

Embodying a Virtual Agent in a Self-Driving Car: A Survey-Based Study on User Perceptions of Trust, Likeability, and Anthropomorphism

Clarisse Lawson-Guidigbe, Univ. Bordeaux, CNRS, Bordeaux INP-ENSC, IMS, Talence, France

Nicolas Louveton, Université de Poitiers, Université François-Rabelais de Tours, CNRS, CeRCA, Poitiers, France*

 <https://orcid.org/0000-0002-9021-2589>

Kahina Amokrane-Ferka, IRT SYSTEMX, Palaiseau, France

Benoît Le Blanc, Univ. Bordeaux, CNRS, Bordeaux INP-ENSC, IMS, Talence, France

Jean-Marc André, Univ. Bordeaux, CNRS, Bordeaux INP-ENSC, IMS, Talence, France

ABSTRACT

This article considers the visual appearance of a virtual agent designed to take over the driving task in a highly automated car, to answer the question of which visual appearance is appropriate for a virtual agent in a driving role. The authors first selected five models of visual appearance thanks to a picture sorting procedure ($N = 19$). Then, they conducted a survey-based study ($N = 146$) using scales of trust, anthropomorphism, and likability to assess the appropriateness of those five models from an early-prototyping perspective. They found that human and mechanical-human models were more trusted than other selected models in the context of highly automated cars. Instead, animal and mechanical-animal ones appeared to be less suited to the role of a driving assistant. Learnings from the methodology are discussed, and suggestions for further research are proposed.

KEYWORDS

Embodied Virtual Agents, Highly Automated Driving, Perception of Anthropomorphism, Perception of Trust, Visual Agent Appearance

INTRODUCTION

Virtual agents represent a shift toward more natural interactions with technology, for example, Siri, Google Assistant, Alexa, and Cortana. In everyday life, they enable quick access to web content and other functionalities through computers, mobile devices, and even our cars. Using virtual agents for in-car interactions may benefit drivers, enhancing cockpit features (e.g., music and navigation) with novel user interfaces (Weng et al., 2016). Recent implementations have come in various forms, such

DOI: 10.4018/IJMHCI.330542

*Corresponding Author

This article published as an Open Access article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0/>) which permits unrestricted use, distribution, and production in any medium, provided the author of the original work and original publication source are properly credited.

as mobile apps (e.g., Apple CarPlay and Google Assistant Driving Mode), in-car dedicated devices (e.g., Echo Auto, Chris, and Garmin Speak; Lugano, 2017), and proprietary personal agents from car manufacturers (e.g., Honda, BMW, Mercedes, and Volkswagen; Majji & Baskaran, 2021). With the rise of highly automated cars, new concepts of virtual agents (Ajitha & Nagra, 2021; Okamoto & Sano, 2017) are emerging with new roles. Beyond delivering complex messages about driving and road safety, particularly during handover procedures, these virtual agents embody the artificial intelligence of the automated system. In this context, the perception of trust toward such agents plays an important role in adopting this new technology.

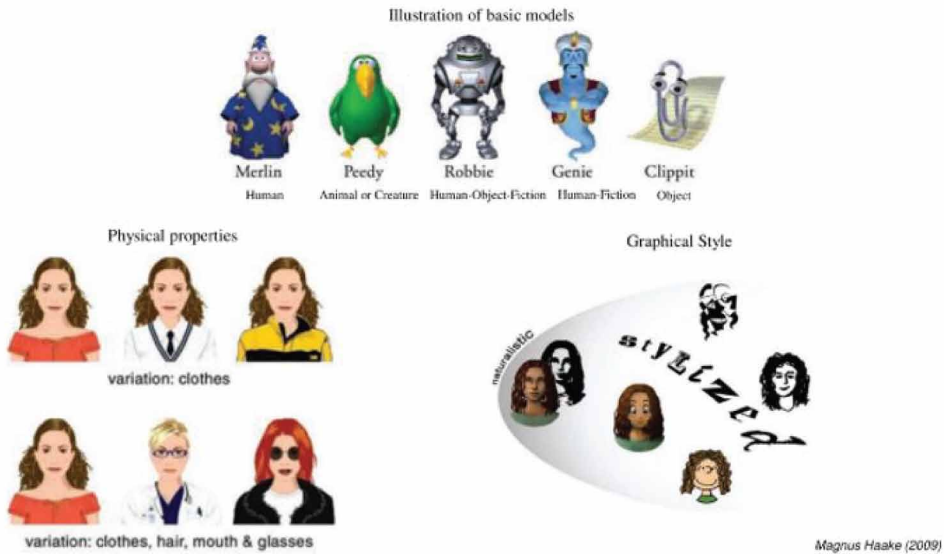
Previous research showed that providing a visual appearance, specifically an anthropomorphic one, might be a way of ensuring trust toward these agents (Meng et al., 2021; Ekman et al., 2018; Waytz et al., 2014). Similar to human-to-human relationships, in which people tend to judge others' social warmth, honesty, trustworthiness, and intellectual competence based on facial appearance or attire (Zebrowitz & Montepare, 2008; Willis & Todorov, 2006; Smith et al., 2018), human – virtual-agent interactions are greatly influenced by visual appearance. Research revealed that the visual appearance of an embodied virtual agent tends to have an impact on user first impressions (ter Stal et al., 2020), perceptions of competence, intelligence, trust, and acceptance. Several studies have shown that pedagogical agents' visual attributes, including realism, gender, and ethnicity, significantly affect student self-efficacy and motivation (Baylor & Kim, 2004; Gulz & Haake, 2006; McDonnell et al., 2012; Ashby Plant et al., 2009).

Further, it appears that the appropriate appearance for a virtual agent may vary with the use case or context of use. For example, agents with realistic human-like looks tend to be preferred for health-related roles, whereas cartoon-like human looks are preferred for social interactions (Ring et al., 2014). Also, the level of realism in human-like representations encourages positive subjective user ratings (Yee et al., 2007). Zoomorphic agents tend to be associated with entertainment, education, and therapy-related roles, and agents with more mechanical looks are expected to carry out security-related tasks (Dingjun et al., 2010; Kalegina et al., 2018). Parmar et al. (2018) provided evidence that a virtual health agent designed with attire fitting the role (e.g., a white coat) was perceived to be more trustworthy, reassuring, and persuasive than an agent whose attire was not role appropriate (e.g., a casual outfit).

Strohmann et al. (2019) attempted to provide visual design recommendations for in-car virtual agents, but the choice of a visual appearance for a virtual agent in a driving role can still be challenging. This is mainly due to the almost unlimited number of visual design options when creating a virtual agent for a new application and specific type of user. Haake and Gulz (2009) proposed a comprehensive framework for this design space (see Figure 1). It defines three high-level aspects of static visual characteristics: basic model, physical properties, and graphical style. *Basic model* refers to the constitution of an embodied agent and the conceptual entity related to this constitution (human, animal or creature, inanimate object, fantasy, or a combination of these). *Physical properties* refers to design elements (e.g., face shape; eye, hair, and skin color; clothes, and accessories) that carry the most variety and suggest cultural, social, and psychological characteristics. Finally, *graphical style* (realistic or stylized) refers both to the degree of realism and details applied in a visual design.

The aspect of graphical style has received a lot of attention in visual design of a virtual agent research (Duffy, 2003; Billard, 2005; Blow et al., 2006; Kopp et al., 2005). It is inspired by the work of McCloud (1993), where the author presents an in-depth theory of comic design space and gives clear definitions of abstract, iconic, and realistic visual design. Abstract design does not necessarily relate to a real-world entity (e.g., a Jackson Pollock painting); on the contrary, realistic design resembles actual things, like a photograph of a person, animal, or object; and finally, iconic design (e.g., cartoons) are an abstraction of real appearances, but still perceived as resembling something real (Manning, 1998). Trust is a key factor for technology adoption in safety-critical contexts such as in driving. While prolonged visual interactions might not be suitable for manual driving (i.e., because of competing attentional resources; Wickens, 2008), highly automated cars allow for building trust

Figure 1. Pedagogical embodied agents design framework (adapted from Haake & Gulz, 2009)



through visual interfaces such as visually embodied agents. Thus, it is necessary to investigate the relationship between visual design choices and perception of trust.

Trust is a complex psychological concept, although it is a key factor in safety-critical human–system collaboration (Hertzum et al., 2002). Overall, throughout the literature, trust has been defined as a belief (Knack, 2001), an expectation (Robinson, 1996; Chouk & Perrien, 2004), an attitude (Ajzen & Fishbein, 2000; Castelfranchi & Falcone, 2000) or a judgment (Gambetta, 2000; Madhok, 1995; Lorenz, 1999). It might be described as a behavioral intention (Mayer et al., 1995; Corritore et al., 2003) or as an actual behavior (Currall & Judge, 1995; Lewis & Weigert, 1985). It has also been defined as a part of an individual personality that drives one to believe another party is reliable (Muir, 1994). Trust measurement is very variable across studies (Hancock et al., 2011; K. E. Schaefer et al., 2014; Helldin et al., 2013; Häuslschmid et al., 2017). Researchers often use self-reported scales as a primary measurement tool along with behavioral tasks and physiological measurements (French et al., 2018).

Some researchers have used a single-question measurement (Haesler et al., 2018; Rossi et al., 2018; Shu et al., 2018;), or created multiple items scales specifically for their studies (Nordheim et al., 2019; Byrne & Marín, 2018; Linnemann et al., 2018; Waytz et al., 2014; Garcia et al., 2015; Holthausen et al., 2020). Others re-used scales from psychological research (Martelaro et al., 2016; Sebo et al., 2019; Herse et al., 2018). Standardized trust scales have also been developed to fit a large variety of use cases and allow comparisons across studies (Jian et al., 2000; K. Schaefer, 2016; Yagoda & Gillan, 2012). Finally, recent measurement tools have focused on integrating social dimensions into trust scales (Malle & Ullman, 2021; Law & Scheutz, 2021).

In the case of anthropomorphism, measurement tools are also diverse. Anthropomorphism is a cognitive process in which human form or characteristics such as intelligence or capacity of emotions are projected onto a non-human object such as a computer or robot. Single (MacDorman, 2006) or two item scales have been used (Aggarwal & McGill, 2007) to measure anthropomorphism. Powers and Kiesler (2006) proposed a multiple item scale that is often used for comparisons across studies (Bartneck et al., 2009; Häuslschmid et al., 2017; Ekman et al., 2018). Other scales have been developed to suit specific research (van der Woerd & Haselager, 2019; Waytz et al., 2014; Ruijten

et al., 2019). Other studies use behavioral measurement (Minato et al., 2005) and scales have been developed specifically for intelligent personal assistants (Moussawi & Koufaris, 2019).

Additionally, the concepts of mindful (Araujo, 2018) and mindless (Kim & Sundar, 2012) anthropomorphism have been measured. *Mindful anthropomorphism* is defined as a “sincere belief that the object has human characteristics” (Nass & Moon, 2000, p. 93; Powers & Kiesler, 2006; Bartneck et al., 2009). In other words, after interacting with the machine, the user judges it as possessing qualities such as intelligence or the ability to be aware of its environment. Conversely, *mindless anthropomorphism* is described as an automatic response of the human to an anthropomorphic stimulus that causes it not only to attribute human characteristics but also to sense the presence of the machine from which the stimulus originates (Nass & Moon, 2000; Kim & Sundar, 2012). The difference between the two types of anthropomorphism is the perception of *social presence*, defined as the “feeling of being in the presence of another person” (Burgoon et al., 2001, p. 2). Hence, it is important to assess the trustworthiness of a virtual agent in the role of a driving assistant in a highly automated car (such as in SAE Level 4; SAE International, 2014). Higher trust for driving assistants may increase highly automated driving system acceptance or improve safety. A first step toward this goal is to identify basic models that encourage the perception of trust toward a virtual agent. For these reasons, we present a study that compares different basic models to potential users and assesses their relative trustworthiness, likability, and anthropomorphism.

In our study, we chose to use the *zero-acquaintance approach* (Vartanian et al., 2012; ter Stal et al., 2020). It consists of evaluating virtual agents based on static image prototypes of virtual agents. It consists of evaluating virtual agents based solely on static images, without any actual interaction between the agent and user. While high-fidelity and interactive prototypes are more representative, such a low-fidelity approach is particularly relevant early in the design process. Being able to predict user trust with this approach may save resources for researchers and design practitioners. Here, the images were sorted into categories by a group of participants using a picture sorting procedure. Selected categories were included in an online survey which immersed respondents into the context of autonomous vehicles, virtual agents, and driving situations, using text and pictures to illustrate context. Finally, the categories of basic models were rated in terms of trust, likability, and anthropomorphism. The scale from Waytz et al. (2014) was used as it was specifically developed for the context of highly automated driving and included measurement for both trust and anthropomorphism.

MATERIALS AND METHODS

Participants

A total of 146 participants (60% of them were females, mean age = 36.92, $SD = 14.04$) completed the online questionnaire. They were recruited through a public link to the questionnaire shared on social media and mailing lists (academic and professional). The first page of the questionnaire explained the context and goals of the study and asked the participant to give informed consent. A total of 99% reported France as their current residency country and 91% as the country where they have lived the most; 69% of participants had a driver's license for over five years; 41% drive every day and 31% once a week; 58% declared using at least one driving assistance system. When asked about their prior usage of virtual assistants from mainstream consumer electronics (such as those of Google and Apple), 29% of our participants declared using those; 85% of our participants graduated from a level above or equal to a baccalaureate level (high school or higher education diploma).

Material

This work focused on the basic model dimension from Haake's and Gulz's design space (2009). Five basic models were evaluated using a French-translated version of the 16-item scale from Waytz et al. (2014). This scale was used because it was specifically developed for the context of highly automated

driving and includes items related to anthropomorphism, liking, and trust. Static high-fidelity images from an online database were used, and, to avoid evaluation of a specific image instead of a basic model category, four images were given for each basic model. The four example images were selected using a picture sorting procedure (described in more detail below). This procedure was performed with an independent group of participants to ensure that basic models were built from user perception and understanding.

Picture Sorting Procedure

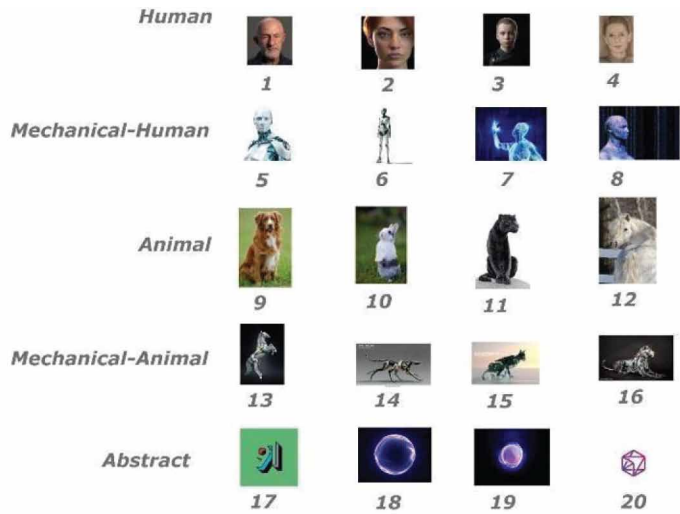
Picture sorting (Lobinger & Brantner, 2020) is a variant of the card sorting technique that helps to study a user's mental model and understand how they categorize items. The aim of this procedure was to create basic model categories and find pictures that illustrate these categories. Multiple illustration pictures are necessary for each basic model to avoid the interpretation of a single image instead of the model as a concept. Firstly, we created 10 category labels using two aspects (basic model and graphical style) of the framework proposed by Haake and Gulz (2009). Then, each label was written on the Pinterest website search bar (one at a time), and pictures were selected from the prompted results, based on their quality. Depending on the category label, results from the website were variable. Hence, the number of pictures selected for each category is different. The complete selection was as follows: *realistic human* (5), *stylized human* (11), *realistic animal* (9), *stylized animal* (7), *mechanical human* (5), *mechanical animal* (5), *mini mechanical* (5), *abstract* (9), *inanimate object* (5), and *fantasy* (21). Ultimately, 83 pictures were collected from the Pinterest website. The pictures were printed and placed on a large table, with the model labels placed on an adjacent table (see Figure 2).

A total of 19 participants (eight females, mean age = 28.10) recruited from a research institute and from a university campus (Paris, France) participated in the card sorting. The card sorting was individual and took place at IRT SystemX (Paris, France). Firstly, each participant was presented with the definition of every label and with questions to ensure they understood what each label meant. Then, they were asked to sort the pictures using the provided labels. A picture could be categorized to one or many labels or stay uncategorized. A picture was considered as representative of a basic model when it was placed in that model category by more than 70% of participants (Fincher & Tenenberg, 2005). Following the results from the picture sorting procedure, five basic models had at least four representative pictures: *realistic human*, *mechanical human*, *realistic animal*, *mechanical animal*, and *abstract*. The word *realistic* was removed for simplification (see Figure 3). These basic models were selected for the questionnaire. All the other basic models did not have enough representatives' pictures to be selected for the questionnaire.

Figure 2. Picture sorting results (after participant sorting)



Figure 3. Basic models categories used in the online survey



Survey Procedure

The questionnaire was designed using LimeSurvey. From the publicly shared link, participants landed on the introductory page of the questionnaire. Information was gathered about participant driving behavior and experience as well as prior knowledge of virtual agent technology. To increase the relevance of participant ratings, we provided them with driving scenarios including potential interaction with a virtual agent. Participants were immersed in such scenarios using a picture and an accompanying text (see Figure 4).

Several questions about their perception and intention were asked (e.g., “Will you let the assistant drive the car from the garage to the front door?”) to help immerse participants in the automated driving context. They also had to project themselves into the scenario by answering a question about their intended behavior (“When the automated mode is on and the virtual assistant is in charge of the driving task, what are you most likely to do in this situation?”). The immersive procedure is achieved thanks to presenting a risky driving situation (see Figure 5).

Once this immersive phase was completed, participants went through the trust rating procedure for each of the five basic models selected, thanks to our picture sorting results. For each basic model (represented by four examples; see Figure 6), participants completed the French-translated version

Figure 4. Normal automated driving situation with little road traffic



Figure 5. Risky automated driving situation with high traffic density and fast reaction required by the presence of an ambulance in the rearview mirror

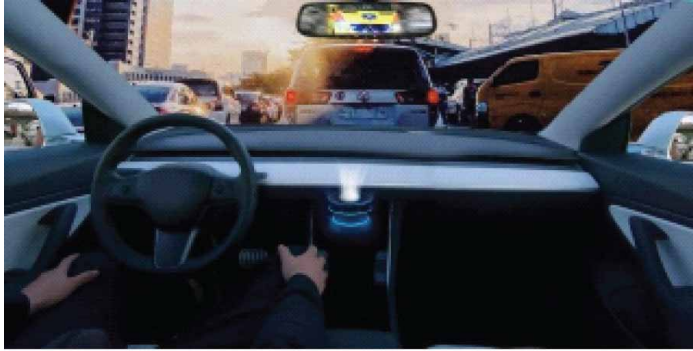
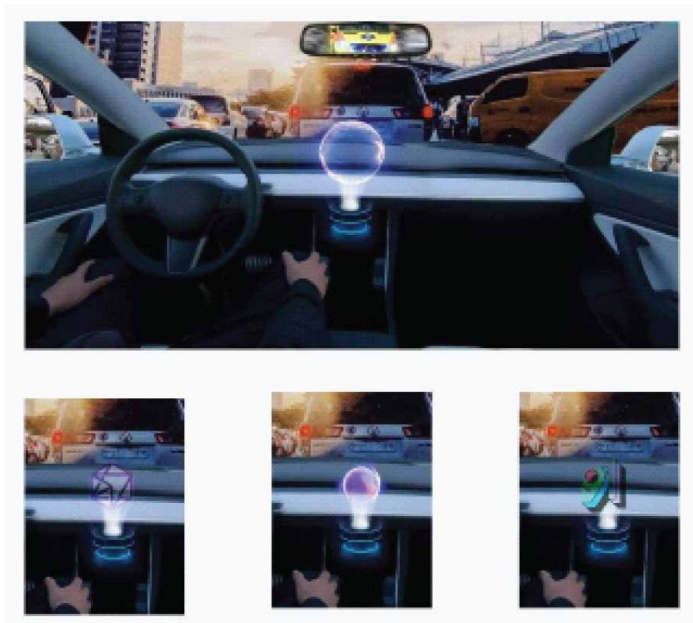


Figure 6. Model presentation in the survey with examples from the abstract model category, integrated into the central cockpit space



of the Waytz et al. (2014) scale that assessed trust, likability, and perceived anthropomorphism. General feedback and socio-demographic data were collected at the end of the questionnaire. The questionnaire took about 20 to 25 minutes to complete.

RESULTS

To calculate the perceived anthropomorphism score, four items (1–4 in the translated scale) were averaged to form a single composite ($\alpha = .92$). For likability, three items (5–7) were averaged ($\alpha = .96$) and eight items (8–16) were averaged for trust ($\alpha = .97$). Table 1 shows median scores, standard deviations, and quartiles for each basic model on the three measures of trust, likability,

Table 1. Median ratings and (std. dev.) for each model; every item in the scale is ranked on a Likert scale (From 0 = Not at all to 10 = Very much)

		Human	Animal	Mechanical human	Mechanical animal	Abstract
Anthropomorphism	Median	5.750	5.000	6.375	5.000	5.750
	Std. dev.	2.413	2.532	2.459	2.701	2.365
	25 th percentile	4.250	2.500	4.750	2.563	4.250
	50 th percentile	5.750	5.000	6.375	5.000	5.750
	75 th percentile	7.500	6.250	8.000	6.750	7.000
Liking	Median	5.000	4.000	5.333	4.667	5.000
	Std. dev.	2.958	2.866	2.995	2.969	2.861
	25 th percentile	2.000	1.750	3.000	2.000	3.333
	50 th percentile	5.000	4.000	5.333	4.667	5.000
	75 th percentile	7.250	6.000	7.667	6.917	7.333
Trust	Median	5.611	4.444	5.833	5.000	5.278
	Std. dev.	2.438	2.556	2.432	2.626	2.472
	25 th percentile	4.111	2.278	4.444	2.889	3.778
	50 th percentile	5.611	4.444	5.833	5.000	5.278
	75 th percentile	7.333	5.972	7.528	6.444	7.111

and anthropomorphism. For each of these scales, the higher the score, the better the judgment of participants toward a specific model.

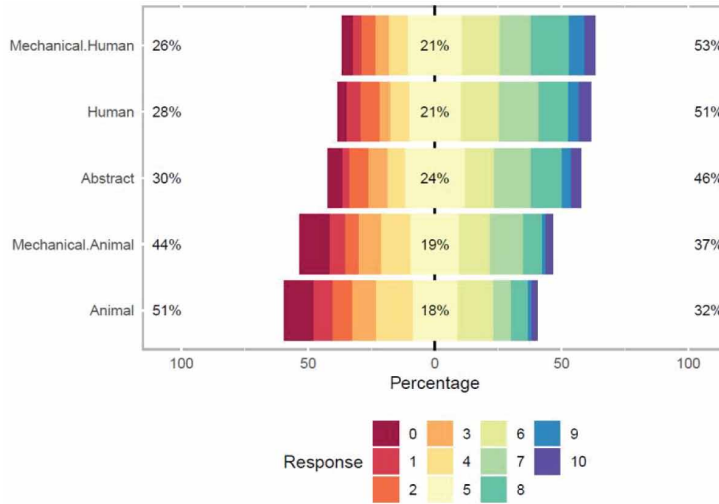
Trust

The mechanical-human and human categories appeared to be the most trusted, with 53% and 51% of ratings above neutrality (see Figure 7). The abstract category was less trusted, with 46% of ratings above neutrality. In contrast, mechanical-animal and animal categories had a very low trust rating: 44% and 51% below neutrality, respectively. It is noteworthy that neutral ratings were from 18% to 24% across categories. These differences appeared to be statistically significant as shown by the Friedman non-parametric test (Chi-squared value = 73.889; $p < 0.05$). An additional post hoc pair-wise comparison test (Wilcoxon test) with Bonferroni correction (see Table 2) showed that the animal and mechanical-animal categories were significantly different from all other categories but not different from each other.

Table 2. The p-values (significance threshold: $p < 0.05$) for the post hoc pair-wise comparison of trust

	Human	Animal	Mechanical-human	Mechanical-animal
Animal	0.0004*	-		
Mechanical-Human	1.0000	< 0.001*	-	
Mechanical-Animal	0.0281*	1.0000	0.0040*	-
Abstract	1.0000	0.0041*	1.0000	0.1664

Figure 7. Diverging stacked chart representing trust ratings from Likert scales across basic model categories: The central percentage represents neutral judgments, while the percentages at the right and left are positive and negative judgments, respectively



Liking

The mechanical-human, abstract, and human categories appeared to be the most liked (see Figure 8), with relatively similar ratings: 47%, 47%, 43% above neutrality, respectively. Animal and mechanical animal were the least liked, with 36% and 31% of ratings above neutrality, respectively. Neutral ratings ranged from 11% to 18% across categories. These differences were statistically significant as shown by a Friedman non-parametric test (Chi-square value = 47.788; $p < 0.05$). An additional post hoc pair-wise comparison test (Wilcoxon test) with Bonferroni correction (see Table 3) showed the Animal model was significantly different from all other models.

Figure 8. Diverging stacked chart representing likability ratings from Likert scales across basic model categories: The central percentage represents neutral judgments, while the percentages at the right and left are positive and negative judgments, respectively

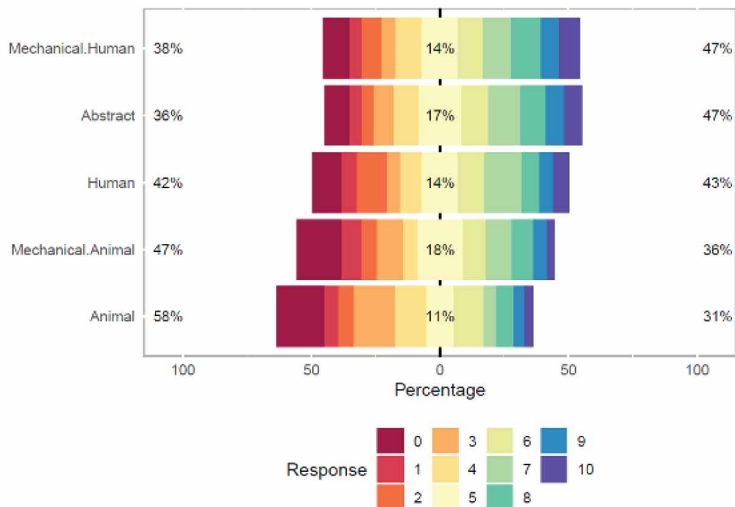


Table 3. The *p*-values (significance threshold: $p < 0.05$) for the post-hoc pair-wise comparison of liking, with Bonferroni correction

	Human	Animal	Mechanical-Human	Mechanical-Animal
Animal	0.2219	-		
Mechanical human	1.0000	0.0032*	-	
Mechanical animal	1.0000	1.0000	0.1086	-
Abstract	1.0000	0.0031*	1.0000	0.1308

Anthropomorphism

The human and abstract models were judged as the most anthropomorphic (see Figure 9) with 58% and 55% of ratings above neutrality, respectively. They were followed by the mechanical-animal (47%), mechanical-human (45%), and animal (38%) categories. Neutral ratings ranged from 13% to 21% across categories. These differences were statistically significant, as shown by a Friedman non-parametric test (Chi-square value = 109.42; $p < 0.05$). An additional post hoc pair-wise comparison test (Wilcoxon test) with Bonferroni correction (see Table 4) showed that anthropomorphism ratings for the animal and mechanical-animal categories were significantly different from all other models but not different from each other.

Cross-Indicator Correlations

When correlating the three different scales together we found that all correlation were high and statistically significant. Which means that when a user is ranking a category high (or low) on a scale it also ranks high (or low, respectively) on other scales. More precisely, it means that anthropomorphism is positively correlated with liking ($r(728) = 0.89$; $p < 0.05$), and with trust ($r(728) = 0.90$; $p < 0.05$). Trust and liking are also highly and positively correlated or each other ($r(728) = 0.90$; $p < 0.05$).

Figure 9. Diverging stacked chart representing anthropomorphism ratings from Likert scales across basic model categories: The more the stacked chart is shifted to the right, the more liked the basic model

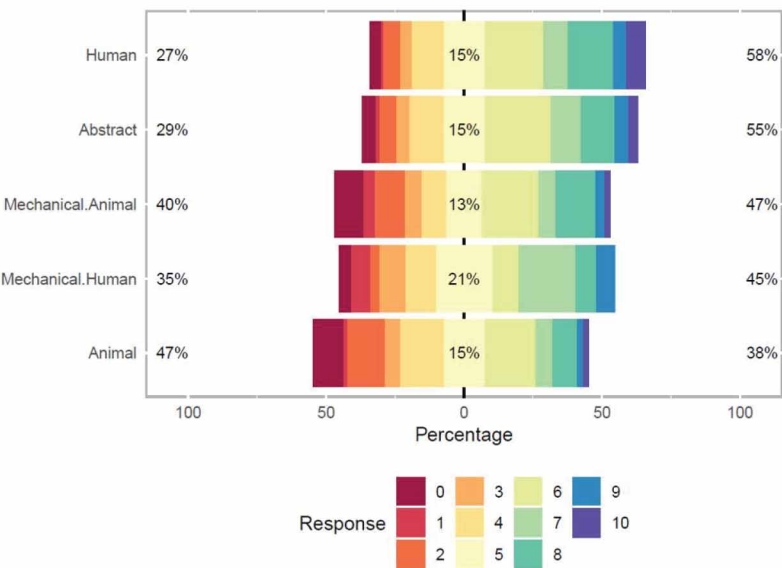


Table 4. The p -values (significance threshold: $p < 0.05$) for the post hoc pair-wise comparison of anthropomorphism, with Bonferroni correction

	Human	Animal	Mechanical human	Mechanical animal
Animal	0.00032*	-		
Mechanical human	1.00000	< 0.001*	-	
Mechanical animal	0.03276*	1.00000	0.00032*	-
Abstract	1.00000	0.00515*	0.27101	0.25675

Discussion

This study aimed to assess the trustworthiness, perceived anthropomorphism, and likability of five visual appearance models for a virtual driving agent in the context of highly automated driving. We used an early-prototyping survey-based methodology to make early design decisions and explore potential-user attitudes. Rather than observing a high rating in the level of trust, our goal was to find the most appropriate visual representation for the role under study. Our findings point to mechanical-human and human models as the most appropriate in terms of both trust and likability. Human and abstract models were judged as the most anthropomorphic. In contrast, animal and mechanical-animal models were least liked, trusted, and anthropomorphic, pointing to a low relevance of those visual representations in the context under study. This confirms our expectations that some visual appearances are more appropriate than others for the role of driving (Verberne et al., 2015). Specifically, animal appearances are not associated with safety-related tasks (Nomura et al., 2008) such as driving. Further, results confirm that anthropomorphism is a key factor in trust perception toward virtual agents (Waytz et al., 2014; P. Ruijten et al., 2018; Lemoine & Cherif, 2012; Christoforakos et al., 2021).

An unexpected result is that the abstract model was perceived as more anthropomorphic than the mechanical-human model. Firstly, participants might have more experience with agents designed as an abstract model (such as Alexa, Google Assistant, and Siri) than mechanical-human ones. Indeed, there is evidence that these intelligent personal agents are highly anthropomorphized by users (Kuzminykh et al., 2020). Secondly, abstract visual representations might give more room for imagination. Thus, it might be easier for participants to project human qualities onto them than onto mechanical-human models. Also, it is worth noting that on the three scales, positive scores (above neutral) were never very high, and neutral scores corresponded to a significant part of participant ratings. This may come from the lack of interactivity within our setup. These latter findings call for further reflection.

This study would benefit from additional replication and extension. First, further research might implement selected embodiment models into a more ecological driving environment, such as driving simulation. Indeed, such an experimental setup will allow for studying the impact on trust of actual interaction with the virtual agent. Interactivity might have possible side effects on driving performance and safety. Also, in this study, we did not explore the effect of individual differences on the trust evaluation of specific basic models. Indeed, human factor literature points to the effect of individual differences on trust in automation (Körber & Bengler, 2014; Matthews et al., 2019). Studies showed an effect of prior knowledge and context (Khastgir et al., 2018; Lee & See, 2004); the effect of cognitive abilities in managing automation unreliability (Rovira et al., 2017); and even neurobiological attributes that modulate overtrust in computerized decision aid (Parasuraman et al., 2012). A replication of this study with a controlled effect of individual differences might bring more understanding of how educational and professional background modulates trust in mainstream product automation.

Our findings might also need to be reproduced across different countries and cultures. Indeed, most of our respondents were nationals or residents of the country where this study was realized (i.e., France). Also, most of them held a high school degree or above (higher education) which does

reflect graduation level statistics for younger generations in France (Institut national de la statistique et des études économiques, 2021) and in Europe (EuroStat, 2021). However, we cannot conclude whether our findings may generalize to countries of different cultural backgrounds or individuals with lower graduation levels.

Finally, while our study targeted a high automation level (SAE Level 4; SAE International, 2014), prolonged visual interactions are unsuitable in a manual driving context as driving is mainly a visual-manual task. Face-embodied interactions are suitable either when the car is stopped or in automated driving mode. This might allow for building a personalized and trustworthy relationship with the automation system. In case of a handover situation, the virtual driving assistant might switch from a visual embodiment to a voice-only embodiment to limit interference with the visual channel (Wickens, 2008). This raises important questions about the impact of prior visually embodied calibration on how trust is enduring variability of the situation (automated driving to manual), of the communication channel (visual to auditory), and the possible side effects on actual situational awareness and driving performance.

CONCLUSION

Our rapid-prototyping approach based on a card sorting procedure and a survey, allowed us to select embodiment models that are likely to elicit trust toward a virtual driving assistant in a highly automated driving context. According to our data, human and mechanical-human models are more likely to be trusted in a highly automated driving context than animal and mechanical-animal ones. Follow-up research work is needed to investigate the extension of these results in an ecological environment, such as in driving simulation, individual differences that might impact the perception of trustworthiness in visual appearance models, and the impact of visual embodiment on trust calibration and robustness.

ACKNOWLEDGMENT

This work has been supported by the French government under the France 2030 program, as part of the SystemX Technological Research Institute within the CAB project.

REFERENCES

- Aggarwal, P., & McGill, A. (2007). Is that car smiling at me? Schema congruity as a basis for evaluating anthropomorphized products. *The Journal of Consumer Research*, 34(4), 468–479. doi:10.1086/518544
- Ajitha, P. V., & Nagra, A. (2021). An overview of artificial intelligence in automobile industry: A case study on Tesla cars. *Solid State Technology*, 64(2), 11.
- Ajzen, I., & Fishbein, M. (2000). Attitudes and the attitude–behavior relation: Reasoned and automatic processes. *European Review of Social Psychology*, 11(1), 1–33. doi:10.1080/14792779943000116
- Araujo, T. (2018). Living up to the chatbot hype: The influence of anthropomorphic design cues and communicative agency framing on conversational agent and company perceptions. *Computers in Human Behavior*, 85, 183–189. doi:10.1016/j.chb.2018.03.051
- Ashby Plant, E., Baylor, A. L., Doerr, C. E., & Rosenberg-Kima, R. B. (2009). Changing middle-school students' attitudes and performance regarding engineering with computer-based social models. *Computers & Education*, 53(2), 209–215. doi:10.1016/j.compedu.2009.01.013
- Bartneck, C., Kulić, D., Croft, E., & Zoghbi, S. (2009). Measurement instruments for the anthropomorphism, animacy, likeability, perceived intelligence, and perceived safety of robots. *International Journal of Social Robotics*, 1(1), 71–81. doi:10.1007/s12369-008-0001-3
- Baylor, A. L., & Kim, Y. (2004). *Pedagogical agent design: The impact of agent realism, gender, ethnicity, and instructional role*. Intelligent Tutoring Systems., doi:10.1007/978-3-540-30139-4_56
- Blow, M., Dautenhahn, K., Appleby, A., Nehaniv, C., & Lee, D. (2006). The art of designing robot faces: Dimensions for human-robot interaction. *Proceedings of the First ACM SIGCHI–SIGART Conference on Human-Robot Interaction*, (pp. 331–332). ACM. doi:10.1145/1121241.1121301
- Burgoon, J. K., Twitchell, D. P., Jensen, M. L., Meservy, T. O., Adkins, M., Kruse, J., & Younger, R. E. (2009). Detecting concealment of intent in transportation screening: A proof of concept. *IEEE Transactions on Intelligent Transportation Systems: A Publication of the IEEE Intelligent Transportation Systems Council*, 10(1), 103–112. IEEE. 10.1109/TITS.2008.2011700
- Byrne, K., & Marín, C. A. (2018). Human trust in robots when performing a service. *Proceedings of the 2018 IEEE 27th International Conference on Enabling Technologies: Infrastructure for Collaborative Enterprises*. IEEE. doi:10.1109/WETICE.2018.00009
- Castelfranchi, C., & Falcone, R. (2000, January 1). Trust is much more than subjective probability: Mental components and sources of trust. *Proceedings of the 33rd Annual Hawaii International Conference on System Sciences*. IEEE. doi:10.1109/HICSS.2000.926815
- Chouk, I., & Perrien, J. (2004). Les facteurs expliquant la confiance du consommateur lors d'un achat sur un site marchand: Une étude exploratoire. *Décision Marketing*, 35, 75–86. doi:10.7193/DM.035.75.86
- Christoforakos, L., Gallucci, A., Surmava-Große, T., Ullrich, D., & Diefenbach, S. (2021). Can robots earn our trust the same way humans do? A systematic exploration of competence, warmth, and anthropomorphism as determinants of trust development in HRI. *Frontiers in Robotics and AI*, 8, 640444. doi:10.3389/frobt.2021.640444 PMID:33898531
- Corritore, C., Kracher, B., & Wiedenbeck, S. (2003). On-line trust: Concepts, evolving themes: A model. *International Journal of Human-Computer Studies*, 58(6), 737–758. doi:10.1016/S1071-5819(03)00041-7
- Currall, S. C., & Judge, T. A. (1995). Measuring trust between organizational boundary role persons. *Organizational Behavior and Human Decision Processes*, 64(2), 151–170. doi:10.1006/obhd.1995.1097
- Dingjun, L., Rau, P.-L., & Li, Y. (2010). A cross-cultural study: Effect of robot appearance and task. *International Journal of Social Robotics*, 2(2), 175–186. doi:10.1007/s12369-010-0056-9
- Duffy, B. (2003). Anthropomorphism and the social robot. *Robotics and Autonomous Systems*, 42(3–4), 177–190. doi:10.1016/S0921-8890(02)00374-3

Ekman, F., Johansson, M., & Sochor, J. (2018). Creating appropriate trust in automated vehicle systems: A framework for HMI design. *IEEE Transactions on Human-Machine Systems*, 48(1), 95–101. doi:10.1109/THMS.2017.2776209

EuroStat. (2021). *Population by educational attainment level, sex and age (%)* [Dataset EDAT_LFS_9903]. https://ec.europa.eu/eurostat/databrowser/view/edat_lfs_9903/default/bar

French, B., Duenser, A., & Heathcote, A. (2018). *Trust in automation: A literature review*.

Gambetta, D. (2000). Trust: Making and breaking cooperative relations. *The British Journal of Sociology*, 13. Advance online publication. doi:10.2307/591021

Garcia, D., Kreutzer, C., Badillo-Urquiola, K., & Mouloua, M. (2015, August 2). *Measuring trust of autonomous vehicles: A development and validation study*. Springer. 10.1007/978-3-319-21383-5_102

Gulz, A., & Haake, M. (2006). Design of animated pedagogical agents: A look at their look. *International Journal of Human-Computer Studies*, 64(4), 322–339. doi:10.1016/j.ijhcs.2005.08.006

Haake, M., & Gulz, A. (2009). A look at the roles of look and roles in embodied pedagogical agents: A user preference perspective. *International Journal of Artificial Intelligence in Education*, 19, 39–71.

Haesler, S., Kim, K., Bruder, G., & Welch, G. (2018). Seeing is believing: Improving the perceived trust in visually embodied Alexa in augmented reality. *2018 IEEE International Symposium on Mixed and Augmented Reality Adjunct*, (pp. 204–205). IEEE. doi:10.1109/ISMAR-Adjunct.2018.00067

Hancock, P. A., Billings, D. R., Schaefer, K. E., Chen, J. Y. C., de Visser, E. J., & Parasuraman, R. (2011). A meta-analysis of factors affecting trust in human-robot interaction. *Human Factors*, 53(5), 517–523. doi:10.1177/0018720811417254 PMID:22046724

Häuslschmid, R., von Bülow, M., Pflöging, B., & Butz, A. (2017). Supporting trust in autonomous driving. *Proceedings of the 22nd International Conference on Intelligent User Interfaces*, (pp. 319–329). ACM. doi:10.1145/3025171.3025198

Helldin, T., Falkman, G., Riveiro, M., & Davidsson, S. (2013, October 30). Presenting system uncertainty in automotive UIs for supporting trust calibration in autonomous driving. *Proceedings of the 5th International Conference on Automotive User Interfaces and Interactive Vehicular Applications*. ACM. doi:10.1145/2516540.2516554

Herse, S., Vitale, J., Ebrahimi, D., Tonkin, M., Ojha, S., Sidra, S., Johnston, B., Phillips, S., Gudi, S. L. K. C., Clark, J., Judge, W., & Williams, M.-A. (2018). Bon appetit! Robot persuasion for food recommendation. *Companion of the 2018 ACM–IEEE International Conference on Human–Robot Interaction*, (pp. 125–126). IEEE. doi:10.1145/3173386.3177028

Hertzum, M., Andersen, H. H., Andersen, V., & Hansen, C. B. (2002). Trust in information sources: Seeking information from people, documents, and virtual agents. *Interacting with Computers*, 14(5), 575–599. doi:10.1016/S0953-5438(02)00023-1

Holthausen, B., Wintersberger, P., Walker, B., & Riener, A. (2020, September 21). Situational trust scale for automated driving (STS-AD): Development and initial validation. *Proceedings of the 12th International Conference on Automotive User Interfaces and Interactive Vehicular Applications*. ACM. doi:10.1145/3409120.3410637

Institut national de la statistique et des études économiques. (2021). *Proportion de bacheliers dans une génération selon la filière et le sexe données Annuelles de 1997 à 2020*. Institut national de la statistique et des études économiques. <https://www.insee.fr/fr/statistiques/2383585>

Jian, J.-Y., Bisantz, A. M., & Drury, C. G. (2000). Foundations for an Empirically determined scale of trust in automated systems. *International Journal of Cognitive Ergonomics*, 4(1), 53–71. doi:10.1207/S15327566IJCE0401_04

Kalegina, A., Schroeder, G., Allchin, A., Berlin, K., & Cakmak, M. (2018). Characterizing the design space of rendered robot faces. *Proceedings of the 2018 ACM–IEEE International Conference on Human–Robot Interaction*, (pp. 96–104). ACM. doi:10.1145/3171221.3171286

- Khastgir, S., Birrell, S., Dhadyalla, G., & Jennings, P. (2018, July). Effect of knowledge of automation capability on trust and workload in an automated vehicle: A driving simulator study. *Proceedings of the International Conference on Applied Human Factors and Ergonomics*, (pp. 410–420). Springer.
- Kim, Y., & Sundar, S. S. (2012). Anthropomorphism of computers: Is it mindful or mindless? *Computers in Human Behavior*, 28(1), 241–250. doi:10.1016/j.chb.2011.09.006
- Knack, S. (2001). Trust, associational life and economic performance. In J. F. Helliwell (Ed.), *The contribution of human and social capital to sustained economic growth and well-being* (pp. 172–202). Enquiries Centre - Human Resources Development Canada
- Kopp, S., Gesellensetter, L., Krämer, N. C., & Wachsmuth, I. (2005). A conversational agent as museum guide—design and evaluation of a real-world application. In *Intelligent Virtual Agents: 5th International Working Conference*. Springer Berlin Heidelberg.
- Körber, M., & Bengler, K. (2014). Potential individual differences regarding automation effects in automated driving. *Proceedings of the XV International Conference on Human–Computer Interaction*. ACM. doi:10.1145/2662253.2662275
- Kuzminykh, A., Govindaraju, N., Avery, J., & Lank, E. (2020). Genie in the bottle: Anthropomorphized perceptions of conversational agents. *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, (pp. 1–13). ACM. doi:10.1145/3313831.3376665
- Law, T., & Scheutz, M. (2021). Trust: Recent concepts and evaluations in human-robot interaction. In C. S. Nam & J. B. Lyons (Eds.), *Trust in human–robot interaction* (pp. 27–57). Academic Press. doi:10.1016/B978-0-12-819472-0.00002-2
- Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50–80. doi:10.1518/hfes.46.1.50.30392 PMID:15151155
- Lemoine, J.-F., & Cherif, E. (2012). Comment générer de la confiance envers un agent virtuel à l’aide de ses caractéristiques? Une étude exploratoire. *Management et avenir*, 58, 169. 10.3917/mav.058.0169
- Lewis, J. D., & Weigert, A. (1985). Trust as a social reality. *Social Forces*, 63(4), 967–985. doi:10.2307/2578601
- Linnemann, G. A., & Jucks, R. (2018). “Can I trust the spoken dialogue system because it uses the same words as I do?”: Influence of Lexically aligned spoken dialogue systems on trustworthiness and user satisfaction. *Interacting with Computers*, 30(3), 173–186. doi:10.1093/iwc/iwy005
- Lobinger, K., & Brantner, C. (2020). Picture-sorting techniques: Card-sorting and Q-sort as alternative and complementary approaches in visual social research. *Sage (Atlanta, Ga.)*, 309–321. Advance online publication. doi:10.4135/9781526417015.n19
- Lorenz, E. (1999). Trust, contract and economic cooperation. *Cambridge Journal of Economics*, 23(3), 301–315. doi:10.1093/cje/23.3.301
- Lugano, G. (2017). Virtual assistants and self-driving cars. *Proceedings of the 15th International Conference on ITS Telecommunications*, (pp. 1–5). IEEE. doi:10.1109/ITST.2017.7972192
- MacDorman, K. (2006). *Subjective ratings of robot video clips for human likeness, familiarity, and eeriness: An exploration of the uncanny valley*. Semantic Scholar. <https://www.semanticscholar.org/paper/Subjective-Ratings-of-Robot-Video-Clips-for-Human-%2C-MacDorman/9bd36d63aac878782184217d73e1fda9b2603bb5>
- Madhok, A. (1995). Revisiting multinational firms’ tolerance for joint ventures: A trust-based approach. *Journal of International Business Studies*, 26(1), 117–137. doi:10.1057/palgrave.jibs.8490168
- Majji, K. C., & Baskaran, K. (2021). Artificial intelligence analytics: Virtual assistant in UAE automotive industry. In V. Suma, J. I.-Z. Chen, Z. Baig, & H. Wang (Eds.), *Inventive systems and control* (pp. 309–322). Springer. doi:10.1007/978-981-16-1395-1_24
- Malle, B. F., & Ullman, D. (2021). A multidimensional conception and measure of human-robot trust. In C. S. Nam & J. B. Lyons (Eds.), *Trust in human–robot interaction* (pp. 3–25). Academic Press., doi:10.1016/B978-0-12-819472-0.00001-0

- Manning, A. D. (1998). Understanding comics: The invisible art. *IEEE Transactions on Professional Communication*, 41(1), 66–69. doi:10.1109/TPC.1998.661632
- Martelaro, N., Nneji, V. C., Ju, W., & Hinds, P. (2016). Tell me more: Designing HRI to encourage more trust, disclosure, and companionship. *Proceedings of the 2016 11th ACM–IEEE International Conference on Human–Robot Interaction*. IEEE. doi:10.1109/HRI.2016.7451864
- Matthews, G., Lin, J., Panganiban, A. R., & Long, M. D. (2019). Individual differences in trust in autonomous robots: Implications for transparency. *IEEE Transactions on Human-Machine Systems*, 50(3), 1–11. doi:10.1109/THMS.2019.2947592
- Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20(3), 709–734. doi:10.2307/258792
- McCloud, S. (1993). *Understanding comics: The invisible art*. Kitchen Sink Press.
- McDonnell, R., Breidt, M., & Bühlhoff, H. H. (2012). Render me real? Investigating the effect of render style on the perception of animated virtual humans. *ACM Transactions on Graphics*, 31(4), 91. Advance online publication. doi:10.1145/2185520.2185587
- Meng, F., Cheng, P., Yao, J., & Wang, Y. (2021). Intelligent agents in cars: Investigating the effects of anthropomorphism and physicality of agents in driving contexts. In T. Z. Ahram & C. S. Falcão (Eds.), *Advances in usability, user experience, wearable and assistive technology* (pp. 685–691). Springer International Publishing. doi:10.1007/978-3-030-80091-8_81
- Minato, T., Shimada, M., Itakura, S., Lee, K., & Ishiguro, H. (2005). Does gaze reveal the human likeness of an android? *Proceedings of the Fourth International Conference on Development and Learning*, (pp. 106–111). IEEE. doi:10.1109/DEVLRN.2005.1490953
- Moussawi, S., & Koufaris, M. (2019). Perceived intelligence and perceived anthropomorphism of personal intelligent agents: Scale development and validation. *Proceedings of the 2019 Hawaii International Conference on System Sciences*. IEEE. doi:10.24251/HICSS.2019.015
- Muir, B. M. (1994). Trust in automation: I. Theoretical issues in the study of trust and human intervention in automated systems. *Ergonomics*, 37(11), 1905–1922. doi:10.1080/00140139408964957
- Nass, C., & Moon, Y. (2000). Machines and mindlessness: Social responses to computers. *The Journal of Social Issues*, 56(1), 81–103. doi:10.1111/0022-4537.00153
- Nomura, T., Suzuki, T., Kanda, T., Han, J., Shin, N., Burke, J., & Kato, K. (2008). What people assume about humanoid and animal-type robots: Cross-cultural analysis between Japan, Korea, and the United States. *International Journal of HR; Humanoid Robotics*, 5(1), 25–46. doi:10.1142/S0219843608001297
- Nordheim, C. B., Folstad, A., & Bjorkli, C. A. (2019). An initial model of trust in chatbots for customer service: Findings from a questionnaire study. *Interacting with Computers*, 31(3), 317–335. doi:10.1093/iwc/iwz022
- Okamoto, S., & Sano, S. (2017). Anthropomorphic AI agent mediated multimodal interactions in vehicles. *Proceedings of the Ninth International Conference on Automotive User Interfaces and Interactive Vehicular Applications Adjunct*, (pp. 110–114). IEEE. doi:10.1145/3131726.3131736
- Parasuraman, R., de Visser, E. J., Lin, M. K., & Greenwood, P. M. (2012). Dopamine beta hydroxylase genotype identifies individuals less susceptible to bias in computer-assisted decision making. *PLoS One*, 7(6), e39675. Advance online publication. doi:10.1371/journal.pone.0039675 PMID:22761865
- Parmar, D., Ólafsson, S., Utami, D., & Bickmore, T. (2018). Looking the part: The effect of attire and setting on perceptions of a virtual health counselor. *Proceedings of the 18th International Conference on Intelligent Virtual Agents*, (pp. 301–306). IEEE. doi:10.1145/3267851.3267915
- Powers, A., & Kiesler, S. (2006). The advisor robot: Tracing people’s mental model from a robot’s physical attributes. *Proceedings of the First ACM SIGCHI–SIGART Conference on Human–Robot Interaction*, (pp. 218–225). ACM. doi:10.1145/1121241.1121280
- Ring, L., Utami, D., & Bickmore, T. (2014). The right agent for the job? In T. Bickmore, S. Marsella, & C. Sidner (Eds.), *Intelligent Virtual Agents* (pp. 374–384). Springer International Publishing. doi:10.1007/978-3-319-09767-1_49

- Robinson, S. L. (1996). Trust and breach of the psychological contract. *Administrative Science Quarterly*, 41(4), 574–599. doi:10.2307/2393868
- Rossi, A., Dautenhahn, K., Koay, K., & Walters, M. L. (2018). The impact of peoples' personal dispositions and personalities on their trust of robots in an emergency scenario. *Paladyn : Journal of Behavioral Robotics*, 9(1), 137–154. doi:10.1515/pjbr-2018-0010
- Rovira, E., Pak, R., & McLaughlin, A. (2017). Effects of individual differences in working memory on performance and trust with various degrees of automation. *Theoretical Issues in Ergonomics Science*, 18(6), 573–591. doi:10.1080/1463922X.2016.1252806
- Ruijten, P., Terken, J., & Chandramouli, S. (2018). Enhancing trust in autonomous vehicles through intelligent user interfaces that mimic human behavior. *Multimodal Technologies and Interaction*, 2(4), 62. doi:10.3390/mti2040062
- Ruijten, P. A. M., Haans, A., Ham, J., & Midden, C. J. H. (2019). Perceived human-likeness of social robots: Testing the Rasch model as a method for measuring anthropomorphism. [z]. *International Journal of Social Robotics*, 11(3), 477–494. doi:10.1007/s12369-019-00516-z
- Schaefer, K. (2016). Measuring trust in human robot interactions: Development of the “Trust Perception Scale—HRI”. In R. Mittu, D. Sofge, A. Wagner, & W. F. Lawless (Eds.), *Robust intelligence and trust in autonomous systems* (pp. 191–218)., doi:10.1007/978-1-4899-7668-0_10
- Schaefer, K. E., Billings, D. R., Szalma, J. L., Adams, J. K., Sanders, T. L., Chen, J. Y., & Hancock, P. A. (2014). *A meta-analysis of factors influencing the development of trust in automation: implications for human-robot interaction*. DTIC. <https://apps.dtic.mil/docs/citations/ADA607926>
- Sebo, S. S., Krishnamurthi, P., & Scassellati, B. (2019). “I don’t believe you”: Investigating the effects of robot trust violation and repair. *Proceedings of the 2019 14th ACM–IEEE International Conference on Human–Robot Interaction*, (pp. 57–65). ACM. doi:10.1109/HRI.2019.8673169
- Shu, P., Min, C., Bodala, I., Nikolaidis, S., Hsu, D., & Soh, H. (2018). Human trust in robot capabilities across tasks. *Companion of the 2018 ACM–IEEE International Conference on Human–Robot Interaction*, (pp. 241–242). ACM. doi:10.1145/3173386.3177034
- Smith, J. K., Liss, M., Erchull, M. J., Kelly, C. M., Adragna, K., & Baines, K. (2018). The relationship between sexualized appearance and perceptions of women’s competence and electability. *Sex Roles*, 79(11), 671–682. doi:10.1007/s11199-018-0898-4
- Strohmman, T., Siemon, D., & Robra-Bissantz, S. (2019). Designing virtual in-vehicle assistants: Design guidelines for creating a convincing user experience. *AIS Transactions on Human-Computer Interaction*, 11(2), 54–78. doi:10.17705/1thci.00113
- ter Stal, S., Tabak, M., op den Akker, H., Beinema, T., & Hermens, H. (2020). Who do you prefer? The Effect of age, gender and role on users’ first impressions of embodied conversational agents in eHealth. *International Journal of Human-Computer Interaction*, 36(9), 881–892. doi:10.1080/10447318.2019.1699744
- van der Woerd, S., & Haselager, P. (2019). When robots appear to have a mind: The human perception of machine agency and responsibility. *New Ideas in Psychology*, 54, 93–100. doi:10.1016/j.newideapsych.2017.11.001
- Vartanian, O., Stewart, K., Mandel, D., Pavlovic, N., McLellan, L., & Taylor, P. (2012). Personality assessment and behavioral prediction at first impression. *Personality and Individual Differences*, 52(3), 250–254. doi:10.1016/j.paid.2011.05.024
- Verberne, F., Ham, J., & Midden, C. (2015). Trusting a Virtual driver that looks, acts, and thinks like you. *Human Factors*, 57(5), 1–15. doi:10.1177/0018720815580749 PMID:25921302
- Waytz, A., Heafner, J., & Epley, N. (2014). The mind in the machine: Anthropomorphism increases trust in an autonomous vehicle. *Journal of Experimental Social Psychology*, 52, 113–117. doi:10.1016/j.jesp.2014.01.005
- Weng, F., Angkitittrakul, P., Shriberg, E. E., Heck, L., Peters, S., & Hansen, J. H. L. (2016). Conversational in-vehicle dialog systems: The past, present, and future. *IEEE Signal Processing Magazine*, 33(6), 49–60. doi:10.1109/MSP.2016.2599201

Wickens, C. D. (2008). Multiple resources and mental workload. *Human Factors*, 50(3), 449–455. doi:10.1518/001872008X288394 PMID:18689052

Willis, J., & Todorov, A. (2006). First impressions: Making up your mind after a 100 ms exposure to a face. *Psychological Science*, 17(7), 592–598. doi:10.1111/j.1467-9280.2006.01750.x PMID:16866745

Yagoda, R., & Gillan, D. (2012). You want me to trust a robot? The development of a human–robot interaction trust scale. *International Journal of Social Robotics*, 4(3), 235–248. doi:10.1007/s12369-012-0144-0

Yee, N., Bailenson, J., & Rickertsen, K. (2007). A meta-analysis of the impact of the inclusion and realism of human-like faces on user experiences in interfaces. *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, (pp. 1–10). IEEE. doi:10.1145/1240624.1240626

Zebrowitz, A., & Montepare, J. M. (2008). First impressions from facial appearance cues. In N. Ambady & J. J. Skowronski (Eds.), *First Impressions* (pp. 171–204). Guilford Publications.

Clarisse Lawson-Guidigbe currently works as a UX Researcher in a french company. Her research work during her Ph.D. thesis focused on anthropomorphic design and its impact on user trust in a self-driving environment.

Nicolas Louveton is a lecturer in psychology and cognitive ergonomics at the University of Poitiers, at the Research Center on Cognition and Learning. He holds a PhD in human movement sciences and human factors. His research and teaching are revolving around human-centered design, visual search of information in graphic displays, collaboration with industrial robots or autonomous cars using virtual and augmented reality. He worked on many projects, involving both academic and industrial partners in several European countries.

Kahina Amokrane-Ferka is Senior Research Engineer in Human Machine Interaction. She is currently a system architect at IRT SystemX, France. Her education includes computer science degree for UMMTO University, Algeria, in 2005, and then Master research degree from Evry Val d'Essonne University, France, in 2006. She received Ph.D. degree in Information Technologies and System from University of Technology of Compiègne in 2010, then worked as academic researcher for 4 years. Her research interests include knowledge management, serious gaming, human factors, simulation, intelligent systems, system engineering, system modeling and human system integration. She has also taught courses in Algorithmic and system modeling at the UTC and Central Supélec University.

Benoît Le Blanc is professor at the National Engineering School of Cognitics (ENSC - Bordeaux INP), specialized in artificial intelligence and applied cognitive sciences. His research activities concern knowledge management, decision systems and human trust in artificial intelligence systems. He receives a PhD (1992) and a French accreditation to conduct research activities (2002) from University of Bordeaux. He is at the editorial board of the scientific journal *Hermes* (CNRS) in Human Sciences and he is President of the French learned society for Artificial Intelligence (AFIA).

Jean-Marc André is a Senior University Professor (Phd in Environmental Biology). He heads the Cognitics group and the Cognitics and Human Engineering (CIH) team of the IMS UMR CNRS 5218 laboratory. He is the scientific director of the École Nationale Supérieure de Cognitique (ENSC) of the Institut Polytechnique de Bordeaux. For the last few years he has been the scientific leader of several research projects in the fields of cognitive engineering: management of cognitive resources, human-system interfaces and collective cognition.