.429$, 95% CI [2.67–4.40]), in comparison to the baseline ($M = 4.47$, $SE = 0.429$, 95% CI [3.60–5.33]), ($\beta = 0.933$, $SE = 0.416$; $t(28) = 2.244$, $p = 0.033$). However, no significant difference in trust was found for participants of the second group (*recommendation-type advice*) after interacting with the robot ($\beta = 0.533$, $SE = 0.416$; $t(28) = 1.282$, $p = 0.210$). These results highlighted a type of protective role of robot-based recommendations on the *trust of older adults in robots for prescription medicines*. More simply, when participants experienced robot-based advice for supplements, we found that this experience negatively influenced their trust for medicine when the robot simply informed them about supplements (i.e., information-type). As shown in Fig. 6, the trust of older adults in robot-based advice for medicines was medium with values around the central point of the scale before the interaction with the robot. After receiving information-type advice on supplements, their trust in receiving the same type of advice on medicines was lower, with values below the central point of the scale. This might indicate that trusting a robot with information-type advice on supplements is not sufficient to invoke trust in the robot for more sensitive tasks such as prescription medicines for more severe health conditions. However, when the robot recommended further supplements (i.e., recommendation-type), the level of trust in the robot did not change compared to the baseline (pre-interaction with the robot). We speculate that older adults might have a better overall perception of a proactive robot (i.e., one that gives recommendations), compared to a reactive one (i.e., which strictly follows instructions) in giving useful advice in health-related contexts, which resulted in somewhat resilience in their trust.

7.3 Associations (Primary and Secondary Results)

Intention to use Robots (RQ2, RQ3 and RQ2', RQ3')

Our research questions also concerned the strength of the association between—trust toward robots for *supplements* and the closer proxy of actual behaviour, i.e., *intention to use robots at home* and *intention to use robots in facilities* (e.g., pharmacies) in the future. As shown in Fig. 8, Spearman's Rank correlation coefficients revealed the strongest association between the *trust in robots for medicines* and *intention to use robots in facilities*, while the lowest association was found between *trust in robots for supplements* and *intention to use robots at home*. All correlations were significant, at $p \leq .001$ except for Trust (Supplements)—Home use ($p = 0.008$) and for Trust (Medicines)—Home use ($p = 0.001$).

The results indicate two main outcomes:

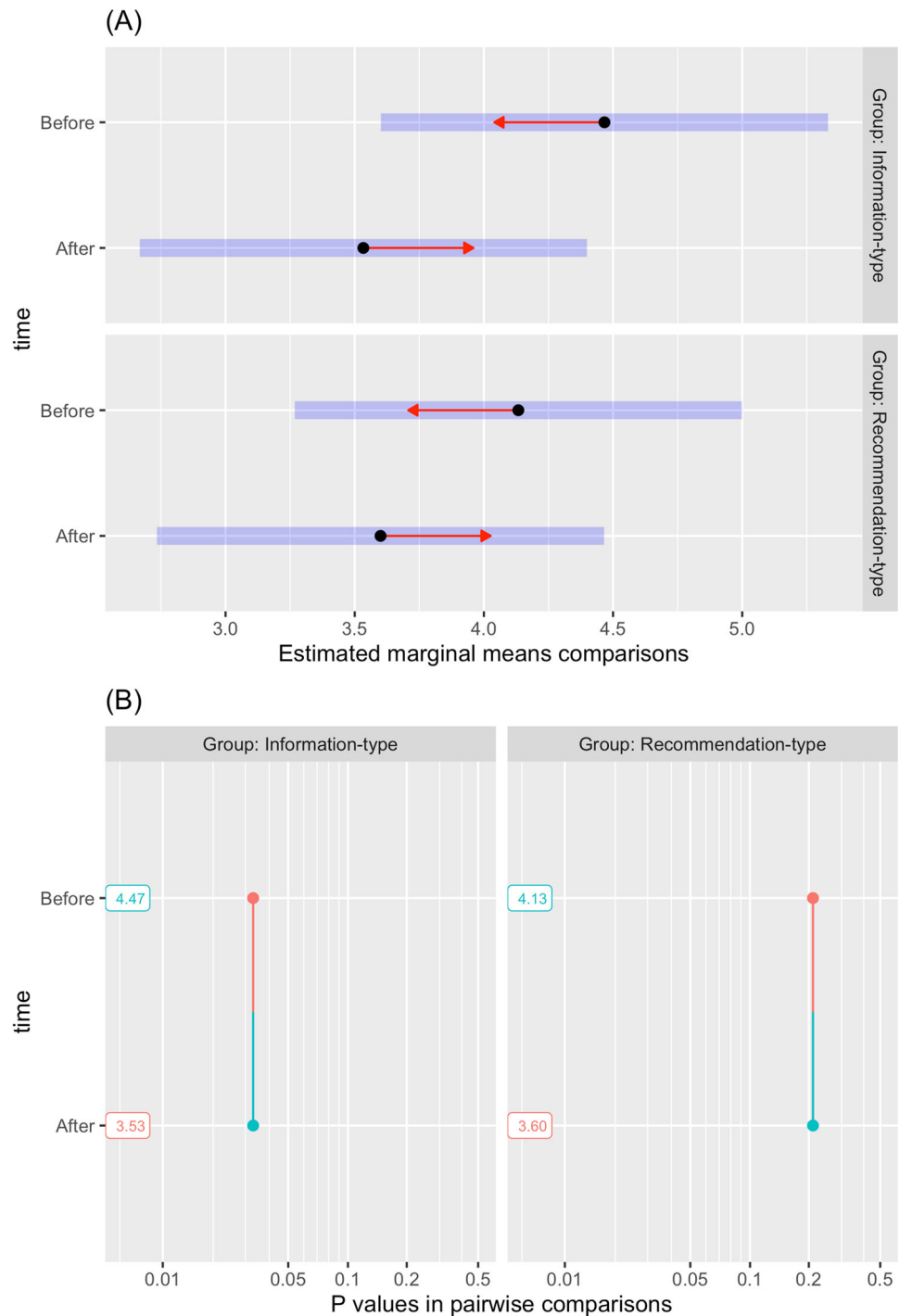
1. Given a certain level of trust in the robot-based advice for supplements, its association with the intention of older adults to use robots in a human-controlled environment, such as a health facility is higher (0.64) than that of welcoming robots in their own homes (0.45). Similarly, the correlation between trust in robot-based advice for medicines and intention to use robots in facilities (0.73) was stronger than its correlation with the intention to use robots at home (0.56). The strength of these associations can be interpreted as: low < 0.6 ; moderate ≥ 0.6 ; strong ≥ 0.8 .
2. When older adults trust the robot in a more sensitive task, such as medication for severe illness compared to milder health supplements, their intention to use a robot in their homes is higher. Instead, trusting a robot in a task with rather negligible consequences of robot-based advice, such as in the case of supplements, does not strongly indicate that older adults intend to use robots at home (lowest association, Fig. 8).

These results should be read with caution since we have simply explored bivariate correlations. However, they may provide useful insights for future studies when planning the estimation of the causal relationship between the variables.

8 Discussion

The aim of the present study was to investigate whether older adults would trust a humanoid robot giving them advice on over-the-counter supplements and vitamins. From literature insights, we identified two desirable attributes of advisers: anthropomorphism and perceived competence (or task fit). Here, we intended not to investigate the effect of either attribute on people's advice-taking behaviour. Instead,

Fig. 7 Plot A shows the estimated marginal means (EMMs) comparisons in two groups, as a function of time. The degree to which arrows overlap reflects the significance of the comparison of the two estimates; the bars represent the confidence interval of the EMMs. Plot B shows a visual representation of the *p*-values in the pairwise comparisons test. Note. Since we used a balanced design and the model did not include covariates, marginal means based on the statistical model are equal to the descriptive means



we aimed to explore the preference for the type of robot-based advice in health-related contexts (information-type vs. recommendation type), assuming that the sensitivity of the context would play a role in which advice was accepted and trusted more. However, we accounted for the anthropomorphism and competence of the adviser to strengthen the adviser–advisee synergy, thus we used a humanoid robot as an adviser (i.e., anthropomorphism) and labelled it as a “trusted

expert” by a human (i.e., competence). For the latter, we used a cover story of a fictional human dispenser who gave the participants a wellbeing questionnaire to complete before the experiment, based on which the dispenser recommended two health supplements. The participants were blind to the choice of supplements and were asked to consult the robot about their use. Our intention to include a fictional dispenser was motivated by studies suggesting that people are more likely to

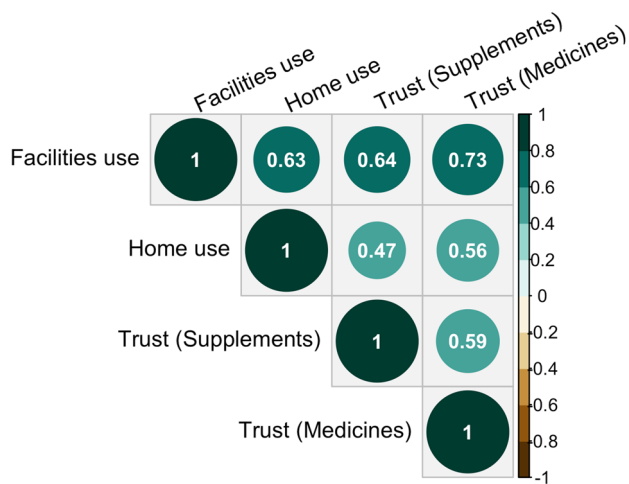


Fig. 8 Spearman's Rank correlation matrix. The size of the circles indicates the strength of the association between the variables; coefficients are also included inside the circles (refer to the text for the corresponding p -values)

follow instructions given by an official representative than an unofficial one [54]. We hypothesised that participants would prefer to receive only *information-type* advice from the robot on the supplements, trusting the robot more when it simply described the supplements given by the dispenser and answer questions on them. We expected this pattern to change when the robot gave *recommendation-type* advice, i.e., it gave advice in favour of an additional supplement not provided by the dispenser, based on the participants' indicated health conditions on the wellbeing questionnaire. Specifically, we expected higher trust in the robot for the participants that received information-type advice compared to the baseline (before interacting with the robot) and no changes in trust for participants that received recommendation-type advice. This hypothesis followed theoretical literature insights that information-type advice is often preferred [15] and given the sensitive nature of the experimented task. To understand people's response to recommendation-type advice, we also measured their post-decision compliance with the advice considering it as a potential behavioural consequence of their trust in the robot. We expected that higher trust would be accompanied by higher levels of compliance. Our exploratory study represents a first attempt to understand the perception, trust and intention to use robot advisers (a) in an older adult population and (b) in a close-to real-life health context with more critical implications.

The scope of our study was not to directly control for human-based advice. Instead, we compared a proactive versus a reactive robot behaviour mediated indirectly by the human factor, i.e., the information-type advice condition implied a condition in which the robot explicitly followed instructions from a human dispenser, whereas the recommendation-type advice mismatched that of the

human. Our experimental findings showed that the robot-based advice was equally trusted by older adults whether the robot simply informed them about the supplements or recommended additional supplements. The preliminary non-parametric permutation ANOVAs revealed that the groups showed no differences in the overall level of trust in robots prior to the interaction. Moreover, we also verified that the differences in the trust variable between the two experimental groups were not indirectly affected by the participants' perceived easiness of interacting with the robot, i.e., participants of both groups perceived the robot as equally easy to interact with. Hence, our hypothesis H1 was only partially supported. Our analysis revealed that the trust of older adults in the robot before the interaction was moderate and remained such even after they received advice from the robot on health supplements. While we cannot affirm with certainty that older adults manifest a generally positive attitude in receiving robot-based advice, it appears that they might not necessarily ill-favour it. In particular, when measuring their compliance with the recommendation-type advice, we found low levels of negative feelings after the advice, although with higher variation in response compared to the *actual acceptance of the advice* measure. The latter produced results beyond the central point of the scale, indicating that older adults may be willing to take and use the advice of the robot, even when it arises from the robot's self-recommendations. These results suggest that higher trust is accompanied by higher levels of compliance, as per our hypothesis 2. We exclude that any cognitive anchoring bias (overreliance on initial information offered) might have occurred, given that participants were not given a choice of adviser to run the experiment with [71], i.e., compliance arose from the experimented conditions. We used a realistic laboratory experiment, having participants interact with a real robot adviser. This is important given that authentic human–robot interactions can lead to differences in perceptions of the robot adviser's characteristics, for example, emotional trust or expertise [35].

Here we considered a scenario involving health-related advice using supplements and vitamins. While certainly, this is a relatively sensitive task, the investigation of other tasks of higher complexity and sensitive characteristics would reveal useful and solid conclusions. Task difficulty influences the utilisation of advice and perceived expertise (both self and that of the adviser) [16, 72]. To this aim, we sought to investigate whether there would be any consequential behaviours for older adults in the context of a similar but more complex task, such as robot-based advice on medication including those for severe illnesses. Although we did not experiment with this task directly given ethical reasons, we expected the close context of supplements/vitamins to prime people's response toward prescription medicine given an anchoring effect [55]. In the lack of empirical investigations, our study may provide some predictions in this context. The pattern of our results

suggested that indeed task difficulty might influence trust in the advice received (and/or in the adviser's characteristics) and so does the type of advice, consistent with previous literature. For example, after receiving information-type advice on supplements, participants' trust in the robot-based advice for medicines dropped. In contrast, the trust of participants who received recommendation-type advice on supplements was not negatively affected compared to the pre-interaction baseline, remaining beyond the central point of the scale. We speculate that receiving further recommendations from the robot might have led to a better perception of the robot adviser's expertise, which resulted in trust resilience or some type of protective role of trust. It is worth noting, however, that our insights were drawn from the inference that individuals may rely on their experience with the robot's suggestion of vitamins to evaluate their trust in its ability to prescribe medicine. It is possible that this information may not be the sole determining factor in their trust, which requires further careful investigations. Moreover, our results confirm literature insights that task difficulty and type require significantly different adviser characteristics to fit that task. Hence, robot-based advisers should be studied extensively and empirically in several meaningful contexts before making conclusive assumptions about their use. Nonhuman advisers in other applications, such as domestic, social, health, or military must be informed based on contextual factors that influence compliance with them [73].

Another finding of our study was that older adults are generally willing to use robots in the future at home or in health facilities. With trust being positively associated with the intention to use robots, these findings strengthen our previous measures of trust in robot-based advice. Moreover, they provide insights for future studies to study the causal relationship between the two variables. Our results indicated once more that the type of task is important even in people's intention to use robots. For example, trusting robot-based advice for supplements is of lesser reason to use robots domestically, compared to trusting a robot for more sensitive tasks like medicines, which increases people's intention to use one at home. However, older adults would yet slightly prefer to use robot advisers in a more controlled environment like health facilities. The results should be read with caution since we only explored bivariate correlations. Our findings offer valuable insights into the relationship between trust in older adults and the use of robots for health-related advice. However, to further validate these results, it is necessary to conduct additional research with larger sample sizes and higher statistical power.

In conclusion, our exploratory study showed that the older population is willing to receive health advice from robots when these have low-risk consequences, while more studies are needed to clarify what would happen when high-risk consequences are involved. Our findings are promising and

motivate conducting larger studies to understand in which way high-risk advice should be communicated by the robot to increase the trust and compliance of older adults. Moreover, future studies could explore the differences between the type of advice (i.e., information-based vs recommendation-based) and type of adviser (robot vs human) in health-related contexts within older adults. This would give a better understanding of the role of robots in the trust toward health-related advice. For example, in future work, we aim to investigate other experimental conditions, such as comparing against multiple robot advisers of different characteristics, like degree of anthropomorphism, personality traits, level of confidence, level of expertise and competence. In turn, the results of this study should be evaluated more broadly in a non-vulnerable population (i.e., adults) and experiment with real-world conditions for high-sensitivity contexts. Finally, in our study we presumed that labelling the robot as expert by an official human representative (i.e., dispenser) would impact its perceived competence and credibility. Thus, it is not clear how our participants would respond to robot-based advice if the robot was not labelled as expert by humans. In this sense, future work should fill this gap by investigating the role of the robot title in robot-based advice.

Funding The work described in this paper was funded by the Interreg 2 Seas Mers Zeeën AGE'In project (2S05-014).

Data Availability This manuscript has no associated data.

Declarations

Conflict of interest The authors have no competing interests to disclose. The contents of the manuscript have been unanimously agreed by all co-authors. We certify that the submission is original work and not submitted for consideration elsewhere.

Ethical Approval This study involved human subjects in its research. Written consent was taken from all participants. The study was conducted in accordance with the Declaration of Helsinki and approved by the Faculty Research Ethics and Integrity Committee of the University of Plymouth (Project ID 3162; date of approval; 30 August 2021).

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

References

- Abdi J, Al-Hindawi A, Ng T, Vizcaychipi MP (2018) Scoping review on the use of socially assistive robot technology in elderly care. *BMJ Open* 8(2):e018815
- Henschel A, Laban G, Cross ES (2021) What makes a robot social? a review of social robots from science fiction to a home or hospital near you. *Curr Robot Reports* 2(1):9–19
- United Nations Department of Economic and Social Affairs, P.D. World Population Prospects 2022: Summary of Results; United Nations Department of Economic and Social Affairs, P.D.: New York, NY, USA, 2022. Available at: <https://www.un.org/development/desa/pd/content/World-Population-Prospects-2022>
- de Graaf MM, Ben Allouch S, Van Dijk JA (2019) Why would I use this in my home? A model of domestic social robot acceptance. *Hum Comput Interact* 34(2):115–173
- Piçarra N, Giger JC (2018) Predicting intention to work with social robots at anticipation stage: Assessing the role of behavioral desire and anticipated emotions. *Comput Hum Behav* 86:129–146
- Rossi S, Conti D, Garramone F, Santangelo G, Staffa M, Varrasi S, Di Nuovo A (2020) The role of personality factors and empathy in the acceptance and performance of a social robot for psychometric evaluations. *Robotics* 9(2):39
- Van Swol LM (2009) The Effects of Confidence and Advisor Motives on Advice Utilization. *Commun Res* 36(6):857–873. <https://doi.org/10.1177/0093650209346803>
- Bonaccio S, Dalal RS (2006) Advice Taking and Decision-Making: An Integrative Literature Review, and Implications for the Organizational Sciences. *Organ Behav Hum Decis Process* 101(2):127–151
- Van Swol LM (2011) Forecasting Another's Enjoyment versus Giving the Right Answer: Trust, Shared Values, Task Effects, and Confidence in Improving the Acceptance of Advice. *Int J Forecast* 27(1):103–120
- Rousseau DM, Sitkin SB, Burt RS, Camerer C (1998) Not so different after all: A cross-discipline view of trust. *Acad Manag Rev* 23(3):393–404
- Komiak SYX, Benbasat I (2006) “The Effects of Personalization and Familiarity on Trust and Adoption of Recommendation Agents”, *MIS Quarterly* (30:4). University of Minnesota, Management Information Systems Research Center, pp 941–960
- Mayer, R. C., Davis, J. H., and Schoorman, F. D. 1995. “An Integrative Model of Organizational Trust,” *The Academy of Management Review* (20:3), p. 709.
- Kim, H., Benbasat, I., and Cavusoglu, H. 2017. “Online Consumers’ Attribution of Inconsistency Between Advice Sources,” in *Thirty Eighth International Conference on Information Systems*, Seoul, pp. 1–10.
- Schultze T, Rakotoarisoa A-F, Schulz-Hardt S (2015) Effects of Distance between Initial Estimates and Advice on Advice Utilization. *Judgm Decis Mak* 10(2):144–171
- Dalal RS, Bonaccio S (2010) What types of advice do decision-makers prefer? *Organ Behav Hum Decis Process* 112(1):11–23. <https://doi.org/10.1016/j.obhdp.2009.11.007>
- Gino F, Moore DA (2007) Effects of task difficulty on use of advice. *J Behav Decis Mak* 20(1):21–35
- Anthes, G. 2017. “Artificial Intelligence Poised to Ride a New Wave,” *Communications of the ACM* (60:7), ACM, pp. 19–21.
- Goodhue, D. L. 1995. “Understanding User Evaluations of Information Systems,” *Management Science* (41:12), INFORMS , pp. 1827–1844.
- Goodhue, D. L., and Thompson, R. L. 1995. “Task-Technology Fit and Individual Performance,” *MIS Quarterly* (19:2), p. 213.
- Liu C (2010) Human-machine trust interaction: A technical overview. *International Journal of Dependable and Trustworthy Information Systems* 1:61–74. <https://doi.org/10.4018/jdtis.2010100104>
- Muir BM, Moray N (1996) Trust in automation: Part II. Experimental studies of trust and human intervention in a process control simulation. *Ergonomics* 39:429–460. <https://doi.org/10.1080/00140139608964474>
- Madhavan P, Wiegmann DA (2007) Similarities and differences between human-human and human-automation trust: An integrative review. *Theor Issues Ergon Sci* 8:277–301. <https://doi.org/10.1080/14639220500337708>
- Smith MA, Allaham MM, Wiese E (2016) Trust in automated agents is modulated by the combined influence of agent and task type. *Proceedings of the Human Factors and Ergonomics Society Annual Meeting* 60:206–210. <https://doi.org/10.1177/1541931213601046>
- Goetz, J., Kiesler, S., & Powers, A. 2003. Matching robot appearance and behavior to tasks to improve human-robot cooperation. In H. F. M. Vander Loos & K. Yana (Eds.), *Proceedings of RO-MAN 2003: The 12th IEEE International Workshop on Robot and Human Interactive Communication* (Vol. 19, pp. 55–60). Piscataway, NJ: IEEE.
- Pak R, Fink N, Price M, Bass B, Sturre L (2012) Decision support aids with anthropomorphic characteristics influence trust and performance in younger and older adults. *Ergonomics* 55(9):1059–1072
- De Visser EJ, Monfort SS, McKendrick R, Smith MA, McKnight PE, Krueger F, Parasuraman R (2016) Almost human: Anthropomorphism increases trust resilience in cognitive agents. *J Exp Psychol Appl* 22(3):331
- Gray K, Wegner DM (2012) Feeling robots and human zombies: Mind perception and the uncanny valley. *Cognition* 125:125–130. <https://doi.org/10.1016/j.cognition.2012.06.007>
- Koh YJ, Sundar SS (2010) Effects of specialization in computer, web sites, and web agents on e-commerce trust. *Int J Hum Comput Stud* 68:899–912. <https://doi.org/10.1016/j.ijhcs.2010.08.002>
- Shaefer, K. E., Billings, D. R., Szalma, J. L., Adams, J. K., Sanders, T. L., Chen, J. Y. C., & Hancock, P. A. 2014. A meta-analysis of factors influencing the development of trust in automation: Implications for human-robot interaction (Report No. ARL-TR-6984). Aberdeen, MD: U.S. Army Research Laboratory.
- Heerink, M., Krose, B., Evers, V. and Wielinga, B., 2007, June. Observing conversational expressiveness of elderly users interacting with a robot and screen agent. In *2007 IEEE 10th International Conference on Rehabilitation Robotics* (pp. 751–756). IEEE.
- Barakova EI, De Haas M, Kuijpers W, Irigoyen N, Betancourt A (2018) Socially grounded game strategy enhances bonding and perceived smartness of a humanoid robot. *Connect Sci* 30(1):81–98
- Hertz N, Wiese E (2019) Good advice is beyond all price, but what if it comes from a machine? *J Exp Psychol Appl* 25(3):386
- Ghazali AS, Ham J, Barakova E, Markopoulos P (2020) Persuasive robots acceptance model (PRAM): roles of social responses within the acceptance model of persuasive robots. *Int J Soc Robot* 12:1075–1092
- Babel, F., Kraus, J., Miller, L., Kraus, M., Wagner, N., Minker, W. and Baumann, M., 2021. Small talk with a robot? The impact of dialog content, talk initiative, and gaze behavior of a social robot on trust, acceptance, and proximity. *International Journal of Social Robotics*, pp.1–14.
- Tauchert, C. and Mesbah, N., 2019, December. Following the Robot? Investigating Users’ Utilization of Advice from Robo-Advisors. In *ICIS*.
- Al-Tae MA, Kapoor R, Garrett C, Choudhary P (2016) Acceptability of robot assistant in management of type 1 diabetes in children. *Diabetes Technol Ther* 18(9):551–554

37. Ghazali AS, Ham J, Barakova E, Markopoulos P (2018) The influence of social cues in persuasive social robots on psychological reactance and compliance. *Comput Hum Behav* 87:58–65
38. Langedijk RM, Ham J (2021) More than advice: The influence of adding references to prior discourse and signals of empathy on the persuasiveness of an advice-giving robot. *Interact Stud* 22(3):396–415
39. Ritschel, H., Janowski, K., Seiderer, A. and André, E., 2019. Towards a robotic dietitian with adaptive linguistic style.
40. Ren X, Guo Z, Huang A, Li Y, Xu X, Zhang X (2022) Effects of social robotics in promoting physical activity in the shared workspace. *Sustainability* 14(7):4006
41. McColl D, Nejat G (2013) Meal-time with a socially assistive robot and older adults at a long-term care facility. *Journal of Human-Robot Interaction* 2(1):152–171
42. Law, T. and Scheutz, M., 2021. Trust: Recent concepts and evaluations in human-robot interaction. *Trust in human-robot interaction*, pp.27–57.
43. Bargain O, Aminjonov U (2020) Trust and compliance to public health policies in times of COVID-19. *J Public Econ* 192:104316
44. Chancey ET, Bliss JP, Yamani Y, Handley HAH (2017) Trust and the Compliance-Reliance Paradigm: The Effects of Risk, Error Bias, and Reliability on Trust and Dependence. *Hum Factors* 59(3):333–345. <https://doi.org/10.1177/0018720816682648>
45. Roesler, E., Naendrup-Poell, L., Manzey, D., & Onnasch, L. (2022). Why context matters: the influence of application domain on preferred degree of anthropomorphism and gender attribution in human–robot interaction. *International Journal of Social Robotics*, 1–12. <https://doi.org/10.1007/s12369-021-00860-z>
46. Giorgi I, Tiroto FA, Hagen O, Aider F, Gianni M, Palomino M, Masala GL (2022) Friendly But Faulty: A Pilot Study on the Perceived Trust of Older Adults in a Social Robot. *IEEE Access* 10:92084–92096
47. Ahn J, Kim J, Sung Y (2021) AI-powered recommendations: the roles of perceived similarity and psychological distance on persuasion. *Int J Advert* 40(8):1366–1384
48. Birnbaum MH, Stegner SE (1979) Source credibility in social judgment: Bias, expertise, and the judge's point of view. *J Pers Soc Psychol* 37(1):48
49. DeBono KG, Harnish RJ (1988) Source expertise, source attractiveness, and the processing of persuasive information: A functional approach. *J Pers Soc Psychol* 55(4):541
50. Wilson EJ, Sherrell DL (1993) Source effects in communication and persuasion research: A meta-analysis of effect size. *J Acad Mark Sci* 21(2):101–112
51. Sundar SS, Nass C (2000) Source orientation in human-computer interaction: Programmer, networker, or independent social actor. *Commun Res* 27(6):683–703
52. Moon Y, Nass C (1996) How “real” are computer personalities? Psychological responses to personality types in human-computer interaction. *Commun Res* 23(6):651–674
53. Nass C, Fogg BJ, Moon Y (1996) Can computers be teammates? *Int J Hum Comput Stud* 45(6):669–678
54. Bickman L (1974) The social power of a uniform 1. *J Appl Soc Psychol* 4(1):47–61
55. Furnham A, Boo HC (2011) A literature review of the anchoring effect. *J Socio-Econ* 40(1):35–42
56. NAO the humanoid and programmable robot, Aldebaran Robotics, retrieved from <https://www.aldebaran.com/en/nao>. Last accessed October 2022.
57. Amirova A, Rakhymbayeva N, Yadollahi E, Sandygulova A, Johal W (2021) 10 years of human-NAO interaction research: a scoping review. *Front Robot A I*:8
58. Smarr CA, Mitzner TL, Beer JM, Prakash A, Chen TL, Kemp CC, Rogers WA (2014) Domestic robots for older adults: attitudes, preferences, and potential. *Int J Soc Robot* 6(2):229–247
59. Redmon J, Divvala S, Girshick R, Farhadi A (2016) You only look once: Unified, real-time object detection. In: *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 779–788.
60. Chen H, Liu X, Yin D, Tang J (2017) A survey on dialogue systems: recent advances and new frontiers. *ACM SIGKDD Explorations Newsl* 19(2):25–35
61. Devlin J, Chang MW, Lee K, Toutanova K (2018) Bert: pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*.
62. Minutolo A, Damiano E, De Pietro G, Fujita H, Esposito M (2022) A conversational agent for querying Italian Patient Information Leaflets and improving health literacy. *Comput Biol Med* 141:105004
63. Vlasov V, Mosig JE, Nichol A (2019) Dialogue transformers. *arXiv preprint arXiv:1910.00486*.
64. Barr DJ, Levy R, Scheepers C, Tily HJ (2013) Random effects structure for confirmatory hypothesis testing: Keep it maximal. *J Mem Lang* 68(3):255–278. <https://doi.org/10.1016/j.jml.2012.11.001>
65. Wainwright PE, Leatherdale ST, Dubin JA (2007) Advantages of mixed effects models over traditional ANOVA models in developmental studies: a worked example in a mouse model of fetal alcohol syndrome. *Dev Psychobiol J Int Soc Dev Psychobiol* 49(7):664–674. <https://doi.org/10.1002/dev.20245>
66. Schielzeth H, Dingemanse NJ, Nakagawa S, Westneat DF, Allegate H, Teplitsky C, Reale D, Dochtermann NA, Garamszegi LZ, Araya-Ajoy YG (2020) Robustness of linear mixed-effects models to violations of distributional assumptions. *Methods Ecol Evol* 11(9):1141–1152. <https://doi.org/10.1111/2041-210X.13434>
67. Kherad-Pajouh S, Renaud O (2015) A general permutation approach for analyzing repeated measures ANOVA and mixed-model designs. *Stat Pap* 56(4):947–967. <https://doi.org/10.1007/s00362-014-0617-3>
68. Frossard J, Renaud O (2021) Permutation tests for regression, ANOVA, and comparison of signals: the permuco package. *J Stat Softw* 99:1–32. <https://doi.org/10.18637/jss.v099.i15>
69. Bates D, Mächler M, Bolker B, Walker S (2014) Fitting linear mixed-effects models using lme4. *arXiv preprint arXiv:1406.5823*. <https://doi.org/10.18637/jss.v067.i01>.
70. Lenth R. 2022. *_emmeans: Estimated Marginal Means, aka Least-Squares Means_*. R package version 1.7.5. <https://cran.r-project.org/package=emmeans>.
71. Madhavan P, Wiegmann DA (2005) Cognitive anchoring on self-generated decisions reduces operator reliance on automated diagnostic aids. *Hum Factors* 47:332–341. <https://doi.org/10.1518/0018720054679489>
72. Ehrlinger J, Johnson K, Banner M, Dunning D, Kruger J (2008) Why the unskilled are unaware: further explorations of (Absent) self-insight among the incompetent. *Organizational Behav Human Dec Process* (105:1), NIH Public Access, pp 98–121
73. Hancock PA, Billings DR, Schaefer KE, Chen JYC, de Visser EJ, Parasuraman R (2011) A meta-analysis of factors affecting trust in human-robot interaction. *Hum Factors* 53:517–527. <https://doi.org/10.1177/0018720811417254>

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Dr Ioanna Giorgi is a Lecturer in Artificial Intelligence at the School of Computing, University of Kent. She received a PhD degree in Computer Science from the University of Manchester, U.K. She was a Research Fellow in Computational Intelligence and Robotics at the

University of Plymouth, in the EU Interreg 2 Seas Mers Zeeën “AGE Independently” (AGE’IN) Project. Dr Giorgi specialises in brain-inspired architectures to explore the role of human language for high-level cognitive modelling in developmental and cognitive robotics. She is currently involved in several industry-focused projects as a co-Investigator to analyse the content and use of natural language across practical applications. Her research interests include Human-language-facilitated neuro-robotics, Multilingual Cognition, Human-Robot Interaction (HRI) and Trust, and Cognitive and Social Robotics.

Dr Aniello Minutolo is a researcher and member of the “Knowledge and Language Engineering” research group at the Institute for High-Performance Computing and Networking of the National Research Council of Italy (ICAR-CNR). He received his M.Sc. in Computer Science Engineering from the University of Naples “Federico II”, and a PhD degree in Information Technology Engineering from the University of Naples “Parthenope”. Since 2018, he has been a contract professor of Computer Science at the University of Naples “Federico II”. Since 2022, he has been a contract professor of Web Technologies at the telematic University “Giustino Fortunato”. His current research interests include Artificial Intelligence, Dialog Systems, Natural Language Processing, Knowledge management, Modelling and Reasoning. He has been involved in many national and European projects. He has been on the program committee of many international conferences and workshops and, moreover, is currently a member of the editorial board of the Internet of Things journal by Elsevier as well as some other international journals.

Dr Francesca A. Tiroto received her PhD in Social and Environmental Psychology at the University of Plymouth, UK (2022). Her doctoral project focused on social identity processes involved in the social acceptance of deep geothermal energy in Cornwall, Wales and Scotland. She worked as a Research Assistant in human–robot interaction with the School of Engineering, Computing and Mathematics, University of Plymouth. Previously, she worked as a Data Scientist, Teaching and Research Assistant. She is currently a Research Associate at Cardiff University’s Centre for Climate Change and Social Transformations, working on the project “Public understanding of and behavioural factors in indoor air quality”. She is interested in how people think and behave in group settings, in different types of pro-environmental behaviour, and in how people interact with different technologies.

Oksana Hagen received a B.Sc. degree in electrical and computer engineering from the National Chiao Tung University, Taiwan and an M.Sc. degree in computer vision and robotics as a joint European degree (Vibot) from Heriot-Watt University (U.K.), the University of Burgundy (France), and the University of Girona (Spain). She is currently pursuing a PhD degree in computing with the University of Plymouth, U.K. She joined ITRI CCMA, Taiwan, in 2011, as a Software Engineering Intern, and later as a Research Engineer with the Speech Processing Laboratory. She worked in the field of medical imaging at the University of Girona, Spain and at the Softbank Robotics Europe AI laboratory as a Robotics Researcher under Marie Curie ESR ITN APRIL Fellowship. In 2021, she joined the AGE’IN Project at the University of Plymouth, as a Researcher, in the field of social robotics.

Dr Massimo Esposito is currently a senior researcher at the Institute for High-Performance Computing and Networking of the National Research Council of Italy (ICAR-CNR). He received an M.Sc. degree (cum laude) in computer science engineering from the University of Naples Federico II, in March 2004, and a PhD degree in information technology engineering from the University of Naples, Parthenope, in

April 2011. Since 2012, he has been a Contract Professor of Informatics with the Faculty of Engineering, University of Naples Federico II. Since 2022, he has been a contract professor of Artificial Intelligence for Health Systems at the telematic University “Giustino Fortunato”. Since 2020, he has been the Leader of the “Language and Knowledge Engineering” Group at ICAR-CNR. He has been involved in different national and European projects. He is the author of more than 140 peer-reviewed papers in international journals and conference proceedings. His current research interests include artificial intelligence, natural language processing, knowledge engineering, and conversational systems. He is a member of the editorial board of some international journals. He has been on the program committee of many international conferences and workshops.

Dr Mario Gianni is currently a Senior Lecturer in Computer Science at the University of Liverpool. Previously, he was an Associate Professor of robotics with the School of Engineering, Computing and Mathematics, University of Plymouth, where he was Co-PI of the AGE’IN Project (Interreg 2 Seas Mers Zeeën, European Regional Development Fund). He was PI of an EPSRC Project and of an Innovate U.K. and was involved in several European Projects. He specialises in the design and development of models and methods combining AI, machine learning, control theory, and optimization to solve complex problems in robotics. He has been granted by the Southwest Creative Technology Network for developing a Smart AI-Assisted Robot Swarm for the autonomous drawing of human sketches through speech. He has published over 40 academic papers and served as a reviewer and Committee Member in several International Journals and Conferences in robotics, including SSRR, Autonomous Robots, IROS, ICRA, and the Journal of Field Robotics.

Dr Marco Palomino is an Associate Professor in Big Data and Information Systems at the University of Plymouth. He received a B.Sc. degree (Hons.) in mathematics from the National Autonomous University of Mexico, and a PhD degree in computer science from Downing College, University of Cambridge, in 2007. He started his academic career after working as a Software Consultant in London. Dr Palomino received a grant as co-PI funded by the Interreg 2 Seas 2014–2020 “AGE Independently (AGE’IN)” EU Project for the University of Plymouth (3 years, €371K). He has significant expertise in information retrieval and has pursued interests in web mining. He has worked with leading academics from Exeter, Bristol, the Environment Agency, Lloyd’s of London, the NHS, and the STFC Research Council. He is currently the Chair of the BCS Southwest Committee and an Honorary Fellow of the UK National Health Service (NHS).

Dr Giovanni Masala is a Senior Lecturer in Computer Science and leader of the Cognitive Robotics lab at the University of Kent. He is also a Visiting Researcher with the University of Plymouth, Plymouth, U.K. Dr Masala has published more than 90 peer-reviewed papers on topics in AI and robotics, Natural Language Processing, data security and Machine Learning for biomedical data classification. In 2019 Dr Masala obtained a grant as PI funded by the Interreg 2 Seas 2014–2020 “AGE Independently (AGE’IN)” EU Project for the University of Plymouth to explore the role of technologies on senior care and several other grants from Innovate UK. Dr Masala is also a Senior Member of the IEEE Computational Intelligence Society and Associate Editor in Frontiers in Robotics and AI and for the International Journal of Environmental Research and Public Health.