

The role of anthropomorphic, $\hat{x}$ enocentric, intentional, and social (A $\hat{X}$ IS) robotics in human-robot interaction

Anshu Saxena Arora^{a,*}, Amit Arora^a, K. Sivakumar^b, Vasyl Taras^c

^a School of Business and Public Administration, University of the District of Columbia, 4200 Connecticut Ave NW, Washington, DC 20008, USA

^b Department of Marketing, Lehigh University, 621 Taylor Street, Bethlehem, PA 18015, USA

^c Bryan School of Business and Economics, University of North Carolina at Greensboro, 516 Stirling Street, Greensboro, NC 27412, USA

ARTICLE INFO

Keywords:

Anthropomorphism
X $\hat{e}$ nocentrism
Intentionality
Sociality
A $\hat{X}$ IS robotics
Human-robot interaction

ABSTRACT

This research explores the socio-cognitive mechanisms of human intelligence through the lens of anthropomorphic, $\hat{x}$ enocentric, intentional, and social (A $\hat{X}$ IS) robotics. After delving into three pivotal A $\hat{X}$ IS concepts – robotic anthropomorphism, intentionality, and sociality – the study examines their impact on robot likeability and successful human-robot interaction (HRI) implementation. The research introduces the concept of robotic $\hat{x}$ enocentrism (represented by perceived inferiority and social aggrandizement) as a new global dimension in social robotics literature, positioning it as a higher-order concept that moderates the impact of pivotal independent variables on robot likeability. Analyzing a sample of 308 respondents in global cross-cultural teams, the study confirms that pivotal A $\hat{X}$ IS robotics concepts foster positive robot likeability and successful HRI implementation for both industrial and social robots. Perceived inferiority negatively moderated the relationship between anthropomorphism and robot likeability, but it was a positive moderator between intentionality and robot likeability. However, social aggrandizement did not act as a significant boundary condition. Sociality remains unaffected by the moderating influence of $\hat{x}$ enocentrism. The study concludes by outlining future research directions for A $\hat{X}$ IS robotics.

1. Introduction

When [artificial intelligence] is focused on augmenting humans rather than mimicking them, then humans retain the power to insist on a share of the value created. What's more, augmentation creates new capabilities and new products and services, ultimately generating far more value than merely humanlike AI (Brynjolfsson, 2022, p. 272, p. 272).

Human-robot interaction (HRI) is a multidisciplinary field of study that draws research insights from psychology, sociology, computer science, technology, engineering, mathematics, business, and other disciplines (Arora & Arora, 2020; Chang & Kim, 2022; Mahdi, Saleh, Shariff, & Dautenhahn, 2020; Marchesi, Spatola, & Wykowska, 2021). HRI is relevant to a wide range of domains, including manufacturing, healthcare, education, and entertainment. As robots become more sophisticated and integrated into our lives, it is important to develop a deep understanding of how humans interact with robots. This understanding is needed to design robots that are safe, easy to use, and beneficial to society.

Social robots have sparked controversies in ethical, legal, and social spheres, though their paramount importance for the global society is unquestionable (Ullrich & Diefenbach, 2017). Although human-human interaction (HHI) differs substantially from HRI, researchers have tried to map similarities between HHI and HRI to develop social robots that can successfully interact with humans (Schellen & Wykowska, 2019). Social robots are artificial embodied agents that react to natural social stimuli (Chang & Kim, 2022; Wiese, Metta, & Wykowska, 2017; Wykowska, Chaminade, & Cheng, 2016) and have the social and psychological capabilities of social cognition (Schellen & Wykowska, 2019). HRI researchers have paid particular attention to humanoid robots because these robots are deeply engrained in human psychology (Li, Terfurth, Woller, & Wiese, 2022; Wiese et al., 2017; Wykowska et al., 2016) as artificial humanlike agents capable of forming lasting relationships with humans by displaying anthropomorphism (i.e., embodying human characteristics; Zogaj, Mähner, Yang, & Tschulin, 2023; Arora, Parnell, & Arora, 2022; Kaplan, Sanders, & Hancock, 2019; Woods et al., 2007) resulting in a successful HRI implementation.

* Corresponding author.

E-mail addresses: anshu.arora@udc.edu (A.S. Arora), amit.arora@udc.edu (A. Arora), k.sivakumar@lehigh.edu (K. Sivakumar), v_taras@uncg.edu (V. Taras).

<https://doi.org/10.1016/j.chbah.2023.100036>

Received 5 July 2023; Received in revised form 2 December 2023; Accepted 3 December 2023

Available online 7 December 2023

2949-8821/© 2023 The Author(s). Published by Elsevier Inc. This is an open access article under the CC BY-NC-ND license (<http://creativecommons.org/licenses/by-nc-nd/4.0/>).

A ‘successful HRI implementation’ is characterized by a social robot engaging consumers through a meaningful user experience. Key components of such an engagement include a careful consideration of the use context, a focused purpose integrated into the robot design and use, and an interaction flow relying on both verbal and non-verbal communication (Dang & Liu, 2023; Kim & Im, 2023; Mikkelsen & Rehm, 2022, pp. 950–955). In real-world scenarios, successfully implemented robots are generally more socially acceptable or artificial agents that are successfully integrated with society and complete their tasks (as needed) without having the sense of the presence of any foreign agent among humans. An example of a “successfully implemented robot” is PackBot, which is a range of military robots developed by iRobot - an international robotics company designed for search and rescue activities. These robots were deployed in Iraq and Afghanistan; they were also used to search the debris at the World Trade Center in 2001 and the Fukushima Nuclear plant disaster.¹

Some examples of anthropomorphized humanoids include: (a) Pepper Robot, an internationally renowned humanoid robot developed by SoftBank Robotics, is used as a humanoid companion that communicates intuitively with color-changing lights in its eyes and the tablet on its torso. It operates in Pepper Parlor cafés in Japan, taking customer orders and interacting with customers at tables (Chang & Kim, 2022); (b) Ocean One, a bimanual underwater humanoid robot created by the Stanford Robotics Lab, explores coral reefs; (c) ATLAS humanoid, developed by Boston Dynamics with funding from the United States Defense Advanced Research Projects Agency, displays agility and maneuvering capabilities for navigating difficult terrains through range sensing, stereo vision, and other sensors; (d) NAO humanoid, developed by the French robotics company Aldebaran Robotics, has been successful working with children diagnosed with autism and other learning disorders; and (e) Sophia, developed by roboticist David Hanson and Hanson Robotics, was the first robot to be granted citizenship by the Kingdom of Saudi Arabia (Fernandes, 2022, pp. 51–64). In addition, schools and other public institutions have recently begun using robots to minimize the spread of COVID-19 by cleaning surfaces with ultraviolet radiation, sanitizing floors, scanning schoolchildren for fevers, enforcing mask-wearing, spraying antimicrobial gases and disinfectants in outdoor public spaces, and taking on jobs considered dangerous for humans (Mims, 2020).

In HRI situations, robots exhibit anthropomorphism through their looks and/or voice (e.g., Asimo, Kirobo Mini, Pepper, Nao, Sophia, MIT robots Kismet and Cog). They also display intentionality, which refers “more narrowly to adopting a strategy in predicting and explaining others’ behavior with reference to mental states” (Schellen & Wykowska, 2019, p. 139). In addition, the concept of xenocentrism (Arora & Arora, 2020), or “the belief that what is foreign is best, that our own lifestyle, products, or ideas are inferior to others” (Eshleman, Cashion, & Basirico, 1993, p. 109), adds value to social robotics. We anticipate direct linkages and relationships among anthropomorphism, intentionality, and sociality to robot likeability and HRI implementation, and we propose that xenocentrism moderates these relationships. To distinguish and differentiate the moderating effects of xenocentrism (considered a higher-order variable) from the direct effects of anthropomorphism, intentionality, and sociality on robot likeability, herein, we refer to this higher order construct as “ $\hat{x}$ enocentrism.” In line with the social-collaborative robotics domain, we collectively refer to these phenomena as anthropomorphic, $\hat{x}$ enocentric, intentional, and social ($\hat{A}\hat{X}$ IS) robotics in the HRI context.

Our paper makes several contributions to address the dearth of research in social robotics and HRI. First, we enhance the understanding of social-collaborative robotics by comprehensively considering levels of

$\hat{A}\hat{X}$ IS in social robotics and their subsequent influence on robot likeability and successful HRI implementation. Second, we add to the limited research addressing how these socio-behavioral relationships are associated with $\hat{A}\hat{X}$ IS characteristics in robots (Ciardo, De Tommaso, & Wykowska, 2022; Kaplan et al., 2019; Lajante, Tojib, & Ho, 2023; Letheren, Kuhn, Lings, & Pope, 2016; Marchesi et al., 2021; Woods et al., 2007). We delve into each of the constructs using the *computers are social actors* (CASA) paradigm as an overarching theory, the uncanny valley effect, and the sociality, effectance, and elicited agent knowledge (SEEK) model theories. Third, we investigate robotic $\hat{x}$ enocentrism as a higher-order social robotics construct comprising perceived inferiority and social aggrandizement through the lenses of robotic anthropomorphism, robotic intentionality, and robotic sociality. This construct and its impact have not been previously examined in social robotics and HRI research. Furthermore, we analyze the linkages of robotic $\hat{x}$ enocentrism to robot likeability and HRI implementation (Arora & Arora, 2020). We strive to address the following questions: How do robots mimic humanlike characteristics, and how do their implicit characteristics of sociality and intentionality arouse robot likeability that leads to a successful HRI? Our research explores these questions and fills the research gaps through an in-depth examination of $\hat{A}\hat{X}$ IS robotics that focuses on robot likeability and overall successful HRI implementation.

This article consists of four sections. First, we define and describe $\hat{A}\hat{X}$ IS robotics as a part of social-collaborative robotics in HRI. Second, we examine how robotic $\hat{A}\hat{X}$ IS are interrelated in the social-collaborative robotics and HRI context, propose our conceptual framework, and develop a series of hypotheses. Third, we test our conceptual framework by analyzing data from 308 respondents using the moderating effects of $\hat{A}\hat{X}$ IS robotics on robot likeability and HRI implementation. Fourth, we discuss our study’s theoretical and practical implications, limitations, and future research directions.

2. Conceptual background

2.1. Social robotics and HRI

Humanoid social robots are human-made technologies with physical (e.g., NAO, Pepper, Zora) or digital forms (e.g., voice assistants such as Siri and Alexa, chatbots) that bear some resemblance to humans, whether bodily or through anthropomorphic/humanlike features (Di Dio et al., 2020; Fox & Gambino, 2021; Liu & Sundar, 2018). Social robotics is an emerging research field, and HRI research is still in its infancy. One of the significant theories applied to HRI is the CASA (computers are social actors) paradigm. CASA framework was derived from Reeves and Nass’s (1996) media equation, and it refers to the phenomenon of humans conditioned to react mindlessly to technology, thereby treating “technology” as yet another “social being” and mimicking HHI in HRI situations (Fox & Gambino, 2021; Gambino, Fox, & Ratan, 2020). Although social robotics research may question whether CASA applies equally to both HHI and HRI situations, elements of CASA perspectives are clearly present in both HHI and HRI, and modern-day HRI is aimed to simulate HHI, especially when the robots demonstrate social and anthropomorphic cues. As technologies in artificial intelligence (AI) and robotics advance, social robots continue to become increasingly sophisticated, thus blurring the boundaries of HHI into HRI despite complexity and cost constraints (Fox & Gambino, 2021; Nicolas & Agnieszka, 2021; Song & Kim, 2022).

2.2. $\hat{A}\hat{X}$ IS robotics

The HRI literature defines social-collaborative robots as robots that are (1) socially evocative (relying on CASA and anthropomorphic features), (2) socially situated (relying on social and environmental cues), (3) sociable (relying on active human engagement and social cognition),

¹ <https://interestingengineering.com/innovation/a-brief-history-of-military-robots-including-autonomous-systems>.

Table 1

Definitions of key terms/concepts used in the current research.

Concept	Definition	Theories Used	Sources
AXIS Robotics	The current research proposes a mix of robotics concepts of anthropomorphism, $\hat{x}$ enocentrism, intentionality, and sociality in the HRI context.	The CASA paradigm means that humans are conditioned to react mindlessly to technology, thereby treating technology as another social being and mimicking HHI in HRI situations.	Kim and Im (2023) Fox and Gambino (2021) Gambino et al. (2020)
Robotic Anthropomorphism	Robotic anthropomorphism entails exhibiting human characteristics in robots (e.g., mimicking human emotions in robots, facial and voice recognition, and exhibiting walking/dancing behaviors). Humans anthropomorphize robots by simulating and exhibiting emotional associations with artificial human agents.	The uncanny valley effect is a dip in positive perception, whereby the robotic likeability increases as the robot becomes more humanlike and then drops if the robot becomes too humanlike.	Dang and Liu (2023) Chung, Kang, and Jun (2023) Arora, Arora, Jentjens, McIntyre, and Sepehri (2022) Schuetz and Venkatesh (2020) Damiano and Dumouchel (2018) Turtle (2017) Wiese et al. (2017) Aly & Tapus (2016) Zlotowski, Proudfoot, Yogeewaran, and Bartneck (2015) Hesslow (2012) Bartneck, Kulić, Croft, and Zoghbi (2009) Epley, Waytz, and Cacioppo (2007) Sung et al. (2007) Hesslow (2002) Wilson (2002) Chartrand and Bargh (1999) Mori (1970)
Robotic $\hat{x}$enocentrism	$\hat{x}$ enocentrism is “a psychological attribute which implies a biased view ... one who is xenocentric sees faults where none exist” (Kent & Burnight, 1951, pp. 256–57).	System justification theory uses the psychological process by which existing social arrangements are legitimized at the expense of personal and group interests	Klüber and Onnasch (2022) Arora and Arora (2020) Balabanis and Diamantopoulos (2016) Eshleman et al. (1993) Zhou and Belk (2004)
Robotic Intentionality	Robotic intentionality, or intentional stance, in AXIS robots, involves activating brain regions related to mentalizing and social cognition. This leads to positive or negative human behaviors toward robots and high-level decision-making in the HRI context.	Uncanny Valley Effect theory	Marchesi et al. (2021) Nicolas and Agnieszka (2021) Spatola and Wudarczyk (2021) Schellen and Wykowska (2019) Hesslow (2012) Wiese et al. (2017) Özdem et al. (2017) Bartneck et al. (2009)
Robotic Sociality	Humans are motivated to seek alternative ways (e.g., using robots as social companions) through the social reconnection hypothesis.	The SEEK model theory predicts that humans like to interact with technology (e.g., social robots) when they are motivated to be effective social agents and/or when they lack a sense of social connection to other humans	Leo-Liu (2023) Christoforakos and Diefenbach (2022) Marchesi et al. (2021) Kwok, Grisham, and Norberg (2018) Damiano and Dumouchel (2018) Gaudiello, Zibetti, Lefort, Chetouani, and Ivaldi (2016) Bartneck et al. (2009) Epley et al. (2007) DeWall and Baumeister (2006)
Robot Likeability	Favorable attitudes and behaviors toward social robots. Research shows that there is more likeability between robots and humans, especially individuals with ASD, as anthropomorphic characteristics increase. Anthropomorphism amplifies the big five human personality traits in individuals with ASD and other learning/cognitive disabilities.	Uncanny Valley Effect theory	Li, Guo, Wang, Chen, and Ham (2023) Klüber and Onnasch (2022) Chang and Kim (2022) Li et al. (2022) Marchesi et al. (2021) Arora, Fleming, Arora, Taras, and Xu (2021) Bartneck et al. (2009)

(continued on next page)

Table 1 (continued)

Concept	Definition	Theories Used	Sources
Successful and Positive HRI Implementation	In $\hat{\text{A}}\hat{\text{X}}\text{IS}$ robotics, when robots mimic human emotions and portray humanlike characteristics (both implicitly and explicitly) either through their humanlike face, voice, and other external features or through their internal and implicit characteristics of intelligence and intentional mindset, they arouse positive human behaviors toward robots and lead to a positive implementation for robots in human spheres.	Uncanny Valley Effect theory	Dang and Liu (2023) Li et al. (2023) Mikkelsen and Rehm (2022) Christoforakos and Diefenbach (2022) Chang and Kim (2022) Li et al. (2022) Arora, Parnell, and Arora (2022) Marchesi et al. (2021) Bartneck et al. (2009)

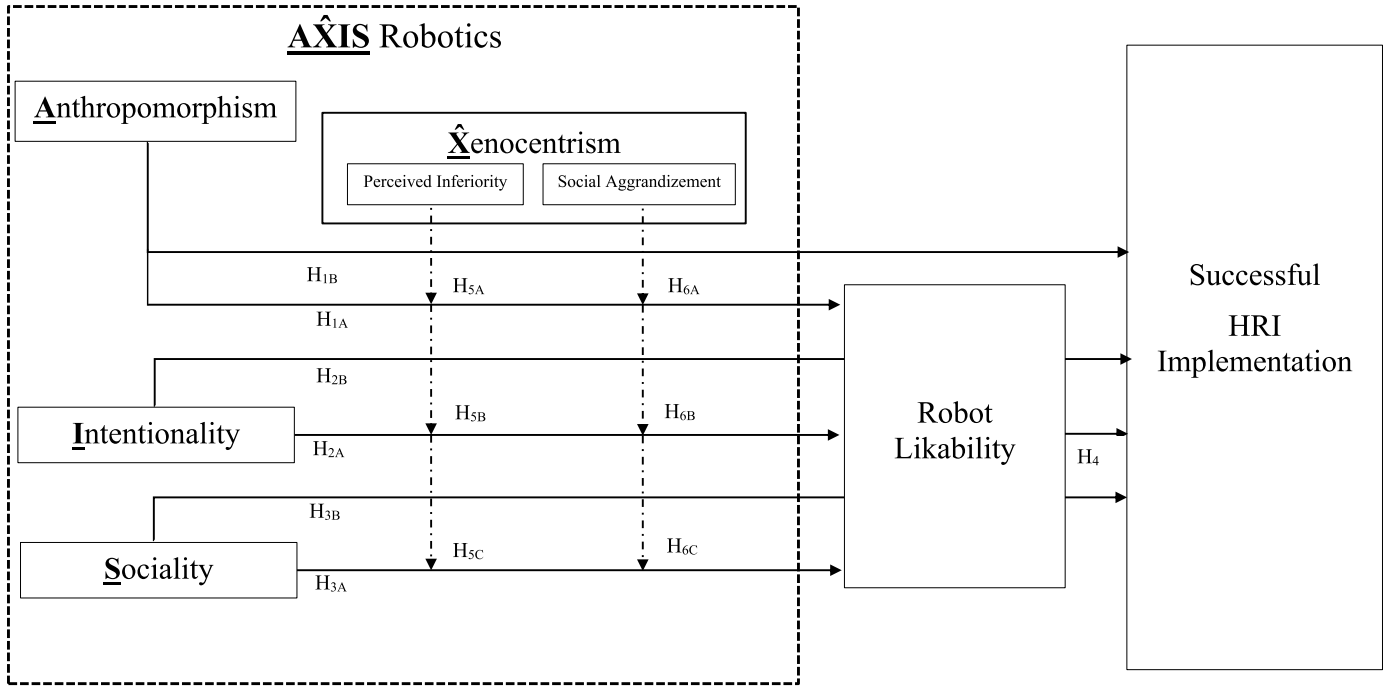


Fig. 1. A relational view of the $\hat{\text{A}}\hat{\text{X}}\text{IS}$ robotics-likeability-HRI success framework.

(4) socially intelligent (relying on human social intelligence and human cognition), and (5) socially interactive (relying on peer-to-peer HRI) (Grundke, Stein, & Appel, 2023; Mahdi et al., 2020). $\hat{\text{A}}\hat{\text{X}}\text{IS}$ robots refer to social-collaborative robots (e.g., robot interrogators, therapeutic robots, physical assistance robots, Wizard-of-Oz (WoZ), industrial robots with human interaction capabilities) that employ principles of socio-cognitive intelligence by ensuring that social robots collaborate, “follow and exhibit socially acceptable behaviors, and understand the societal and ethical consequences of their interactions in the sociocultural context in which they operate” (Arora & Arora, 2020, p. 4; see also Leo-Liu, 2023; Spatola & Wudarczyk, 2021). Collaborative robots (cobots) are generally flexible, easily programmable, capable of working with humans in social contexts and sharing workplaces, and offer organizations the ability to successfully implement human-robot interaction (HRI) experiences and applications (Cherubini et al., 2016; Schou et al., 2018; Kopp et al., 2021). Humanoid robots exhibit socio-cognitive intelligence by modeling humanlike intelligence through human cognition, including decision-making, perception, reasoning, and problem-solving skills (Arora, Parnell, & Arora, 2022; Li et al., 2022; Lieto, Chella, & Frixione, 2017). We argue that social robotics should be defined and examined through $\hat{\text{A}}\hat{\text{X}}\text{IS}$ robotics, which covers social, cognitive, and collaborative robotics. Table 1 illustrates key concepts of $\hat{\text{A}}\hat{\text{X}}\text{IS}$ robotics, definitions, theories used, and sources/references.

Fig. 1 exemplifies relationships among the $\hat{\text{A}}\hat{\text{X}}\text{IS}$ constructs that lead to robot likeability and successful HRI outcomes along with moderating effects, as described in H₁–H₆. Previous research in social robotics and HRI has not investigated robotic $\hat{\text{x}}$ enocentrism. To address this gap, we explore this concept in depth and through the lenses of the three pivotal concepts of robotic anthropomorphism, intentionality, and sociality, as well as the impact of these concepts on robot likeability and successful HRI implementation. Robotic $\hat{\text{x}}$ enocentrism can be considered a higher-order concept that comprises perceived inferiority and social aggrandizement, and it moderates the linkages of robotic anthropomorphism, intentionality, and sociality to robot likeability and HRI implementation (Arora & Arora, 2020). In $\hat{\text{A}}\hat{\text{X}}\text{IS}$ robotics, when robots mimic human emotions and portray humanlike characteristics (both implicitly and explicitly)—whether through their humanlike faces, voices, or other external features or their mental, internal, and implicit characteristics of intelligence, sociality, and intentional mindset—they arouse robot likeability moderated by robotic $\hat{\text{x}}$ enocentrism, which in turn leads to a successful HRI implementation.

2.2.1. Anthropomorphism

Our study’s first independent, pivotal variable/construct is robotic anthropomorphism. Anthropomorphism is “the human tendency to attribute human traits to non-human entities” (Damiano & Dumouchel,

2018, p. 2; see also Chung et al., 2023; Li et al., 2023; Epley et al., 2007; Zlotowski et al., 2015), and robotic anthropomorphism refers specifically to robots exhibiting human characteristics (e.g., mimicking human emotions, facial and voice recognition, exhibiting walking/dancing behaviors). Humans anthropomorphize robots by simulating and exhibiting emotional associations with artificial human agents (Arora, Parnell, & Arora, 2022; Schuetz & Venkatesh, 2020; Sung et al., 2007; Turkle, 2017; Zogaj et al., 2023). Robotic anthropomorphism can facilitate strong social relationships, interactions, and exchanges between humans and robots. Robot designers and roboticists are aware that human likeness to robots increases with the addition of more anthropomorphic robotic features, and companies such as Furhat Robotics and SoftBank Robotics make and design their robots using the principles of anthropomorphism. The principle of robotic anthropomorphism uses socio-cognitive processes based on perception, action, and emotion and emphasizes thinking similar to interaction with the external environment (Hesslow, 2002, 2012, Wilson, 2002; Aly & Tapus, 2016, p. 193; see also Chartrand & Bargh, 1999).

Mori's (1970) uncanny valley effect can be described as the phenomenon of robotic likeability increasing up to the point at which the robot becomes highly humanlike or humanoid and then drops when the robot becomes too humanlike (Arora et al., 2021; Spatola & Wudarczyk, 2021). The uncanny valley effect curve proves that although humans tend to anthropomorphize robots, they accept the robot's humanlike attributes (or anthropomorphism) only to a certain extent, beyond which human affinity/likeness for robots decreases exponentially. According to Marchesi et al. (2021), anthropomorphism and how humans differ in their tendencies to anthropomorphize robots can be explained by three factors: the display of an agent's physical characteristics activating knowledge and heuristics related to humans, fulfillment of humans' sociality needs, and human personality traits. The uncanny valley effect is critical for robot design and development in terms of the robot's anthropomorphic appearance and behavior, considering the primary aim for robot development is human likeability, social interaction, and responsiveness toward robots. In other words, "an increase of human likeness raises a robot's likeability until the resemblance becomes nearly perfect" (Damiano & Dumouchel, 2018, p. 2), a point that can be designated as a social threshold between human likeness and familiarity toward robots leading to a successful HRI implementation. The more anthropomorphic a robot is (until the point of social threshold, as highlighted in the uncanny valley effect curve), the more likable the robot is, and the more likelihood of a successful HRI implementation for humans (Chung et al., 2023; Damiano & Dumouchel, 2018; Kim & Im, 2023; Marchesi et al., 2021). Therefore, we posit the following:

H1A, 1B. Robotic anthropomorphism is positively associated with robot likeability and successful HRI implementation.

2.2.2. *Ẋenocentrism*

Robotic Ẋenocentrism refers to a psychological attitude that implies a biased and favorable view of social robots and uses the system justification theory (Balabanis & Diamantopoulos, 2016; Hesslow, 2012), in which consumers prefer social robots that reflect social power because of a social justification motive (Arora & Arora, 2020; Shepherd, Chartrand, & Fitzsimons, 2015). In social robotics and HRI, Ẋenocentric robots are liked and preferred due to two traits of Ẋenocentrism: (1) perceived inferiority (i.e., ingroup derogation whereby people negatively stereotype and undervalue themselves and fail to appreciate domestic products and brands), which is related to AẊIS robotics in that humans perceive robots to be "superior" to humans (in other words, humans tend to perceive themselves as "inferior" to robots and AI) due to the concept of foreignness, and (2) social aggrandizement (i.e., outgroup favoritism, whereby people prefer foreign goods more than domestic ones to enhance perceived social status), which is related to AẊIS robotics in that robots are perceived as more intelligent than humans; come from a different, foreign world; and have a different composition

(Arora & Arora, 2020; Balabanis & Diamantopoulos, 2016).

Herein, we treat Ẋenocentrism as a moderator rather than an independent, pivotal concept because it is a higher-order construct measured through perceived inferiority and social aggrandizement. Defined as "the belief that what is foreign is best, that our own lifestyle, products, or ideas are inferior to those of others" (Eshleman et al., 1993, p. 109; see also Kent & Burnight, 1951), its key attribute is negative stereotypical perceptions of one's own group or perceived inferiority, whereby in the context of humans versus robots, robots are always perceived to be more intelligent. These perceptions are coupled with favoritism toward outgroups or social aggrandizement, whereby in the context of humans versus robots, humans prefer robots with reliable information and technology aid (e.g., using robot vacuum cleaners in homes; asking Siri, Alexa, or Google for news, weather conditions, or help with navigating when driving; robots being used in hospitals for performing intricate surgeries). Ẋenocentrism is studied in conjunction with the system justification theory, defined as "the psychological process by which existing social arrangements are legitimized, even at the expense of personal and group interests" (Jost & Banaji, 1994, p. 2), whereby humans accept and legitimize robots as superior to themselves due to outgroup favoritism and their alleged inferiority in the HRI context.

2.2.3. *Intentionality*

Our study's second independent, pivotal variable/construct is robotic intentionality, or intentional stance, defined as activating brain regions related to mentalizing and social cognition leading to positive or negative human behaviors toward robots and high-level decision-making in HRI (Candrian & Scherer, 2022; Schellen & Wykowska, 2019). Human thinking employs the brain's "perception, action, and emotion. The mental model covertly simulates actions and their associated perceptual effects" (Vanderelst & Winfield, 2018, p. 57). Socio-cognitive processes result in an intentional stance or intentionality, in which "humans engage socially, wisely, and utilize their social cognition and information processing with robots with the assumption that their interaction partner has a brain resulting in highly efficient [HRIs]" (Arora & Arora, 2020, p. 8; see also Hesslow, 2012; Özdem et al., 2017; Schellen & Wykowska, 2019; Wiese et al., 2017).

In the HRI context, intentionality is related to perceived robotic intelligence that is governed by humans' (state-of-mind) social associations and mental states, resulting in favorable behaviors toward robots (Arora & Arora, 2020; Christoforakos & Diefenbach, 2022; Dennett, 1971, 1988, 1997; Marchesi et al., 2019; Schellen & Wykowska, 2019). While using humanlike robots has resulted in positive outcomes such as increased feelings of familiarity or ease in working with robots (Saupé & Mutlu, 2015), researchers have also identified consumers' adverse feelings (Esterwood & Robert, 2023; Mori, MacDorman, & Kageki, 2012) toward robots due to uncanny valley effect theory. Studies show that when robots work with children diagnosed with cognitive disorders (e.g., autism), successful outcomes occur in these HRI implementation scenarios because the children enjoy interacting with robots more than humans due to overall robot likeability, situatedness, embodiment, and consistency (Dautenhahn, Werry, Salter, & Boekhorst, 2003; Ferrara & Hill, 1980; Hudson & Lewis, 2020; Huijnen, Lexis, Jansens, & de Witte, 2019). Therefore, we posit the following:

H2A, 2B. Robotic intentionality is positively associated with robot likeability and successful HRI implementation.

2.2.4. *Sociality*

The third independent, pivotal variable/construct concept, robotic sociality, refers to the phenomenon that people's tendency to anthropomorphize is related to the fundamental need for sociality, acknowledged in the sociality, effectance, and elicited agent knowledge (SEEK) model theory (Christoforakos & Diefenbach, 2022; Epley et al., 2007). Humans prioritize their social well-being (Ormel, Lindenberg, Steverink, & Verbrugge, 1999). If their social needs are unsatisfied, they are

motivated to seek alternative ways (e.g., using robots as social companions) through the social reconnection hypothesis (DeWall & Baumeister, 2006; Piçarra & Giger, 2018). We argue that social needs act as drivers to search for social cues and fulfill human needs to anthropomorphize animals/non-living objects by attributing humanlike characteristics to them (Christoforakos & Diefenbach, 2022). Kwok et al. (2018) relate anxious attachment with social tendencies and prove that these tendencies are moderately positively related. Humans with intense isolation levels tend to exhibit stronger robot sociality and likeability due to their increased social needs, and successful HRI implementation outcomes are likely due to the sociality needs of humans through HRI. Therefore, we posit:

H3A, 3B. Robotic sociality is positively associated with (a) robot likeability and (b) successful HRI implementation.

Humans prefer A $\hat{X}$ IS robots because they are more humanlike (often humanoid), sophisticated, artificially intelligent foreign beings. Robot likeability plays a significant role and leads to a successful HRI implementation (Arora et al., 2021; Christoforakos & Diefenbach, 2022; Klüber & Onnasch, 2022; Li et al., 2022; Marchesi et al., 2021). Thus, we posit the following:

H4. Robot likeability is positively associated with successful HRI implementation.

2.2.5. Anthropomorphism, intentionality, and sociality as independent constructs/variables

As proposed in our A $\hat{X}$ IS robotics-likeability-HRI success framework (Fig. 1), anthropomorphism, intentionality, and sociality are considered as independent variables because they can be manipulated and measured independently. Robotic anthropomorphism has been studied in conjunction with intentional stance/intentionality in robots, focusing on people's attribution of varying degrees of anthropomorphism and intentionality to robot actions (Bossi et al., 2020; Hegel, Krach, Kircher, Wrede, & Sagerer, 2008; Marchesi et al., 2019; Thellman & Ziemke, 2021). Dennett (1971, 1988, 1997) examines intentional stance or intentionality through various human attitudes (toward robots) for predicting (robot) behavior in HRI implementation. For the intentionality construct, social robots may be programmed to act as if they adopt (or not) an intentional stance, which may result in increased likeability and an efficient HRI implementation, even though these humanoid social robots do not possess beliefs and desires as they are interpreted in folk psychology (Schellen & Wykowska, 2019).

Studies show conflicting relationships regarding robotic anthropomorphism and intentionality. One stream of research posits that anthropomorphism results in intentionality—in other words, “a higher tendency to anthropomorphize is associated with a higher tendency to adopt the intentional stance to explain the behavior of a humanoid robot” (Marchesi et al., 2019; Marchesi et al., 2021, p. 4). However, another stream posits that anthropomorphism and intentional stance/intentionality are separate, independent constructs because the adoption of an intentional stance relies primarily on neuropsychological processes of predicting human behavior toward robots regarding mental states that are independent of humans' ability to anthropomorphize robots (Schellen & Wykowska, 2019; Wiese et al., 2017; Wykowska et al., 2016). We align with the latter view, even though anthropomorphism and intentionality both focus on the human psyche and traits.

Robots demonstrate sociality through their capability to interact with humans through trust-building during HRI situations by following social acceptance norms (Gaudiello et al., 2016; Ribino, Lodato, & Infantino, 2018). Previous robotics research has explored robotic sociality in conjunction with the attribution of anthropomorphic qualities to robots and robotic agents (e.g., Christoforakos & Diefenbach, 2022; Epley et al., 2007; Epley, Akalis, Waytz, & Cacioppo, 2008; Niemyjska & Drat-Ruszczak, 2013). Studies have focused on comparing anthropomorphic and non-anthropomorphic technology agents and assessing

Table 2
Hypotheses, theories used, and rationale.

Hypotheses	Theories Used	Rationale
H _{1A} , 1B. Anthropomorphism is positively associated with robot likeability and successful HRI implementation.	The CASA framework, derived from Reeves and Nass's (1996) media equation, suggests that humans apply stereotypes and norms to computers and technology, further assigning personality traits and making inferences as if the computers were human, even though they understand that computers are not human	The more anthropomorphic a robot is (until the point of social threshold, as highlighted in the uncanny valley effect curve), the more likable the robot is, and anthropomorphism will result in a successful HRI implementation for humans (Damiano & Dumouchel, 2018; Marchesi et al., 2021).
H _{2A} , 2B. Intentionality is positively associated with robot likeability and successful HRI implementation.	The uncanny valley effect is a dip in positive perception, whereby the robotic likeability increases as the robot becomes more humanlike and then drops if the robot becomes too humanlike.	Intentionality (a.k.a., intentional stance or intentional mindset) is related to perceived robotic intelligence that is governed by humans' (state-of-mind) social associations and mental states, resulting in favorable behaviors toward robots (Arora & Arora, 2020; Christoforakos & Diefenbach, 2022; Dennett, 1971, 1997; Schellen & Wykowska, 2019).
H _{3A} , 3B. Sociality is positively associated with robot likeability and successful HRI implementation.	The SEEK model theory predicts that humans like to interact with technology (e.g., social robots) when motivated to be effective social agents and/or lacking a sense of social connection to other humans.	Sociality acts as a driver to search for social cues and fulfill human needs to anthropomorphize animals/non-living objects by attributing humanlike characteristics to these objects (Christoforakos & Diefenbach, 2022). Humans with intense isolation levels tend to exhibit stronger sociality, robot likeability due to their increased social needs, and successful HRI implementation outcomes due to the sociality needs of humans through HRI (Kwok et al., 2018).
H ₄ . Robot likeability is positively associated with successful HRI implementation.	The CASA paradigm means that humans are conditioned to react mindlessly to technology, thereby treating technology as another social being and mimicking HHI in HRI situations.	Humans prefer A $\hat{X}$ IS robots since they are more humanlike (often humanoids), sophisticated, artificially intelligent foreign beings. Robot likeability plays a significant role in a successful HRI implementation (Arora & Arora, 2020; Arora et al., 2021; Christoforakos & Diefenbach, 2022; Li et al., 2022; Marchesi et al., 2021).
H _{5A} , 5B, 5C. $\hat{X}$ enocentrism (through perceived inferiority) moderates the relationships among (a) anthropomorphism, (b)	System justification theory uses the psychological process by which existing social	Robotic $\hat{X}$ enocentrism is a boundary condition among the pivotal concepts of anthropomorphism, intentionality, sociality,

(continued on next page)

Table 2 (continued)

Hypotheses	Theories Used	Rationale
intentionality, and (c) sociality and robot likeability.	arrangements are legitimized at the expense of personal and group interests	and robot likeability. While humanlike robots have resulted in positive outcomes such as increased feelings of familiarity or ease in working with robots (Sauppé & Mutlu, 2015), researchers have also identified adverse (negative) feelings (Mori, 1970; Mori et al., 2012) of consumers toward robots. Social robotics research has never used $\hat{\text{x}}\text{enocentrism}$ as a construct in measuring a successful HRI implementation, so the direction of moderating effects is unknown.
H6A, 6B, 6C. $\hat{\text{x}}\text{enocentrism}$ (through social aggrandizement) moderates the relationships among (a) anthropomorphism, (b) intentionality, and (c) sociality and robot likeability.		

their sociality needs (on intentional and behavioral levels), but no causality has been established (Christoforakos & Diefenbach, 2022; Mourey, Olson, & Yoon, 2017). Thus, we investigate sociality as an independent variable, such as anthropomorphism and intentionality, in our $\hat{\text{A}}\hat{\text{X}}\hat{\text{I}}\hat{\text{S}}$ robotics-likeability-HRI success framework (Fig. 1).

2.2.6. $\hat{\text{x}}\text{enocentrism}$ and robot likeability

Robotic $\hat{\text{x}}\text{enocentrism}$ is closely related to the concept of foreignness and robotic intelligence in $\hat{\text{A}}\hat{\text{X}}\hat{\text{I}}\hat{\text{S}}$ robots. Humans perceive robots as foreign beings who are “intelligent” (more intelligent than humans). Robotic intelligence refers to displaying both AI and socio-cognitive intelligence (Bartneck et al., 2009; Hesslow, 2012). Social robots face the challenge of behaving intelligently in HRI situations requiring high levels of social cognition and human decision-making skills. Although AI simulations in robots work well in experimental design methods and situations, when social robots are deployed in the complex world of everyday users, their limitations will become known to their users because they interact with these users over a time span of years rather than a few minutes or seconds (Bartneck et al., 2009). Bartneck et al. (2009, p. 78) propose a series of questionnaires to measure the users’ perception of robots called “Godspeed” because “it is intended to help creators of robots on their development journey.”

This research is the first of its kind and makes several contributions to theory and practice in $\hat{\text{A}}\hat{\text{X}}\hat{\text{I}}\hat{\text{S}}$ robotics. The most significant contribution is the identification of robotic $\hat{\text{x}}\text{enocentrism}$ as a boundary condition among the pivotal concepts of anthropomorphism, intentionality, and sociality leading to robot likeability. Drawing from system justification theory, Balabanis and Diamantopoulos (2016) conceptualized the ‘consumer xenocentrism’ construct as a combination of two dimensions: ‘perceived inferiority’ and ‘social aggrandizement’. They developed (and validated) a new scale (the C-XENSCALE) for understanding consumer attraction toward foreign things (social robots in this research) for measuring consumers’ xenocentric tendencies (p. 58). While some studies show that humanlike robots have resulted in positive outcomes such as increased feelings of familiarity or ease in working with robots (Sauppé & Mutlu, 2015), others identify consumers’ adverse feelings (Esterwood & Robert, 2023; Mori, 1970; Mori et al., 2012) toward robots. Social robotics research has never used $\hat{\text{x}}\text{enocentrism}$ (through perceived inferiority and social aggrandizement) to assess a successful HRI implementation; thus, the direction of moderating effects is unknown. With limited research available in the field of HRI, we rely on empirical results for assessing the impact of these moderating effects of $\hat{\text{x}}\text{enocentrism}$ in social robotics. Thus, we offer the following hypotheses

(considering positive moderation effects for the purpose of this research):

H5A, 5B, 5C. The impact of perceived inferiority (first dimension of robotic $\hat{\text{x}}\text{enocentrism}$) of humans vis-à-vis robots on robot likeability is stronger (and positive) when the robot exhibits (a) anthropomorphism, (b) intentionality, and (c) sociality.

H6A, 6B, 6C. The impact of social aggrandizement (second dimension of robotic $\hat{\text{x}}\text{enocentrism}$) between humans and robots on robot likeability is stronger (and positive) when the robot exhibits (a) anthropomorphism, (b) intentionality, and (c) sociality.

Table 2 summarizes the hypotheses, theories used, and rationales.

3. Methodology

To collect the data to test the conceptual framework, we used the X-Culture project (www.X-Culture.org), a large-scale international business collaboration and consulting project that employs approximately 5000 participants per academic semester. The participants are business students and working professionals from over 150 universities in 50 countries in six continents. The project is run twice a year. Working in global virtual teams, typically six to seven people per team, each from a different country, the project participants spend the semester solving real-life challenges presented by client companies, typically involving market research, competition analysis, product design, developing a marketing strategy, and completing other tasks related to identifying market expansion opportunities for the client company. The project participants rely on the same online collaboration and communication tools commonly used in the corporate world, such as Google Docs, Dropbox, Zoom, Slack, etc.

3.1. Sample

In total, 308 respondents who participated in the 2021 project completed the questionnaires. The average age of the respondents was 23.3 years, ranging from 18 to over 50, and 39% were male. The majority of the respondents had at least some work experience (average of 3.2 years), and many (30.1%) were employed at the time of the project. Some even ran their own business or held managerial positions (5.1%). The global virtual teams participating in the project submitted weekly deliverables, and all project respondents completed weekly progress surveys. The average response rate was 97.2%, resulting in a sample size of 308 useable, fully completed questionnaires, which provided data on the respondent’s demographics, cultural background, values, and attitudes.

To ensure that respondents understood the field of industrial and social robotics, we required X-Culture respondents to watch three videos (2–3 min each) of social robots in industrial, personal, and social-collaborative situations before being exposed to the final questionnaire. Research in social sciences and interpersonal communication has revealed that messages/communications can be made more persuasive and compliant by cuing humans’ involvement with objects and behaviors (Clark, 1998; Cleveland, Kalamas, & Laroche, 2005). Thus, to help the respondents understand the field of social-collaborative robotics, we used video messages/advertisements as cues to ensure they grasped social behaviors in varying HRI situations. After multiple exposures to industrial and social robots through videos, respondents received an electronic web-based questionnaire with questions focusing on two robots: (1) KUKA Industrial Robot,² which manufacturing companies use for automation and digitization, turnkey production facilities, and smart software solutions, and (2) PARO Seal Therapeutic Robot,³ a personal

² <https://www.kuka.com/en-us/about-kuka/>.

³ <http://www.parorobots.com/>.

Table 3

Overview of indicators and measures of reliability and validity.

Constructs and indicators		Sources	Outer loadings	
			Point estimation	t-value
Anthropomorphism ($\alpha = .902$, AVE = .719, CR = .927)		Godspeed Questionnaires (Bartneck et al., 2009)		
Anthro1	Please rate your impression about the PARO/KUKA robot (5-point scale where 1 = extreme left choice and 5 = extreme right choice): Fake–Natural		.852	46.696
Anthro2	Machinelike–Humanlike		.867	44.867
Anthro3	Unconscious–Conscious		.821	33.456
Anthro4	Artificial–Lifelike		.891	53.491
Anthro5	Moving rigidly–Moving elegantly		.804	30.606
Intentionality ($\alpha = .723$, AVE = .676, CR = .846)		Marchesi et al.'s (2019) InStance Test (IST); Schellen and Wykowska (2019)		
Int1	Please rate your impression about the PARO/KUKA robot (5-point scale where 1 = mechanistic/less intentional and 5 = mentalistic/intentional): The robot makes soothing sounds/noises versus the robot is enjoying attention. (Mechanistic–Mentalistic)		.793	11.329
Int2	The robot looks at me when I talk to the robot versus the robot reduces my stress. (Mechanistic–Mentalistic)		.851	14.714
Sociality ($\alpha = .851$, AVE = .626, CR = .893)		Epley et al. (2007); Eyssel et al. (2011) An extension of Godspeed Questionnaires		
Soc1	Please rate your impression about the PARO/KUKA robot (5-point scale where 1 = extreme left choice and 5 = extreme right choice): Non-social–Social		.763	28.614
Soc2	Non-Trustworthy–Trustworthy		.817	29.424
Soc3	Non-Communicative–Communicative		.757	23.837
Soc4	Non-Interactive–Interactive		.828	27.527
Soc5	Non-engaged/Non-reciprocal–Engaged/Reciprocal		.790	23.979
Perceived Inferiority ($\alpha = .843$, AVE = .617, CR = .887)		C-XENSCALE: Balabanis and Diamantopoulos (2016)		
PerInf1	I prefer 'robots' over humans because robots represent 'foreignness' (with different composition from humans) as compared to humans.		.853	13.337
PerInf2	Robots are better in quality than humans.		.890	14.015
PerInf3	I trust robots over humans performing jobs and services.		.789	12.082
PerInf4	Robots outperform humans in major activities/tasks.		.814	11.438
PerInf5	Even though humans deliver good quality jobs/services, yet robots are far better than humans.		.529	3.925
Social Aggrandizement ($\alpha = .896$, AVE = .669, CR = .909)		C-XENSCALE: Balabanis and Diamantopoulos (2016)		
SocAgg1	Using (foreign) robots enhances my self-esteem.		.684	2.218
SocAgg2	People who don't use/buy robots are less regarded by others.		.844	3.266
SocAgg3	I prefer (foreign) robots over humans as most of my friends prefer robots as well.		.804	3.148
SocAgg4	Buying/using robots makes me trendier.		.847	3.450
SocAgg5	I use/purchase robots to differentiate myself from others.		.894	3.659
Robot Likeability ($\alpha = .912$, AVE = .741, CR = .934)		Godspeed Questionnaires (Bartneck et al., 2009)		
Like1	Please rate your impression about the PARO/KUKA robot (5-point scale where 1 = extreme left choice and 5 = extreme right choice): Dislike–Like		.781	31.112
Like2	Unfriendly–Friendly		.860	45.415
Like3	Unkind–Kind		.888	51.620
Like4	Unpleasant–Pleasant		.900	64.357
Like5	Awful–Nice		.869	31.997
HRI Implementation ($\alpha = .869$, AVE = .792, CR = .920)		Godspeed Questionnaires (Bartneck et al., 2009)		
HRI1	Please rate your impression about the PARO/KUKA robot (5-point scale where 1 = extreme left choice and 5 = extreme right choice): Inert–Interactive		.901	61.661
HRI2	Stagnant–Lively		.907	70.653
HRI3	Failure–Success		.862	47.483

assistant social robot intended to help humans reduce anxiety, depression, and loneliness while also stimulating, collaborating and engaging with people who are living with dementia (Pu, Moyle, & Jones, 2020). We avoided using a humanoid robot (e.g., NAO Robot) due to a potential likeability bias that could be generated vis-à-vis the industrial KUKA robot and the non-humanoid PARO robot.

Admittedly, the present sample comprises students, and we acknowledge certain concerns about the generalizability of the findings because students differ from the general population in their demographic characteristics, particularly age. However, the threat to the validity and generalizability of the findings is likely minimal. The fact that students are typically younger is of little concern if the maturation effect is not expected to influence the effects studied significantly. The respondents, the project settings, and the inter-member differences were real, and the work design was closely reminiscent of the real business world. Within a year or two, all the project respondents would be part of the labor force, and many already were; we have no reason to believe that their attitudes toward robots would drastically change at that time. Therefore, we consider the threat that the present study's findings would not generalize to the real-world consumer population minimal. However, we acknowledge that we can generalize our findings only to the younger to working-age populations. Our results may not apply to seniors, whose attitudes to robots might differ, and our sample does not fully capture such differences.

3.2. Measures

3.2.1. Survey instrument

To design and validate an appropriate survey instrument, we extensively reviewed the literature to identify scales used in past research. We adopted or adapted established scales from past literature to measure anthropomorphism, intentionality, sociality, xenocentrism, robot likeability, and successful HRI implementation as a part of the AXIS robotics-likeability-HRI success conceptual framework. Table 3 lists all constructs and scales used in this research.

Godspeed questionnaires using 5-point semantic differential scales measured robotic anthropomorphism, robot likeability, and HRI implementation (Bartneck et al., 2009). We measured sociality using Godspeed scales describing the need and desire to establish social connections with others through robots (Epley et al., 2007; Eyssel, Kuchenbrandt, & Bobinger, 2011). We used Marchesi et al.'s (2019) InStance Test (IST), a novel tool that assesses the adoption of the intentional stance, wherein two sentences are displayed as possible descriptions, and respondents choose the sentence that best fits the scenario description. For our purposes, one sentence refers to a mental state (i.e., intentional stance), and the other refers to a mechanistic explanation (i.e., design stance) of the HRI scenario. We used sentences mentioning mental states to measure intentionality. We used the consumer xenocentrism scale (C-XENSCALE) to measure xenocentrism (Balabanis & Diamantopoulos, 2016) in robots, with robots treated as foreign (intelligent) beings. All seven constructs (including xenocentrism represented as two constructs: $\hat{X}$ eno A (perceived inferiority) and $\hat{X}$ eno B (social aggrandizement) in the conceptual model constitute latent variables requiring indirect measurement (Bagozzi & Phillips, 1982; Churchill, 1979).

Because the constructs in our research reflect (i.e., cause) their indicators, we specified them as reflective (Diamantopoulos, Riefler, & Roth, 2008; Diamantopoulos & Winklhofer, 2001). We selected indicators of all constructs from existing literature and according to academic and practitioner evidence. We conducted Harman's single-factor test (Podsakoff, MacKenzie, Lee, & Podsakoff, 2003; Podsakoff & Organ, 1986) to allay concerns about common method variance. In the exploratory factor analysis, the first factor explained 15% of the variance, the last factor explained 6% of the variance, and no single factor accounted for the majority (i.e., 50%) of the variance; thus, we consider

it unlikely that common method bias is of serious concern in this study (Podsakoff & Organ, 1986).

3.2.2. Content validity

Content validity aims to analyze whether the scales in the survey questionnaire fully represent the domain being researched (Bollen, 1989). We pretested the scales with experienced managers and researchers to ensure content validity. Two industry professionals directly involved in robotics research and implementing robotics curricula in schools and universities in the United States reviewed our survey questionnaire. They pointed out ambiguities and provided suggestions to improve the survey instrument. In addition, four scholarly experts (outside the authors) also reviewed the questionnaire for clarity, structure, and representativeness.

Based on the feedback of industry professionals and academic experts, we carefully considered the order of presenting questions in our questionnaire to ensure both face validity and content validity. The sequence of questions was strategically designed to flow logically and intuitively, enhancing the face validity by making the questionnaire appear sensible and relevant to participants. This ordering was crucial to maintain participant engagement and reduce response fatigue, which can significantly impact the quality of data collected. We arranged the questions to gradually transition from general to more specific items, thereby providing a coherent narrative that aligns with our research objectives.

3.3. Data collection

We used a web-based survey to collect the data to test the proposed hypotheses. Web surveys are becoming increasingly popular across various research streams (Seepana, Huq, & Paulraj, 2021; Statsenko & Corral de Zubielqui, 2020) because of such advantages as quicker and higher response rates. In addition, web-based surveys allow for the collection of valuable information about the respondents' survey completion process (Griffis, Goldsby, & Cooper, 2003).

3.4. Analytical procedure

We validated our measures and tested our hypotheses using partial least squares (PLS), specifically SmartPLS software version 3.2.8 (Ringle, Wende, & Becker, 2015). Partial least squares structural equation modeling (PLS-SEM) methodology is preferred for our research for several reasons. First, we focus on theory development and prediction (Matthews, Hair, & Matthews, 2018). Second, PLS-SEM is the preferred statistical method to analyze a model with higher-order constructs (Hair, Hult, Ringle, & Sarstedt, 2022; Manley, Williams, & Hair, 2022), such as $\hat{X}$ enocentrism. Third, PLS-SEM employs a fixed point or component-based least squares estimation procedure to obtain parameter estimates. Fourth, PLS uses a series of interdependent ordinary least squares regressions to minimize residual variances. Fifth, it places minimal demands on data in terms of measurement scales, sample size, and distributional assumptions (Chin, 1998; Fornell & Bookstein, 1982; Wold, 1982), which makes it a preferable approach compared to covariance-based maximum likelihood methods (e.g., LISREL) when examining data for which the sample size is relatively small (Bagozzi, Yi, & Singh, 1991). Finally, PLS is a more conservative modeling approach that tends to underestimate path coefficients (Dijkstra, 1983), reducing the likelihood of Type I errors in hypothesis testing (Bagozzi et al., 1991).

To test our model's indicator reliability, we employed a bootstrapping procedure with 2000 randomized samples taken from the original sample (Henseler, Ringle, & Sinkovics, 2009). The results of the analysis are available in Table 3. All estimates of outer loadings exceed the minimum recommended value of 0.7 and also exhibit sufficiently high *t*-values. We also assessed the convergent validity of all constructs. All loadings were greater than 0.7, which implies that all indicators

Table 4
Correlations between constructs.

Construct	Anthro	HRI Implementation	Intentionality	Likeability	Sociality	Perceived Inferiority	Social Aggrandize-ment
Anthropomorphism	.848						
HRI Implementation	.618	.890					
Intentionality	.524	.561	.594				
Likeability	.581	.753	.534	.861			
Sociality	.613	.718	.427	.604	.791		
Perceived Inferiority	.167	.172	.229	.180	.211	.786	
Social Aggrandizement	.066	.026	.053	-.092	.081	.372	.818

Notes: The square root of AVE on the diagonal is in boldface.

Table 5
Path coefficients and R^2 of the structural model.

Relationship	Path coefficients		Hypotheses	
	Point estimate	t-Value		
<i>Likeability (R² = .510)</i>				
Anthro → Likeability	.260	4.175	H _{1A}	Supported
Intention → Likeability	.206	3.641	H _{2A}	Supported
Sociality → Likeability	.356	6.227	H _{3A}	Supported
<i>HRI Implement (R² = .691)</i>				
Anthro → HRI Implement	.096	2.037	H _{1B}	Supported
Intention → HRI Implement	.104	2.469	H _{2B}	Supported
Sociality → HRI Implement	.350	6.697	H _{3B}	Supported
Likeability → HRI Implement	.438	7.662	H ₄	Supported
<i><u>Likeability (Moderating effects)</u></i>				
Anthro × Perceived Inferiority	−.208	2.428	H _{5A}	Supported (-ve relationship)
Intentionality × Perceived Inferiority	.158	2.631	H _{5B}	Supported
Sociality × Perceived Inferiority	−.006	0.092	H _{5C}	Not supported
Anthro × Social Aggrandizement	.031	0.428	H _{6A}	Not supported
Intentionality × Social Aggrandizement	−.106	1.472	H _{6B}	Not supported
Sociality × Social Aggrandizement	.048	0.719	H _{6C}	Not supported

share more variance with their constructs than with error variances (Chin, 1998). We assessed Cronbach's alpha (α) and composite reliability (CR) for construct reliability. Table 3 shows that the α values for all constructs are above the cutoff value of 0.7 (Cronbach, 1951; Litwin, 1995). The same applies to all CR values that exceed the recommended cutoff value of 0.6 (Bagozzi & Yi, 1988; Henseler et al., 2009). The average variance extracted (AVE) values (Table 3) are above the threshold value of 0.5 (Fornell & Larcker, 1981; Henseler et al., 2009), thus establishing convergent validity. We also assessed discriminant validity; Table 4 shows the correlations between the latent variables and the square roots of the AVE on the diagonal, indicating that the AVE's square roots are greater than the correlations among the latent variable scores in all cases. Thus, we can conclude that no construct shares more variance with another, thereby establishing discriminant validity (Fornell & Larcker, 1981; Henseler et al., 2009). Furthermore, to test the prediction relevance of the model, we applied a blindfolding procedure with an omission distance of 5 (Henseler et al., 2009). All resulting Q^2 values were positive, thus establishing sufficient predictive power of the structural model (Geisser, 1975; Stone, 1974).

3.5. Results of analysis

Table 5 presents the results from the evaluation of the structural model (Fig. 1). They show that the R^2 values of the endogenous latent variables (robot likeability [$R^2 = 0.510$] and successful HRI

implementation [$R^2 = 0.691$]) are substantial according to Chin's (1998) specifications. Overall, the results of our analysis indicate a good model fit with sufficient predictive power. We tested the significance of the relationships among the latent variables using the associated t -statistics obtained from PLS bootstrapping. As the results shown in Table 5 indicate, 9 of the 13 hypotheses are confirmed, all of which are significant at $p < 0.01$ level.

4. Discussion and implications

In an HHI situation, social cues (e.g., language and emotional displays) are critical for coordination and communication. Similarly, in an HRI context, social/emotional cues that are integrated into robot design can improve long-term collaboration between robots and humans (Fischer, 2019). Social-collaborative robots use collaboration as social actors through implicit and explicit mechanisms of communicating information efficiently (Admoni, Dragan, Srinivasa, & Scassellati, 2014; Fischer, Jensen, Suvei, & Bodenhagen, 2016). This research uses the CASA paradigm, which equates HHI with HRI by mimicking interpersonal interactions and relationships (Fox & Gambino, 2021), as the theoretical framework for describing $\hat{A}\hat{X}\hat{I}\hat{S}$ robotics.

As CASA posits that robots with anthropomorphic and emotional cues can support higher forms of social and collaborative interactions with humans (Eyssel & Hegel, 2012; Fox & Gambino, 2021), we conceptualize and define $\hat{A}\hat{X}\hat{I}\hat{S}$ robots as social-collaborative robots that use principles of socio-cognitive intelligence for human-robot collaboration and exhibit anthropomorphic, $\hat{x}$ enocentric, intentional, and social behaviors in the sociocultural HRI context in which they operate (Arora & Arora, 2020). A real-world example is Vanderbilt University's ASK NAO program developed for the NAO robot (a French humanoid robot by Aldebaran Robotics) to interact with students diagnosed with autism spectrum disorder (ASD) and other learning disorders. This NAO robot (with ASK NAO software) can be examined as an $\hat{A}\hat{X}\hat{I}\hat{S}$ robot with socio-cognitive-collaborative intelligence operating in a sociocultural HRI situation aimed at ASD students globally. This example provides some face validity to our conceptual framework and empirical results.

The robotics industry is currently experiencing a major transformation (Ballestar, García-Lazaro, Sainz, & Sanz, 2022). Industrial robots are increasingly used in businesses and organizations as precision instruments, mimicking the capabilities of skilled human labor, and repeating a handful of tasks thousands of times over (without showing signs of fatigue). Recently, interest in non-industrial (social) robots has also emerged: businesses are moving away from purely industrial robots to ones that are more social, autonomous, collaborative, and easily trainable (i.e., less programming intensive) (Sanneman, Fourie, & Shah, 2021). Robotics companies and roboticists are studying demand patterns for co-robots or cobots (safe, flexible, vision-enabled, and easily trainable robotic assistants) helping and collaborating with humans in their social spheres (Cherubini et al., 2016; Schou et al., 2018; Kopp et al., 2021). This research used industrial KUKA and social PARO robots in social HRI settings. KUKA has developed a robot called the youBot, which can be used for education and research. It includes a mobile, two-fingered (with five degrees of freedom) plug-and-play robotic arm.

The seal robot PARO is an advanced Japanese interactive robot administered to patients in environments such as hospitals and extended care facilities.

We found that all pivotal concepts of AXIS robotics (anthropomorphism, intentionality, and sociality) lead to positive robot likeability and a successful HRI implementation for both KUKA and PARO robots. These findings corroborate our research perspective illustrated through the CASA framework that humans anthropomorphize robots by simulating and exhibiting social-emotional-collaborative-cognitive associations (Marchesi et al., 2021) before developing affinity/liking toward them. According to CASA framework and uncanny valley effect theories, intentionality (i.e., the tendency to exhibit mentalizing and social cognition behavior) can lead to positive human behavior and likeability of robots. Using the SEEK model theory, we propose that sociality (i.e., the tendency to act as a social/emotional agent) in robots can lead to developing likeability of robots and facilitating HRI implementation (Damiano & Dumouchel, 2018). According to our findings, implementing a social HRI in practice (and ensuring its success) depends on the pivotal AXIS robotics concepts. Anthropomorphic, intentional, and social robots can help generate a successful HRI implementation.

Our results further reveal that using an intentional mentalistic description for robots (e.g., attributing human mental capabilities to robots for activating neural representations), and highlighting a social (or emotional/attachment) cue will result in an efficient and successful HRI implementation that benefits humans involved in HRI situations by engendering feelings of joy, accomplishment, excitement, and enjoyment, and even improved health outcomes (resulting from health care robotic implementations). Relatedly, prior research shows that robot likeability can happen within seconds in HRI situations, and the impression of likeability significantly influences HRI implementation (Bartneck et al., 2009; Kaplan et al., 2019). We found that robot likeability is positively associated with successful HRI implementation: the more anthropomorphic (humanlike), intentional, and social a robot is, the more likable it is, and thus, there is a stronger possibility of a successful HRI implementation. However, the uncanny valley effect should be considered when designing and developing robots, paying careful attention to the point of inflection at which a dip in likeability occurs when the robots appear or behave in too human a manner (Arora et al., 2021; Mori et al., 2012). Robotic designers and roboticists employ the uncanny valley effect in many real-world instances. For example, Soul Machines, a New Zealand company, created Ava (a digital-human avatar and a virtual assistant) for Autodesk Inc. (an American multinational software corporation that makes software products and services for the architecture, engineering, construction, manufacturing, media, education, and entertainment industries). Ava appears remarkably lifelike; she is designed to analyze human facial expressions and voices of joy, sadness, anger, and frustration. She can generate an emotional reaction on her face in return for the received human expression. Because of these AI capabilities, Soul Machines intentionally created Ava with purple eyes so that she is not perceived as too human (or too real), thus avoiding the uncanny valley effect.

We encountered plausible results while examining the moderating effects of $\hat{\text{x}}\text{enocentrism}$ on the linkages between robot likeability and anthropomorphism, intentionality, and sociality, as well as between successful HRI implementation and these pivotal concepts. The concept of $\hat{\text{x}}\text{enocentrism}$ can be explained through the system justification theory, which highlights social aggrandizement (i.e., the phenomena of outgroup favoritism for enhancing perceived social status) and perceived inferiority (i.e., the phenomena of ingroup derogation, particularly among members of low-status groups) (Balabanis & Diamantopoulos, 2016). Diverse domains such as consumer behavior, human resources management, marketing, organizational behavior, corporate social responsibility, and business ethics research have globally implemented system justification theory (e.g., DiTomaso, 2015; Fujimoto, Härtel, & Azmat, 2013; Li & Agrawal, 2014; Shepherd et al.,

2015); herein, we expand the use of this theory for social robotics and HRI research. We employed the C-XENSCALE to measure $\hat{\text{x}}\text{enocentrism}$ and how it moderates the relationships between pivotal concepts, robot likeability, and successful HRI implementation. We find that humans perceive robots to be more intelligent and superior (Rampersad, 2020), thus confirming that robotic $\hat{\text{x}}\text{enocentrism}$ exists. Robots are used in manufacturing, science, surgery, and performing/delivering services. Robots work efficiently and effectively; for example, they fly planes more safely than humans and perform household cleaning tasks better than humans, and driverless cars are better than human drivers who may be distracted by their cellphones or driving under the influence.

When delving deeper into our results, we found that social aggrandizement (outgroup favoritism) did not act as a significant boundary condition for the relationships between pivotal concepts and robot likeability but that perceived inferiority (ingroup derogation) was a significant negative moderator between anthropomorphism and robot likeability. In addition, perceived inferiority was a positive (significant) moderator in the relationship between intentionality and robot likeability. Sociality was not affected due to the moderating influence of $\hat{\text{x}}\text{enocentrism}$. These findings demonstrate that humans perceive robots as superior because they see themselves as inferiors or of low status and not because they think that robots come from a foreign (better) world and have a different composition than humans.

An important finding is that robotic $\hat{\text{x}}\text{enocentrism}$ moderates the linkage between anthropomorphism and robot likeability. Due to the uncanny valley effect, there is a significant negative moderating relationship of $\hat{\text{x}}\text{enocentrism}$ on the linkage between anthropomorphism and robot likeability. This finding confirms robot designers and roboticists' insistence on having imperfections in robots to prevent the occurrence of high levels of anthropomorphic feelings in consumers during HRI. In contrast, we found a significant, positive moderating relationship of $\hat{\text{x}}\text{enocentrism}$ on the linkage between intentionality and robot likeability, meaning that the more intentional the robot is, the more likable it is.

Computers are social actors (CASA) paradigm was utilized as an overarching theory (along with the uncanny valley effect and the sociality, effectance, and elicited agent knowledge (SEEK) model theories) proposing that social robots display high social cognition and decision-making skills. These social robots can behave intelligently in complex HRI situations (e.g., military operations: biological, radiological, chemical, and nuclear detection; battle-space awareness and environmental sensing; precision targeting and precision strike; counter-improvised explosive device capabilities). During these situations, humans tend to use mentalistic descriptions more than mechanistic descriptions for social robots. Thus, perceived robotic intentionality is higher, resulting in positive robot likeability when moderated by robotic $\hat{\text{x}}\text{enocentrism}$. The HRI research on human facial responses and animatronics is crucial considering the importance of robotic intentionality and human mental states. Currently, robots can portray and mimic human emotions, but in the future, the ability to respond to human emotions appropriately will take innovations in AXIS robotics to the next level.

5. Limitations and future research directions

The study has some limitations, which point to further research directions. First, the research involved students and working professionals who examined HRI situations through video messages rather than in-person HRI encounters. However, we would not have been able to garner 308 survey responses from a global audience using an in-person study. Therefore, future research should focus on in-person HRI encounters. Second, our AXIS robotics research has added a new global dimension/concept of robotic $\hat{\text{x}}\text{enocentrism}$ (Balabanis & Diamantopoulos, 2016) to social robotics. While this concept enriches the field of social-collaborative robotics, it is still in its infancy and needs

Table 6
Future research directions for A[^]XIS robotics research.

Concept	Theories used	Future research directions and questions
A[^]XIS Robotics	CASA paradigm (overarching framework for A [^] XIS Robotics)	Further exploration of A [^] XIS concepts in relation to each other along with investigating the moderating effects of xenocentrism.
Robotic Anthropomorphism	Uncanny valley effect	a Why does robotic anthropomorphism negatively influence robot likeability in HRI situations when it is moderated by perceived inferiority? b Why does robotic anthropomorphism have no effect on robot likeability when it is moderated by social aggrandizement for either industrial or social robots? c Are there any other significant effects of anthropomorphism on other key A [^] XIS concepts?
Robotic Xenocentrism	System justification theory;	a Why does robotic xenocentrism moderate anthropomorphism negatively with robot likeability? b Why does robotic xenocentrism moderate intentionality positively with robot likeability? c Why does robotic xenocentrism have no moderation effect on sociality with robot likeability for either industrial or social robots?
Robotic Intentionality	Uncanny valley effect theory	a Why does robotic intentionality positively impact robot likeability in HRI situations, when it is moderated by perceived inferiority? b Why does robotic intentionality have no effect (whatsoever) on robot likeability when it is moderated by social aggrandizement for either industrial or social robots? c Are there any other significant effects of intentionality on other key A [^] XIS concepts?
Robotic Sociality	SEEK model theory	a Why does robotic sociality have no effect on robot likeability in HRI situations, when it is moderated by perceived inferiority and/or social aggrandizement for either industrial or social robots? b Are there any other significant effects of sociality on other key A [^] XIS concepts?

further investigation. Further empirical research should focus on implementing robotic xenocentrism (and examining its moderating effects) in social robotics and HRI contexts and situations. Table 6 presents some suggestions and questions that can guide future scholars in developing new research directions.

Our research on A[^]XIS robotics is the first of its kind to use socio-cognitive, management, and international business concepts in the ever-growing field of social robotics and HRI. We hope our study will help designers of the next generation of social-collaborative robots, thus further bridging the gap between humans and robots.

CRedit authorship contribution statement

Anshu Saxena Arora: Conceptualization, Formal analysis, Funding acquisition, Resources, Writing – original draft, Writing – review & editing. **Amit Arora:** Data curation, Investigation, Methodology, Software, Validation, Writing – review & editing. **K. Sivakumar:** Conceptualization, Project administration, Visualization, Writing – original draft, Writing – review & editing. **Vasyl Taras:** Data curation, Methodology, Resources, Validation, Visualization.

Declaration of competing interest

None.

Acknowledgement

The research is supported by the National Science Foundation (NSF) Research Initiation Award (RIA) #2100934 - Social Motivation Approach for Rehabilitation Through Educational Robotics (SMARTER) research and NSF Catalyst Project Award #2106411 - STEM-Business Focused Sustainable Financial Technology - Clean Technology (Sustainable FinTech-CleanTech) at the University of the District of Columbia.

References

Admoni, H., Dragan, A., Srinivasa, S. S., & Scassellati, B. (2014). Deliberate delays during robot-to-human handovers improve compliance with gaze communication. In G. Sagerer, & M. Imai (Eds.), *Hri '14: Proceedings of the 2014 ACM/IEEE international conference on human-robot interaction* (pp. 49–56). Association for Computing Machinery.

Aly, A., & Tapus, A. (2016). Towards an intelligent system for generating an adapted verbal and nonverbal combined behavior in human-robot interaction. *Autonomous Robots*, 40(2), 193–209.

Arora, A. S., & Arora, A. (2020). The race between cognitive and artificial intelligence: Examining socio-ethical collaborative robots through anthropomorphism and xenocentrism in human-robot interaction. *International Journal of Intelligent Information Technologies*, 16(1), 1–16.

Arora, A. S., Arora, A., Jentjens, S., McIntyre, J. R., & Sepehri, M. (2022). Managing social robotics and socio-cultural business norms: Parallel worlds of emerging AI and human virtues. *International marketing and management research series*. Palgrave Macmillan-Springer Nature.

Arora, A. S., Fleming, M., Arora, A., Taras, V., & Xu, J. (2021). Finding “H” in HRI: Examining human personality traits, robotic anthropomorphism, and robot likeability in human-robot interaction. *International Journal of Intelligent Information Technologies*, 17(1), 19–38.

Arora, A. S., Parnell, C., & Arora, A. (2022). Applying service-dominant logic and conversation management principles to social robotics for autism spectrum disorder. *International Journal of Intelligent Information Technologies*, 18(1), 1–19.

Bagozzi, R. P., & Phillips, L. W. (1982). Representing and testing organizational theories: A holistic construal. *Administrative Science Quarterly*, 27(3), 459–490.

Bagozzi, R. P., & Yi, Y. (1988). On the evaluation of structural equation models. *Journal of the Academy of Marketing Science*, 16(1), 74–94.

Bagozzi, R. P., Yi, Y., & Singh, S. (1991). On the use of structural equation models in experimental designs: Two extensions. *International Journal of Research in Marketing*, 8(2), 125–140.

Balabanis, G., & Diamantopoulos, A. (2016). Consumer xenocentrism as determinant of foreign product preference: A system justification perspective. *Journal of International Marketing*, 24(3), 58–77.

Ballestar, M. T., García-Lazaro, A., Sainz, J., & Sanz, I. (2022). Why is your company not robotic? The technology and human capital needed by firms to become robotic. *Journal of Business Research*, 142, 328–343.

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.

Bollen, K. A. (1989). *Structural equations with latent variables*. Wiley.

Bossi, F., Willemse, C., Cavazza, J., Marchesi, S., Murino, V., & Wykowska, A. (2020). The human brain reveals resting state activity patterns that are predictive of biases in attitudes toward robots. *Science Robotics*, 5(46), Article eabb6652.

Brynjolfsson, E. (2022). The Turing trap: The promise and peril of human-like artificial intelligence. *Daedalus*, 151(2), 272–287.

Candrian, C., & Scherer, A. (2022). Rise of the machines: Delegating decisions to autonomous AI. *Computers in Human Behavior*, 134, Article 107308.

Chang, W., & Kim, K. K. (2022). Appropriate service robots in exchange and communal relationships. *Journal of Business Research*, 141, 462–474.

Chartrand, T., & Bargh, J. (1999). The chameleon effect: The perception-behavior link and social interaction. *Personality and Social Psychology*, 76(6), 893–910.

Chin, W. W. (1998). The partial least squares approach to structural equation modeling. In G. A. Marcoulides (Ed.), *Modern methods for business research* (pp. 295–336). London: Erlbaum.

Christoforakos, L., & Diefenbach, S. (2022). Technology as a social companion? An exploration of individual and product-related factors of anthropomorphism. *Social Science Computer Review*. <https://doi.org/10.1177/08944393211065867>

- Chung, H., Kang, H., & Jun, S. (2023). Verbal anthropomorphism design of social robots: Investigating users' privacy perception. *Computers in Human Behavior*, Article 107640.
- Churchill, G. A. (1979). A paradigm for developing better measures of marketing constructs. *Journal of Marketing Research*, 16, 64–73.
- Ciarlo, F., De Tommaso, D., & Wykowska, A. (2022). Joint action with artificial agents: Human-likeness in behaviour and morphology affects sensorimotor signaling and social inclusion. *Computers in Human Behavior*, 132, Article 107237.
- Clark, H. H. (1998). Communal lexicons. In K. Malmkjær, & J. Williams (Eds.), *Context in language learning and language understanding* (pp. 63–87). Cambridge University Press.
- Cleveland, M., Kalamas, M., & Laroche, M. (2005). Shades of green: Linking environmental locus of control and pro-environmental behaviors. *Journal of Consumer Marketing*, 22(4), 198–212.
- Cronbach, L. (1951). Coefficient alpha and the internal structure of tests. *Psychometrika*, 16, 297–334.
- Damiano, L., & Dumouchel, P. (2018). Anthropomorphism in human-robot co-evolution. *Frontiers in Psychology*, 9, 468.
- Dang, J., & Liu, L. (2023). Do lonely people seek robot companionship? A comparative examination of the loneliness-robot anthropomorphism link in the United States and China. *Computers in Human Behavior*, 141, Article 107637.
- Dautenhahn, K., Werry, I., Salter, T., & Boekhorst, R. T. (2003). Towards adaptive autonomous robots in autism therapy: Varieties of interactions. In *Proceedings 2003 IEEE international symposium on computational intelligence in robotics and automation. Computational intelligence in robotics and automation for the new millennium (cat. No. 03EX694)* (pp. 577–582). IEEE, 2.
- Dennett, D. C. (1971). Intentional systems. *Journal of Philosophy*, 68(4), 87–106. <https://doi.org/10.2307/2025382>
- Dennett, D. C. (1988). Précis of the intentional stance. *Behavioural and Brain Sciences*, 11(3), 495–505. <https://doi.org/10.1017/S0140525X00058611>
- Dennett, D. C. (1997). *Consciousness in human and robot minds*. Oxford University Press.
- DeWall, C. N., & Baumeister, R. F. (2006). Alone but feeling no pain: Effects of social exclusion on physical pain tolerance and pain threshold, affective forecasting, and interpersonal empathy. *Journal of Personality and Social Psychology*, 91(1), 1–16.
- Di Dio, C., Manzi, F., Peretti, G., Cangelosi, A., Harris, P. L., Massaro, D., et al. (2020). Shall I trust you? From child-robot interaction to trusting relationships. *Frontiers in Psychology*, 11, 469.
- Diamantopoulos, A., Riefler, P., & Roth, K. (2008). Advancing formative measurement models. *Journal of Business Research*, 61, 1203–1218.
- Diamantopoulos, A., & Winklhofer, H. M. (2001). Index construction with formative indicators: An alternative to scale development. *Journal of Marketing Research*, 37, 269–277.
- Dijkstra, T. (1983). Some comments on maximum likelihood and partial least squares methods. *Journal of Econometrics*, 22(1/2), 67–90.
- DiTomaso, N. (2015). Racism and discrimination versus advantage and favoritism: Bias for versus bias against. *Research in Organizational Behavior*, 35, 57–77.
- Epley, N., Akalis, S., Waytz, A., & Cacioppo, J. T. (2008). Creating social connection through inferential reproduction: Loneliness and perceived agency in gadgets, gods, and greyhounds. *Psychological Science*, 19(2), 114–120.
- Epley, N., Waytz, A., & Cacioppo, J. T. (2007). On seeing human: A three-factor theory of anthropomorphism. *Psychological Review*, 114(4), 864–886.
- Eshleman, C., Cashion, B., & Basirico, L. (1993). *Sociology: An introduction* (3rd ed.). BVT Publishing.
- Esterwood, C., & Robert, L. P., Jr. (2023). Three Strikes and you are out!: The impacts of multiple human-robot trust violations and repairs on robot trustworthiness. *Computers in Human Behavior*, 142, Article 107658.
- Eyssel, F., & Hegel, F. (2012). (S)he's got the look: Gender stereotyping of robots. *Journal of Applied Social Psychology*, 42, 2213–2230.
- Eyssel, F., Kuchenbrandt, D., & Bobinger, S. (2011). Effects of anticipated human-robot interaction and predictability of robot behavior on perceptions of anthropomorphism. In *Proceedings of the 6th international conference on human-robot interaction* (pp. 61–68). IEEE.
- Fernandes, J. V. (2022). Robot citizenship and gender (in) equality: The case of Sophia the robot in Saudi Arabia. *Janus: E-journal of international relations. Thematic dossier the Middle East. Local dynamics, regional actors, global challenges*.
- Ferrara, C., & Hill, S. D. (1980). The responsiveness of autistic children to the predictability of social and nonsocial toys. *Journal of Autism and Developmental Disorders*, 10(1), 51–57.
- Fischer, K. (2019). Why collaborative robots must be social (and even emotional) actors. *Techné: Research in Philosophy and Technology*, 23(3), 270–289.
- Fischer, K., Jensen, L. C., Suvel, S. D., & Bodenhagen, L. (2016). Between legitimacy and contact: The role of gaze in robot approach. In *2016 25th IEEE international symposium on robot and human interactive communication (RO-MAN)* (pp. 646–651). IEEE.
- Fornell, C., & Bookstein, F. L. (1982). Two structural equation models: LISREL and PLS applied to consumer exit-voice theory. *Journal of Marketing Research*, 19, 440–452.
- Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18, 39–50.
- Fox, J., & Gambino, A. (2021). Relationship development with humanoid social robots: Applying interpersonal theories to human-robot interaction. *Cyberpsychology, Behavior, and Social Networking*, 24(5), 294–299.
- Fujimoto, Y., Härtel, C. E., & Azmat, F. (2013). Towards a diversity justice management model: Integrating organizational justice and diversity management. *Social Responsibility Journal*, 9(1), 148–166.
- Gambino, A., Fox, J., & Ratan, R. A. (2020). Building a stronger CASA: Extending the computers are social actors paradigm. *Human-Machine Communication*, 1, 71–85.
- Gaudiello, I., Zibetti, E., Lefort, S., Chetouani, M., & Ivaldi, S. (2016). Trust as indicator of robot functional and social acceptance. An experimental study on user conformation to iCub answers. *Computers in Human Behavior*, 61, 633–655.
- Geisser, S. (1975). A predictive approach to the random effect model. *Biometrika*, 61, 101–107.
- Griffis, S. E., Goldsby, T. J., & Cooper, M. (2003). Web-based and mail surveys: A comparison of response, data, and cost. *Journal of Business Logistics*, 24(2), 237–258.
- Grundke, A., Stein, J. P., & Appel, M. (2023). Improving evaluations of advanced robots by depicting them in harmful situations. *Computers in Human Behavior*, 140, Article 107565.
- Hair, J. F., Hult, G. T. M., Ringle, C., & Sarstedt, M. (2022). *A primer on partial least squares structural equation modeling (PLS-SEM)* (3rd ed.). SAGE Publications.
- Hegel, F., Krach, S., Kircher, T., Wrede, B., & Sagerer, G. (2008). Understanding social robots: A user study on anthropomorphism. In *RO-MAN 2008—the 17th IEEE international symposium on robot and human interactive communication* (pp. 574–579). IEEE.
- Henseler, J., Ringle, C. M., & Sinkovics, R. R. (2009). The use of partial least squares path modeling in international marketing. *Advances in International Marketing*, 20, 277–319.
- Hesslow, G. (2002). Conscious thought as simulation of behaviour and perception. *Trends in Cognitive Sciences*, 6(6), 242–247.
- Hesslow, G. (2012). The current status of the simulation theory of cognition. *Brain Research*, 1428, 71–79.
- Hudson, C., & Lewis, L. (2020). Social robots as reinforcement in applied behavior analysis. In *Companion of the 2020 ACM/IEEE international conference on human-robot interaction* (pp. 266–268). IEEE.
- Huijnen, C. A., Lexis, M. A., Jansens, R., & de Witte, L. P. (2019). Roles, strengths, and challenges of using robots in interventions for children with autism spectrum disorder (ASD). *Journal of Autism and Developmental Disorders*, 49(1), 11–21.
- Jost, J. T., & Banaji, M. R. (1994). The role of stereotyping in system-justification and the production of false consciousness. *British Journal of Social Psychology*, 33(1), 1–27.
- Kaplan, A. D., Sanders, T., & Hancock, P. A. (2019). The relationship between extroversion and the tendency to anthropomorphize robots: A bayesian analysis. *Frontiers in Robotics and AI*, 5, 135.
- Kent, D. P., & Burnight, R. G. (1951). Group centrism in complex societies. *American Journal of Sociology*, 57(3), 256–259.
- Kim, J., & Im, I. (2023). Anthropomorphic response: Understanding interactions between humans and artificial intelligence agents. *Computers in Human Behavior*, 139, Article 107512.
- Klüber, K., & Onnasch, L. (2022). Appearance is not everything-Preferred feature combinations for care robots. *Computers in Human Behavior*, 128, Article 107128.
- Kwok, C., Grisham, J. R., & Norberg, M. M. (2018). Object attachment: Humanness increases sentimental and instrumental values. *Journal of Behavioral Addictions*, 7(4), 1132–1142.
- Lajante, M., Tobji, D., & Ho, T. I. (2023). When interacting with a service robot is (not) satisfying: The role of customers' need for social sharing of emotion. *Computers in Human Behavior*, Article 107792.
- Leo-Liu, J. (2023). Loving a "defiant" AI companion? The gender performance and ethics of social exchange robots in simulated intimate interactions. *Computers in Human Behavior*, 141, Article 107620.
- Letheren, K., Kuhn, K. A. L., Lings, I., & Pope, N. K. L. (2016). Individual difference factors related to anthropomorphic tendency. *European Journal of Marketing*, 50(5/6), 973–1002.
- Li, X., & Agrawal, N. (2014). Thriving in a sinking system: When does a threatening world promote meaningful behaviors. In J. Cotte, & S. Wood (Eds.), *Advances in consumer research* (pp. 575–577). Association for Consumer Research.
- Lieto, A., Chella, A., & Frixione, M. (2017). Conceptual spaces for cognitive architectures: A lingua franca for different levels of representation. *Biologically Inspired Cognitive Architectures*, 19, 1–9.
- Li, M., Guo, F., Wang, X., Chen, J., & Ham, J. (2023). Effects of robot gaze and voice human-likeness on users' subjective perception, visual attention, and cerebral activity in voice conversations. *Computers in Human Behavior*, 141, Article 107645.
- Li, Z., Terfurth, L., Woller, J. P., & Wiese, E. (2022). Mind the machines: Applying implicit measures of mind perception to social robotics. In *2022 17th ACM/IEEE international conference on human-robot interaction (HRI)* (236–245). IEEE.
- Litwin, M. S. (1995). *How to measure survey reliability and validity*. SAGE Publications.
- Liu, B., & Sundar, S. S. (2018). Should machines express sympathy and empathy? Experiments with a health advice chatbot. *Cyberpsychology, Behavior, and Social Networking*, 21(10), 625–636.
- Mahdi, H., Saleh, S., Shariff, O., & Dautenhahn, K. (2020). Creating myjay: A new design for robot-assisted play for children with physical special needs. In *International conference on social robotics* (pp. 676–687). Cham: Springer.
- Manley, S. C., Williams, R. I., Jr., & Hair, J. F., Jr. (2022). Enhancing TQM's effect on small business performance: A PLS-SEM exploratory study of TQM applied with a comprehensive strategic approach. *The TQM Journal*. <https://doi.org/10.1108/TQM-10-2021-0299>. Sep. 9.
- Marchesi, S., Ghiglini, D., Ciarlo, F., Perez-Osorio, J., Baykara, E., & Wykowska, A. (2019). Do we adopt the intentional stance toward humanoid robots? *Frontiers in Psychology*, 10, 450.
- Marchesi, S., Spatola, N., & Wykowska, A. (2021). The mediating role of anthropomorphism in adopting the intentional stance towards humanoid robots. *Working paper*. University of Manchester.
- Matthews, L., Hair, J. O. E., & Matthews, R. (2018). PLS-SEM: The holy grail for advanced analysis. *Marketing Management Journal*, 28(1), 1–13.

- Mikkelsen, R. V., & Rehm, M. (2022). *Creating informative and successful human robot interaction in a real-world service environment. 2022 31st IEEE international Conference on Robot and human interactive communication (RO-MAN)*. IEEE.
- Mims, C. (2020). Reporting for coronavirus duty: Robots that go where humans fear to tread. *The Wall Street Journal* (April 4) <https://www.wsj.com/articles/reporting-for-coronavirus-duty-robots-that-go-where-humans-fear-to-tread-11585972801>.
- Mori, M. (1970). The uncanny valley: The original essay by Masahiro Mori. *IEEE Spectrum*, 6.
- Mori, M., MacDorman, K. F., & Kageki, N. (2012). The uncanny valley [from the field]. *IEEE Robotics and Automation Magazine*, 19(2), 98–100.
- Mourey, J. A., Olson, J. G., & Yoon, C. (2017). Products as pals: Engaging with anthropomorphic products mitigates the effects of social exclusion. *Journal of Consumer Research*, 44(2), 414–431.
- Nicolas, S., & Agnieszka, W. (2021). The personality of anthropomorphism: How the need for cognition and the need for closure define attitudes and anthropomorphic attributions toward robots. *Computers in Human Behavior*, 122, Article 106841.
- Niemyska, A., & Drat-Ruszczak, K. (2013). When there is nobody, angels begin to fly: Supernatural imagery elicited by a loss of social connection. *Social Cognition*, 31(1), 57–71.
- Ormel, J., Lindenberg, S., Steverink, N., & Verbrugge, L. M. (1999). Subjective well-being and social production functions. *Social Indicators Research*, 46(1), 61–90.
- Özdem, C., Wiese, E., Wykowska, A., Müller, H., Brass, M., & Van Overwalle, F. (2017). Believing androids—fMRI activation in the right temporo-parietal junction is modulated by ascribing intentions to non-human agents. *Social Neuroscience*, 12(5), 582–593.
- Piçarra, N., & Giger, J. C. (2018). Predicting intention to work with social robots at anticipation stage: Assessing the role of behavioral desire and anticipated emotions. *Computers in Human Behavior*, 86, 129–146.
- Podsakoff, P. M., MacKenzie, S. B., Lee, J.-Y., & Podsakoff, N. P. (2003). Common method biases in behavioral research: A critical review of the literature and recommended remedies. *Journal of Applied Psychology*, 88, 879–893.
- Podsakoff, P. M., & Organ, D. W. (1986). Self-reports in organizational research: Problems and prospects. *Journal of Management*, 12, 531–544.
- Pu, L., Moyle, W., & Jones, C. (2020). How people with dementia perceive a therapeutic robot called PARO in relation to their pain and mood: A qualitative study. *Journal of Clinical Nursing*, 29(3–4), 437–446.
- Rampersad, G. (2020). Robot will take your job: Innovation for an era of artificial intelligence. *Journal of Business Research*, 116, 68–74.
- Reeves, B., & Nass, C. (1996). *The media equation: How people treat computers, television, and new media like real people and places*. Cambridge.
- Ribino, P., Lodato, C., & Infantino, I. (2018). Handling robot sociality: A goal-based normative approach. *AIC*, 2018, 59–73.
- Ringle, C. M., Wende, S., & Becker, J.-M. (2015). *SmartPLS 3*. SmartPLS GmbH. <http://www.smartpls.com>.
- Sanneman, L., Fourie, C., & Shah, J. A. (2021). The state of industrial robotics: Emerging technologies, challenges, and key research directions. *Foundations and Trends in Robotics*, 8(3), 225–306.
- Sauppé, A., & Mutlu, B. (2015). The social impact of a robot co-worker in industrial settings. In *Proceedings of the 33rd annual ACM conference on human factors in computing systems* (pp. 3613–3622). Association for Computing Machinery.
- Schellen, E., & Wykowska, A. (2019). Intentional mindset toward robots—open questions and methodological challenges. *Frontiers in Robotics and AI*, 5, 139.
- Schuetz, S., & Venkatesh, V. (2020). The rise of human machines: How cognitive computing systems challenge assumptions of user-system interaction. *Journal of the Association for Information Systems*, 21(2), 460–482.
- Seepana, C., Huq, F. A., & Paulraj, A. (2021). Performance effects of entrepreneurial orientation, strategic intent and absorptive capacity within cooperative relationships. *International Journal of Operations & Production Management*, 41(3), 227–250.
- Shepherd, S., Chartrand, T. L., & Fitzsimons, G. J. (2015). When brands reflect our ideal world: The values and brand preferences of consumers who support versus reject society's dominant ideology. *Journal of Consumer Research*, 42(1), 76–92.
- Song, C. S., & Kim, Y. K. (2022). The role of the human-robot interaction in consumers' acceptance of humanoid retail service robots. *Journal of Business Research*, 146, 489–503.
- Spatola, N., & Wudarczyk, O. A. (2021). Ascribing emotions to robots: Explicit and implicit attribution of emotions and perceived robot anthropomorphism. *Computers in Human Behavior*, 124, Article 106934.
- Statsenko, L., & Corral de Zubielqui, G. (2020). Customer collaboration, service firms' diversification and innovation performance. *Industrial Marketing Management*, 85, 180–196.
- Stone, M. (1974). Cross-validated choice and assessment of statistical predictions. *Journal of the Royal Statistical Society*, 36, 111–147.
- Sung, G. H., Hywel-Jones, N. L., Sung, J. M., Luangsa-Ard, J. J., Shrestha, B., & Spatafora, J. W. (2007). Phylogenetic classification of Cordyceps and the clavicipitaceous fungi. *Studies in Mycology*, 57, 5–59.
- Thellman, S., & Ziemke, T. (2021). The perceptual belief problem: Why explainability is a tough challenge in social robotics. *ACM Transactions on Human-Robot Interaction (THRI)*, 10(3), 1–15.
- Turkle, S. (2017). How computers change the way we think. In *Law and society approaches to cyberspace* (pp. 3–7). Routledge.
- Ulrich, D., & Diefenbach, S. (2017). Truly social robots—understanding human-robot interaction from the perspective of social psychology. In *Proceedings of the 12th international joint conference on computer vision, imaging and computer graphics theory and applications (VISIGRAPP 2017)* (pp. 39–45). Science and Technology Publications.
- Vanderelst, D., & Winfield, A. (2018). An architecture for ethical robots inspired by the simulation theory of cognition. *Cognitive Systems Research*, 48, 56–66.
- Wiese, E., Metta, G., & Wykowska, A. (2017). Robots as intentional agents: Using neuroscientific methods to make robots appear more social. *Frontiers in Psychology*, 8, 1663.
- Wilson, M. (2002). Six views of embodied cognition. *Psychonomic Bulletin & Review*, 9(4), 625–636.
- Wold, H. (1982). Soft modeling: The basic design and some extensions. In K. G. Joreskog, & H. Wold (Eds.), *Systems under indirect observation, part II* (pp. 1–54). North-Holland Publishing Co.
- Woods, S., Dautenhahn, K., Kaouri, C., te Boekhorst, R., Koay, K. L., & Walters, M. L. (2007). Are robots like people? Relationships between participant and robot personality traits in human-robot interaction studies. *Interaction Studies*, 8(2), 281–305.
- Wykowska, A., Chaminade, T., & Cheng, G. (2016). Embodied artificial agents for understanding human social cognition. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 371(1693), Article 20150375.
- Zhou, N., & Belk, R. W. (2004). Chinese consumer readings of global and local advertising appeals. *Journal of Advertising*, 33(3), 63–76.
- Zlotowski, J., Proudfoot, D., Yogeeswaran, K., & Bartneck, C. (2015). Anthropomorphism: Opportunities and challenges in human-robot interaction. *International Journal of Social Robotics*, 7(3), 347–360.
- Zogaj, A., Mähner, P. M., Yang, L., & Tscheulin, D. K. (2023). It's a Match! The effects of chatbot anthropomorphization and chatbot gender on consumer behavior. *Journal of Business Research*, 155, Article 113412.