

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Younger and Older Speakers' Use of Linguistic Redundancy with a Social Robot

Permalink

<https://escholarship.org/uc/item/9zt616fw>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 43(43)

Authors

Saryazdi, Raheleh

Nuque, Joanne

Chambers, Craig

Publication Date

2021

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Younger and Older Speakers' Use of Linguistic Redundancy with a Social Robot

Raheleh Saryazdi (raheleh.saryazdi@mail.utoronto.ca)

Department of Psychology, University of Toronto, 100 St. George Street, Toronto, ON, M5S 3G3

Joanne Nuque (joanne.nuque@mail.utoronto.ca)

Department of Psychology, University of Toronto Mississauga, 3359 Mississauga Road, Mississauga, ON, L5L 1C6

Craig Chambers (craig.chambers@utoronto.ca)

Department of Psychology, University of Toronto, 100 St. George Street, Toronto, ON, M5S 3G3

Abstract

Social robotics has shown expansive growth in areas related to companionship/assistance for older adults. Critically, everyday interactions with artificial agents often involve spoken language in the context of a shared visual environment. Therefore, language interfaces for these applications must account for the distinctive nature of visually-situated communication revealed by psycholinguistic studies. In traditional frameworks, "rational" speakers were thought to avoid redundancy, yet human-human communication research shows that both younger and older speakers include redundant information (e.g., color adjectives) in descriptions to facilitate listeners' visual search. However, this "cooperative" use of redundant expressions hinges on beliefs about listeners' perception (e.g., "pop-out" nature of human color processing). We explored the incidence and nature of younger and older speakers' redundant descriptions for a robot partner in different visual environments. Whereas both age groups produced redundant descriptions, there were important age differences for when these descriptions occurred and for the properties encoded in them.

Keywords: Informativity; Social Robotics; Aging; Human-Robot Communication

Introduction

The increase in the global aging population has led to a significant rise in research on new technologies to promote aging in place (living independently in one's residence as long as possible; Kim et al., 2017; Mois & Beer, 2020). These technologies (e.g., smart homes, robots) often rely on spoken language interfaces, and as such there is growing interest in ensuring they are equipped to manage *natural and contextualized* patterns of language use. The present research explores this issue in a context where speakers generate instructions for a social robot. The core question is whether younger and older speakers produce redundant adjectives that reflect implicit beliefs about a robot addressee's perceptual abilities and capacity for incremental interpretation. This can reveal features of younger and older adults' language use with a robot partner and also extends our understanding of human-robot communication.

As background, one of the key applications for social robots that support healthy aging is to provide instructions for everyday tasks (e.g., cooking, clothing selection, managing a daily schedule; Begum et al., 2013). Critically, these

instructions often involve real objects in the immediate visual environment or depicted objects on a screen. One key consideration for spoken language interfaces used in social robots should therefore be the distinctive nature of *visually-situated* spoken communication, as demonstrated in psycholinguistic work on human-human communication. This work has shown that *the shared visual environment* alters the character of communicative behavior, such as the extent to which human speakers adhere to conventions governing the amount of information they express to achieve communicative goals. In particular, when referring to the visual here-and-now, speakers tend to include redundant information (e.g., color adjectives) to facilitate listeners' real-time visual search for the intended referent. Recent evidence also suggests that older adults tend to provide more redundant (and lexically diverse) descriptions than younger adults (Healey & Grossman, 2016; Saryazdi, Bannon, et al., 2019; although see Long et al., 2020). The use of redundant information reflects certain implicit beliefs on the speaker's part regarding the perceptual abilities of their human partner. What is not yet known is whether younger and older speakers generate the same implicit beliefs with a robot partner and use these beliefs to tailor descriptions in different visual contexts. Thus, the aim of the present study is to explore age-related differences in speakers' use of linguistic redundancy with a robot partner in various visual contexts.

Linguistic Redundancy

In face-to-face communication, people often refer to objects in the immediate environment. In accordance with Grice's Maxim of Quantity (Grice, 1975), rational speakers should provide referential descriptions that contain "only the right amount" of information, not too much nor too little. In practice, this is often reflected in speakers' inclusion or omission of *modifiers*. For example, in a context containing two different-colored cups, using a modified description (*blue cup*) would be considered optimal. However, in a context where there is only one cup present, the same description would technically contain too much information. In information-theoretic approaches to language, referring expressions that include too much information are called *overspecified*, whereas those that provide too little information are called *underspecified*. Although an attempt

to tread the line between too much and too little information explains many aspects of language behavior, an increasing body of work has shown that younger and older adults often spontaneously overspecify descriptions for objects in the immediate visual environment (Engelhardt et al., 2006; Long et al., 2020; Pechmann, 1989; Saryazdi, Bannon, et al., 2019). This apparent violation to Grice's Maxim of Quantity has been a topic of much debate. It has been argued that the reason people overspecify is not that they are failing to be rational, but because they are instead providing information they implicitly believe would help facilitate their conversational partners' comprehension (Jara-Ettinger & Rubio-Fernandez, 2020; Rubio-Fernandez, 2016, 2019).

In line with these findings, it has been shown that the rate of overspecification varies as a function of visual complexity. Visual complexity is often manipulated by increasing the number of overall objects in the field of view and/or the number of items from the same category as the named target (often referred to as competitors). Evidence reveals that speakers have a higher tendency to overspecify in complex visual contexts compared to simple contexts (e.g., Elsner et al., 2018; Healey & Grossman, 2016). Further, object characteristics that are high in perceptibility and can easily be used to discriminate a target from other objects in the crowded environment are more likely to be included in an object's description. Overspecification is therefore intended to aid listeners' visual search, which is why it tends to occur more in complex visual environments when it is not as easy to efficiently locate an object. This in turn suggests that overspecification is in fact a sign of cooperativeness on the part of the speaker (Rubio-Fernandez, 2016, 2019). In line with this, a number of studies have shown that color, often being the most salient characteristic of an object, commonly occurs in overspecified descriptions produced by both younger (Belke & Meyer, 2002; Tarenskeen et al., 2015) and older adults (Saryazdi, Bannon, et al., 2019). Redundant color adjectives also tend to not impair listeners' comprehension (Tourtour et al., 2019).

Additional observations can be made regarding patterns of overspecification in the speech of older adults in particular. Some studies have shown that older adults exhibit an increased tendency for overspecification relative to younger adults when referring to objects in the immediate environment, as well as greater lexical diversity in the choice of modifiers used (Healey & Grossman, 2016; Saryazdi, Bannon, et al., 2019). One explanation for why older adults overspecify more than younger adults is that they are more diligent in ensuring success in listeners' referential identification, and thus being cooperative. This builds on earlier work suggesting that redundancy is a type of communicative strategy (Long et al., 2020; Saryazdi, Bannon, et al., 2019). Another explanation is that older adults' overspecification is caused by age-related declines in cognitive abilities (Hasher & Zacks, 1988; Park et al., 2002) that give rise to word finding difficulties as well as increased verbosity (Arbuckle et al., 2000; James et al., 1998). Therefore, redundancy might be a strategy to "buy more

time" to retrieve the object's name. However, analysis of this question suggests that this is an unlikely explanation (Saryazdi, Bannon, et al., 2019). The question of interest here is whether age-related differences in overspecification also persist in human-robot communication.

Present Study

The present study examines the descriptions produced by younger and older speakers as they provide instructions to a social robot (Furhat Robotics, Sweden). Of interest was whether speakers would generate overspecified descriptions that putatively help their robot addressee to quickly locate the target object. As mentioned earlier, the implicit beliefs that underlie this behavior hinge on perceptual phenomena (the "pop-out" nature of certain visual properties) and human listeners' tendency to engage in rapid incremental language interpretation. Do these beliefs guide how speakers design descriptions for a robot, which might not be ascribed the same abilities or behaviors?

We used a game-like task involving a "busy" visual context with eight items. In addition, on critical trials, we manipulated whether the target object was the only object in the display bearing certain properties (e.g., the only pink object), or whether another item in the display also bore the properties in question. This influenced the extent to which a redundant modifier would be considered helpful in rapidly narrowing listeners' attention to the intended object. To illustrate, in the competitor present condition, the target object (e.g., pink, open box, *see Figure 1*) shared properties with another object (e.g., pink, open door), which we refer to as the property-matched competitor. In the competitor absent condition, the competitor was replaced by an unrelated item from a different object category that did not share properties with the target object (e.g., a clownfish, neither pink nor 'open'). According to classic Gricean theory, participants should not overspecify in either condition because the target is unique in its category. However, the work described earlier suggests that, in "busy" visual contexts (such as the one we use here), overspecification will likely occur as a result of speakers' cooperativeness in aiding object identification. Although overspecification in both the competitor present and absent condition would narrow attention to fewer objects (one or two), we would expect the attentional narrowing to be stronger in the latter case, when the modifier picks out a single object. One question is whether this will still be the case when the listener is a robot.

One possibility is that speakers might rarely overspecify, suggesting they believe the robot's language processing and visual-perceptual abilities are more limited than a human's. Alternatively, they may routinely produce overspecified descriptions, typically involving objects' color properties (as in human-human communication). This would suggest speakers believe the robot has similar perceptual abilities as a human and is therefore able to benefit from color's "pop out" feature. It is also possible that older adults may overspecify more than their younger counterparts, particularly in the competitor absent condition. This idea is

motivated by research showing older adults are more likely to spontaneously attribute human characteristics to a robot (Alimardani & Qurashi, 2019; Pak et al., 2020). Thus, as for a human partner, a redundant modifier would be most beneficial for rapid identification of the intended target when the modifier strongly restricts visual search. However, if older speakers simply show a general trend toward overspecification, this may instead suggest their inclusion of a modifier arises from age-related patterns described earlier, such as increased word finding difficulties.

Method

Participants

Participants were 20 younger ($M_{age} = 18.95$, $SD = 2.14$) and 13 older adults ($M_{age} = 76.31$, $SD = 5.23$). (Data collection was interrupted due to COVID-19.) Younger adults were recruited from a student participant pool and received partial course credit or \$12/hour. Older adults were recruited from a volunteer database and received \$15/hour. All participants reported normal or corrected-to-normal vision and were screened for visual acuity and color blindness.

Materials

On each trial, participants were presented with a set of eight images in a 2 x 4 grid format. Images were displayed at 375 x 375 pixels on a 24 in. Dell touchscreen monitor (Figure 1). On critical trials, the target item was always a unique object in its category, requiring no additional information for identification beyond a one-word label (e.g., *box*). Each target item was chosen so it could be naturally referred to with a color and state modifier (e.g., *pink box* or *open box*). However, the use of these or any other modifier types (e.g., size, shape, location) on critical trials would reflect linguistic redundancy (overspecification), given that, as mentioned, the name of the object on its own would be sufficient. To explore the incidence of linguistic redundancy and whether this behavior varied depending on the visual context, we included a manipulation varying the presence/absence of an object in a different category that matched the depicted state/color of

the target object (e.g., open pink door in the presence of the open pink box target). This manipulation helps clarify whether overspecification is limited to cases where the information in the modifier uniquely identifies the target (e.g., pink box is the only pink object present) compared to when this information is less helpful because it applies to more than one candidate (e.g., where two pink objects are present, only one of which is the box). In summary, as shown on Figure 1, each critical trial contained a unique target, a modifier-matched (in competitor present condition only, right panel) or an unrelated object (in competitor absent condition, left panel), with the six remaining items all unrelated objects that did not share the same color/state as the target object. Filler trials were included to counteract strategies that might arise due to the critical trials. For example, some fillers had another same category object requiring a modified description for identification (e.g., *the grey sock* in the presence of a yellow sock).

Participants were asked to provide instructions to a social robot so that it could identify the target objects. The robot used was a Gen2 Furhat model which possessed human-like face and voice features (Figure 2). In total there were 50 distinct object arrays (2 practice, 24 critical, and 24 filler trials). For each object array, participants would provide two instructions, the first always referring to the critical object and the second to another object (24 critical and 72 non-critical descriptions produced).

Procedure

Participants were led to a sound booth and completed a consent form, language questionnaire, and color blindness test. Next, they were seated in front of the monitor positioned between themselves and the robot, whose eyes remained closed as the experimenter explained the task. Participants were asked to provide the robot verbal instructions regarding which objects to select on the screen. Participants were given a standing spiral-bound booklet. Each page (which corresponded to a single trial) showed two grid squares labelled as "1" and "2" within a 2 x 4 grid. This denoted the targets' locations for the first and second instructions,

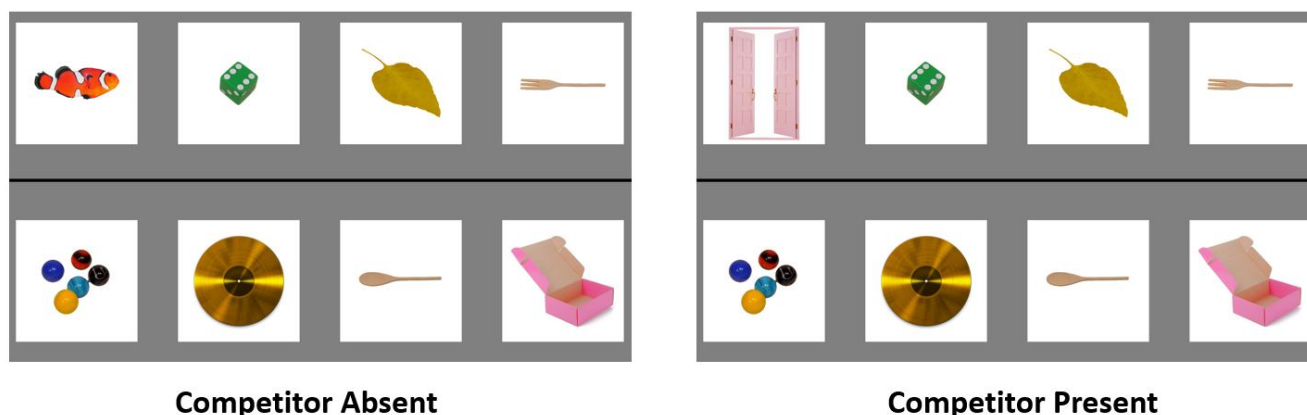


Figure 1: Sample experimental trial in the competitor absent and competitor present conditions (target item = box).

respectively. The participant related this information to items displayed on the screen located between them and the robot.

Participants were asked to phrase their instructions in the form "Select the...[object]" and to refrain from using any location terms (e.g., object in the bottom right). The experimenter led participants through two practice items (which the participant subsequently repeated with the robot) to confirm that they understood the task and could locate the target objects given the booklet's information. The experimenter then left the sound booth. Control of the robot's speech and behaviors was achieved via a Wizard-of-Oz interface managed by the experimenter, who stayed outside the sound booth during the task. The experimenter could see the booth interior via a window and could hear, via closed-circuit audio monitoring, the sounds occurring inside. A customized Graphic User Interface (GUI) was made to allow for control over the robot's verbal and nonverbal reactions to the participant's instructions. An "attend-to-user" feature (whereby the robot's head and gaze reoriented automatically to follow the user's face) was enabled to create a sense of social engagement during the task. Before beginning the experiment, the robot, which was already facing the participant, opened its eyes and provided an introduction.

At the start of each trial, the display screen reminded participants to flip to the correct trial in their booklet and to touch a green arrow on the monitor to proceed. After the participant provided an instruction, the robot scanned the screen and gave one of four responses confirming its recognition of the object (e.g., "Yes, I found it"). A red box appeared around the intended image, accompanied by a tone. The robot asked for clarification if participants did not follow task instructions correctly or provided the robot unclear instructions. The task took around 20 minutes to complete.

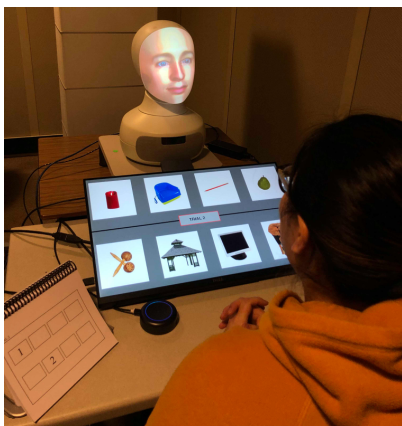


Figure 2: Experimental setup.

Coding Procedure

Research assistants blind to the experimental conditions transcribed participants' critical descriptions. They were coded in terms of 1) the type of noun used, namely whether the label was considered a basic level (e.g., shirt), subordinate (e.g., dress-shirt), or superordinate term (e.g., clothing); 2)

whether the noun was accompanied by a modifier; 3) the type of modifier, whether it referred to the object's color (e.g., green shirt), state (e.g., wrinkled shirt), size (e.g., large shirt), location (e.g., shirt on the left), or other (e.g., men's shirt); 4) a further coding for the "other" category to index what alternative modifiers participants used; 5) the location of the modifier(s): prenominal (e.g., green shirt), postnominal (e.g., shirt on the left), or both (e.g., green shirt on the left).

Results

All analyses were conducted with R statistical package version 3.6.1 (R Core Team, 2019). Logistic mixed effect analyses were performed using the lme4 package version 1.1-21 (Bates et al., 2015) and statistical significance was assessed with the lmerTest package version 3.1-0 (Kuznetsova et al., 2017). Unless otherwise specified, each model included age (younger adult = 1 vs. older adult = -1), competitor condition (competitor present = 1, competitor absent = -1), and their interaction as fixed effects. The random effects were intercept terms for participants and items, as well as by-participant slope for the competitor condition, and by-item slope terms for age, competitor, and their interaction. Any incomplete audio recordings or cases involving incorrect reference were excluded from analyses. The item pail (color = blue, state = full) was also excluded because younger adults excessively referred to this item as a "bucket of sand". It was not evident how this noun phrase should be coded, namely whether bucket is the name or a type of classifier specifying the containment of the actual object. Note that the inclusion/exclusion of this item does not change the pattern of results observed.

Noun Type

Some items showed very high name agreement, with almost all participants referring to the depicted object with the same label (e.g., candle). Other items were referred to with more variable names (e.g., pills - vitamins). As mentioned, nouns were coded as basic level, subordinate, or superordinate. This initial check was motivated by recent evidence that both younger and older adults often use subordinate terms instead of modified basic terms when possible (e.g., *Dalmatian, sunflower*; Saryazdi, Bannon, et al., 2019). Reflecting the fact that the materials in the present study were not specifically designed to elicit subordinate terms, only 15% of descriptions produced by older adults and 13% of descriptions produced by younger adults contained a subordinate term. Further, this was mostly driven by a few items, such as pants (e.g., *jeans*), pills (e.g., *vitamins*), and bag (e.g., *purse*). Given that instances of subordinate terms (and superordinate terms) were low and item-specific, they were not excluded from the main analysis.

Modifier Use

To examine the incidence of linguistic redundancy, each trial was coded for whether the description contained at least one modifier (yes = 1, no = 0). Because the fully specified statistical model did not converge, we removed the

interaction term for the by-item random slope. The results revealed a significant interaction of Age x Competitor, $Estimate = 0.22$, $SE = 0.11$, $z = 2.06$, $p = .039$ and no other significant effects. To explore this interaction, we conducted two separate follow-up analyses for the competitor absent and present conditions. Whereas there were no differences between younger and older adults' use of modifiers in the competitor absent condition, $Estimate = 0.28$, $SE = 0.32$, $z = 0.87$, $p = .385$, younger adults provided significantly more modifiers than older adults in the competitor present condition, $Estimate = 0.77$, $SE = 0.35$, $z = 2.20$, $p = .028$. Interestingly, as shown in Figure 3, the numerical pattern reveals that older adults reduce their use of modifiers when there was a modifier-matched competitor present, in comparison to when the competitor was absent.

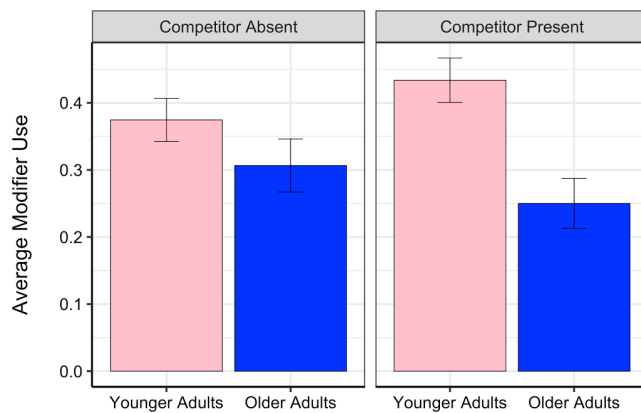


Figure 3: Average incidence of modifier use.

Modifier Type

Next, we explored the type of modifiers used by participants. As mentioned, color has shown to be the most common type of linguistic redundancy in referential communication studies involving objects in the visual environment. This was borne out in the current study where 62% of modifiers were color terms. To explore whether there were age-related differences in the use of color and whether this varied as a function of the competitor condition, we first subsetting the dataset to include only trials with at least one modifier. Next, we coded whether the first modifier used by the participant was color (color = 1, everything else = 0). Although there were cases where participants used more than one modifier, there was not a sufficient number of these trials to include in the analysis. The results of a logistic mixed-effect analysis based on the first modifier showed no significant differences in the use of color as a function of age and competitor condition (all p 's > .78). However, some interesting patterns were observed across both age groups. As shown on Table 1, unlike younger adults, older adults were almost as likely to use non-color modifiers (~46%) as much as they used color (~54%). This was not driven by the incidence of state adjectives (e.g., open/closed), as the rate of state modifiers was similar across both age groups (~21%). Although state adjectives are often not considered in studies of referential production, the target

items in the present study were specifically selected so they could be referred to with a state modifier, thus increasing the likelihood that participants would use these adjectives.

Table 1: Breakdown of modifier type by age group (%).

Modifier Type	Younger	Older
Color	65	54
State	21	21
Size	2	4
Location	0	1
Other	12	20

Older adults also tended to use modifiers categorized as "other" more than younger adults. We explored several subcategories within this class to test for possible patterns, namely attributes (e.g., *skinny jeans*), function (e.g., *packing box*), material (e.g., *glass jar*), count (e.g., *pair of shoes*), and gender (e.g., *men's shirt*). Interestingly, younger adults never used gender type adjectives compared to older adults. Younger adults were also more likely to refer to items in terms of their material. In fact, this was the most common type of "other" modifier used in the second modifier position by younger adults. Finally, consistent with previous research, older and younger adults used mostly prenominal modifiers in their descriptions (83% and 98%, respectively). Note that excluding postnominal modifiers from the earlier color analysis does not change the pattern of results observed.

Discussion

The present study examined the incidence and nature of overspecification in descriptions produced by younger and older adults when communicating with a social robot. We found that, as in human-human communication, speakers in both age groups often produced overspecified descriptions when referring to objects for a robot partner. In fact, approximately a third of all descriptions were overspecified. Recent work has suggested that speakers produce these descriptions in "busy" visual contexts (such as ours) because modifiers make it easier for listeners to efficiently identify the intended target. Some accounts also predict that the rate of overspecification should be higher when the redundant modifier limits visual attention to only a *single* object. This is because visual search would be most optimized by redundancy in the context where modifier information uniquely identifies the target item. However, we found no effect of the competitor manipulation that varied whether the target shared the denoted property with another object. Instead, there was an interaction involving age, where both groups performed similarly when the competitor was absent, yet older adults were significantly less likely to overspecify when the competitor was present. The behavior of older adults is consistent with earlier claims that overspecification should be more likely to occur when the property allows for the rapid narrowing of visual attention.

One possible explanation for this pattern is that older adults are more consciously diligent in their design of descriptions for the robot, providing the most optimal information and performing more rationally. Given the evidence that older adults implicitly expect robots to have the same perceptual abilities as humans (Saryazdi, Nuque, et al., 2019), it is possible that older adults produced redundant descriptions more often in the competitor absent condition because they believed the robot would reap the greatest benefit of overspecification in this context, in the same way as human listeners. However, what is somewhat intriguing is that the *overall* rate of overspecification was lower with older adults. It is interesting to note that, although earlier work by Saryazdi, Bannon, et al. (2019) showed that older adults provided comparatively more redundant modifiers, this pattern was not observed when producing descriptions for a computer addressee in the no-contrast condition. At a minimum, this outcome rules out the possibility that older adults' greater tendency to overspecify (as reported in past studies) results from patterns of age-related decline in cognitive and linguistic abilities. Instead, this suggests that older adults implement different strategies depending on the situational context. We also speculate that the overall greater likelihood of overspecification in younger adults' descriptions in the present study still reflects "cooperativeness" on the part of the speakers, but arises from a different set of beliefs. That is, younger adults' behavior may in fact reflect a specific assumption about the robot's capabilities. Recall that the robot scanned the screen before confirming it found the denoted object. Younger adults may be more likely to spontaneously view this behavior as reflecting advanced image processing on the part of the robot, which may have encouraged their use of providing more redundant descriptions. Alternatively, younger adults who arguably have more experience with technology and are thus more likely to understand its limitations, may have provided redundant descriptions as a means of helping the robot as much as possible. Additional research is needed to explore whether these patterns of results hold in a larger sample size of participants, as well as with other types of artificial agents.

It is also important to note that people tend to ascribe more humanlike characteristics to robots that look and behave like a human (Briggs & Scheutz, 2014; Torrey et al., 2010). In the current study, our robot possessed a humanlike face, spoke with a humanlike voice, and engaged in joint attention, as is common in face-to-face referential communication. These humanlike features may have increased the social engagement participants experienced, in turn encouraging them to provide redundant information to the robot addressee as a way of facilitating its comprehension (in the same way they would do for a human listener). It would be interesting to explore whether the patterns found in the current study would persist if the robot listener lacked humanlike features.

We also found that, when modifiers were produced, both age groups were most likely to use an adjective denoting the target object's color, which aligns with findings from previous human-human communication studies (Rubio-

Fernandez, 2016; Saryazdi, Bannon, et al., 2019). Interestingly, the incidence of state modifiers was both comparatively high and quite similar across age groups. This is noteworthy as many studies have not coded this property or have not used materials where state information could be readily encoded (although see Parker & Heller, 2019). Further, older adults were more likely to use modifier types in our "other" category, consistent with previous findings (Saryazdi, Bannon, et al., 2019). We also found differences in the incidence of subtypes in this category. For example, younger adults never used gendered modifiers (e.g., "lady's purse"), but older adults did (see also Saryazdi, Bannon, et al., 2019). This may be because gendered distinctions in clothes and other categories may not be salient to younger adults.

In summary, the present study reveals that overspecification occurs spontaneously in human-robot communication, similar to patterns observed in human-human communication. This is noteworthy because the factors argued to make overspecification "rational" or "cooperative" hinge on a speaker's beliefs about their addressee's perceptual abilities and tendency to interpret spoken language incrementally. These are beliefs that might not be spontaneously extended to a robot listener. However, the current results suggest participants do implicitly ascribe these attributes to the robot. We also found age-related differences in the rate and nature of overspecification, which may reflect younger and older adults' different approaches to being "cooperative" in the communicative task. Additional research is needed to pinpoint the nature of these differences, which may be important for the design and development of tailored language interfaces for advanced technologies. Together with studies of human-human communication, the current results highlight the ways in which the production of redundant descriptions depends on the age of the speaker and the distinct nature of the visually-situated environment.

Acknowledgments

Funding for this research was provided by Discovery Grant 458164 to Craig Chambers from the Natural Sciences and Engineering Research Council (Canada) and doctoral awards to Raheleh Saryazdi from the Natural Sciences and Engineering Research Council and the Pfizer/Queen Elizabeth II Prize in Science and Technology. The authors thank Dhrumil Patel, Jan Dizon, Simran Kanda, Kelly Qi, Caren Khalaf, and Jarrod Servilla for their assistance.

References

- Alimardani, M., & Qurashi, S. (2019). Mind perception of a sociable humanoid robot: A comparison between elderly and young adults. *Iberian Robotics Conference*, 96–108.
- Arbuckle, T. Y., Nohara-LeClair, M., & Pushkar, D. (2000). Effect of off-target verbosity on communication efficiency in a referential communication task. *Psychology and Aging*, 15(1), 65–77.
- Bates, D., Mächler, M., Bolker, B. M., & Walker, S. C. (2015). Fitting linear mixed-effects models using lme4.

- Journal of Statistical Software*, 67(1).
- Begum, M., Wang, R., Huq, R., & Mihailidis, A. (2013). Performance of daily activities by older adults with dementia: The role of an assistive robot. *IEEE International Conference on Rehabilitation Robotics*, 1–8.
- Belke, E., & Meyer, A. S. (2002). Tracking the time course of multidimensional stimulus discrimination: Analyses of viewing patterns and processing times during “same”-“different” decisions. *European Journal of Cognitive Psychology*, 14(2), 237–266.
- Briggs, G., & Scheutz, M. (2014). How robots can affect human behavior: Investigating the effects of robotic displays of protest and distress. *International Journal of Social Robotics*, 6(3), 343–355.
- Elsner, M., Clarke, A., & Rohde, H. (2018). Visual complexity and its effects on referring expression generation. *Cognitive Science*, 42, 940–973.
- Engelhardt, P. E., Bailey, K. G. D., & Ferreira, F. (2006). Do speakers and listeners observe the Gricean Maxim of Quantity? *Journal of Memory and Language*, 54(4), 554–573.
- Grice, P. H. (1975). Logic and conversation. In P. Cole & J. L. Morgan (Eds.), *Syntax and semantics, Vol. 3, speech acts* (Vol. 3, pp. 41–58). Academic Press.
- Hasher, L., & Zacks, R. T. (1988). Working memory, comprehension, and aging: A review and a new view. *The Psychology of Learning and Motivation*, 22, 193–225.
- Healey, M. L., & Grossman, M. (2016). Social coordination in older adulthood: A dual-process model. *Experimental Aging Research*, 42(1), 145–164.
- James, L. E., Burke, D. M., Austin, A., & Hulme, E. (1998). Production and perception of “verbosity” in younger and older adults. *Psychology and Aging*, 13(3), 355–367.
- Jara-Ettinger, J., & Rubio-Fernandez, P. (2020). *The social basis of referential communication: Speakers construct reference based on listeners’ expected visual search*. PsyArXiv.
- Kim, K. il, Gollamudi, S. S., & Steinhubl, S. (2017). Digital technology to enable aging in place. *Experimental Gerontology*, 88, 25–31.
- Kuznetsova, A., Brockhoff, P. B., & Christensen, R. H. B. (2017). lmerTest package: Tests in linear mixed effects models. *Journal of Statistical Software*, 82(13), 1–26.
- Long, M., Rohde, H., & Rubio-Fernandez, P. (2020). The pressure to communicate efficiently continues to shape language use later in life. *Scientific Reports*, 10(1), 1–13.
- Mangold, R., & Pobel, R. (1988). Informativeness and instrumentality in referential communication. *Journal of Language and Social Psychology*, 7, 181–191.
- Mois, G., & Beer, J. M. (2020). Robotics to support aging in place. In R. Pak, E. J. de Visser, & E. B. T. Rovira (Eds.), *Living with robots: Emerging issues on the psychological and social implications of robotics* (pp. 49–74). Academic Press.
- Pak, R., Crumley-Branyon, J. J., de Visser, E. J., & Rovira, E. (2020). Factors that affect younger and older adults’ causal attributions of robot behaviour. *Ergonomics*, 63(4), 421–439.
- Park, D. C., Lautenschlager, G., Hedden, T., Davidson, N. S., Smith, A. D., & Smith, P. K. (2002). Models of visuospatial and verbal memory across the adult life span. *Psychology and Aging*, 17(2), 299–320.
- Parker, M., & Heller, D. (2019). Overspecifying state information in the production of referring expressions. *CUNY Conference on Human Sentence Processing*.
- Pechmann, T. (1989). Incremental speech production and referential overspecification. *Linguistics*, 27(1), 89–110.
- R Core Team. (2019). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing. <http://www.r-project.org/>
- Rubio-Fernandez, P. (2016). How redundant are redundant color adjectives? An efficiency-based analysis of color overspecification. *Frontiers in Psychology*, 7, 1–15.
- Rubio-Fernandez, P. (2019). Overinformative speakers are cooperative: Revisiting the Gricean Maxim of Quantity. *Cognitive Science*, 43(11).
- Saryazdi, R., Bannon, J., & Chambers, C. G. (2019). Age-related differences in referential production: A multiple-measures study. *Psychology and Aging*, 34(6), 791–804.
- Saryazdi, R., Nuque, J., & Chambers, C. G. (2019). Communicative perspective taking with artificial agents. *Embodied Situated Language Processing (ESLP) and Attentive Listener in the Visual World (AttLis) Joint Conference*.
- Tarenskeen, S., Broersma, M., & Geurts, B. (2015). Overspecification of color, pattern, and size: Salience, absoluteness, and consistency. *Frontiers in Psychology*, 6, 1703.
- Torrey, C., Fussell, S. R., & Kiesler, S. (2010). What robots could teach us about perspective taking. In E. Morsella (Ed.), *Expressing Oneself / Expressing One’s Self: Communication, Cognition, Language, and Identity* (pp. 93–106). Psychology Press / Taylor & Francis.
- Tourtour, E. N., Delogu, F., Sikos, L., & Crocker, M. W. (2019). Rational over-specification in visually-situated comprehension and production. *Journal of Cultural Cognitive Science*, 3(2), 175–202.