

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Two's company but six is a crowd: emergence of conventions in multiparty communication games

Permalink

<https://escholarship.org/uc/item/3gd9j28x>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 44(44)

Authors

Boyce, Veronica

Hawkins, Robert

Goodman, Noah

et al.

Publication Date

2022

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

Two's company but six is a crowd: emergence of conventions in multiparty communication games

Veronica Boyce (vboyce@stanford.edu) **Robert D. Hawkins** (rdhawkins@princeton.edu)
Department of Psychology
Stanford University
Princeton Neuroscience Institute
Princeton University

Noah D. Goodman (ngoodman@stanford.edu) **Michael C. Frank** (mcfrank@stanford.edu)
Departments of Psychology and Computer Science
Stanford University
Department of Psychology
Stanford University

Abstract

From classrooms to dinner parties, many of our everyday conversations take place in larger groups where speakers address multiple listeners at once. Such multiparty settings raise a number of challenges for classical theories of communication, which largely focus on dyadic interactions. In this study, we investigated how speakers adapt their referring expressions over time as a function of the feedback they receive from multiple parties. We collected a large corpus of multiparty repeated reference games (98 games, 390 participants, 116K words) where speakers designed referring expressions for groups of 1 to 5 listeners. Larger groups tended to use more words total and to introduce more new words; nonetheless, most groups were able to converge to more efficient conventions regardless of the number of listeners.

Keywords: Pragmatics; Communication; Reference game; Convention; Reduction; Efficiency

Introduction

Verbal communication is an integral part of our daily lives. We coordinate schedules with partners, socialize with friends over board games, learn and teach in seminar classes, and listen to podcasts. Communicative environments range in size from one-on-one dialogue to broadcast communication to large groups, but the goal of efficient communication is shared across these (Branigan, 2006; Ginzburg & Fernandez, 2005; Traum, 2004). Shared referring expressions are a necessity for efficient communication; a thing or an idea needs some sort of name that the interlocutors will jointly understand. In many cases, there are widely shared conventionalized expressions for objects or ideas, but in other cases, spontaneous ad-hoc expressions must be invented.

The formation of these new reference expressions is well-studied in dyadic contexts and has been a case study for efficient communication more broadly. But these dynamics may be different in larger groups, which are less studied. Our current work builds on the dyadic reference game tradition by extending it to larger groups.

Clark & Wilkes-Gibbs (1986) established an experimental method for studying the emergence of new referring expressions that has now become standard (building on Krauss & Weinheimer, 1964, 1966). Two participants see the same set of tangram figures; the speaker describes each figure in turn so the listener can select the target from the set of figures. The speaker and listener repeat this process with the same images over a series of blocks. Early descriptions are long and make

reference to multiple features in the figure, but in later iterations, shorthand conventional names for each figure emerge; this shortening of utterances is called ‘reduction’.

Recently, online participant recruitment and web-based experiments have made it possible to study this convergence in larger populations (Haber et al., 2019; Hawkins, Frank, & Goodman, 2020). In Hawkins et al. (2020), 83 pairs completed a similar iterated reference experiment where they communicated via a chat box. Speakers reduced their utterances, producing fewer words per image in later blocks than in earlier blocks, in line with results from face-to-face, oral paradigms.¹

How does this process proceed in multi-party communication? In a dyad, speakers can tailor their utterances to the one listener, but in large groups, speakers must balance the competing needs of different listeners (Schober & Clark, 1989; Tolins & Fox Tree, 2016). These effects likely vary by both the knowledge state of and communication channels available to the listeners (Fox Tree & Clark, 2013; Horton & Gerrig, 2002, 2005). Prior work has focused on manipulating knowledge states by adding new listeners to established groups.

In this context, one approach for speakers is to ‘aim low’ and produce utterances tailored to the least knowledgeable listener (Yoon & Brown-Schmidt, 2018). For instance, in Yoon & Brown-Schmidt (2014), speakers developed conventions with one listener but then used longer descriptions with a new listener. Another strategy for speakers is to integrate across listeners and balance efficiency with informativeness by ‘aiming in the middle’. In Yoon & Brown-Schmidt (2019), speakers communicating to a mixed group of 3 experienced listeners and 1 naive listener used shorter utterances and made fewer accommodations than they did in groups with a greater fraction of naive listeners. Both of these strategies predict that larger groups will be slower to converge than smaller groups.

Disagreements about how to conceptualize referents can also slow groups down. In Weber & Camerer (2003), pairs of participants played a reference game with the same image sets before a listener switched groups and joined a different pair, making a group of three. The addition of the new listener slowed both listeners down for multiple rounds. When a listener switched groups, they brought preconceptions about

¹We use “speaker” and “listener” to refer to the roles describing and selecting targets, regardless of communication modality.

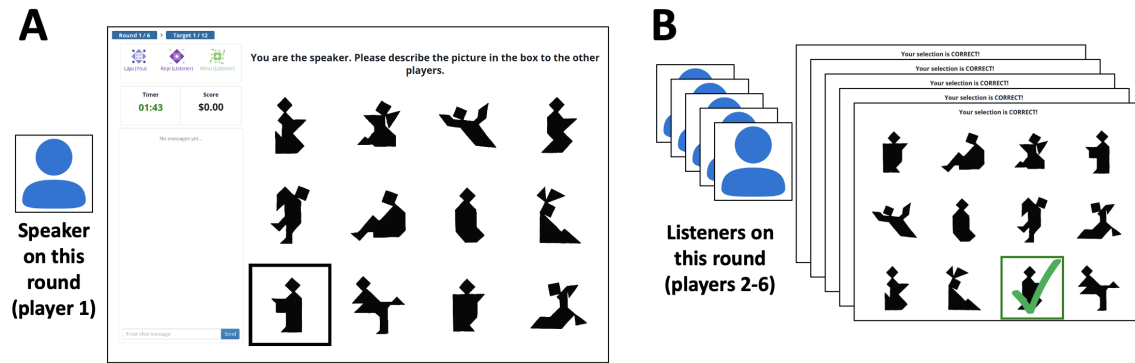


Figure 1: All participants saw all 12 tangram images. (A) Speaker’s view during selection phase. (B) During the feedback stage, speakers saw what figure each person chose, but listeners only learned if their selection was correct or incorrect. Listeners were not shown what other listeners chose.

how the pictures should be described which conflicted with how the speaker was used to describing the images. This result predicts that, with more perspectives in play, larger groups may have more difficulty agreeing on common conceptualizations.

In general, listeners expect speakers to maintain conventions and stick to descriptions that were similar to successful descriptions. However, listeners were not surprised to hear different descriptions of a familiar object if it came from a new speaker who had just entered the room (Metzing & Brennan, 2003). It’s unclear what this finding predicts about new speakers who are present as fellow listeners during prior blocks – will listeners expect them to maintain conventions?

Work on multi-party communication has focused on the addition of a new person into a pair or group that had built up some shared representations. Our present work complements this prior work by examining the effect of group size during the process of convention formation. We extend the dyadic repeated reference game paradigm of Hawkins et al. (2020) to games for 2–6 players who rotate between speaker and listener roles. This paradigm allows us to confirm that these findings in dyads extend to larger groups: that accuracy and speed will increase across blocks (question 1) and that speakers will reduce their utterances (produce fewer words) in later blocks (question 2). Additionally, we will be able to test for trends across group size, allowing us to ask whether smaller groups use shorter utterances and reduce faster than larger groups (question 3) and how conventions emerge in larger groups (question 4). In sum, these analyses will fill a gap in the literature by providing a basic characterization of how convention-formation and communication occurs in larger groups.

Methods

Building on the methods of Hawkins et al. (2020), we used Empirica (Almaatouq et al., 2020) to create real-time multi-player reference games. In each game, one of the players started as the speaker who saw an array of tangrams with one

highlighted (Figure 1A) and communicated which figure to click to the other players (listeners). After the speaker had identified each of the 12 images in turn, the speaker role rotated to another player and the process repeated with the same images. In total, there were 6 blocks, giving each player at least one chance to be the speaker. We recorded what participants said in the chat, as well as who selected what image and how long they took to make their selections.² We report how we determined our sample size, all data exclusions, all manipulations, and all measures in the study.³

Participants

We recruited participants between May and July 2021 using the Prolific platform; participants had all self-reported as fluent native English speakers on Prolific’s demographic pre-screen. Participants were paid \$7 for 2-player games, \$8.50 for 3-player games, \$10 for 4-player games, and \$11 for 5- and 6-player games (with the intention of a \$10 hourly rate), in addition to up to \$2.88 in performance bonuses. A total of 390 people each participated in one game.

Materials

We used the 12 tangram images used by Hawkins et al. (2020) and Clark & Wilkes-Gibbs (1986) (see Figure 1). These images were displayed in a grid with order randomized for each participant (thus descriptions such as “top left” were ineffective as the image might be in a different place on the speaker’s and listeners’ screens). The same images were used every block.

Procedure

We implemented the experiment using Empirica, a Javascript-based platform for running real-time interactive experiments online (Almaatouq et al., 2020). From

²Code to run the experiment, as well as data and analysis code are available at https://osf.io/qdvbr/?view_only=47aebfde243f405e9c42a45cacb697d2.

³Our preregistrations are at https://osf.io/cn9f4/?view_only=7fdacd698b24465cbl8699050af5bfc and https://osf.io/rpz67?view_only=5284203e2b644fc5ac39cf3e723b9a7e.

Players	Partial	Complete
2	4	15
3	2	18
4	2	19
5	3	17
6	6	12

Table 1: Number of games run for each player count.

Prolific, participants were directed to our website where they navigated through a self-paced series of instruction pages explaining the game. Participants had to pass a quiz to be able to play the game. They were then directed to a “waiting room” screen until their partners were ready.

Once the game started, participants saw screens like Figure 1A. Each trial, the speaker described the highlighted tangram image so that the listeners could identify and click it. All participants were free to use the chat box to communicate, but listeners could only click once the speaker had sent a message. Once a listener clicked, they could not change their selection. There was no signal to the speaker or other listeners about who had already made a selection.

Once all listeners had selected (or a 3-minute timer ran out), participants were given feedback (Figure 1B). Listeners learned whether they individually had chosen correctly or not; listeners who were incorrect were not told the correct answer. The speaker saw which tangram each listener had selected, but listeners did not. This feedback regime is different from Hawkins et al. (2020) where listeners were shown what the right answer was during feedback. We made this change to prevent listeners from learning conventions purely as a memorized mapping between utterance and correct answer.

Listeners got 4 points for each correct answer; the speaker got points equal to the average of the listeners’ points. These points translated into performance bonus at the end of the experiment.

In each block, each of the 12 tangrams was indicated to the speaker once. The same person was the speaker for an entire block, but participants rotated roles between blocks. Thus, over the course of the 6 blocks, participants were speakers 3 times in 2-player games, twice in 3-player games, once or twice in 4 and 5-player games, and once in 6-player games. Rotating the speaker was chosen to keep participants more equally engaged (the speaker role is more work), and to give a more robust test for reduction and convention.

After the game finished, participants were given a survey asking for optional demographic information and feedback on their experience with the game.

Data pre-processing and exclusions

Participants could use the chat box freely, which meant that the chat transcript contained some non-referential language. The first author skimmed through the chat transcripts, tagging utterances that did not refer to the current tangram. These were primarily pleasantries (“Hello”), meta-

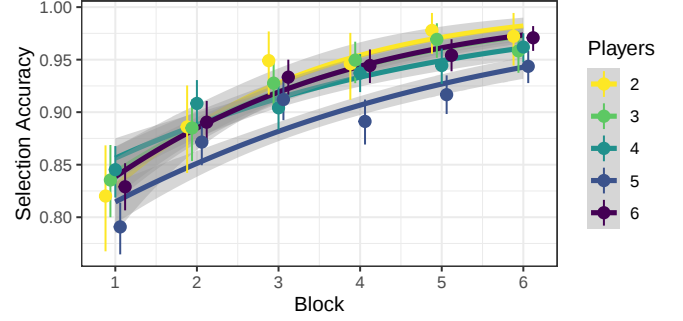


Figure 2: Players’ accuracy at correctly selecting the target figure by block and group size. Accuracy increases across blocks.

commentary about how well or fast the task was going, and confirmations or denials (“ok”, “got it”, “yes”, “no”). We exclude these utterances from our analyses. Note that chat lines sometimes included non-referential words in addition to words referring to the tangrams (“ok, so it looks like a zombie”, “yes, the one with legs”); these lines were retained intact.

Our intended sample size was 20 complete games in each group size, but we ended up with fewer due to games not filling or participants disconnecting early (Table 1). We excluded incomplete blocks from analyses, but included complete blocks from partial games.

Results

Accuracy and Speed

Our first question was whether accuracy and speed increased across groups of different sizes.

Accuracy is high and increasing. Most individuals were accurate in their selections, with accuracy rising across blocks (Figure 2). In a logistic model of accuracy⁴, participants are more accurate in later blocks (block: Est=0.38, CrI=[0.25, 0.5]), and there was no strong effect of group size on accuracy (numPlayers: Est=-0.02, CrI=[-0.08, 0.03]) or interaction between block and group size (block:numPlayers: Est=-0.01, CrI=[-0.04, 0.01]).

Participants speed up in later blocks. Participants selected images faster in later blocks (Figure 3), although there was wide variability. In a linear model of selection time⁵, participants got faster across blocks (block: Est=-10.03, CrI=[-11.03, -9.03]) and were slightly slower in larger games (numPlayers: Est=1.03, CrI=[0.4, 1.66]). This speed up is consistent with prior work by Weber & Camerer (2003) which used speed as the dependent measure. Wide variability in selection time meant that especially for larger groups, there was a wide spread in how long it took groups to complete the experiment.

⁴correct.num~ block × numPlayers This and all subsequent regression models were run in brms with weakly regularizing priors.

⁵time~ block × numPlayers

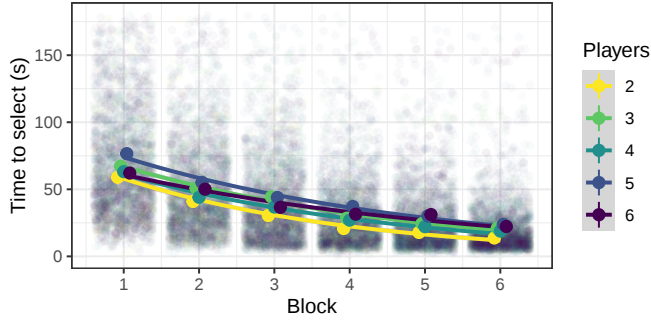


Figure 3: How long listeners took to select a figure in seconds by block and group size. Listeners selected images faster in later blocks. Only times for correct responses are shown.

Reduction

Our second question was whether speakers reduce their referring expressions in larger groups.

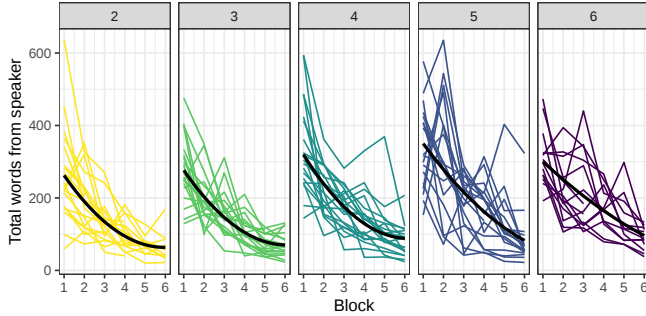


Figure 4: Number of words from speaker (total, across all 12 figures) in a block. Each colored line is one group, the overall trend is shown in black. Across group size, the number of words decreases as conventions emerge, but convention formation is not a smooth process, and there is variability between speakers.

Speakers’ utterances reduce in length. As shown in Figure 4, the number of words produced by speakers decreases over the course of rounds, both in aggregate and for many individual groups. Nonetheless, in some groups, a later speaker may be more verbose than an earlier speaker. Speakers make longer utterances in early blocks that reduce to shorter utterances in later blocks. From a linear model⁶, the effect of being one block later is -3.35 (CrI=[-4.58, -2.13]) words.

Listeners rarely talk. Listeners often don’t talk much, but are more likely to ask questions or make clarification in early blocks. In a linear regression for the number of words each listener said,⁷ there was an effect of block (block: Est=-0.48, CrI=[-0.79, -0.18]), but no clear effect of game size (numPlayers: Est=0.2, CrI=[-0.13, 0.51]).

⁶words_i ~ block × numPlayers + (block|tangram) + (1|playerId) + (1|tangram_group) + (block|gameId)

⁷words_i ~ block × numPlayers + (block|tangram) + (1|playerId) + (1|tangram_group) + (block|gameId)

Effects of group size on conventions

Our third question was whether smaller groups would use fewer words or reduce faster than larger groups.

Larger groups say more. The overall effect of having more players in a group is 1.67 (CrI=[0.68, 2.71]) words from the speaker per trial per additional player. There is no clear interaction between block and group size (block:numPlayers: Est=-0.1, CrI=[-0.39, 0.18]). Larger groups saying more is consistent with predictions from audience design that with more listeners to accommodate, the speaker may use multiple conceptualizations, either initially as a hedge or in response to listener clarifications.

Speaker experience does not fully explain group size effects. One potential concern is that group size correlates with whether the speaker has had the speaker role before (smaller groups repeat speakers more). To address this confound, we coded for whether the speaker has been speaker in an earlier block⁸. Repeat speakers do use fewer words (speaker.repeat: Est=-8.55, CrI=[-10.41, -6.79]), but there are still effects of group size (numPlayers: Est=1.63, CrI=[0.58, 2.66]) and block (block: Est=-5.26, CrI=[-6.84, -3.69]). The effects of block and repeat speaker are subadditive (block:speaker.repeat: Est=3.2, CrI=[2.65, 3.78]), and there is minimal interaction between block and group size (block:numPlayers: Est=0.08, CrI=[-0.22, 0.41]).

Development of conventions

Our final question was how conventions emerge in larger groups.

Speakers who don’t know the convention reduce less. In our games (which had limited feedback), listeners who got a tangram wrong didn’t have a way of knowing what the right answer was unless they asked for clarification in the chat. If a speaker got a tangram wrong as a listener in the previous block, they may not have known the conventional description that went with it, and thus were unlikely to follow the convention. If we assume that reduction is a sign of convention development, then speakers should say more words when they got the tangram wrong the previous block. We added prior errors as an additional predictor to our regression predicting number of words and found that speakers said more words for tangrams after they were incorrect (numPlayers: Est=2.18, CrI=[0.93, 3.44]).

Smaller groups reduce and stabilize conventions sooner. Another angle to look at conventions is to take the speaker’s utterances in the last block as the “convention”, and look at how far back they started. We took the contentful words said by the speaker in the last block and looked at how many of them were used to describe that tangram in prior

⁸words_i ~ block × numPlayers + block × speaker.repeat + (block|tangram) + (1|playerId) + (1|tangram_group) + (block|gameId)

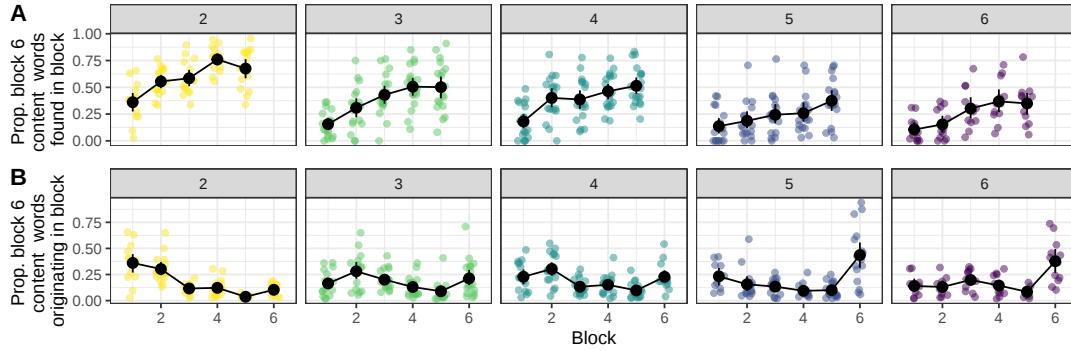


Figure 5: Comparison of the content words in the final (6th) block to words in previous blocks. (A) The proportion 6th block content words found in each prior block. (B) The proportion 6th block content words that were first used by the speaker in each block. Smaller games have higher overlap and a greater proportion of words originating earlier.

rounds.⁹ Then we calculated the proportion overlap between the last block utterance and earlier blocks, shown in Figure 5A. In a linear model of the overlap between an earlier block and the last block¹⁰, later blocks have more overlap (block: Est=0.1, CrI=[0.08, 0.12]). Blocks with the same speaker as the last round have more overlap (same_speaker: Est=0.08, CrI=[0.05, 0.1]); this pattern is visible in the peaks for blocks 2 and 4 in the 2 player games. Larger groups have less overlap between blocks (numPlayers: Est=-0.07, CrI=[-0.09, -0.04]), but there is no interaction between blocks and game size (block:numPlayers: Est=0, CrI=[-0.01, 0]). One potential confound is that in smaller games, players spend more time in the speaker role; however, there is still more overlap in smaller games even to blocks with a different speaker.

Another way to measure conventions is to look at when these words were first introduced by the speaker (Figure 5B). A greater fraction of 6th block words were new in 5 or 6 player games compared with smaller games, whereas most words used in 2-player games originated in the 1st or 2nd blocks. Overall, conventions reduced and stabilized sooner in smaller groups, perhaps because fewer people need to implicitly agree on them.

Groups varied in their strategies and reduction. While most groups did form conventions for most tangrams, it's illustrative to look at a case where a group did not. Table 2 shows the transcript of a 4-person group for a specific figure where they described it geometrically every round, leading to long and not very informative descriptions. Nearly all the figures have diamond heads, so this isn't a distinguishing feature, yet it is described. This illustrates the variability between groups, but also why conventions might be useful.

Table 2: Excerpt from a group that did not reduce very much. The speaker for each round is marked with (S). Figure under discussion is row 3, column 3 in Figure 1A.

Block	Person	Text
1	A(S)	Diamond on top. Body with no real arms or legs. The body is shaped like a boot with the diamond on top.
2	C	Is the boot pointed left or right?
2	B(S)	diamond on top, large body beneath it. Left is a straight line all the way down, small variations on the right to the main body
3	C(S)	Diamond in center on top. Left side straight, right side carved out like a vase.
4	D(S)	Diamond head, flat topped body, straight on the left side with two triangles pointing out on the left
5	D(S)	*on the right
5	A(S)	Diamond on top. Left side is straight, right side is obstructed, looks like a boot
5	B	what do you mean by obstructed?
5	A(S)	The left side of the body is right, right side has bents in it
6	B(S)	Diamond on top of a long large body/rectangle. Left side is complete, right side has bits missing

⁹Contentful words were defined as all words in a referential message with a part of speech identified by the Spacy (<http://spacy.io>) tagger as being a noun, verb, or adjective; that were not on the Spacy stop word list; and that did not have a lemma in the set ['look', 'like', 'body', 'person', 'man', 'guy'], a list generated as being extremely common vocabulary across tangrams.

¹⁰overlap ~ block × numPlayers + same_speaker + (1|gameId) + (1|target)

A different 4-person group had a member who during the first block shared the idea that the task would be easier if they explicitly gave “codenames” to the figures. The transcript for this group and one of the tangrams is shown in Table 3. Of note, multiple speakers forget the assigned codename, demonstrating that meta-knowledge doesn't always help. This group also describes the figure in relation to another already-named figured. Nonetheless, the group success-

Table 3: Excerpt from a group that explicitly gave nicknames to the figures. The speaker for each round is marked with (S). Tangram under discussion is row 1, column 4 in Figure 1A.

Block	Person	Text
1	A(S)	[...] yes, the legs are like a zig zag
	C	CODE name ZIGZAG
2	A(S)	There are no legs upwards
	B(S)	okay so similar to begger guy but no foot pointing up
	B(S)	its like a zigzag
	B(S)	i forgot the code name
	D	zigzag yea
	A	The one standing with knees bent
	B(S)	yeah
	B(S)	standing
	C	Yeah zigzag
	C(S)	The begger with no foot coming out from the left
3	B	zigzag
	C(S)	zigzag it is
	C(S)	sorry i forgot
4	D(S)	zigzag
5	A(S)	zigzag
6	B(S)	beggar guy
	B(S)	zigzag

fully conventionalizes on a couple reduced names for this figure: “zigzag” and “beggar”. This dual-naming of figures from multiple conceptual angles contributed by different speakers also occurs in other games.

Discussion

The emergence of conventions has been a key case study for communication more broadly. Yet this issue has – for the most part – been studied only in dyadic communication. While some studies have examined aspects of convention formation in larger groups (e.g., Yoon & Brown-Schmidt, 2014; Yoon & Brown-Schmidt, 2019), basic descriptive work has not yet investigated how group size changes the dynamics of interaction in a standard referential communication task, in part because such tasks can be difficult to administer to larger groups. Taking advantage of a new online multi-player experiment platform, we ran repeated reference games with groups of 2–6 players and characterized the nature of group performance.

Consistent with dyadic games, listeners’ selection accuracy increased over blocks at the same time as listeners sped up their selections (question 1). Crucially, speakers reduced the length of their descriptive utterances as they conventionalized on concepts for each image (question 2). Because speakers rotated, this reduction finding is robust: not only did speakers say less in later repetitions than they themselves said earlier, speakers later in the order said less than speakers earlier in

the rotation. This reduction varied with group size; smaller groups used shorter utterances, but group size did not significantly interact with block (question 3). The trajectory of reduction also depended on whether the current speaker correctly identified the tangram in the prior block and whether the current speaker was new to being speaker. This pattern is consistent with both the ‘aim low’ and ‘aim middle’ hypotheses from previous work (Yoon & Brown-Schmidt, 2014; Yoon & Brown-Schmidt, 2019).

What was specifically different across group sizes? Smaller groups showed more agreement in how each tangram was identified across blocks (question 4), coming to consensus earlier: Their overlap between descriptions in the first 5 blocks to the final block was higher, and words in the final block tended to originate earlier. The greater diversity in how tangrams were described in larger groups could be explained by slower convergence to a convention or parallel competing conceptualizations favored by different speakers. Larger groups have more people for the speaker to communicate to, but also more people who might interrupt with questions, and more people who have opinions about what each image looks like. Bigger groups differ from smaller groups in a number of ways, however, and disentangling these differences is an area for future work.

Group interactions are rich, and this experiment is necessarily a schematic simplification with a number of limitations. Real-life situations vary widely in who the interlocutors are, their relationships, their goals, and their environment (Carletta, Garrod, & Fraser-Krauss, 1998; Fay, Garrod, & Carletta, 2000). Our participants were a convenience sample of Prolific workers who were strangers to each other; thus we miss richness that could come from prior relationships or shared community. Reference is only one goal out of many possible communicative goals, and the tangram images are artificial. We provided less feedback than previous studies such as Hawkins et al. (2020); this regime imitates situations where interlocutors can’t show each other examples, but it’s not representative of all communicative environments. Further, our text-based online paradigm meant that participants’ individual identities were not especially salient. In sum, communication takes place in a plethora of situations; our experiment provides some insights, but also misses many complexities that should be a focus of further experiments.

The experimental paradigm presented here could be a valuable tool to disentangle the mechanisms of group size and determine which design parameters are relevant to reduction. Luckily, with an online implementation, recruiting for and running experiments is feasible, and thus it will be possible to iterate on this experiment to determine how far the patterns generalize. While much is left to be explored, this initial data set provides a rich corpus of how humans adapt language dynamically to communicate.

References

Almaatouq, A., Becker, J., Houghton, J. P., Paton, N.,

- Watts, D. J., & Whiting, M. E. (2020). Empirica: A virtual lab for high-throughput macro-level experiments. *arXiv:2006.11398 [Cs]*. Retrieved from <http://arxiv.org/abs/2006.11398>
- Branigan, H. (2006). Perspectives on multi-party dialogue. *Research on Language and Computation*, 4(2), 153–177.
- Carletta, J., Garrod, S., & Fraser-Krauss, H. (1998). Placement of Authority and Communication Patterns in Workplace Groups: The Consequences for Innovation. *Small Group Research*, 29(5), 531–559. <http://doi.org/10.1177/1046496498295001>
- Clark, H. H., & Wilkes-Gibbs, D. (1986). Referring as a collaborative process. *Cognition*.
- Fay, N., Garrod, S., & Carletta, J. (2000). Group Discussion as Interactive Dialogue or as Serial Monologue: The Influence of Group Size. *Psychol Sci*, 11(6), 481–486. <http://doi.org/10.1111/1467-9280.00292>
- Fox Tree, J. E., & Clark, N. B. (2013). Communicative Effectiveness of Written Versus Spoken Feedback. *Discourse Processes*, 50(5), 339–359. <http://doi.org/10.1080/0163853X.2013.797241>
- Ginzburg, J., & Fernandez, R. (2005). Action at a distance: The difference between dialogue and multilogue. *Proceedings of DIALOR*, 9.
- Haber, J., Baumgärtner, T., Takmaz, E., Gelderloos, L., Bruni, E., & Fernández, R. (2019). The PhotoBook Dataset: Building Common Ground through Visually-Grounded Dialogue. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics* (pp. 1895–1910). Florence, Italy: Association for Computational Linguistics. <http://doi.org/10.18653/v1/P19-1184>
- Hawkins, R. D., Frank, M. C., & Goodman, N. D. (2020). Characterizing the dynamics of learning in repeated reference games. *arXiv:1912.07199 [Cs]*. Retrieved from <http://arxiv.org/abs/1912.07199>
- Horton, W. S., & Gerrig, R. J. (2002). Speakers' experiences and audience design: Knowing when and knowing how to adjust utterances to addressees. *Journal of Memory and Language*, 18.
- Horton, W. S., & Gerrig, R. J. (2005). The impact of memory demands on audience design during language production. *Cognition*, 96(2), 127–142. <http://doi.org/10.1016/j.cognition.2004.07.001>
- Krauss, R. M., & Weinheimer, S. (1964). Changes in reference phrases as a function of frequency of usage in social interaction: A preliminary study. *Psychon Sci*, 1(1-12), 113–114. <http://doi.org/10.3758/BF03342817>
- Krauss, R. M., & Weinheimer, S. (1966). Concurrent feedback, confirmation, and the encoding of referents in verbal communication. *Journal of Personality and Social Psychology*, 4(3), 343–346. <http://doi.org/10.1037/h0023705>
- Metzing, C., & Brennan, S. E. (2003). When conceptual pacts are broken: Partner-specific effects on the comprehension of referring expressions. *Journal of Memory and Language*, 49(2), 201–213. [http://doi.org/10.1016/S0749-596X\(03\)00028-7](http://doi.org/10.1016/S0749-596X(03)00028-7)
- Schober, M. F., & Clark, H. H. (1989). Understanding by addressees and overhearers. *Cognitive Psychology*, 21(2), 211–232. [http://doi.org/10.1016/0010-0285\(89\)90008-X](http://doi.org/10.1016/0010-0285(89)90008-X)
- Tolins, J., & Fox Tree, J. E. (2016). Overhearers Use Addressee Backchannels in Dialog Comprehension. *Cogn Sci*, 40(6), 1412–1434. <http://doi.org/10.1111/cogs.12278>
- Traum, D. (2004). Issues in Multiparty Dialogues. In G. Goos, J. Hartmanis, J. van Leeuwen, & F. Dignum (Eds.), *Advances in Agent Communication* (Vol. 2922, pp. 201–211). Berlin, Heidelberg: Springer Berlin Heidelberg. http://doi.org/10.1007/978-3-540-24608-4_12
- Weber, R. A., & Camerer, C. F. (2003). Cultural Conflict and Merger Failure: An Experimental Approach. *Management Science*, 49(4), 16.
- Yoon, S. O., & Brown-Schmidt, S. (2014). Adjusting conceptual pacts in three-party conversation. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 40(4), 919–937. <http://doi.org/10.1037/a0036161>
- Yoon, S. O., & Brown-Schmidt, S. (2018). Aim Low: Mechanisms of Audience Design in Multiparty Conversation. *Discourse Processes*, 55(7), 566–592. <http://doi.org/10.1080/0163853X.2017.1286225>
- Yoon, S. O., & Brown-Schmidt, S. (2019). Audience Design in Multiparty Conversation. *Cognitive Science*, 43(8), e12774. <http://doi.org/10.1111/cogs.12774>