



A Practical Guide to Conversation Research: How to Study What People Say to Each Other

Citation

Yeomans, Michael, Katelynn Boland, Hanne K. Collins, Nicole Abi-Esber, and Alison Wood Brooks. "A Practical Guide to Conversation Research: How to Study What People Say to Each Other." *Advances in Methods and Practices in Psychological Science* 6, no. 4 (October–December 2023).

Published version

<https://doi.org/10.1177/25152459231183919>

Link

<https://nrs.harvard.edu/URN-3:HUL.INSTREPOS:37377717>

Terms of use

This article was downloaded from Harvard University's DASH repository, and is made available under the terms and conditions applicable to Other Posted Material (LAA), as set forth at

<https://harvardwiki.atlassian.net/wiki/external/NGY5NDE4ZjgzNTc5NDQzMGIzZWZhMGFIOWI2M2EwYTg>

Accessibility

<https://accessibility.huit.harvard.edu/digital-accessibility-policy>

Share Your Story

The Harvard community has made this article openly available.
Please share how this access benefits you. [Submit a story](#)

A Practical Guide to Conversation Research: How to Study What People Say to Each Other

Michael Yeomans¹, F. Katelynn Boland², Hanne K. Collins³,
Nicole Abi-Esber⁴, and Alison Wood Brooks³

¹Business School, Imperial College London, London, England; ²Columbia Business School, Columbia University, New York, New York; ³Harvard Business School, Harvard University, Cambridge, Massachusetts; and

⁴London School of Economics, London, England

Advances in Methods and
Practices in Psychological Science
October-December 2023, Vol. 6, No. 4,
pp. 1–38
© The Author(s) 2023
Article reuse guidelines:
sagepub.com/journals-permissions
DOI: 10.1177/25152459231183919
www.psychologicalscience.org/AMPPS



Abstract

Conversation—a verbal interaction between two or more people—is a complex, pervasive, and consequential human behavior. Conversations have been studied across many academic disciplines. However, advances in recording and analysis techniques over the last decade have allowed researchers to more directly and precisely examine conversations in natural contexts and at a larger scale than ever before, and these advances open new paths to understand humanity and the social world. Existing reviews of text analysis and conversation research have focused on text generated by a single author (e.g., product reviews, news articles, and public speeches) and thus leave open questions about the unique challenges presented by interactive conversation data (i.e., dialogue). In this article, we suggest approaches to overcome common challenges in the workflow of conversation science, including recording and transcribing conversations, structuring data (to merge turn-level and speaker-level data sets), extracting and aggregating linguistic features, estimating effects, and sharing data. This practical guide is meant to shed light on current best practices and empower more researchers to study conversations more directly—to expand the community of conversation scholars and contribute to a greater cumulative scientific understanding of the social world.

Keywords

natural language processing, text analysis, conversation, social interaction, open science

Received 5/31/22; Revision accepted 6/6/23

Conversation is one of the most pervasive of all human behaviors—people talk to each other all the time, all over the world (Dunbar et al., 1997). Most interpersonal relationships develop through a series of conversations over time—time spent talking and not talking, together and apart. Although a frequent and familiar task, each conversation is complex—it requires (and enables) people to coordinate their behavior and beliefs about the world (Clark et al., 2011; Jaques et al., 2019; Misyak et al., 2014; Rossignac-Milon et al., 2021). Conversations are consequential, allowing people to pursue a wide array of informational and relational goals (Yeomans et al., 2022) in the short term and over the long term—spanning each individual conversation and longer-term relationships (Fitzsimons & Finkel, 2018). Indeed, the

amount and quality of social interaction is one of the most enduring predictors of human well-being (H. K. Collins et al., 2022; Diener & Seligman, 2002; Epley & Schroeder, 2014; Mehl et al., 2010; Quoidbach et al., 2019; Sun, Harris, & Vazire, 2020).

It is no surprise that researchers are increasingly interested in studying conversations, the contextual factors that surround them, and the short- and long-term effects of having them. This practical guide argues for the relevance of this work now, the benefits and challenges

Corresponding Author:

Michael Yeomans, Business School, Imperial College London, London, England
Email: m.yeomans@imperial.ac.uk



Creative Commons NonCommercial CC BY-NC: This article is distributed under the terms of the Creative Commons Attribution-NonCommercial 4.0 License (<https://creativecommons.org/licenses/by-nc/4.0/>), which permits noncommercial use, reproduction, and distribution of the work without further permission provided the original work is attributed as specified on the SAGE and Open Access pages (<https://us.sagepub.com/en-us/nam/open-access-at-sage>).

researchers should expect from studying conversations, and how to analyze conversation data, pair transcripts with surveys, and share results as the field moves toward a cumulative science of conversation (see Figs. 1 and 2).

Why Now?

At least three developments have enabled a recent boom in conversation research. First, conversations have become increasingly mediated through technology as a consequence of the Digital Revolution and Information Age of the 20th century and the social media era of the 21st century (Rainie & Wellman, 2012), shifts that were accelerated during the COVID-19 pandemic (M. Nguyen et al., 2020). These mediated communication technologies allow for the recording of text, audio, and/or video and thus preserve a rich source of conversation data for analysis. Second, there have been many advances in natural language processing (NLP)—an interdisciplinary subfield at the intersection of linguistics, computer science, and artificial intelligence that seeks to learn, parse, and understand human-language content using quantitative techniques (Hirschberg & Manning, 2015). This field develops computational tools that turn raw conversations into behavioral data—words into numbers—especially at scale (Jurafsky & Martin, 2017). Finally, the value of larger-scale analyses has been underscored by the recent revolution in research practices (Nelson et al., 2018). Taken together, these cultural and methodological developments offer wide promise for the study of conversation across a variety of academic disciplines.

Measuring Conversation

Although conversations are common and consequential, they are also complicated—no two are identical. Researchers have dealt with the complexity of conversation with a wide range of approaches aimed to simplify and isolate different aspects of a conversation. In exchange for simplicity, these approaches can make conversations less natural and more abstract. For example, researchers often study dialogue indirectly by having participants talk to a trained confederate, respond to hypothetical vignettes, make evaluations of carefully selected transcript segments, recall a previous conversation from memory, or offer holistic evaluations of a conversation after it is over. These approaches constitute creative and generative ways to study conversations and were particularly useful when conversation technology was nascent.

These approaches allow researchers to simplify and study conversations, but they also suffer from several well-known biases. For instance, confederate simulations rely on faithful execution of researchers' instructions; hypothetical and recall methods suffer from errors in

forecasting and memory; self-report measures suffer from social-desirability bias, hindsight bias, and demand effects; and experimenter-generated stimuli remove the conversational context in which they would occur in the real world. Conversation is a complex, contextual, and improvisational environment, and these kinds of simplifications can result in a misunderstanding between the assumed, perceived, and actual goals and psychological experiences of the speakers (Stokoe, 2021).

On the other hand, many researchers have taken on the daunting task of studying natural, contextualized conversational behavior, beginning with study of “ordinary language” as early as the mid-20th century (e.g., Garfinkel, 1956; Goffman, 1981; Heritage, 2008; Pomerantz, 1990; Sacks et al., 1974; Schegloff, 1968; Schegloff & Sacks, 1973; Stivers et al., 2010; Stivers & Sidnell, 2012). This work has typically prioritized attention to descriptive detail in natural settings by scrutinizing isolated portions of transcripts at the expense of scalability and controlled measures of outcomes and effects. Furthermore, linguistic inquiry often assumes rationality on the part of speakers (e.g., Goodman & Frank, 2016; Grice, 1975; Misyak et al., 2014) and infers intent on the basis of outcomes. This assumption can be limiting. People constantly deviate from rational behavior (Kahneman, 2002), so it is important to measure both intentions and outcomes to see whether speakers are making wise choices, enacting good behaviors, or making mistakes (Yeomans et al., 2022).

More recent work has conceptualized conversation as a diagnostic window into variables such as health status, personality, and well-being (e.g., G. Collins et al., 2018; H. K. Collins et al., 2022; Conner & Mehl, 2015; de Barbaro, 2019; Jaidka et al., 2020; Mehl & Pennebaker, 2003; Robbins et al., 2011). This simplification abstracts away from the details of each particular conversation and focuses instead on person-level variables. The focus is on what speakers' behavior says about themselves, rather than the effects of their behavior on their partner, and ignores the specific goals and outcomes of individual conversations.

Many prior articles have compared conversation behavior with context and outcome data (e.g., Weingart et al., 2004; Word et al., 1974). But this work usually relies on human annotation to quantify conversational behavior from recordings or transcripts. Although these insights are useful, they are costly to scale and often do not give a transparent or interpretable definition of how a measure is calculated (see below for more on this point). Likewise, speakers are often asked to quantify the content of the conversation themselves using retrospective survey measures. Again, these measures are convenient but opaque and suffer from the same self-report and memory biases of other survey methods.

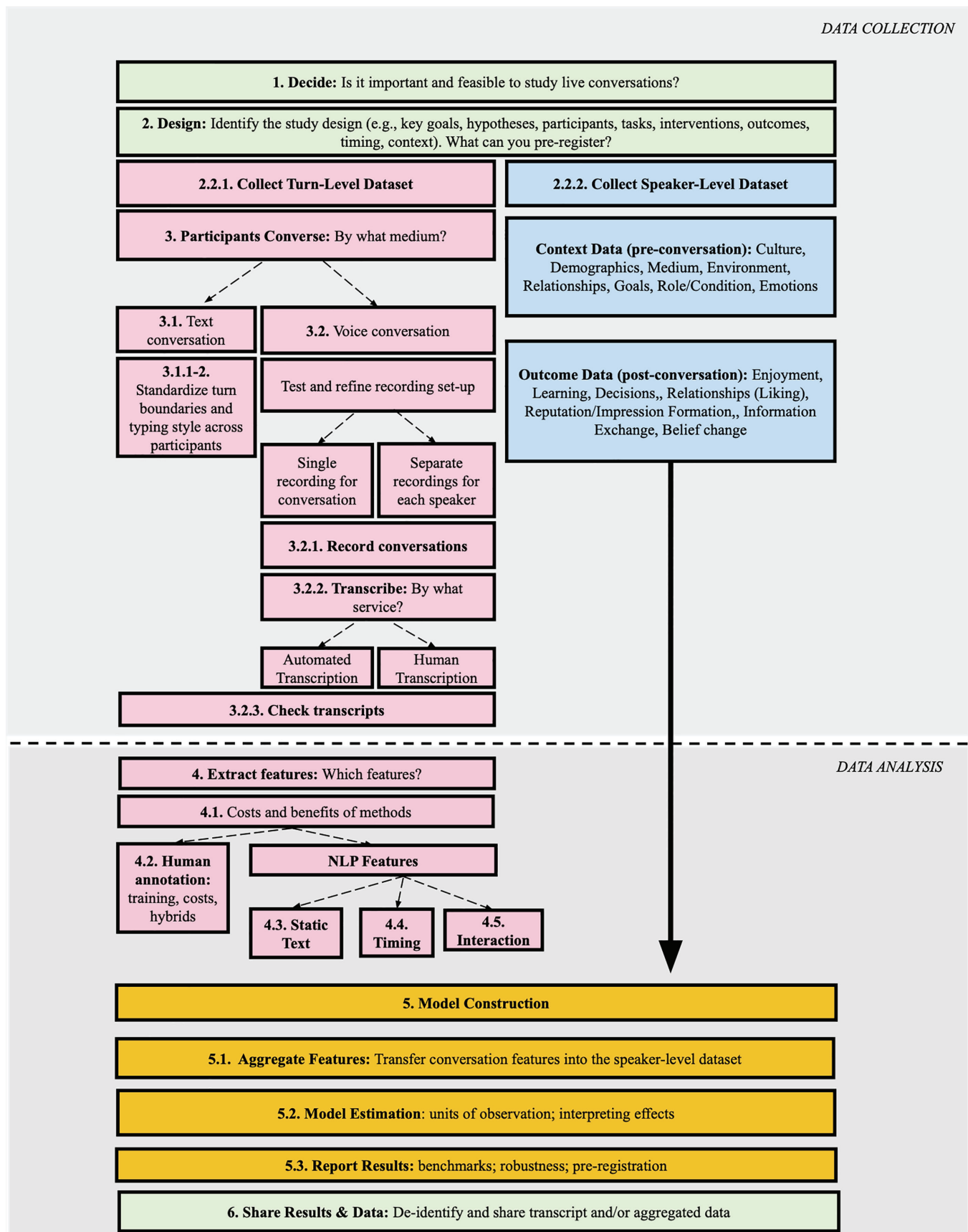


Fig. 1. A workflow for researchers to collect and analyze conversation data.

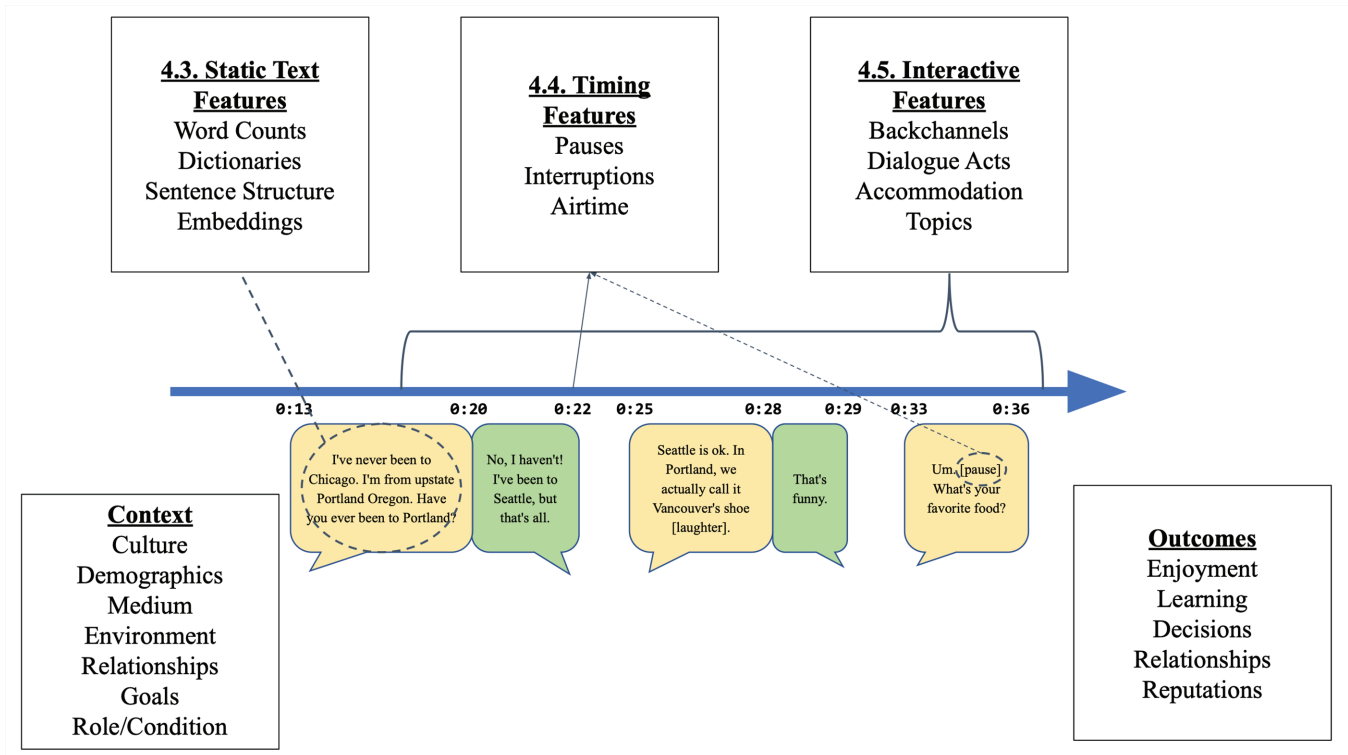


Fig. 2. A map of data sources for conversation research. From the transcript itself, features can be extracted using static text methods and relying on time stamps and interactivity. These conversation features are then compared with preconversation context variables and post-conversation outcomes.

In this article, we highlight how recent technological advances provide researchers with novel capabilities to combine the best aspects of these research approaches and directly measure conversation behavior in more natural contexts at scale. The tools for conversation science are rapidly improving—both for recording conversations and for analyzing them—leading to an emerging boom of conversation research in a wide range of contexts across a wide range of academic disciplines (for a review, see Table 1). Modern workflows have made it easier than ever for researchers to combine detailed transcript analysis with algorithmic tools to scale up their insights and obtain robust measures of context and outcome variables surrounding conversational choices.

In this practical guide, we aim to make data, tools, and methods more accessible to a wider group of researchers by describing common challenges that face behavioral researchers who wish to study conversations and suggesting approaches that address those challenges. This review is aimed at researchers across disciplines who are looking to incorporate conversation-research methodology into their work for the first time or expand on a body of conversation research by incorporating new methods and techniques.

The Scope of This Article: A Focus on Transcripts

Conversations can include a wide array of psychological and behavioral content, including verbal features—what words are uttered, by whom, and in what order—and nonverbal features—tone of voice, gesture, posture, facial expressions, and so on (via visual and/or audio inputs). We focus primarily on verbal content for three key reasons. First, every conversation includes verbal content, whereas nonverbal cues are not present in many conversations (e.g., emails and phone calls). Second, verbal content presents common challenges for conversation research no matter what other types of cues are also present. The decisions, beliefs, and consequences that stem from the verbal content of conversation are only beginning to be rigorously understood. Finally, although nonverbals can inflect the meaning of the words spoken, it is the words themselves that form most of the meaning: They define the topics of conversation and what is being said about them. Indeed, verbal content has an overwhelming effect on how nonverbals are interpreted (Lapacko, 1997).

For these reasons, the scope of this article focuses on the aspects of conversations that can be captured in a

Table 1. A Nonexhaustive List of Recent Research That Analyzes Transcript Data Across Behavioral Domains, Conducted Across Academic Disciplines

Behavioral domain	Article title	Application	Citation
Negotiations	“Communicating With Warmth in Distributive Negotiations Is Surprisingly Counterproductive”	The authors trained a natural-language-processing algorithm to quantify the difference between how people enact warm and friendly versus tough and firm communication styles in a distributive negotiation.	Jeong et al., 2019
	“Communication and Bargaining Breakdown: An Empirical Analysis”	The authors used text analysis to show that repeat players learn how to use communication in bargaining and that the messaging strategies of experienced sellers are correlated with successful bargaining.	Backus et al., 2020
	“Setting the Stage for Negotiations: How Superordinate Goal Dialogues Promote Trust and Joint Gain in Negotiations Between Teams”	The authors used structured dialogues to identify the boundary conditions in negotiations that shape when superordinate goal dialogues are most likely to increase joint gain and when they will not be effective.	Swaab et al., 2021
Work emails	“Social Networks Under Stress”	The authors analyzed instant messages among the decision-makers in a large hedge fund and their network of outside contacts to investigate the link between price shocks, network structure, and change in the affect and cognition of decision-makers in the network.	Romero et al., 2016
	“Alignment at Work: Using Language to Distinguish the Internalization and Self-Regulation Components of Cultural Fit in Organizations”	The authors developed a measure of cultural fit based on linguistic alignment and used this measure to find that patterns of alignment in the first 6 months of employment are predictive of individuals’ downstream outcomes, especially involuntary exit.	Doyle et al., 2017
Work meetings	“Virtual Communication Curbs Creative Idea Generation”	The authors randomly assigned work teams to conduct team meetings in person or on Zoom and studied how that affected idea generation and decision quality.	Brucks & Levav, 2022
Interviews	“Tie-Breaker: Using Language Models to Quantify Gender Bias in Sports Journalism”	The authors proposed a language-model-based approach to quantify differences in questions posed to female vs. male athletes and applied it to tennis postmatch interviews.	Fu et al., 2016
Entrepreneurial pitches	“Pitching a Business Idea to Investors: How New Venture Founders Use Micro-Level Rhetoric to Achieve Narrative Plausibility and Resonance”	The authors analyzed micro-level arguments underpinning pitch narratives of entrepreneurs who joined a business incubator and discerned four rhetorical strategies that these entrepreneurs used to achieve narrative plausibility and resonance.	van Werven et al., 2019
	“Actions Speak Louder Than Words: How Figurative Language and Gesturing in Entrepreneurial Pitches Influences Investment Judgments”	The authors identified distinct pitching strategies entrepreneurs use involving different combinations of verbal tactics and gesture and examined the impact of these strategies on investors’ propensity to invest.	Clarke et al., 2019
Quarterly earnings calls	“Manager-Analyst Conversations in Earnings Conference Calls”	The authors conducted sentiment analysis to look at how well the questions asked (and their associated answers) predict changes in stock prices following quarterly earnings calls by publicly traded companies.	Chen et al., 2018
	“Disclosure Sentiment: Machine Learning vs. Dictionary Methods”	The authors found that machine-learning methods are better at detecting disclosure sentiment than dictionary methods in 10,000 filings and earnings calls.	Frankel et al., 2022

(continued)

Table 1. (continued)

Behavioral domain	Article title	Application	Citation
Medical conversations	“Miscommunication in Doctor–Patient Communication”	The authors used conversation analysis to explore the effectiveness of medical treatment and shared understanding between patient and clinician in the context of psychiatric consultations.	McCabe & Healey, 2018
	“Naturalistically Observed Sighing and Depression in Rheumatoid Arthritis Patients: A Preliminary Study.”	This study tested the degree to which naturalistically observed sighing in daily life is a behavioral indicator of depression and reported physical symptoms in patients with rheumatoid arthritis.	Robbins et al., 2011
Police investigations	“I’m Not Gonna Hit A Lady’: Conversation Analysis, Membership Categorization and Men’s Denials of Violence Towards Women”	The authors used British police interrogation materials and conversation analysis to shed light on the location and design of and responses to suspects’ “category-based denials” that they are not “the kind of men who hit women.”	Stokoe, 2010
	“Language From Police Body Camera Footage Shows Racial Disparities in Officer Respect”	The authors presented a systematic analysis of officer body-worn camera footage using computational linguistic techniques to automatically measure the respect level that officers display to community members.	Voigt et al., 2017
Courtrooms	“On Racial Diversity and Group Decision Making: Identifying Multiple Effects of Racial Composition on Jury Deliberations”	The authors examined the effects of racial diversity on group decision-making and extended previous findings that racial issues, in the form of jury selection questions, increase leniency toward a Black defendant on trial.	Sommers, 2006
	“Echoes of Power: Language Effects and Power Differences in Social Interaction”	The authors proposed an analysis framework based on linguistic coordination that they then use to study how conversational behavior can reveal power relationships in discussions among Wikipedians and arguments before the United States Supreme Court.	Danescu-Niculescu-Mizil et al., 2012
	“Justice, Interrupted: The Effect of Gender, Ideology, and Seniority at Supreme Court Oral Arguments”	The authors studied how the Justices of the United States Supreme Court compete to have influence at oral argument by examining the extent to which the Justices interrupt each other and how advocates interrupt the Justices, contrary to the rules of the Court.	Jacobi & Schweers, 2017
Central bank meetings	“Transparency and Deliberation Within the FOMC: A Computational Linguistics Approach”	The authors used computational linguistics algorithms to explore the effect of transparency on monetary policymakers’ deliberations.	Hansen et al., 2018
Voter turnout drives	“Unacquainted Callers Can Predict Which Citizens Will Vote Over and Above Citizens’ Stated Self-Predictions”	The authors used conversation analysis to find that strangers can use nonverbal signals to improve predictions of follow through on self-reported intentions.	Rogers et al., 2016
Game shows	“Malleable Lies: Communication and Cooperation in a High Stakes TV Game Show”	The authors conducted an empirical analysis that showed that statements that carry an element of conditionality or implicitness are associated with a lower likelihood of cooperation and confirmed that malleability is a good criterion for judging the credibility of cheap talk.	Turmunkh et al., 2019
Government debates	“Asking Too Much? The Rhetorical Role of Questions in Political Discourse”	The authors used an unsupervised methodology for extracting surface motifs that recur in questions and for grouping them according to their latent rhetorical role.	Zhang et al., 2017

(continued)

Table 1. (continued)

Behavioral domain	Article title	Application	Citation
Online forums	"No Country for Old Members: User Lifecycle and Linguistic Change in Online Communities"	The authors proposed a framework for tracking linguistic change in online communities and for understanding how specific users react to these evolving linguistic norms.	Danescu-Niculescu-Mizil et al., 2013
	"Winning Arguments: Interaction Dynamics and Persuasion Strategies in Good-Faith Online Discussions"	The authors used discussions from an online community on Reddit to study and understand the mechanisms behind persuasion.	Tan et al., 2016
	"Tracking Group Identity Through Natural Language Within Groups"	The authors developed and validated a language-based metric of group identity strength and demonstrated its potential in tracking identity processes over time in Reddit communities.	Ashokkumar & Pennebaker, 2022
Classrooms	"Investigating How Student's Cognitive Behavior in MOOC Discussion Forums Affect Learning Gains"	The authors adopted a content-analysis approach to analyze students' cognitively relevant behaviors in a massive open online course (MOOC) discussion forum and further explored the relationship between the quantity and quality of that participation with their learning gains.	X. Wang et al., 2015
	"The Civic Mission of MOOCs: Engagement Across Political Differences in Online Forums"	The authors collected measures of students' political ideology and observed student behavior in the course discussion boards to find that students hold diverse political beliefs, participate equitably in forum discussions, directly engage with students holding opposing beliefs, and converge on a shared language rather than talking past one another.	Yeomans, Stewart, et al., 2018
Academic seminars	"Women's Visibility in Academic Seminars: Women Ask Fewer Questions Than Men"	The authors quantified women's visibility through the question-asking behavior of academics at seminars using observations and an online survey.	Carter et al., 2018
	"Gender and the Dynamics of Economics Seminars"	The authors collected data on every interaction between presenters and their audience in hundreds of research seminars, summer conferences, and job-market talks across most leading economics departments to find that women presenters are treated differently than their male counterparts.	Dupas et al., 2021
Speed dates	"It's Not You, It's Me: Detecting Flirting and Its Misperception in Speed-Dates"	The authors created a flirtation-detection system that uses paralinguistic, dialogue, and lexical features to detect a speaker's intent to flirt on a speed-date with up to 71.5% accuracy, outperforming both the baseline and the human interlocutors.	Ranganath et al., 2009
	"It Doesn't Hurt to Ask: Question-Asking Increases Liking"	The authors trained a natural-language-processing algorithm as a "follow-up question detector" and applied it to speed-dating data to find that speed daters who ask more follow-up questions during their dates are more likely to elicit agreement for second dates from their partners, a behavioral indicator of liking.	K. Huang et al., 2017
Customer service calls	"Conversational Dynamics: When Does Employee Language Matter?"	The authors investigated the warmth-competence trade-off in customer-service agents. They found that warm language is most common during the beginning and ends of successful calls compared with the middle of those calls.	Y. Li et al., 2022
Door-to-door campaigns	"Durably Reducing Transphobia: A Field Experiment on Door-to-Door Canvassing"	The authors showed that a single 10-min conversation that actively encouraged taking the perspective of others markedly reduces prejudice for at least 3 months.	Broockman & Kalla, 2016

transcript. This includes conversations conducted through sound and through writing.¹ Transcript data primarily include all words and phrases uttered by the speakers, the relative order and timing in which they are produced, and who produced them. In addition, transcript data can include some paralinguistic features (e.g., laughter, backchannel feedback like “yea” or “uh huh,” and disfluencies like “um” or “uhhh”). Likewise, written conversation sometimes includes features intended to represent nonverbal information (e.g., emojis or emoticons).

However, this scope excludes data that are present in many types of conversations. Primarily, this excludes paralinguistic (acoustic) information, such as the tone, pitch, and volume of voice, and visual nonverbal information, such as the speakers’ facial expressions, hand gestures, and body posture. We also focus almost exclusively on monolingual English conversations because the complexities of conversing in two or more languages simultaneously are too manifold for us to address properly herein. Most common NLP tools are available in many languages, although in cases in which researchers are studying dialogue from underresourced languages or other complex sources (e.g., slang, jargon, multiple languages at once), they may want to rely more heavily on expert human annotation.

In cases in which these other sources of information are relevant to the research question, we urge researchers to take a more tailored approach rather than rely only on our simplified workflow. For example, we note that these other types of conversational content can be added or annotated within transcript data (which we address in Capturing Conversation Data below) and can be quite important in some cases.

As a final statement of scope, we have also avoided research on dialogue generation (i.e., building models that can converse autonomously, sometimes called “dialogue agents,” “dialogue systems,” or “chatbots”). Transcripts of conversations that include bots can be analyzed in essentially the same way as conversations that include only humans, but building the chatbots themselves is more arduous. A practical reason to avoid this topic is because it is an especially fast-moving field. For example, between our initial submission of this article and its final acceptance, ChatGPT was released (OpenAI, 2022), followed by a rush of similarly impressive language-generation models. Although the future remains uncertain, we do anticipate that novel and enhanced models will emerge and become available in the coming years and will become increasingly important in the field of conversation analysis.

At this point in time, effective chatbots in the real world tend to be task-specific (e.g., customer-service phone trees, smart-home assistants) or serve narrow

roles, such as a conversation facilitator for human speakers (e.g., Adamson et al., 2014; Traeger et al., 2020). When chatbots participate in broader conversations, they often have problems with listening, consistency, factuality, and other basic skills, although this may improve in the near future (M. Huang et al., 2020).

Leveraging the Predictable Structure of Conversation

Conversation is constructed jointly by (at least) two people, each of whom has their own independent goals, preferences, beliefs, perceptions, traits, and choices, often intertwined in an interdependent relational system (Fitzsimons & Finkel, 2018; Yeomans et al., 2022). In light of the difficult coordination puzzle that conversation presents, it is a wonder that humans manage to communicate at all. Remarkably, people do figure out how to understand each other (Goodman & Frank, 2016; Grice, 1975; Misyak et al., 2014). In fact, the predictable and intuitive structure of conversation—a pattern humans learn to recognize and produce from a very young age—facilitates information flow between speakers. The raw data of conversation are carefully structured by the participants themselves. For example, conversation partners alternate turns, jointly establish topics as a common frame of reference, and ask and answer questions (Pickering & Garrod, 2004; Schegloff, 2007).

However, from a researcher’s perspective, conversational transcript data are difficult to analyze quantitatively and involve many steps (see Fig. 1). First, sounds sometimes need to be converted into words (e.g., “uh-huh” or “[laughter]”). Then, all words need to be arranged in sentences and turns. Once the transcript is generated, researchers will notice how conversation data are high-dimensional—no two conversations are exactly alike. Within one conversation, every possible turn branches into an exponentially large decision tree containing what could be said next in quick, recursive cycles across multiple speakers. Although researchers can take advantage of the predictable aspects of conversational structure, they must also sift through the exponential complexity—they must make many judgment calls to determine which features are counted and how to do so from raw text.

The distinctiveness of dialogue versus single-voice text

For our purposes, conversations consist of dialogue² generated between two or more people over a series of turns. This definition primarily serves to distinguish conversations from documents authored from a single perspective, including speeches, essays, newspaper and magazine articles, books, product reviews, legal documents, and

social media posts. Although the “great bulk” of language use is conversational (Levinson, 2016), single-voice documents have been the dominant source material in applied text analysis, and many review articles in related fields have focused only on single-voice documents (e.g., Benoit, 2020; Berger et al., 2020; Boyd & Schwartz, 2021; Dehghani & Boyd, 2022; Gentzkow et al., 2019; Grimmer & Stewart, 2013; Hansen & Ash, 2023; Hirschberg & Manning, 2015; Jackson et al., 2022; Pennebaker et al., 2003). Many of the techniques developed for single-voiced documents are also useful for studying conversations. However, the differences between single-voiced text and dialogue motivate different ways that researchers should capture and analyze conversation data.

First, unlike single-voiced text, conversations include multiple interchanging contributors. Each person’s contribution to the full conversation must be disambiguated (e.g., Who said what?). Second, conversations are generated on the spot and responsively, which puts a special priority on understanding the sequence of what is said, when it is said, and how it relates to adjacent conversational turns. Third, conversations are usually less thoroughly edited than single-voice documents, often because the turns are spontaneously composed. This means conversation entails looser sentence structure and breakdowns in the coordination of common ground, including more interruptions, cross-talk, silence, repairs, repetitions, misarticulations, clarifications, backchannels, conflicts, slurs, and jargon (Fox Tree, 2010). Lack of editing means conversations tend to have more spelling and grammatical errors and disfluencies (e.g., “umm,” “uh-huh”). Fourth, conversation often covers many topics and goals (Cooney et al., 2020; Yeomans et al., 2022), whereas most single-voiced documents focus on one or a small number of topics (e.g., product reviews or news articles). These complications of conversation pose many novel challenges (and opportunities) for researchers, even for researchers familiar with text analysis of single-voiced text data.

Managing conversation data sets: analyzing turn-level and speaker-level data simultaneously

Managing conversation data requires researchers to handle two distinct data sets: a turn-level data set (to examine conversational behavior) and a speaker-level data set (to compare that conversational behavior with pre-conversation or postconversation data, such as individual differences, experimental conditions, or outcomes). Thus, conversation analysis calls for analytical software (e.g., R or Python) that allows researchers to efficiently manage multiple data sets at once.

Turn-level data set: the contents of conversation. Most conversations can be discretized into a series of

speaker “turns,” much like a screenplay or script. These data can be represented as a transcript in which each row contains information about a single turn—specifically, who was speaking, the words spoken during the turn, and time stamps indicating when the turn started and ended. This data structure requires the turn-level data set to have unique identifiers for every conversation or group (e.g., Group 1, 2, 3, 4 . . .), every turn in each conversation (e.g., Turn Number 1, 2, 3, 4, 5 . . .), and each speaker in the group (e.g., Speaker A, B, C . . .). We provide an example of a turn-level data set in Table 2.

In general, the boundaries of each turn are determined by the time during which a single speaker is talking. Every new turn will involve a different speaker than the turn prior. Linguists distinguish the concept of a turn from that of an “utterance,” defined as a single continuous expression by a speaker. A turn can be composed of multiple utterances. For example, speakers could send several messages in a row before their partner responds. In that case, as a simplifying assumption, researchers typically collapse multiple consecutive utterances from a single speaker into a single turn.

Speaker-level data set: data from outside the conversation. In a speaker-level data set, there is a unique row for each speaker. Each conversation will have multiple rows (one for each speaker in that conversation), and speakers who joined multiple conversations will have multiple rows (one for each conversation). The unique identifiers for the conversation (or group) and speaker included in the turn-level data set can be used to connect the speakers’ conversational behaviors to the speaker-level data set, which also contains the conversation (or group) and speaker identifiers in addition to other variables recorded before the conversation (e.g., random assignment, time of day, context, demographics) and after the conversation (e.g., self-reported survey items, negotiated outcomes). We provide an example of a speaker-level data set in Table 3.

Many researchers will conduct their final analyses in the speaker-level data set because many research questions focus on variation at the person level or context level. When this is the case, the turn-level data set is used to generate measures of conversational behaviors (e.g., the number of questions or interruptions), which are then summarized at the person-level data set and tallied in the speaker-level data set (e.g., Speaker A in Group 4 asked 41 questions, five hedges, and interrupted three times during the conversation). We provide further detail on this topic in Model Construction below.

Capturing Conversation Data

There are considerable challenges involved in coercing conversation data into the data sets described above, and they vary based on modality. We focus on the two

Table 2. Example of a Turn-Level Data Set

Group ID	Turn	Start time	End time	Speaker ID	Text	Question	Laughter	Word count
1	1	0:00:01	0:00:03	A1	Hey, how are you? My name is [name] but my friends call me [name].	1	0	14
1	2	0:00:04	0:00:06	B1	Nice to meet you, [name]. I'm [name]. Where are you from?	1	0	11
1	3	0:00:06	0:00:12	A1	Thanks for asking! I'm from a small town outside of Chicago actually, you probably haven't heard of it. What about you?	1	0	21
1	4	0:00:13	0:00:20	B1	Probably not [laughter]. I've never been to Chicago. I'm from upstate Portland Oregon. Have you ever been to Portland?	1	1	19
1	5	0:00:20	0:00:22	A1	No, I haven't! I've been to Seattle, but that's all.	0	0	10
1	6	0:00:25	0:00:28	B1	Seattle is ok. In Portland, we actually call it Vancouver's shoe [laughter].	0	1	12
1	7	0:00:28	0:00:29	A1	That's funny.	0	0	3
1	8	0:00:33	0:00:36	B1	Um. [pause]. What's your favorite food?	1	0	5
1	9	0:00:37	0:00:55	A1	Hmm. That's a hard question. [pause] I really like all different foods. I made this really good stew the other day that I think might be the best thing I've eaten lately. But I'm always partial to a good hamburger.	0	0	39
1	10	0:00:56	0:00:59	B1	Cool. What was in your stew?	1	0	6

Note: The column labels are “Group ID,” used to distinguish between conversations; “Turn,” an index for each turn in the conversation in order; “Start Time” and “End Time,” indicating the time span of the turn; “Speaker ID,” indicating the speaking participant; “Text,” what was said in the turn; “Question,” a code for whether the turn contained a question; “Laughter,” code for whether the turn contained laughter; and “Word Count,” a count of words spoken during the turn.

most common conversational modalities, in which words are either written as text or spoken out loud. Each of these major modalities presents unique challenges and opportunities for speakers and researchers (e.g., M. Berry, 2013; Boland et al., 2022; Meredith & Stokoe, 2014; Oba & Berger, in press).

In either case, the fixed cost of structuring a conversation data set is not trivial. Once it is done, a good data set can benefit many subsequent research projects (and, possibly, many different researchers). Thus, we encourage researchers to explore whether it is possible to pilot test their research ideas in data sets from past research, including in archives purpose built for conversation data (e.g., Chang et al., 2020; Liberman & Cieri, 1998; Miller et al., 2017; Reece et al., 2022). For similar reasons, we

also encourage researchers to share their own data after they have structured it (see Data Sharing below).

Text-only conversations

Research on text-only conversation has proliferated in part because of the availability of text data, which are easy to record and store. It is often produced in massive Internet forums (e.g., Wikipedia, Twitter) or in catalogued archives (e.g., newspaper articles, books, legal documents, earnings calls) in which records are public and accessible to researchers (Hirschberg & Manning, 2015) or scraped using one of many available software tools that can scrape text content from webpages. In addition, people often have records of conversations

Table 3. Example of Speaker-Level Data Set With Round-Robin Design

Group ID	Speaker ID	Partner ID	Age	Gender	Partner gender	Condition	Liking	Partner liking	Questions	Laughter	Turns	Word count
1	A1	B1	24	1	2	1	5	6	2	1	5	87
2	A1	B2	24	1	1	1	2	7	3	1	4	60
3	A1	B3	24	1	2	1	7	6	1	0	3	54
1	B1	A1	34	2	1	1	6	5	4	2	5	53
2	B1	A2	34	2	1	1	6	2	0	3	6	102
3	B1	A3	34	2	2	1	5	4	3	1	7	131
1	A2	B2	57	1	1	2	2	5	0	0	2	45
2	A2	B3	57	1	2	2	1	7	1	1	4	75
3	A2	B1	57	1	2	2	2	6	1	0	5	64
1	B2	A2	23	1	1	2	5	2	1	0	4	24
2	B2	A3	23	1	2	2	7	5	3	3	5	33
3	B2	A1	23	1	1	2	7	2	4	2	6	98
1	A3	B3	55	2	2	1	3	4	2	1	3	112
2	A3	B1	55	2	2	1	4	5	5	1	4	33
3	A3	B2	55	2	1	1	5	7	1	2	2	16
1	B3	A3	19	2	2	2	4	3	1	0	3	47
2	B3	A1	19	2	1	2	6	7	0	0	4	87
3	B3	A2	19	2	1	2	7	1	0	1	6	101

Note: The column labels are “Group ID,” used to distinguish between conversations; “Speaker ID” and “Partner ID,” used to distinguish between participants in a conversation; “Age,” the age of the participant; “Gender” and “Partner Gender,” the gender of the participants in the conversation; “Condition,” which represents assignment; “Liking” and “Partner Liking,” self-reported measures; “Questions,” total number of questions the speaker asked in that conversation; “Laughter,” total amount of speaker laughter in that conversation; “Turn,” total number of turns in the conversation; “Word Count,” the word count of the speaker in that conversation.

conducted by chat or email. Accordingly, some researchers use software that allows consenting participants to extract and share their own text or social media conversations (e.g., Stillwell & Kosinski, 2004). Researchers also collect their own text conversations within controlled experiments with emerging technologies such as ChatPlat (www.chatplat.com; K. Huang et al., 2017), iDecisionGames (www.idcisiongames.com), Smartriqs (Molnar, 2019), and survconf (Brodsky et al., 2022).

Text conversation can be easier to analyze than spoken conversation because the words are already transcribed during the conversation itself (by the speakers). The style of conversation conducted via text is also different; compared with voice conversation, text-only conversation tends to be more asynchronous, with more time for cognitive preparation, reflection, and processing within and between turns; clearer sentence structures; and fewer disfluencies (M. Berry, 2013; Meredith & Stokoe, 2014). Still, text-conversation data present unique challenges for researchers.

Turn boundaries. The time course of text-only conversation can be tricky to pinpoint because transcripts often include only one time stamp per turn: when a message is “sent” or “posted.” If the conversation is more synchronous (e.g., instant messages), the lag time between these

stamps may be a useful signal of the time spent reading the last message or composing the next one. If the conversation is more asynchronous (e.g., email), the lag time may not be as informative.

In addition, in text conversations, people can compose their turns simultaneously, which can lead to multiple disjointed threads. When topics overlap, researchers must disentangle them by hand (or else accept some measurement error). Furthermore, most text platforms allow a single person to send multiple messages in a row, essentially replying to themselves. This can be simplified by combining consecutive messages from the same person into discrete, alternating turns—or by considering each message as separate turns.

Standardizing typing. In text-based conversation, people type their own transcripts. Writing style differs across people, cultures, languages, and time, and spelling and grammatical errors are common. There is a range of unique spellings in modern written language, including emojis (e.g., “☺”), variants (e.g., “oh nooo,” “woot!”), representations of sounds (e.g., “jajajaja,” “haha”), and acronyms (e.g., “tbh,” “lmk,” “lol,” “tldr,” “wtf”).

In many analyses, variants are simply ignored, especially if they are rare. However, some research questions might require attention to variants (e.g., grouping

different kinds of typed laughter or unpacking emoji valence to detect emotional sentiment). Clear writing errors can be more pernicious given that most feature-extraction systems rely on correct spelling and grammar. To address this, we strongly recommend that a person looks through each text at least once, perhaps assisted with spell-checking software, to fix obvious errors.

Voice conversations

Research on spoken conversations usually requires additional steps because spoken words are expressed in continuous sound waves that must be discretized into words, sentences, and turns. Some high-stakes audio conversations are routinely transcribed (e.g., interviews, conference calls, government proceedings), and some researchers have examined such documents (e.g., D. S. Berry et al., 1997; Chen et al., 2018; Danescu-Niculescu-Mizil et al., 2012; Hansen et al., 2018). However, the burden of accurately transcribing conversations often falls on researchers themselves. With technological advances, automatic speech recognition (ASR) and speaker disambiguation have improved (Park et al., 2022), but they are still not nearly as good at parsing speech as human transcribers (Errattahi et al., 2018; Meier et al., 2021), and this is likely to remain true for some time. Furthermore, these automated tools are often trained on convenience data samples, so they may be most inaccurate for speakers from underrepresented groups, who may use an accent or vocabulary that is not well represented in the training data (Dehghani et al., 2015).

We urge researchers to put serious effort into assuring data quality both through preparation before the conversations happen and after they have been recorded. Here, we suggest a series of steps and several tips to capture research-quality voice conversations.

Record. Researchers often underestimate the importance of audio-recording quality. This is especially critical when researchers have complete control over the recording protocol (e.g., recording participants speaking to each other inside a behavioral lab). However, there are cases in which researchers have less control, for example, the Electronically Activated Recorder (e.g., Kaplan et al., 2020; Mehl, 2017; Mehl et al., 2001), the Language ENvironment Analysis system (Ganek & Eriks-Brophy, 2016), and other experience-sampling methods that require people to carry microphones with them throughout the day. Furthermore, online experiments may have people conversing through their own home computers, which researchers do not have control over. Nevertheless, each of these protocols involves different considerations and constraints to optimize audio quality. Across all these study designs, we urge researchers to test their recording setup in advance.

High-quality audio recordings will lead to higher-quality transcriptions later. If you are having trouble hearing words when listening to a recording, your transcriber (human or ASR) will certainly struggle. Make sure you can clearly identify what words are being said and by whom. Some of the main factors to consider include microphone quality (e.g., sensitivity, internally generated noise, distortion, and directional characteristics), speaker clarity, background noise, distance from the microphones, and reverb. Ideally, researchers should rely on solutions that do not place a burden on the speakers; for example, a change in microphone placement will be a more reliable fix than asking speakers to enunciate more clearly.

One common decision point for researchers is the number of audio recordings per conversation: Should the entire conversation be captured in one file, or should each individual be recorded separately? A single recording may seem easier to set up but may complicate the analysis later because audio-transcription services often struggle with speaker differentiation, especially when two speakers have similar-sounding voices. With only a single recording, transcribers must determine whether the person talking is (a) different from the previous turn (Did the speaker change?) and (b) the same as any of the previous turns (Has this person spoken before?). This task is especially difficult when speakers have similar speaking styles or vocal registers and as the number of speakers increases. Video recordings can help, although we have found that professional human-transcription services often do not look at videos.

When possible, we recommend collecting separate audio recordings for each speaker. This makes speaker differentiation simple and improves audio quality by moving microphones closer to each speaker. Fortunately, virtual meeting services (e.g., Zoom) record separate audio streams from each computer, which automatically differentiate speakers (if people have their own computer). Some services automatically combine these separate streams into a single turn-by-turn transcript (including Zoom and Microsoft Teams). If separate recordings are set up manually, they must then be combined and sorted into the correct order using the time stamps for each turn.

To connect the speaker-level data to the data collected outside the conversation (e.g., demographics and survey data), each speaker and each conversation must have a unique identifier that can be used to link the turn-level and speaker-level data sets. As a safety measure, researchers may consider reading the conversation identifier out loud at the beginning (or end) of the audio recording and use the identifier as the name of the audio file as well. Likewise, speakers in the conversation should say their unique speaker identifier as one of their first turns in the recording so their voices can be unmistakably matched to their conversation-level and speaker-level data.

We recommend conducting a few test recordings that run through as much of the workflow as possible. The researchers should check to see that the file records well (that the audio is clear and that the spacing of microphones and speakers is appropriate), that it can be played back properly, that it is saved in a format that is compatible with the intended transcription method, and that the researcher can match each recording and each speaker to the metadata. Finally, do not forget to press “record.”

Transcribe. Transcriptions of the audio recordings will form the foundation of the turn-level data set. There are several approaches to generate transcriptions from audio files. Most commonly, researchers pay traditional transcription services, which hire trained humans to type words while they listen to audio recordings. However, this approach is often inadequate (and expensive)—the quality is inconsistent, typos are inevitable, and transcribers use different formatting methods (even within the same company). Some researchers hire research assistants to transcribe. Although this affords more control over formatting, the training can be long, and the work can be arduous and inefficient. Others use automated speech-recognition software. Although software will never be as accurate at recognizing words as the best trained humans, they produce the most precise time stamps, and they deliver consistent formatting and spelling.

We strongly recommend a hybrid approach, combining automated speech-recognition software with trained humans, which is both accurate and cost-effective. First, automatic speech-recognition software can generate a low-cost first-draft transcription, tackling the easiest sections of the transcript quickly and producing transcripts with consistent formatting and reliable time stamps. Then, this initial draft of the transcript can be edited by a human, who can focus time and attention on the more difficult tasks, such as speaker differentiation and correcting any passages with low-quality audio.

It is important to establish consistent formatting conventions early. Many transcription services (human and machine) export their data in text documents (e.g., Microsoft Word, PDF) rather than tabular files (e.g., Microsoft Excel, CSV). However, as long as all files have a consistent format, researchers can write code to parse the text files into an analyzable tabular format. Subtitle file formats (.VTT files) are also common for mapping utterances to time stamps, and these files can be processed into tabular formats automatically in R (Knight, 2023).

There are many automatic transcription services available today (e.g., Otter, Temi, Amberscript, Descript, Trint, Sonix, Happy Scribe, Wreally, Ebby, Scribie; Table A1), and new services and iterations are rapidly emerging (e.g., OpenAI’s Whisper tool, which was released during

the revision process of this article). In September 2020, we systematically tested 10 of the most popular transcription services available. Each service transcribed the same series of audio recordings, and we evaluated the services along the following dimensions: (a) transcription accuracy, (b) speaker differentiation, (c) incorporation of time stamps, (d) user-friendliness, and (e) pricing. We summarize our findings in the Appendix.

This review is not meant to be definitive. Rather, its primary purpose is to demonstrate how researchers might test and compare various transcription tools. Automatic speech-recognition products and services have been rapidly evolving over time. Thus, we strongly encourage readers to conduct their own contemporaneous search at the time they require these services, evaluating their options based on the dimensions we list above. Researchers’ needs may also vary depending on what is best for their projects, so there is not a single best transcription service for everyone. However, we believe a hybrid transcription approach—automated transcription followed by human correction—is and will remain the most cost-efficient way to produce accurate, research-quality transcripts, at least in the near term.

Check. Automated transcription services have become more accurate over time, but they are not perfect (and neither are human transcribers). We strongly recommend asking people to listen to the audio recording while reading through the transcript, fixing any mistakes, and ensuring that formatting conventions are consistent throughout.

For example, transcription services have different policies about how to demarcate inaudible moments. Many will simply skip over this moment and leave a blank, whereas others will flag this with “[inaudible],” sometimes with a time stamp including duration. Our preference is typically to use the “[inaudible]” flag, which can be removed as needed; either way, it is essential to be consistent throughout. Furthermore, there are many paralinguistic features that may be ignored by some transcription services. Common examples of these are laughter (“[laughter],” “[laughs],” or “[laughing]”) and interruptions (“[interruption],” “[interposing],” or “—” at the start of an interrupting turn). Similar approaches are taken for other paralinguistic cues, such as sighing, singing, crying, yelling, whispering, or cross-talk. Research questions should inform your approach: If laughter is important, make sure you annotate it and do so consistently.

Checking transcripts can also uncover errors in the time stamps. One common error is typos from human transcribers—large errors can often be detected in later analyses (e.g., typos often result in negative or very long interturn pauses), although smaller errors also happen. When speakers are recorded separately, their time stamps may be aligned to different benchmarks in each

recording (e.g., if the recordings start at different times). In this case, time stamps must be realigned to a common reference time before the transcripts from each recording are merged.

It is often useful to have human coders fix errors made by the speakers themselves, too, unless those errors are of research interest (e.g., self- and other-initiated repairs are important conversational phenomena). Some examples include the following:

- Include and standardize the spelling of backchannels (e.g., “yeah,” “uh-huh,” “oh”).
- Remove erroneously repeated words (e.g., “I thought you . . . thought you were ready”).
- Include punctuation (e.g., question marks, periods, commas, ellipses).
- Change “gonna,” “sorta,” “dunno,” and so on to “going to,” “sort of,” “don’t know,” and so on.
- Correct misspoken words in cases in which the intended meaning is clear (e.g., “nice to mate you”).

There can be subtle but important differences in meaning among nonstandard variations (e.g., “yes,” “yup,” “yasss”). However, there is a trade-off between specificity and statistical power. In general, differentiation could be reasonable if there is an adequate sample size of each variation and if the distinctions matter for the research questions at hand. Otherwise, it may be best to aim for consistency (e.g., “yes” to study linguistic affirmation broadly).

Although transcript checking can be monotonous, the process can be designed efficiently. We typically find it easier for research assistants to complete all tasks for one document at a time rather than completing one task for all documents before moving to the next task. However, to batch tasks like this, you must plan your checking needs in advance. For more efficiency, error checking can also be batched with human-feature annotation (see Feature-Extraction Objectives below).

Extracting Features From Text

Perhaps the most daunting task for conversation researchers is to decide which features to extract from the transcripts. Each “feature” can be thought of as a measure of one behavior in the transcript (e.g., the number of first-person pronouns, the percentage of words that mention food, the average length of pauses). There are a large (and increasing) number of tools available to researchers for this task, and researchers are presented with a wide array of options, even for measuring the same underlying construct (Schweinsberg et al., 2021; Yeomans, 2021).

We offer a brief review of common techniques with a special focus on the challenges of studying dialogue data (vs. single-voice documents). Although tools for these steps are available in several software environments, we point readers to tools in the R software language. However, we note that Python also has many excellent tools for NLP (ConvoKit, in particular; Chang et al., 2020). Note that both Python and R allow users to manage the two data sets—turn-level and speaker-level data—simultaneously. This means researchers can integrate their feature-extraction code with their analysis code (see Model Construction below).

Feature-extraction objectives

Before we introduce common feature-extraction methods below, we first describe the important dimensions on which these methods can differ. This is important because there is no one “correct” approach. Instead, researchers must choose techniques on the basis of their own idiosyncratic objectives and constraints, which are determined by their skill set, audience, research goals, resources, deadlines, and so on. Each of these dimensions should be considered when choosing a feature-extraction method.

Accuracy. First and foremost, researchers should hope the features they extract from text data are valid, accurate measures of the underlying behavior or belief (Flake & Fried, 2020). Thankfully, accuracy can be evaluated empirically within a validation data set that has labels that can be treated as “ground truth” for comparison. For example, a turn-by-turn measure of question asking should correlate as highly as possible with the true number of questions in each turn.

However, accuracy is not an inherent property of any method—it can be defined only within a particular population of interest. For instance, a model trained to label different types of questions in a doctor’s office may not be as valid for labeling question types in a job interview. Researchers should be explicit about their intended populations and the boundary conditions of their results (Simons et al., 2017). They should also routinely conduct tests of “transfer learning” (Weiss et al., 2016; Yeomans, 2021) by explicitly testing how well their methods perform when they are developed in one context and applied to data from a different context.

Fairness. Bias is a concern shared by both humans and artificial-intelligence (AI) systems. Just as humans are prone to unconscious biases (Greenwald & Banaji, 1995), AI models can exhibit algorithmic bias (Kordzadeh & Ghasemaghahi, 2022). Mitigating this bias is essential to ensuring the accuracy and fairness of research outcomes regardless of the initial source.

Because language models learn about the world from data used to train them, anything that learns from biased language data may unwittingly generate models that reinforce and codify prejudice, stereotypes, or other unsavory aspects of human judgment (Caliskan et al., 2017). And, as is often the case with historical (and present-day) data sets, the speakers in the training data may themselves be biased or prejudiced. Sometimes this bias is the subject of research inquiry itself; however, if the focus is on other aspects of human behavior, this bias can undermine the goals of the research. This is especially true when a model or estimate is used to make decisions that affect real people. Consider, for example, an algorithm used to match job candidates to job postings using similarity to exemplars in past training data. If that training data reflect a past in which some demographic groups (e.g., women, minorities) were excluded or discouraged from leadership roles, then the model on which it is trained may unwittingly reinforce that bias going forward. For example, an algorithm employed for recruitment at Amazon was later shown to be unwittingly discriminating against female applicants because the data it learned from showed that most leaders tended to be male (Dastin, 2018).

The accuracy of a model can thus vary across social groups in ways that may have biased consequences for the outcomes of those group members. Models trained on only one kind of speech, such as data from the most commonly studied sources (e.g., from demographic majority groups, from American-English speech), may be much less accurate when they parse speech from groups that are historically underrepresented, from speakers from non-American countries, or for other reasons not included in the training data (Koenecke et al., 2020). This is an issue for all kinds of slang, jargon, and other language that are contextually—or socially—determined, and this type of language is very common in conversation.

There are no surefire techniques that can ensure a model is unbiased. One approach that has grown more common in recent years is to conduct an “algorithm audit” in which AI systems are evaluated to ensure they work as expected and do so without bias or discrimination (Brown et al., 2021; Koshiyama et al., 2022). Moreover, transfer-learning tests, as described in the Accuracy section, are very useful—by comparing how well a model’s accuracy varies across different populations, researchers can evaluate whether particular groups may be adversely affected. When transfer-learning tests are not possible, researchers should explicitly acknowledge the limitations of their training data so that their tools are not misused by others. To improve the model itself, researchers should try to find training data that best represent the people involved, perhaps even oversampling

from less numerous groups so that they are accounted for in the model. Above all, we recommend not taking model outputs as ground truth; instead, researchers should try to interpret and understand their models as much as possible, evaluate the contents using their own domain expertise, and be as thorough as possible in making sure the model is behaving as expected.

Interpretability. Behavioral scientists are rarely concerned only with prediction accuracy. They also seek to understand and explain how people behave, which means they also need to understand what drives the results of their statistical models. Interpretability allows researchers to scrutinize their models so that they might improve them and think about how well they might generalize to new contexts (Bianchi & Hovy, 2021). Improving interpretability can also improve fairness by allowing users (including regulatory bodies) to evaluate the model’s strengths and failings in detail (Doshi-Velez et al., 2017; Rudin, 2019), and users generally trust models more when they understand them (Gilpin et al., 2018; Yeomans, Shah, et al., 2019). We recommend a similar skepticism from researchers—so-called black-box methods that are not explained should not be relied on to provide scientific insights.

Although interpretability is almost universally desirable, it is difficult to define or quantify it precisely (Lipton, 2018). But generally speaking, models can be made more interpretable along two dimensions. First, the methods themselves should be transparent. Their exact content, code, and training procedure should be shared and benchmarked against related models across diverse contexts (Mitchell et al., 2019). However, transparency is necessary but not sufficient—many modern NLP models are still too complex to scrutinize, even by experts (Bender et al., 2021). More troublingly, this information is often not shared because of expediency and to prioritize individual success over progress as a field (Belz et al., 2021). For example, the DICTION software package provides only broad generalities about how its features are scored or how its formulae were determined and validated (Hart, 2001) even though its license fee is much higher than open-source models that are much more transparent.

In addition to transparency, models can be made more interpretable by generating additional outputs in addition to raw feature scores. One approach is to use the model scores to find excerpts from the dialogue that highlight contrasting levels of a given measure (e.g., high vs. low warmth; follow-up vs. switch question). Often, they can also extract coefficients directly from the model to reveal which features most affect a model’s output (e.g., K. Huang et al., 2017; Voigt et al., 2017). Even when researchers must rely on an uninterpretable model because of their high accuracy (e.g., human annotators

or black-box NLP), they should still try to understand its workings. One approach is to train a simpler model that approximates the predictions of the more complex one and interpret that simpler one instead (Madsen et al., 2021; Ribeiro et al., 2016).

Scalability. Researchers usually need to anticipate the costs of calculating and extracting features at a large scale. All feature-extraction methods involve direct resource costs. These costs come in the form of upfront investment (e.g., learning how to use a new software package or developing an annotation scheme) and in the marginal cost of applying a method to new data (e.g., computation or annotation time). There are other limitations that affect the costs of implementing different methods. For example, when data are proprietary, identifiable, or otherwise sensitive, some methods (e.g., human annotators reading raw text) may come under more intense scrutiny from stakeholders than other, less invasive methods (e.g., computing average turn length).

Complexity. Many of these objectives are related to the complexity of a feature-extraction method, even though complexity is not itself an objective. Complex features tend to be costlier to implement, but this extra effort is typically justified because of improved accuracy, fairness, or interpretability. Conversation is itself complex, so a perfectly accurate feature extractor would have to be correspondingly complex. Instead, researchers often settle on a trade-off between acceptable effort and acceptable accuracy, and this can be done iteratively: Simpler measures can be used first, and if that is insufficient, then more complex measures can be used. To borrow an idiom, before investing in a more complex method, researchers should first consider if “the juice is worth the squeeze.”

Complexity is often related to the scope of information needed from the transcript to identify a single feature, whether it is responsiveness, warmth, question types, expressions of gratitude, disfluency, or interturn pause length. The simplest and most common methods treat a person’s turns as a block of static text, as if they were single-voice documents (see Static-Text Features). This allows researchers to draw on the large tool kit from single-voice document analysis. However, this ignores the features of text that make conversation unique. For example, some features incorporate the time stamps from the transcripts (see Timing Features). Many other features look at consecutive sequences of turns to understand the structure of how speakers are interacting (see Interactive Features). We illustrate these different input scopes in Figure 2.

NLP versus human annotation

Before computational tools were available, researchers traditionally annotated conversations, scoring various

features in transcripts by hand. In theory, any annotation task done by a human could be attempted with an algorithm instead and vice versa. Thus, it is tempting to see NLP as a potential substitute for human labor to automate simple workloads and reduce time spent reading.

However, we argue the opposite: Researchers should consider NLP as a complement to human work. These algorithms make close reading more powerful because they can be used to scale up and interpret human insights. Humans can develop typologies and provide labels to train supervised algorithms. Researchers themselves can read their corpora to guide their intuitions on which algorithms might be the best fit for their data and context.

Advantages of humans. Human and algorithmic feature extraction have contrasting strengths and weaknesses. For example, many conversational phenomena are too complex for current tools to automatically detect with sufficient accuracy. In these cases, trained human annotators usually produce more accurate labels and can be used as the “gold standard” for evaluating NLP performance (Bommasani et al., 2021). Human annotators can use their knowledge about the social context of a conversation to frame their responses, whereas an algorithm typically applies the same scoring rule regardless of context. For example, humans use their knowledge about speakers and context to infer sarcasm, whereas algorithms are typically built to take all of a speaker’s words at face value. Humans are better at understanding nuanced meaning amid social exchange.

Limitations of humans. People can be inconsistent from day to day and between one another—annotators almost always have some amount of disagreement. Furthermore, their thought processes may be hard to know or interpret (Nisbett & Wilson, 1977). Annotators often do not—or cannot—give precise reasons for their judgments. Although the exact protocols used to train the annotators can be shared, this does not guarantee that human annotators followed them or followed them in the same way. Thus, algorithms are not the only black-box feature extractors used in research—humans can be black boxes, too.

Humans can suffer from many of the same problems that algorithms do. Accuracy within and across domains is always a concern. When human annotators perform poorly, it can be hard to know if the task is inherently difficult, human judgment is too subjective, or the annotators are lacking the right training. Human annotators can treat people unfairly because of historical bias and prejudice or inexperience in the domain, among other reasons (Denton et al., 2021). All of the tools available to interpret algorithmic judgments should be used to scrutinize human annotations for unintended biases or blind spots.

Costs of human annotation. The costs of using human annotators are typically higher than using an algorithm. Much of this difference lies in the marginal costs of annotating new data—annotator time scales linearly with the amount of data, whereas the marginal cost of automatically processing more data is trivial once an algorithm is built. However, there are upfront fixed costs for both. For humans, researchers must establish clear definitions and protocols for assigning labels. Annotators then practice until they reach sufficient agreement on training cases. Researchers may revise their protocols during training, as their definitions are applied to edge cases in real data. This process is iterative: drafting a scheme, then testing it individually and via group discussion, revising the scheme, and retesting. These details are usually context-specific, and researchers should work with domain experts to develop their annotation schemes.

Often, researchers try to reduce annotation costs by crowdsourcing label generation to pools of online workers (e.g., from Mechanical Turk). However, crowdsourced workers have their own problems. They are hard to train, do not provide good feedback during protocol development, and can be inattentive. The task must be cleverly allocated across many workers because each one can label only part of the data set (e.g., Benoit et al., 2016; Kiritchenko & Mohammad, 2017). Accuracy concerns are less relevant for simple tasks and can be mitigated in part by averaging over many annotators (although, this reduces their cost advantage).

In general, we have found that if annotation tasks are sufficiently complex, a pair of in-house research assistants can produce more accurate labels than a larger pool of crowdsourced workers. Moreover, in-house annotators can complete the necessary checking and cleaning tasks described above (also see Check section).

Human-algorithm hybrids. As with transcription, a hybrid approach may be useful during feature extraction. Human annotations can be used to train interpretable algorithms that reproduce human judgments. This approach identifies the linguistic features that are driving the humans' judgments. A side benefit to this hybrid approach is that if the resulting algorithm is accurate, it can be directly applied on new data without having to recruit new human annotators. In addition, rough algorithmic approaches can be used as a first pass to focus the efforts of human annotators.

We used this workflow ourselves in K. Huang et al. (2017) when we wanted humans to annotate different question types. First, we applied a simple algorithm to identify turns that included a question (to assist the humans' search through the transcript). Then, human research assistants coded these questions as one of several question types. After the human annotations were

collected, the consensus labels were then fed back into a supervised learning algorithm to train a question-type detector. The final model included both the initial search filter and the supervised model so that it could reproduce the human annotators' judgments at scale. It was trained on 4,209 annotated question turns within 368 conversations from a lab experiment and then applied to an observational data set with 987 conversations and 19,321 question turns.

Static-text features

There are many review articles covering different methods for extracting features from single-voice documents. For brevity, we review the most common methods and focus on why they may function differently in dialogue. These methods treat turn content as though it were from a single-author document, such as a news article. However, individual turns vary wildly in word count. In practice, this means many turns from one speaker are collapsed into a single piece of text (this is discussed in detail in Aggregating Conversation Features below).

Counting words. A common, straightforward approach to analyze text is the “bag of words” approach: Count each word that occurs at least once, ignoring order. This can produce a very large feature set (perhaps thousands of different words in a single conversation). There are many preprocessing steps commonly used to smooth out the raw counts, including reducing words to their stems, expanding contractions, removing rare words, removing common “stop words,” and constructing “n-grams” (two- or three-word phrases).

These techniques improve models, but they should be considered in light of the specific research questions that are being addressed (Denny & Spirling, 2018). Conversation has a lot of stylistic and structural language, which tends to be determined by the more common function words—pronouns (“you,” “they”), adpositions (“to,” “from”), determiners (“the,” “your”), and adverbs (“mostly”). For example, question words (“who,” “what,” “where,” “when,” “why,” “how,” “which”) are essential for determining what types of questions people are asking (K. Huang et al., 2017; Zhang et al., 2017). However, these words tend to get removed by most off-the-shelf stop-word lists, which were typically built for single-voiced text.

Dictionaries. Dictionaries are lists of words generated by expert human annotators that give scores to words that group them into simpler dimensions of meaning. For example, a “food” dictionary would give all the words relating to food (e.g., “pizza,” “broccoli”) a score of 1 and the rest of the words (e.g., “bicycle,” “reading,” “heavenly”) a score of 0.

a score of 0. Other dictionaries assign each word a score on a continuous scale using average ratings (e.g., concreteness; Coltheart, 1981; Warriner et al., 2013). To calculate the summary score for the whole text, the scores of the individual words within it are averaged. For binary dictionaries, this score is the percentage of words that comes from a dictionary.

Dictionaries are common and accessible. The Linguistic Inquiry Word Count (LIWC) is probably the most often used NLP tool in psychology (Tausczik & Pennebaker, 2010) because it requires no special skill to conduct analyses and many features are simple to understand (e.g., first-person pronouns, words about music). Although dictionaries can be quite useful, users should be aware of their limitations. Most obviously, dictionaries (like bag of words) ignore the order of words, sentences, phrases, and topics—how verbal content unfolds in sequence. For example, most dictionaries do not account for negations (“not bad” vs. “bad”) or relative magnitude (“very bad” vs. “bad” vs. “terrible”; although, see Hutto & Gilbert, 2014). Furthermore, the interpretation of dictionary results is often lacking. Although it is tempting to simply take the title of a dictionary at face value, its meaning should be determined from the actual words it contains and the procedure by which it was created and validated. Sometimes these details are not shared publicly.

Furthermore, authors should make sure the dictionary is capturing what is intended in their context by comparing texts from their data with the dictionary’s scores, perhaps starting with texts that get especially high or low scores. Most dictionaries implicitly assume domain-generalizability—that the contained words each have a single, stable meaning (Hamilton et al., 2016). This is not always true in conversation (Boyd & Schwartz, 2021; Eichstaedt et al., 2021; Yeomans, 2021). For example, even something simple such as emotional sentiment (e.g., positive words minus negative words) can fail to measure closely related concepts such as the experience of happiness or well-being of the speaker (Beasley & Mason, 2015; Jaidka et al., 2020; Kross et al., 2019; Sun, Schwartz, et al., 2020) or the nuances of how a business or product is being described (Frankel et al., 2022; Rocklage et al., 2022). Although domain-specific dictionaries can help these concerns (e.g., Loughran & McDonald, 2016), the boundary for what is in- versus out-of-domain is not always clear, and researchers are usually best off conducting their own in-domain validation (Benoit et al., 2019; Yeomans, 2021).

Sentence structure. Modern NLP tools can extract not just the words themselves but also the underlying structure of sentences—that is, the grammatical parsing of sentences into subjects, verbs, objects, modifiers, clauses, and so on. This improves the features extracted from a typical bag-of-words model by making use of structures that

determine meaning—for example, negations (“bad” vs. “not bad”), named entities (“apple” the company vs. the fruit), and homonyms (“like” the positive-valence verb vs. “like” the valence-neutral adposition). Researchers can use pretrained neural-network models (Honnibal & Johnson, 2015; Manning et al., 2014, 2020) to generate grammar tags for each word and then build features using the tagged set.

These tools have been effectively applied to measure markers of politeness from individual turns (Danescu-Niculescu-Mizil, Sudhof, et al., 2013; Voigt et al., 2017; Yeomans et al., 2020; Yeomans, Kantor, & Tingley, 2018). In conversational text, politeness features often succeed at capturing the robust dimensions of how speakers structure their conversational turns—agreement, disagreement, acknowledgment, hedging, gratitude, subjectivity, apologies, greetings, and goodbyes. Models trained on these dimensions have generalized well across multiple domains because they focus on structural and stylistic features rather than the main content features that tend to define a domain (e.g., specific nouns and verbs). Figure 3 provides an example of politeness features extracted from a data set to show the differences in linguistic style that result from a randomized preconversation assignment to condition.

Embeddings. A common approach to detecting semantic content is to use pretrained “embedding spaces” that represent words and sentences as vectors within a space of meaning (e.g., Landauer & Dumais, 1997; Mikolov et al., 2013). Most modern embedding models are extracted from small neural networks trained to estimate which words tend to have the same neighbors (Bhatia et al., 2019). To solve this problem, the inner layer of the neural network groups words with similar meanings close to one another within the space. These embeddings are particularly useful for tasks that involve a similarity calculation—for example, measuring the semantic similarity of two texts (Arora et al., 2017) or improving dictionaries. Rather than using a dictionary to count words in a binary sense (i.e., presence/absence), authors can compute the similarity of a whole document to the dictionary as a continuous measure (e.g., Garten et al., 2018; Sagi & Dehghani, 2014).

Embedding models have several advantages over raw word counts. These models group words with similar meanings into a common dimension, whereas a word-count model treats each word as its own dimension, reducing the feature space considerably. Although word-count models typically remove rare words to simplify the estimation, embedding models are pretrained on large data in which a high frequency of words is seen often enough to be included in the model.

However, embedding spaces are difficult to interpret—the dimensions themselves do not directly correspond to

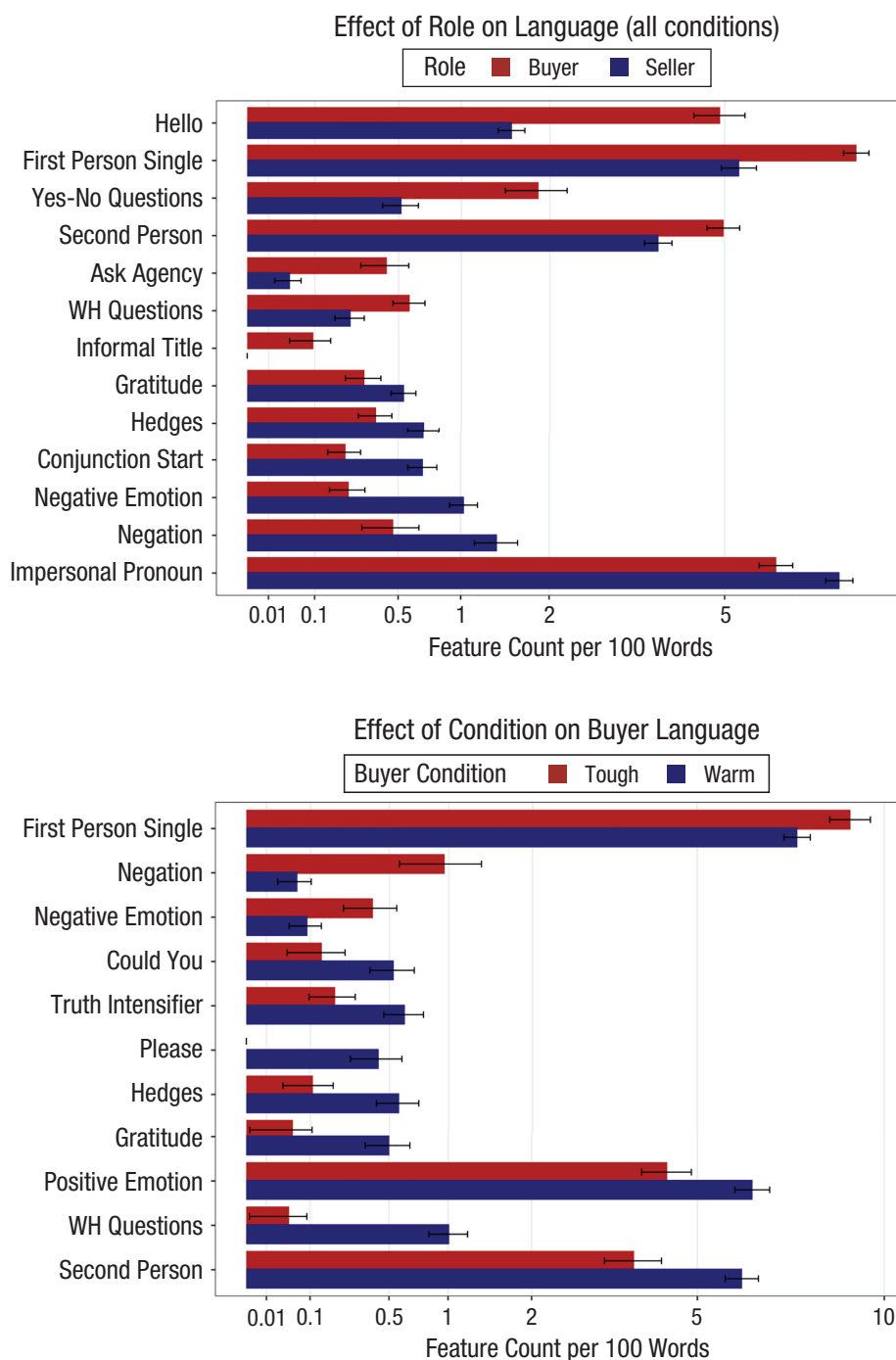


Fig. 3. An example graph showing dialogue features extracted from negotiation transcripts (Jeong et al., 2019) using the *politeness* R package (Yeomans, Kantor, & Tingley, 2018). (Top) Comparison of the feature usage between buyers and sellers. (Bottom) Comparison of the feature usage of buyers instructed to be warm and friendly versus tough and firm. All bars show group means and standard errors. Note that plots show feature counts per 100 words because buyers (especially buyers instructed to be warm) use many more words than sellers.

meaningful concepts, and researchers must use other tools to interpret what the model is doing. In addition, many common pretrained embedding models are mapped

to individual words, which means that they ignore the order of words spoken in conversation and other sources of contextual variation in meanings. Still, newer models

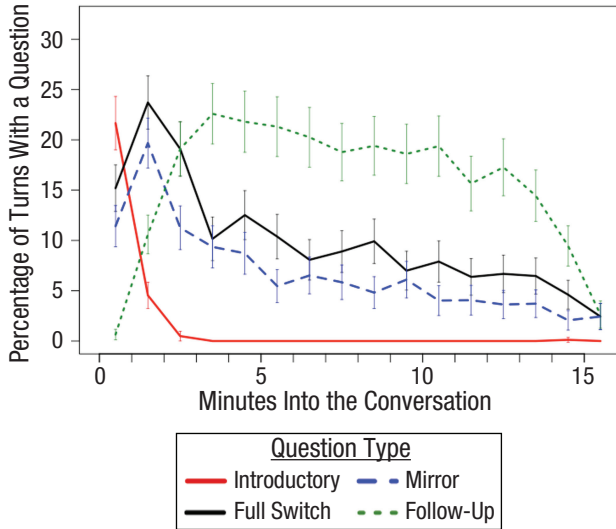


Fig. 4. An example conversational time-series graph showing frequency of question types asked over the course of approximately 300 conversations between strangers (data from K. Huang et al., 2017).

of embeddings can encode entire sentences within an embedding space (e.g., Devlin et al., 2018) and can be fine-tuned to incorporate some contextual differences in meaning if the researchers have enough data. This is a frontier of constant progress in the NLP community.

Timing features

In this section, we review several conversation-specific features that can be derived from time stamps. Many types of conversation features are particularly prevalent in some parts of the conversation (for an example, see Fig. 4). Furthermore, the impact of some features of language may vary in meaning or effect depending on when they are said during a conversation (e.g., Y. Li et al., 2022). The most common use of time stamps is to organize other features of text and to select features from certain parts of the conversation for analysis. This is relevant for causal versus predictive inference (see Model Estimation).

Pauses. Typically, there is some amount of pause between turns, measured as the difference between one turn’s end time stamp and the next turn’s start time stamp. Pauses tend to be longer in asynchronous and text conversations and shorter in synchronous and spoken conversations. Teleconference conversations tend to be somewhere in the middle of the two (Boland et al., 2022). Within a particular data set, pauses of various lengths can be counted as turn-level features (Templeton et al., 2022, 2023). Some researchers simply dichotomize each turn into pause or no pause using a threshold and show that results are robust over a range of thresholds (e.g., Curhan et al., 2022). It is more

difficult to define within-turn pauses, in which people pick up after their own silence, and the relevant time stamps are not included in a turn-level data set. Transcribers (human or algorithmic) can be instructed to indicate a midturn pause as a nonverbal (e.g., “so anyways . . . [pause] did you see them at the wedding?”), which can be counted or removed as needed.

Interruptions. Sometimes speakers do not leave any time in between their turns or even talk over one another. This often happens when the first speaker is interrupted by the second, and this type of interruption is often given a special annotation in transcripts (e.g., a single dash at the beginning or end of a turn) and a zero or negative gap between the end time of the previous turn and the start time of the interruption. The meaning of these interruptions is the subject of scholarly study—as a signal of disrespect or authority in formal settings (H. Z. Li et al., 2004; Mendelberg & Karpowitz, 2016); a sign of excited, enjoyable discourse (H. Z. Li et al., 2004; Yeomans & Brooks, 2023); or a signal that one person was merely filling dead air until the partner was ready to take a turn. The content of the interrupter’s turn also distinguishes different types of interruptions, such as backchannels, questions, and arguments (Shi et al., 2022).

Speaking time. Time stamps can also be used to measure speech patterns over longer periods. For example, speaking time (i.e., “participation” or “airtime”) is commonly measured as the percentage of the total time that is used by a particular speaker. When time stamps are not available, airtime can be approximated using the number of words spoken by each speaker as a percentage of the total words spoken (although this does not account for when no one is speaking). Comparing turn length with the time stamps will give an estimate of the person’s speaking speed (i.e., “cadence”).

Interactive features

Backchannels. During conversation, listeners often insert a brief utterance to signal they understand (e.g., “yeah,” “ok,” “mm-hmm”) while someone else is talking. Different definitions have been used, and it varies according to context (e.g., audio vs. text chat). Typically, backchannels are treated as a single turn within the flow of conversation with zero time gap between the preceding and subsequent turns. This may unnaturally divide the longer turn of the backchannel recipient into two separate turns, which could interfere with sentence-level features. Some researchers have avoided this by considering backchannels as features of the turn receiving the backchannel. Then, each turn has a feature counting the number of backchannels it receives from other speakers (Reece et al., 2022).

Dialogue acts. Most of what is said in conversation imposes a structure on what is said in subsequent turns: asking different types of questions; stating facts, opinions, or feelings; making requests or commands; signaling understanding, agreement, or disagreement; or initiating repair. These “dialogue acts” are essential to understand how speakers are communicating with one another (Bunt et al., 2010; Stolcke et al., 2000). Other theoretical frameworks (e.g., speech acts; Searle, 1965) capture roughly the same idea, which is that conversational turns are usually more than just statements of fact about the world. Rather, they communicate speakers’ intentions and give structure to the response they expect to receive.

Some dialogue acts can be reasonably approximated with features extracted from individual turns by the *politeness* package (e.g., gratitude, apologies, acknowledgment; Yeomans, Kantor, & Tingley, 2018; see Fig. 3). However, many other dialogue acts are difficult to identify without information from other turns. For example, adjacency pairs (e.g., consecutive turns such as question/answer, offer/acceptance, misunderstanding/repair) often demarcate essential decisions in a conversation.

There is no universally accepted, domain-general list of dialogue acts. Instead, the set of relevant dialogue acts will change depending on the conversational context (e.g., the modality of exchange, the goals of the speakers). For example, consider the sequence of formal offers within a negotiation. Specific offers are among the most important dialogue acts, so the impact of measurement error on these features would be considerable. In fact, most negotiation platforms (e.g., iDecisionGames or eBay) require that formal offers be made separately from the unstructured stream of conversation so that the speakers themselves can understand their partners. Algorithms may be able to parse the offers in simple negotiations (Lewis et al., 2017), but if the negotiation involves multiple complicated issues, automatic extraction may not be possible and human annotation may be preferred (Jäckel et al., 2022; Weingart et al., 2004). The same treatment may be necessary for other dialogue in which particular turns have formal significance—for example, voting during a meeting or generating creative ideas (Brucks & Levav, 2022).

Accommodation. One of the most common and reliable results in conversation analysis is accommodation—the tendency of one speaker to mirror the linguistic features of the previous speaker (Giles et al., 1991). Several models of accommodation have been proposed. The most common measure combines the entire transcript of each person separately and then calculates the similarity of those two documents (Ireland et al., 2011). However, this ignores order and directionality (e.g., Which of the speakers is doing the accommodating?). Other models are purpose-built for

conversation and explicitly identify accommodation from one turn to the next (Danescu-Niculescu-Mizil et al., 2011; Demszky et al., 2021; Doyle & Frank, 2016), and this can be aggregated as a feature of one or several turns.

Researchers have considered several feature sets over which accommodation should be measured. Some articles have focused on mirroring of content (e.g., If I talk about my dog, will you talk about your dog?; Babcock et al., 2014; Fusaroli et al., 2012), and others have focused on stylistic categories (e.g., If I use more quantifiers, will you do the same?; Danescu-Niculescu-Mizil et al., 2011) or syntactic structure (e.g., If I use short, clipped sentences, will you?; Boghrati et al., 2018). Other articles have included a wide range of features, combining content and style (Niederhoffer & Pennebaker, 2002; Srivastava et al., 2018). In truth, it is not clear whether conversational style and content can be cleanly separated, and the two often correlate with one another—in essence, some types of content naturally pair with particular styles. This is a subject of ongoing research.

Topics. Conversations are very often broken into discrete topics (e.g., the weather, then work, then cooking, and so on) based on speakers’ varied intentions (Passonneau & Litman, 1993). There are well-known NLP algorithms that focus on extracting topical content from text (i.e., topic modeling). The most common approach, latent Dirichlet allocation, assumes that each text document is a mixture of a small number of topics and that each word’s presence is attributable to one of the document’s topics (Blei et al., 2003; Roberts et al., 2019).

Alas, conversation data are not well suited for topic models built for single-voice text. Topic models focus on the distinctive words that demarcate content and typically remove common words, such as pronouns (e.g., “I,” “you,” “it,” “she,” “they”). However, many turns contain no topic-relevant information (e.g., “Why is that?” could be asked in almost any topic), and most turns are too short to reliably estimate word co-occurrence. Instead, blocks of turns must be segmented into topics for analysis, and dividing dialogue into segments is arguably even harder than assigning a topic to a particular segment (Purver, 2011; although, see Galley et al., 2003; Hearst, 1997; V. A. Nguyen et al., 2014). Furthermore, in both single voice and dialogue, it can be hard to choose the number of topics and interpret the words within each topic (Boyd-Graber et al., 2014; Chang et al., 2009). Still, topic models may be a useful tool for rough exploration and descriptions of the main themes of a body of dialogue.

If the topical structure is important to measure precisely, we suggest researchers avoid relying on an unsupervised algorithm but instead develop their own categories on the basis of their knowledge of the domain

and their exploration. For example, conversations can have a list of preassigned topics, which makes ex-post segmentation much easier (e.g., Yeomans & Brooks, 2023). Many conversations that are repeated often—such as sales calls, customer service, doctor-patient interactions, police interviews, parole hearings—have explicit or accepted dialogue scripts that speakers have been trained to follow as a progression through a series of stages. These scripts can be used to develop domain-specific rules to segment individual transcripts into discrete topics or stages (e.g., Takanobu et al., 2018). This is a subject of ongoing work, and NLP researchers have made progress in tracking topics shifts within dialogue (e.g., Xing & Carenini, 2021; Xu et al., 2021).

Model Construction

Most conversation research does not just examine transcripts. Instead, conversational behavior from transcripts is compared with data from outside the conversation, such as the speakers' gender, when the conversation took place, the terms they negotiated, or how they felt about each other when the conversation ended. This means that feature counts in the turn-level data set (their words) need to be aggregated and merged with the speaker-level data set (other measures outside the conversation). Then a statistical model must be estimated and interpreted. Finally, the results must be reported and benchmarked.

Aggregating conversation features

Although many conversation features are observed at the turn level, other variables of interest may be measured at a higher level, such as at the level of the conversation, individual, dyad, group, organization, or society. Usually, these are measured once per conversation, either as context variables before the conversation (e.g., mood, location, preferences, random assignment to an experimental condition) or outcome variables after it (e.g., enjoyment, learning, negotiated outcomes). However, they can also be measured once per speaker (e.g., demographics) or after multiple conversations in a relationship.

To estimate the links between conversation features and these higher-level measures, turn-level features should be aggregated in some form (e.g., count, average, sum, standard deviation). These aggregations can then be merged to the speaker-level data set using the speaker- and conversation-level unique identifiers (the *dplyr* R package makes this process easier; Wickham et al., 2019).

Aggregation window. Researchers should almost always separate the features of each person in the conversation before analysis (e.g., How many questions did Mary ask?) rather than across the entire transcript (e.g., How many

questions did everyone ask?). This is necessary any time speaker-level variables vary within a conversation, such as occupying different roles, experimental conditions, and demographics.

In addition, researchers may want to aggregate features from only a subset of the conversation. For example, they may remove greetings, off-topic chatter, or final decisions from analysis of a task-focused conversation. In other cases, they may aggregate features only from the beginning of the conversation to focus on each person's behavior before being influenced by the partner's manner of speech or because the meaning of a feature changes at different times (e.g., Y. Li et al., 2022).

Controlling for speaking time. Researchers should be clear about counts versus rates. The total word count of each turn and each conversation is used in many analyses—it is a common and simple benchmark to use for prediction tasks. Other times, feature counts are transformed into feature rates to control for the length of each text (e.g., feature count per minute or per 100 words, which is the default in the LIWC-dictionary approach). Analyses are simplest when word counts are relatively similar across texts. When word-count differences are large, researchers must decide whether the difference is endogenous (i.e., controllable). For example, if someone is studying a mix of 30- and 60-min meetings, then total feature counts would be mainly driven by the prescheduled meeting length. Thus, controlling for the total word count would make it easier to compare language across the two time frames.

Sometimes, total speaking time is an outcome. For example, when people are told to ask more questions, their partner speaks more and enjoys the conversation more (K. Huang et al., 2017). This is not a confound—one reason there is an increase in talking is due to the amount of questions asked. Furthermore, enjoyment early in the conversation can increase talking as the conversation continues. In these cases, it may be better to focus only on the early part of the conversation, before differences in speaking time emerge (Shi et al., 2023). Otherwise, researchers should look at both what and how much is said as two distinct outcomes.

Model estimation

Although a review of the rich existing literature on model estimation (i.e., constructing a statistical model to test a hypothesis) is outside of the scope of this article, we briefly touch on several challenges that are particularly common in conversation research.

Units of observation. Although speakers are given their own row in the speaker-level data set, these are not independent observations. There is often some shared variance

with their partner in the context and outcomes. There is also shared variance when a speaker is present in multiple conversations (e.g., in a round-robin design or when tracking relationships over time) or when outcomes are measured multiple times per conversation (e.g., once per topic). This is commonly addressed by using heteroskedasticity-robust standard errors (e.g., through the *estimat* R package; Zeileis et al., 2020). Researchers who ignore these issues can end up overstating the precision of their estimates and overfit models that are too complex to be estimated well by their data sets (Bertrand et al., 2004; Yeomans, Brooks, et al., 2019).

Interpreting effects. The time course of conversation complicates the interpretation of estimated effects. In particular, we distinguish between “causal” relationships (“What is the effect of X?”), “predictive” relationships (“Will X happen next?”), and “descriptive” relationships (“Did X happen?”). All three have some practical value (J. Kleinberg et al., 2015; Mullainathan & Spiess, 2017), but it is important to know the difference. This is especially difficult in interpersonal interaction because there are many possible third variables that could confound any estimate: Someone’s midconversation behavior could either affect outcomes directly, be correlated with something that affects outcomes, or be an outcome of something that happened earlier in the conversation.

The “gold standard” for causal estimation is a randomized experiment in which at least one speaker is randomly assigned to an intervention that affects some part of the conversational behavior (e.g., try to interrupt a lot vs. try not to interrupt at all) or outcomes the speaker or the speaker’s partner will report (e.g., come with as many ideas as you can vs. choose one idea to pursue). In lieu of experimental control, some empirical approaches can help make causal interpretations more plausible. If speakers have stable conversational tendencies across conversations (e.g., some people always laugh more frequently or have a penchant for arguing), then the random assignment of speakers to their partners can be used as an instrumental variable (Zhang et al., 2020). Researchers have also sharpened their interpretations by focusing on conversation features (as in the Counting Words section) from the beginning of conversations, before speakers are deeply influenced by their partner (e.g., Curhan & Pentland, 2007; Voigt et al., 2017; Zhang et al., 2018). Other common causal inference strategies (e.g., controlling for preconversation variables, matching, event studies) may also be useful (Angrist & Pischke, 2008).

Reporting results

Only a subset of a researcher’s analyses will end up in a final publication. The low cost of additional analyses

can be harnessed to produce a variety of benchmark models, alternative specifications, and robustness checks. Although it is often tempting to report only the positive results, these other analyses are often more useful when they produce negative results because they highlight limitations and boundary conditions.

Although not all of these additional analyses need to make the main body of the article, online appendices often have no word limit. In addition, researchers who share their analysis code and data can encourage their readers to explore alternative models themselves. At the very least, researchers should conduct and report basic sanity checks—for instance, that their results cannot be obtained using simpler text analysis, such as word-count or sentiment analysis.

Benchmarks. Often researchers are focused on a particular variable (e.g., question-asking), and they may want to demonstrate that the variable has a uniquely strong relationship with the outcome of interest. However, because conversation data are complex, there are many potential comparisons that can be constructed.

Instead, researchers should always give context to their focal model with some reasonable set of benchmark models (e.g., Eichstaedt et al., 2021; Yeomans, 2021). For example, computer-science articles routinely include tables comparing the performance of many models on the same data set. Because conversation data are rich, benchmarks could be drawn from contextual data or from other features of the transcript. Another approach to check the importance of a single feature is called an “ablation test.” There, a feature is removed from a more complex model—if the performance of the new model decreases, then the removed feature is considered essential for the original model.

Similar concerns arise when selecting control variables. There are many ways to define a model specification using conversation data, and researchers may find value in estimating alternative models to demonstrate robustness—sometimes called a “multiverse” or “specification curve” analysis (Schweinsberg et al., 2021; Simonsohn et al., 2020). The most reliable results will hold not only across individual specifications within a data set but also across data sets and contexts.

Confirmatory versus exploratory results. The high dimensions of text allow for near-infinite researcher degrees of freedom (Yeomans, 2021). This means the standard concerns about p-hacking, data-dependent modeling choices, and nonreplicability should be especially important for conversation research. Best practices include pre-registering NLP analyses whenever possible—including exact analysis code, detailed information on what data are collected, and how the sample will be determined (Nelson et al., 2018). Likewise, researchers should be wary of

assuming generalizability for models that have been tested in only one data set or one context. However, exploratory results can be tremendously useful (H. K. Collins et al., 2021; D. A. Moore, 2016). Thus, we recommend a balanced approach that prioritizes preregistered results where possible as a complement to (rather than to the exclusion of) well-grounded exploratory work.

When researchers publish results that have not been preregistered, they can still take steps to enhance the credibility of their findings. For example, they can separate validation analyses from their extraction and estimation strategies using cross-validation or split samples within their data set (Poldrack et al., 2020). Although a common default for these validation checks assigns data into training and testing folds randomly, researchers may find added value from nonrandom splits (Weiss et al., 2016). For instance, they could assign data to training and testing at the level of conversations (so that all speakers within a single conversation are all in the same fold together) or the level of speakers (so that when a speaker appears in multiple conversations, all of the speaker's conversations are grouped into the same fold together). This is also relevant when researchers have data across a large time span. For example, researchers who want to forecast stock prices from CEO interviews might train on data from 2010 to 2020 and then test their model on data from 2021 to 2022 so that their model is tested on a simulation of its eventual application: seeing into the future. Other examples might be training and testing on different company types, countries, or CEO characteristics (e.g., gender). These nonrandom splits allow researchers to make stronger claims about the robustness and generalizability of their conclusions.

Data Sharing

Collecting and cleaning conversation data for academic research can be costly in terms of time and money. This can make conversation research prohibitive for early-career scholars and privilege scholars from well-resourced institutions. Moreover, costs may lead individual researchers to be reluctant to share their data with others who did not bear those costs themselves. However, we think this reluctance could be holding conversation science back—it is the costliness of collecting conversation data that makes its sharing especially valuable and productive. The field will be better off if researchers establish norms to share their materials, data, and code openly. We hope to encourage a more cumulative, inclusive, and collaborative research community. To this end, in our own work, we have shared as much of our conversation data as we can. Furthermore, our own research has directly benefited from the generosity of others who were willing to share their data and analyses (e.g., K. Huang et al., 2017; Ranganath et al., 2009).

Open-science practices are important (National Academies of Sciences, Engineering, and Medicine, 2018), and we think they are especially important for conversation science (Reece et al., 2022). First, conversation is so multifaceted that the same data set can be used to answer many research questions, beyond the scope of the initial research question of the researchers who collected the data. Second, the upfront costs of collecting and cleaning large-sample conversation data are immense and may be prohibitive for some researchers. Third, the upfront costs of the analysis are also quite high, so researchers can quickly build on one another's work by publishing reproducible code that can be shared and improved. Finally, individual hypotheses can be more robustly tested if analyses and results can be replicated over multiple data sets that may have been collected in different contexts.

Data privacy

There are barriers to openly sharing data. In our view, the most common and legitimate concern is privacy. Many common privacy issues are exacerbated in conversation research because conversation data sets include identifiable data (Cychosz et al., 2020; Rubinstein & Hartzog, 2016). When conversations are recorded on video and/or audio, these rich media make it easier for subjects to be identified. Furthermore, even the transcripts of conversations can contain revealing details about a person that could be identifiable, either individually or in combination (Sweeney, 2002). These are essential questions for researchers to grapple with, and although there are more extensive treatments of the relevant issues (e.g., Meyer, 2018; Robbins, 2017), we highlight the main concerns.

Preventive measures. The most important step in accounting for privacy is to obtain explicit consent from participants. In practice, we have found that researchers often fail to anticipate future data-sharing needs and are not clear in asking for permission to store and to share deidentified data. Participants and Institutional Review Boards (IRBs) rarely blanch at these requests in consent forms because it is increasingly an essential part of the research process. Furthermore, an explicit warning about sharing may prompt participants not to share anything truly private.

It is worth assessing the importance of individuating information for the research question. For example, if researchers are studying performance during a negotiation simulation in which the particulars are assigned at random in the case materials, then the speakers' true persona (including names, demographics, and location) are irrelevant to many research questions. In these cases, researchers should directly ask participants to

refrain from providing any identifying information before the conversation begins. However, this restriction can interfere with some research questions. Consider two examples—doctor-patient conversations and speed-dating conversations—in which personal information is essential to the goals of the speakers. In these cases, researchers cannot reasonably ask speakers not to share personal information.

Deidentification. It is best practice to anonymize conversation data sets when possible. This is especially important for conversation data because it is open-ended: During a conversation, people can say virtually anything. If data are to be shared for public use (which we encourage), it is essential that the text be completely deidentified. Many feature-extraction techniques automatically remove identifying information. For example, if an *n*-gram model is used and all *n*-grams that occur less than 1% of the time are removed, this will mechanically remove any individuating information (as long as no individual makes up more than 1% of the data).

Anonymizing raw text is more challenging. This can be done manually—by a human coder reading through each transcript and removing any identifiers—or automatically. For example, there are software packages that can deidentify most data by replacing named entities (e.g., specific names, addresses) with generic tags, although no algorithmic method is perfect (B. Kleinberg, 2023; Mendels et al., 2018). Like transcription, the best approach may be hybrid—using an algorithm as a first pass at anonymization followed by a human check to handle the identifiable information most difficult to detect.

Some conversation data are especially difficult to anonymize (e.g., audio or video data). We are not aware of any robust method for automatically deidentifying video or audio data; it may be better to simply focus on sharing transcriptions and turn-level extracted features (metadata) rather than the complete or raw data. Likewise, even transcripts can be difficult to anonymize. For example, a real-estate negotiation will likely reveal identifying features of the property in question, which can then be linked to other public records. In these cases, we still encourage researchers to share the turn-level data set with the text removed, leaving only the unique identifiers and the extracted features. Note, however, that this is not always a guarantee of deidentification. It is possible that text or demographic variables (e.g., gender) could be reconstructed from the feature counts. This is primarily a risk for very elaborate feature extraction (e.g., sentence embeddings), whereas it is exceedingly unlikely to be an issue with simpler features (e.g., counts of pauses or questions).

We encourage researchers to scrutinize the identifiability of the metadata they collect outside the conversation (e.g., demographics). If there is a concern about

these data, they can be deidentified. Common solutions include coarsening variables to broad categories (e.g., reporting age buckets rather than exact age; Samarati & Sweeney, 1998) or perturbing variables by adding noise (e.g., reporting age ± 5 years; Kargupta et al., 2003). This is especially important when researchers combine publicly available text data with nonpublic data, for example, if text from someone's (public) Twitter account is paired with that person's (private) school transcripts. Because the text can be searched, this risks identification of each participant's entire record.

Handling sensitive data. There are unique privacy concerns that arise in many common conversation data settings. Imagine conversations between financial advisors and their clients or between professors and their students. In these cases, researchers must prioritize their responsibilities to protect the rights of the speakers and to uphold the norms of the context in which they were speaking. For example, consent is not always possible to collect from the speakers themselves, and speakers may not be aware of how their data will end up being used.

Many organizations establish their own policies around data sharing. For example, a company may have permission from its users to share data but may not want to make the raw data public because they consider that information proprietary. We strongly encourage researchers to be proactive about this topic when exploring collaborations with outside organizations. Many of the anonymization techniques mentioned above, such as extracting aggregated linguistic features using open-source software (e.g., Yeomans, Kantor, & Tingley, 2018) and using metadata rather than raw transcript data, can be initiated before researchers see any of the data so that no raw text ever leaves the organization.

Depending on their capabilities, organizations may be able to execute analysis code that a researcher writes without ever seeing more than a small example of their internal data. Many feature-extraction algorithms remove identifying information from text (e.g., counts of politeness features). The resulting turn-level feature counts could then be analyzed by researchers and shared publicly along with the code that was used to tally the features.

There are also unique concerns when dealing with text collected from publicly available sources (e.g., social media data or online forums) because there is also a heightened risk that it can be reidentified. If the data set includes metadata that are not publicly available, this creates potential risks for the speakers. For example, if a researcher shares the exact turn-level word embeddings or word counts of entire conversations, that information, although ostensibly anonymized, may be enough to reverse-search and uncover the source of the data. In these cases, researchers may want to increase the

anonymity by adding noise to the extracted feature counts and/or the metadata.

Conclusion

This is an exciting time to be studying conversation, a fundamental activity of the social world. With technological advances, it is becoming easier to collect and analyze large-scale conversation data and to pair turn-level conversation data with speaker-level data containing more traditional survey and behavioral measures. Still, collecting and analyzing text data and combining turn-level and speaker-level data sets present unique challenges. The complexities of this domain provide opportunities for researchers to build a community of inquiry that shares methods, tools, and data and strives for an ever-growing, cumulative science of conversation.

Appendix

Comparative analysis of 10 popular transcription services

There are many automatic-transcription services available, and each service varies in effectiveness on different dimensions. In September 2020, we tested 10 of the most popular transcription services on the market and rated them along five dimensions: (a) transcription accuracy, (b) speaker differentiation, (c) incorporation of time stamps, (d) user-friendliness, and (e) pricing details.

The transcription services that we tested in September 2020 were Otter, Temi, Amberscript, Descript, Trint, Sonix, Happy Scribe, Wreally, Ebby, and Scribie (see Table A1). These services were determined by aggregating the top-rated auto-generated transcription services identified by *The New York Times* (“The Best Transcription Services, 2018), *PC Magazine* (B. Moore, 2018), *TechRadar* (DeMuro & Turner, 2021), and *Poynter* (LaForme, 2018).

Transcription accuracy. To assess transcription accuracy, we first manually transcribed the audio files ourselves. These transcriptions were considered the ground truth. We then generated transcriptions from each service using the same audio files. Each automatic-transcription service that we tested was able to generate the recordings relatively quickly (within 30 min). We then systematically submitted the ground-truth transcriptions along with the matching automatically generated transcriptions to a plagiarism website called CopyLinks. By comparing the two files, we were able to determine the accuracy of the service: The calculated “text accuracy” score represents the percentage of words that overlap between the human-generated and auto-generated transcriptions. The higher the percentage, the better the service was at correctly

transcribing the audio file. Otter and Temi were the most accurate from our testing.

Speaker differentiation. Speaker differentiation is the ability of the transcription service to separate different speakers in the audio recording. Unfortunately, all of the services struggled with auto-generated speaker differentiation, but most allowed you to manually edit and correct this using the service’s main platform. None of the services were able to differentiate by speaker in a consistently accurate way; all required human correction.

Incorporation of time stamps. Incorporation of time stamps refers to whether each transcription service included the time when each speaker began a turn. Each of the top five services that we identified included time stamps at each turn, and most of the services allowed for easy adjustment in the service’s main platform. When speaker differentiation was manually corrected in each service’s platform, most of the time stamps were automatically updated to accurately mark the beginning or end of a turn. Trint was one of the best services at providing accurate time stamps, with the possibility to apply a 120-frame rate.

User-friendliness. User-friendliness refers to how easy each service’s platform was to use and navigate (upload files, edit within the auto-generated transcripts, export final transcripts, etc.). The final data format ended up being especially important. Most transcription services export transcripts as document files (i.e., Microsoft Word, PDF, or text format),³ which all require an additional conversion to a tabular file (i.e., CSV or Microsoft Excel) for analysis. Usually, it is not hard to write cleaning code to parse the Word, PDF, or text files, although this depends on the exact formatting of the transcription service. Subtitle file formats (.VTT files) have also been used by services such as Zoom to map transcribed utterances to time stamps, and there is an R package that can process these files into tabular formats automatically (Knight, 2023).

Pricing details. Pricing details refers to how much each service costs to use. There is quite a bit of variation among the pricing models for each service—some charge a monthly membership fee, and others charge per minute transcribed. In general, those that offered a membership fee tended to have a better overall platform, and those that charged by the minute were more cost-effective.

Approach

We started with all 10 auto-transcription services and five audio files gathered from YouTube. For the first round of testing, we used only audio files that involved two speakers and lasted approximately 10 min. The videos contained a range of accents, ages, and other characteristic differences between the speakers (see Table A2).

Table A1. Transcription Services

Service name	Transcription accuracy	Speaker differentiation	Inclusion of time stamps	User-friendliness	Pricing details
Otter ^a	79.47% match	Must edit speakers in the main service platform (unless paired with Zoom).	Time stamps are present at each turn. It is possible to edit within the platform, and the time stamps will automatically populate based on your edits.	Otter is a great service, especially when it is paired with Zoom (in which it is extremely accurate in speaker differentiation). It also allows for easy edits within the document. The options to export are with a Word document, PDF, text file, SRT, or clipboard. There are a decent number of options in terms of platform capabilities as well.	Offers Otter for individuals, teams, and educators; pricing for a basic plan (free, 600 min per month) with 3 free uploads; premium, and team plans. Also partnered with Zoom.
Temi ^a	75.21% match	Must edit speakers in the main service platform.	Time stamps are present at each turn. It is possible to edit within the platform, and the time stamps will automatically populate based on your edits.	Temi is relatively cost-effective to use, and the output is not only one of the most accurate, but it also allows for easy edits within the document. The options to export are with a Word document, PDF, or text file. There are not as many options in terms of the platform itself, but we think this is definitely a top runner for ease and basic needs.	\$0.25 per audio minute
Amberscript ^a	71.57% match	Must edit speakers in the main service platform. Distinguishing speakers is moderately better when you provide the amount of people in the interaction.	Time stamps are present at each turn. It is possible to edit within the platform, and the time stamps will populate based on your edits.	Amberscript is a good service. It allows for easy edits within the document. The options to export are with a Word document, JSON, or text file. There are not as many options in terms of the platform, but it covers the basics.	Need to create a free account. Pricing is per month (\$25), per hour (\$10), or per minute (\$1.40).
Descript ^a	70.89% match	Must edit speakers in the main service platform.	Time stamps are present at each turn. It is possible to edit within the platform, and the timestamps will populate based on your edits. This service also has the most options regarding time stamps.	Descript is a great service as well. It allows for easy edits within the document and has some of the best capabilities. The options to export are with a Word document or RTF file. You download this service onto your computer, and the platform is pretty straightforward to use.	Offers a one-time trial of 3 transcript hours. Then pricing is \$10 to \$15 per month. Also have a manual/human option.

(continued)

Table A1. (continued)

Service name	Transcription accuracy	Speaker differentiation	Inclusion of time stamps	User-friendliness	Pricing details
Trint ^a	68.13% match	Must edit speakers in the main service platform.	Time stamps are present at each turn. It is possible to edit within the platform, and the time stamps will populate based on your edits. This service also offers the most accurate time stamp (120-frame rate).	Trint is a great service as well. It allows for easy edits within the document. This service has one of the most options to export: Word document, SRT, VTT, STL, EDL, HTML, XML, CSV, or text file. There are also a decent number of options in terms of platform capabilities.	Offers Trint for individuals, teams, and enterprises; pricing for starter (\$48/month or \$60/month depending on payment style; 86 transcriptions per year), advanced, pro, pro team, enterprise.
Sonix	75.86% match	Must edit speakers in the main service platform.	Time stamps are present at each turn. It is possible to edit within the platform, and the time stamps will populate based on your edits.	Sonix is a little bit more accurate than Trint, and it has a similar number of options in terms of platform capabilities. It allows for easy edits within the document and is pretty straightforward. We rank Trint higher because the time stamps tend to be a bit better. This service has one of the most options to export: Word document, PDF, SRT, VTT, SESX, XML, FCPHML, or text file.	Need to create a free account. Pricing is standard (\$10/hr), premium, and enterprise.
Happy Scribe	74.6% match	Must edit speakers in the main service platform.	Time stamps are present at each turn and are pretty close as well (time frame looks high). It is possible to edit within the platform, and the time stamps will populate based on your edits.	HappyScribe is another good option but was simply not as accurate as the other services. It has a similar number of options in terms of platform capabilities. It allows for easy edits within the document and has accurate time stamps. This service has one of the most options to export: Word document, PDF, SRT, VTT, SESX, XML, FCPHML, STL, HTML, XML, JSON, or text file.	Pay for what you need; \$12 per hour for the first 25 hr.
Wreally	79.5% match	Must edit speakers in the main service platform.	Time stamps are present at determined intervals. It is possible to add a time stamp, but they tend to not be as accurate.	Wreally is one of the most accurate services with text comparison, but it is not a great service overall. It allows for edits within the document, but it is not as clear and easy as with other services. The options to export are only with a text file. There are a limited number of options in terms of platform capabilities.	Automatic transcription is \$20/year + \$6 per audio hour

(continued)

Table A1. (continued)

Service name	Transcription accuracy	Speaker differentiation	Inclusion of time stamps	User-friendliness	Pricing details
Ebby	67.58% match	Distinguishing speakers is not great; can edit within the document, but automatic generation is not great.	Time stamps are present at each turn. It is possible to edit within the platform, and the time stamps will populate based on your edits.	Ebby is one of the least accurate services with text comparison, although it does have a decent amount of platform capabilities. It allows for edits within the document. The options to export are Word document, PDF, SRT, VTT, HTML, or text file. This is a decent service but just not as accurate as the others.	Two free previews (not the entire transcript). The service is \$0.10 per audio minute.
Scribie	58.36% match	Distinguishing speakers is not great; can edit within the document, but automatic generation is not great.	Time stamps are present only where you place them. It is possible to add a time stamp, but they tend to not be as accurate.	Scribie is one of the least accurate services with text comparison and does not have that many platform capabilities. It allows for edits within the document with its online editor. The options to export are Word document, PDF, SRT, VTT, ODT, or text file. This is a decent service but just not as accurate as the others.	\$0.10 per minute. Also has a manual option.

^aOne of the top-five services.

Table A2. YouTube Videos

Name	Details	Round	Length	Speakers	YouTube link
Anthony Rizzo on Chicago Cubs Rivalries & Baseball Superstitions While Eating Spicy Wings Hot Ones	Hot Ones is a YouTube series where host Sean Evans asks celebrity guests questions while they eat chicken wings coated in ever-spicier hot sauce.	1	10:17	2	https://www.youtube.com/watch?v=4iSCOtYs_6Q&list=PLAzrgbu8gEMIfV-k5JA89NP3j2JGsFapZ&index=3&t=0s
Christian Woman Interview-Shannon	Interview and portrait of Shannon, a Christian woman in West Virginia who shares her life story.	1	10:36	2	https://www.youtube.com/watch?v=gB6OUSSCdkQ
Penelope Cruz Times Talks with <i>The New York Times</i>	From <i>The New York Times</i> : Logan Hill, <i>The New York Times</i> contributor, talks with Penélope Cruz about her latest film, “Mama,” and her many wide-ranging roles and challenges and opportunities for women in film.	1	9:13	2	https://www.youtube.com/watch?v=PNTGr-dAF0g
Unorthodox: Deborah Feldman’s Escape From Brooklyn to Berlin DW Interview	From <i>DW News</i> : Deborah Feldman tells the story of her life as an ultra-orthodox Jew and her rejection of Hasidic traditions.	1	11:58	2	https://www.youtube.com/watch?v=vvxzIXSAPyw
Shah Rukh Khan, Bollywood Star Journal Interview	From <i>DW News</i> : Interview of Shah Rukh Khan on Bollywood: illusion and reality in India’s Dream factory.	1	9:16	2	https://www.youtube.com/watch?v=J1s1yj4u9I4

(continued)

Table A2. (continued)

Name	Details	Round	Length	Speakers	YouTube link
Suranne Jones' Greatest Weaknesses Are Coffee & Hugh Jackman <i>The Graham Norton Show</i>	<i>The Graham Norton</i> show is a British talk show. Guests on this episode are Suranne Jones, Hugh Jackman, Zack Efron, and Zendaya.	2	10:27	5	https://www.youtube.com/watch?v=rEL3W9xNn9U
From Child Bride to Global Voice for Women's Empowerment, Dr. Tererai Trent Shares Her Journey	From <i>CBS This Morning</i> : Dr. Tererai Trent, the author of <i>The Awakened Woman: Remembering & Reigniting Our Sacred Dream</i> , shares her incredible journey from child bride in Zimbabwe to achieving her doctorate.	2	12:35	2	https://www.youtube.com/watch?v=dDqq-cgG9F4
Pimp and Prostitute Interview - Master J and Little Mama	Soft White Underbelly interview and portrait of "Master J" and "Little Mama" in South Central Los Angeles.	2	16:32	3	https://www.youtube.com/watch?v=1rzBLRp9b34
German Soldier Remembers WW2 Memoirs of WWII #15	From the YouTube series, <i>Memoirs of WWII</i> : Gert Schmitz talks about his life as a German soldier in WWII.	2	13:45	2	https://www.youtube.com/watch?v=3D0Uw-y9mrk
What's It Like Being a Foreigner in Korea? ASIAN BOSS	ASIAN BOSS is a South Korea-based media company, and in this video, they ask people walking by to answer questions.	2	10:49	5+	https://www.youtube.com/watch?v=ldndJslNLgs

After generating the transcriptions of each of the five videos across all 10 platforms, it was clear that some of the transcription services outperformed the others across all dimensions. We narrowed down our list to the top five services and then ran a second round of testing. In this round, we used five different audio files selected following the same criteria as the first round of files except they involved two or more speakers. From this second round of testing, we were then able to decipher our top-five transcription services.

Researchers' choice of transcription service will vary based on their needs and technological improvements and additional services that enter the market.

Transparency

Action Editor: David A. Sbarra

Editor: David A. Sbarra

Author Contribution(s)

Michael Yeomans: Conceptualization; Project administration; Visualization; Writing – original draft.

F. Katelynn Boland: Investigation; Visualization; Writing – review & editing.

Hanne K. Collins: Writing – review & editing.

Nicole Abi-Esber: Writing – review & editing.

Alison Wood Brooks: Funding acquisition; Project administration; Writing – review & editing.

Declaration of Conflicting Interests

The author(s) declared that there were no conflicts of interest with respect to the authorship or the publication of this article.

ORCID iD

Michael Yeomans  <https://orcid.org/0000-0001-5651-5087>

Acknowledgments

This article was much improved by helpful comments on earlier drafts from many other researchers, including (among others) Ken Benoit, Ryan Boyd, Gus Cooney, Morteza Dehghani, Grant Donnelly, Bennett Kleinberg, Andrew Knight, Celia Moore, James Pennebaker, Gillian Sandstrom, Martin Schweinsberg, Lyle Ungar, and Simine Vazire.

Notes

1. For simplicity, we use the term "speaker" to refer generically to all conversation participants throughout.
2. To clarify a common misconception, dialogue refers to any number of speakers—not just two. It is derived from the prefix "dia," meaning through (e.g., diagonal), not the prefix "di," meaning two (e.g., dichotomy).
3. At the time of writing, Trint was the only automatic-transcription service we tested that allowed researchers to export transcriptions as an Excel spreadsheet or CSV file.

References

- Adamson, D., Dyke, G., Jang, H., & Rosé, C. P. (2014). Towards an agile approach to adapting dynamic collaboration support to student needs. *International Journal of Artificial Intelligence in Education*, 24(1), 92–124.
- Angrist, J. D., & Pischke, J. S. (2008). *Mostly harmless econometrics*. Princeton University Press.
- Arora, S., Liang, Y., & Ma, T. (2017, April 24–26). *A simple but tough-to-beat baseline for sentence embeddings* [Conference session]. 5th International Conference on Learning Representations, Toulon, France.
- Ashokkumar, A., & Pennebaker, J. W. (2022). Tracking group identity through natural language within groups. *PNAS Nexus*, 1(2), Article pgac022. <https://doi.org/10.1093/pnas-nexus/pgac022>
- Babcock, M. J., Ta, V. P., & Ickes, W. (2014). Latent semantic similarity and language style matching in initial dyadic interactions. *Journal of Language and Social Psychology*, 33(1), 78–88.
- Backus, M., Blake, T., Pettus, J., & Tadelis, S. (2020). *Communication and bargaining breakdown: An empirical analysis* (No. w27984). National Bureau of Economic Research.
- Beasley, A., & Mason, W. (2015). Emotional states vs. emotional words in social media. In *Proceedings of the ACM Web Science Conference* (pp. 1–10). Association for Computing Machinery.
- Belz, A., Agarwal, S., Shimorina, A., & Reiter, E. (2021). *A systematic review of reproducibility research in natural language processing*. arXiv. <https://doi.org/10.48550/arXiv.2103.07929>
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM conference on fairness, accountability, and transparency* (pp. 610–623). Association for Computing Machinery.
- Benoit, K. (2020). Text as data: An overview. In L. Curini & R. Franzese (Eds.), *The SAGE handbook of research methods in political science and international relations*. Sage. <https://doi.org/10.4135/9781526486387>
- Benoit, K., Conway, D., Lauderdale, B. E., Laver, M., & Mikhaylov, S. (2016). Crowd-sourced text analysis: Reproducible and agile production of political data. *American Political Science Review*, 110(2), 278–295.
- Benoit, K., Munger, K., & Spirling, A. (2019). Measuring and explaining political sophistication through textual complexity. *American Journal of Political Science*, 63(2), 491–508.
- Berger, J., Humphreys, A., Ludwig, S., Moe, W. W., Netzer, O., & Schweidel, D. A. (2020). Uniting the tribes: Using text for marketing insight. *Journal of Marketing*, 84(1), 1–25.
- Berry, D. S., Pennebaker, J. W., Mueller, J. S., & Hiller, W. S. (1997). Linguistic bases of social perception. *Personality and Social Psychology Bulletin*, 23(5), 526–537.
- Berry, M. (2013). Towards a study of the differences between formal written English and informal spoken English. In L. Fontaine, T. Bartlett, & G. O'Grady (Eds.), *Systemic functional linguistics: Exploring choice* (pp. 365–383). Cambridge University Press.
- Bertrand, M., Duflo, E., & Mullainathan, S. (2004). How much should we trust differences-in-differences estimates? *The Quarterly Journal of Economics*, 119(1), 249–275.
- The best transcription services. (2018, October 15). *The New York Times*. <https://www.nytimes.com/wirecutter/reviews/best-transcription-services/>
- Bhatia, S., Richie, R., & Zou, W. (2019). Distributed semantic representations for modeling human judgment. *Current Opinion in Behavioral Sciences*, 29, 31–36.
- Bianchi, F., & Hovy, D. (2021, August). On the gap between adoption and understanding in NLP. In *Findings of the Association for Computational Linguistics: ACL-IJCNLP2021* (pp. 3895–3901). Association for Computational Linguistics
- Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003). Latent Dirichlet allocation. *The Journal of Machine Learning Research*, 3, 993–1022.
- Boghrati, R., Hoover, J., Johnson, K. M., Garten, J., & Dehghani, M. (2018). Conversation level syntax similarity metric. *Behavior Research Methods*, 50(3), 1055–1073.
- Boland, J. E., Fonseca, P., Mermelstein, I., & Williamson, M. (2022). Zoom disrupts the rhythm of conversation. *Journal of Experimental Psychology: General*, 151(6), 1272–1282.
- Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., Bernstein, M. S., Bohg, J., Bosselut, A., Brunskill, E., Brynjolfsson, E., Buch, S., Card, D., Castellon, R., Chatterji, N., Chen, A., Creel, K., Davis, J. Q., Demszy, D., . . . Liang, P. (2021). *On the opportunities and risks of foundation models*. arXiv. <https://doi.org/10.48550/arXiv.2108.07258>
- Boyd, R. L., & Schwartz, H. A. (2021). Natural language analysis and the psychology of verbal behavior: The past, present, and future states of the field. *Journal of Language and Social Psychology*, 40(1), 21–41.
- Boyd-Graber, J., Mimno, D., & Newman, D. (2014). Care and feeding of topic models: Problems, diagnostics, and improvements. In E. M. Airoldi, D. Blei, E. A. Erosheva, & S. E. Fienberg (Eds.), *Handbook of mixed membership models and their applications* (pp. 225–255). Routledge.
- Brodsky, A., Lee, M. J., & Leonard, B. (2022). Discovering new frontiers for dyadic and team interaction studies: Current challenges and an open-source solution—survconf—for increasing the quantity and richness of interactional data. *Academy of Management Discoveries*, 8(3). <https://doi.org/10.5465/amd.2021.0257>
- Broockman, D., & Kalla, J. (2016). Durably reducing transphobia: A field experiment on door-to-door canvassing. *Science*, 352(6282), 220–224.
- Brown, S., Davidovic, J., & Hasan, A. (2021). The algorithm audit: Scoring the algorithms that score us. *Big Data & Society*, 8(1). <https://doi.org/10.1177/2053951720983865>
- Brucks, M. S., & Levav, J. (2022). Virtual communication curbs creative idea generation. *Nature*, 605(7908), 108–112.
- Bunt, H., Alexandersson, J., Carletta, J., Choe, J. W., Fang, A. C., Hasida, K., & Traum, D. (2010, May 17–23). *Towards an ISO standard for dialogue act annotation* [Conference session]. Seventh conference on International Language Resources and Evaluation (LREC'10), Valletta, Malta.
- Caliskan, A., Bryson, J. J., & Narayanan, A. (2017). Semantics derived automatically from language corpora contain human-like biases. *Science*, 356(6334), 183–186.

- Carter, A. J., Croft, A., Lukas, D., & Sandstrom, G. M. (2018). Women's visibility in academic seminars: Women ask fewer questions than men. *PLOS ONE*, 13(9), Article e0202743. <https://doi.org/10.1371/journal.pone.0202743>
- Chang, J., Gerrish, S., Wang, C., Boyd-Graber, J., & Blei, D. (2009, December 7–10). *Reading tea leaves: How humans interpret topic models* [Conference session]. Advances in neural information processing systems 22. Conference on Neural Information Processing Systems, Vancouver, British Columbia, Canada.
- Chang, J. P., Chiam, C., Fu, L., Wang, A., Zhang, J., & Danescu-Niculescu-Mizil, C. (2020, July). ConvoKit: A toolkit for the analysis of conversations. In *Proceedings of the 21th annual meeting of the Special Interest Group on discourse and dialogue* (pp. 57–60). Association for Computational Linguistics.
- Chen, J. V., Nagar, V., & Schoenfeld, J. (2018). Manager-analyst conversations in earnings conference calls. *Review of Accounting Studies*, 23(4), 1315–1354.
- Clark, H., Kirsh, D., Goldin-Meadow, S., & Rogers, Y. (2011, July 20–23). *Interactivity and thought* [Conference session]. CogSci 2011, Boston, Massachusetts, USA.
- Clarke, J. S., Cornelissen, J. P., & Healey, M. P. (2019). Actions speak louder than words: How figurative language and gesturing in entrepreneurial pitches influences investment judgments. *Academy of Management Journal*, 62(2), 335–360.
- Collins, G., Poleski, J., Mehl, M., Tackman, A., Reyes, R., Kraft, A., Russo, J., Kenny, D., Bryan, P., Simons, E., & Casebeer, W. (2018). Building a cognitive profile with a non-intrusive sensor: How speech and sounds map onto our cognitive worlds. *Frontiers in Human Neuroscience*, 12. <https://doi.org/10.3389/conf.fnhum.2018.227.00013>
- Collins, H. K., Hagerty, S. F., Quoidbach, J., Norton, M. I., & Brooks, A. W. (2022). Relational diversity in social portfolios predicts well-being. *Proceedings of the National Academy of Sciences, USA*, 119(43), Article e2120668119. <https://doi.org/10.1073/pnas.2120668119>
- Collins, H. K., Whillans, A. V., & John, L. K. (2021). Joy and rigor in behavioral science. *Organizational Behavior and Human Decision Processes*, 164, 179–191.
- Coltheart, M. (1981). The MRC psycholinguistic database. *The Quarterly Journal of Experimental Psychology, Section A*, 33(4), 497–505.
- Conner, T. S., & Mehl, M. R. (2015). Ambulatory assessment: Methods for studying everyday life. In *Emerging trends in the social and behavioral sciences: An interdisciplinary, searchable, and linkable resource*. Wiley Online Library. <https://doi.org/10.1002/9781118900772.etrds001>
- Cooney, G., Mastroianni, A. M., Abi-Esber, N., & Brooks, A. W. (2020). The many minds problem: Disclosure in dyadic versus group conversation. *Current Opinion in Psychology*, 31, 22–27.
- Curhan, J. R., Overbeck, J. R., Cho, Y., Zhang, T., & Yang, Y. (2022). Silence is golden: Extended silence, deliberative mindset, and value creation in negotiation. *Journal of Applied Psychology*, 107(1), 78–94. <https://doi.org/10.1037/apl0000877>
- Curhan, J. R., & Pentland, A. (2007). Thin slices of negotiation: Predicting outcomes from conversational dynamics within the first 5 minutes. *Journal of Applied Psychology*, 92(3), 802–811.
- Cychosz, M., Romeo, R., Soderstrom, M., Scaff, C., Ganek, H., Cristia, A., Casillas, M., de Barbaro, K., Bang, J. Y., & Weisleder, A. (2020). Longform recordings of everyday life: Ethics for best practices. *Behavior Research Methods*, 52(5), 1951–1969. <https://doi.org/10.3758/s13428-020-01365-9>
- Danescu-Niculescu-Mizil, C., Gamon, M., & Dumais, S. (2011, March). Mark my words! Linguistic style accommodation in social media. In *Proceedings of the 20th international conference on world wide web* (pp. 745–754). Association for Computing Machinery.
- Danescu-Niculescu-Mizil, C., Lee, L., Pang, B., & Kleinberg, J. (2012, April). Echoes of power: Language effects and power differences in social interaction. In *Proceedings of the 21st international conference on world wide web* (pp. 699–708). Association for Computing Machinery.
- Danescu-Niculescu-Mizil, C., Sudhof, M., Jurafsky, D., Leskovec, J., & Potts, C. (2013). A computational approach to politeness with application to social factors. In *51st Annual Meeting of the Association for Computational Linguistics* (pp. 250–259). ACL.
- Danescu-Niculescu-Mizil, C., West, R., Jurafsky, D., Leskovec, J., & Potts, C. (2013, May). No country for old members: User lifecycle and linguistic change in online communities. In *Proceedings of the 22nd international conference on world wide web* (pp. 307–318). Association for Computing Machinery.
- Dastin, J. (2018). Amazon scraps secret AI recruiting tool that showed bias against women. In K. Martin (Ed.), *Ethics of data and analytics* (pp. 296–299). Auerbach Publications.
- de Barbaro, K. (2019). Automated sensing of daily activity: A new lens into development. *Developmental Psychobiology*, 61(3), 444–464.
- Dehghani, M., & Boyd, R. L. (Eds.). (2022). *Handbook of language analysis in psychology*. The Guilford Press.
- Dehghani, M., Khooshabeh, P., Nazarian, A., & Gratch, J. (2015). The subtlety of sound: Accent as a marker for culture. *Journal of Language and Social Psychology*, 34(3), 231–250.
- Demszky, D., Liu, J., Mancenido, Z., Cohen, J., Hill, H., Jurafsky, D., & Hashimoto, T. (2021). *Measuring conversational uptake: A case study on student-teacher interactions*. arXiv. <https://doi.org/10.48550/arXiv.2106.03873>
- DeMuro, J. P., & Turner, B. (2021, March 30). Best transcription services in 2021: Transcribe audio and video into text. *TechRadar*. <https://www.techradar.com/best/best-transcription-services>
- Denny, M. J., & Spirling, A. (2018). Text preprocessing for unsupervised learning: Why it matters, when it misleads, and what to do about it. *Political Analysis*, 26(2), 168–189.
- Denton, E., Díaz, M., Kivlichan, I., Prabhakaran, V., & Rosen, R. (2021). *Whose ground truth? Accounting for individual and collective identities underlying dataset annotation*. arXiv. <https://doi.org/10.48550/arXiv.2112.04554>
- Devlin, J., Chang, M. W., Lee, K., & Toutanova, K. (2018). *Bert: Pre-training of deep bidirectional transformers for language understanding*. arXiv. <https://doi.org/10.48550/arXiv.1810.04805>

- Diener, E., & Seligman, M. E. (2002). Very happy people. *Psychological Science*, 13(1), 81–84.
- Doshi-Velez, F., Kortz, M., Budish, R., Bavitz, C., Gershman, S., O'Brien, D., Scott, K., Schieber, S., Waldo, J., Weinberger, D., Weller, A., & Wood, A. (2017). *Accountability of AI under the law: The role of explanation*. arXiv. <https://doi.org/10.48550/arXiv.1711.01134>
- Doyle, G., & Frank, M. C. (2016, August). Investigating the sources of linguistic alignment in conversation. In *Proceedings of the 54th annual meeting of the Association for Computational Linguistics* (Vol. 1: Long Papers, pp. 526–536). Association for Computational Linguistics.
- Doyle, G., Goldberg, A., Srivastava, S., & Frank, M. C. (2017, July). Alignment at work: Using language to distinguish the internalization and self-regulation components of cultural fit in organizations. In *Proceedings of the 55th annual meeting of the Association for Computational Linguistics* (Vol. 1: Long Papers, pp. 603–612). Association for Computational Linguistics.
- Dunbar, R. I., Marriott, A., & Duncan, N. D. (1997). Human conversational behavior. *Human Nature*, 8(3), 231–246.
- Dupas, P., Modestino, A. S., Niederle, M., & Wolfers, J. (2021). *Gender and the dynamics of economics seminars* (No. w28494). National Bureau of Economic Research.
- Eichstaedt, J. C., Kern, M. L., Yaden, D. B., Schwartz, H. A., Giorgi, S., Park, G., Hagan, C. A., Tobolsky, V. A., Smith, L. K., Buffone, A., Iwry, J., Seligman, M. E. P., & Ungar, L. H. (2021). Closed-and open-vocabulary approaches to text analysis: A review, quantitative comparison, and recommendations. *Psychological Methods*, 26(4), 398–427. <https://doi.org/10.1037/met0000349>
- Epley, N., & Schroeder, J. (2014). Mistakenly seeking solitude. *Journal of Experimental Psychology: General*, 143(5), 1980–1999.
- Errattahi, R., El Hannani, A., & Ouahmane, H. (2018). Automatic speech recognition errors detection and correction: A review. *Procedia Computer Science*, 128, 32–37.
- Fitzsimons, G. M., & Finkel, E. J. (2018). Goal transactivity. In P. A. M. Van Lange, A. W. Kruglanski, & E. T. Higgins (Eds.), *Social psychology: Handbook of basic principles* (3rd ed., pp. 202–221). The Guilford Press.
- Flake, J. K., & Fried, E. I. (2020). Measurement schmeasurement: Questionable measurement practices and how to avoid them. *Advances in Methods and Practices in Psychological Science*, 3(4), 456–465.
- Fox Tree, J. E. (2010). Discourse markers across speakers and settings. *Language and Linguistics Compass*, 4(5), 269–281.
- Frankel, R., Jennings, J., & Lee, J. (2022). Disclosure sentiment: Machine learning vs. dictionary methods. *Management Science*, 68(7), 5514–5532.
- Fu, L., Danescu-Niculescu-Mizil, C., & Lee, L. (2016). *Tie-breaker: Using language models to quantify gender bias in sports journalism*. arXiv. <https://doi.org/10.48550/arXiv.1607.03895>
- Fusaroli, R., Bahrami, B., Olsen, K., Roepstorff, A., Rees, G., Frith, C., & Tylén, K. (2012). Coming to terms: Quantifying the benefits of linguistic coordination. *Psychological Science*, 23(8), 931–939.
- Galley, M., McKeown, K., Fosler-Lussier, E., & Jing, H. (2003). Discourse segmentation of multi-party conversation. In *Proceedings of the 41st annual meeting on Association for Computational Linguistics* (Vol. 1, pp. 562–569). Association for Computational Linguistics.
- Ganek, H., & Eriks-Brophy, A. (2016, November). The Language ENvironment Analysis (LENA) system: A literature review. In *Proceedings of the joint workshop on NLP for Computer Assisted Language Learning and NLP for Language Acquisition* (pp. 24–32). LiU Electronic Press.
- Garfinkel, H. (1956). Conditions of successful degradation ceremonies. *American Journal of Sociology*, 61(5), 420–424.
- Garten, J., Hoover, J., Johnson, K. M., Boghrati, R., Iskiwitsch, C., & Dehghani, M. (2018). Dictionaries and distributions: Combining expert knowledge and large scale textual data content analysis. *Behavior Research Methods*, 50(1), 344–361.
- Gentzkow, M., Kelly, B., & Taddy, M. (2019). Text as data. *Journal of Economic Literature*, 57(3), 535–574.
- Giles, H., Coupland, N., & Coupland, J. (1991). Accommodation theory: Communication, context, and consequence. In H. Giles, J. Coupland, & N. Coupland (Eds.), *Contexts of accommodation: Developments in applied sociolinguistics* (pp. 1–68). Cambridge University Press.
- Gilpin, L. H., Bau, D., Yuan, B. Z., Bajwa, A., Specter, M., & Kagal, L. (2018, October). Explaining explanations: An overview of interpretability of machine learning. In *2018 IEEE 5th International Conference on data science and advanced analytics (DSAA)* (pp. 80–89). IEEE.
- Goffman, E. (1981). *Forms of talk*. University of Pennsylvania Press.
- Goodman, N. D., & Frank, M. C. (2016). Pragmatic language interpretation as probabilistic inference. *Trends in Cognitive Sciences*, 20(11), 818–829.
- Greenwald, A. G., & Banaji, M. R. (1995). Implicit social cognition: Attitudes, self-esteem, and stereotypes. *Psychological Review*, 102(1), 4–27.
- Grice, H. P. (1975). Logic and conversation. In P. Cole & J. L. Morgan (Eds.), *Speech acts* (pp. 41–58). Brill.
- Grimmer, J., & Stewart, B. M. (2013). Text as data: The promise and pitfalls of automatic content analysis methods for political texts. *Political Analysis*, 21(3), 267–297.
- Hamilton, W. L., Clark, K., Leskovec, J., & Jurafsky, D. (2016). Inducing domain-specific sentiment lexicons from unlabeled corpora. In *Proceedings of the conference on empirical methods in natural language processing* (pp. 595–605). Association for Computational Linguistics.
- Hansen, S., & Ash, E. (2023). Text algorithms in economics. *Annual Review of Economics*, 15, 659–688.
- Hansen, S., McMahon, M., & Prat, A. (2018). Transparency and deliberation within the FOMC: A computational linguistics approach. *The Quarterly Journal of Economics*, 133(2), 801–870.
- Hart, R. P. (2001). Redeveloping DICTION: Theoretical considerations. In M. D. West (Ed.), *Theory, method, and practice in computer content analysis* (pp. 43–60). Springer.
- Hearst, M. A. (1997). TextTiling: Segmenting text into multi-paragraph subtopic passages. *Computational Linguistics*, 23(1), 33–64.
- Heritage, J. (2008). Conversation analysis as social theory. In B. S. Turner (Ed.), *The new Blackwell companion to social theory* (pp. 300–320). Wiley.

- Hirschberg, J., & Manning, C. D. (2015). Advances in natural language processing. *Science*, 349(6245), 261–266.
- Honnibal, M., & Johnson, M. (2015, September). An improved non-monotonic transition system for dependency parsing. In *Proceedings of the 2015 conference on empirical methods in natural language processing* (pp. 1373–1378). Association for Computational Linguistics.
- Huang, K., Yeomans, M., Brooks, A. W., Minson, J., & Gino, F. (2017). It doesn't hurt to ask: Question-asking increases liking. *Journal of Personality and Social Psychology*, 113(3), 430–452.
- Huang, M., Zhu, X., & Gao, J. (2020). Challenges in building intelligent open-domain dialog systems. *ACM Transactions on Information Systems (TOIS)*, 38(3), 1–32.
- Hutto, C., & Gilbert, E. (2014). Vader: A parsimonious rule-based model for sentiment analysis of social media text. *Proceedings of the International AAAI Conference on Web and Social Media*, 8(1), 216–225.
- Ireland, M. E., Slatcher, R. B., Eastwick, P. W., Scissors, L. E., Finkel, E. J., & Pennebaker, J. W. (2011). Language style matching predicts relationship initiation and stability. *Psychological Science*, 22(1), 39–44.
- Jäckel, E., Zerres, A., Hemshorn de Sanchez, C. S., Lehmann-Willenbrock, N., & Hüffmeier, J. (2022). NegotiAct: Introducing a comprehensive coding scheme to capture temporal interaction patterns in negotiations. *Group & Organization Management*. Advance online publication. <https://doi.org/10.1177/10596011221132600>
- Jackson, J. C., Watts, J., List, J. M., Puryear, C., Drabble, R., & Lindquist, K. A. (2022). From text to thought: How analyzing language can advance psychological science. *Perspectives on Psychological Science*, 17(3), 805–826. <https://doi.org/10.1177/17456916211004899>
- Jacobi, T., & Schweers, D. (2017). Justice, interrupted: The effect of gender, ideology, and seniority at Supreme Court oral arguments. *Virginia Law Review*, 103, 1379–1496.
- Jaidka, K., Giorgi, S., Schwartz, H. A., Kern, M. L., Ungar, L. H., & Eichstaedt, J. C. (2020). Estimating geographic subjective well-being from Twitter: A comparison of dictionary and data-driven language methods. *Proceedings of the National Academy of Sciences, USA*, 117(19), 10165–10171.
- Jaques, N., Lazaridou, A., Hughes, E., Gulcehre, C., Ortega, P., Strouse, D. J., Leibo, J. Z., & De Freitas, N. (2019). Social influence as intrinsic motivation for multi-agent deep reinforcement learning. *Proceedings of the 36th International Conference on Machine Learning, PMLR* 97, 3040–3049.
- Jeong, M., Minson, J., Yeomans, M., & Gino, F. (2019). Communicating with warmth in distributive negotiations is surprisingly counterproductive. *Management Science*, 65(12), 5813–5837.
- Jia, R., & Liang, P. (2017). *Adversarial examples for evaluating reading comprehension systems*. arXiv. <https://doi.org/10.48550/arXiv.1707.07328>
- Jurafsky, D., & Martin, J. H. (2017). *Speech and language processing* (Vol. 4). Pearson.
- Kahneman, D. (2002). Maps of bounded rationality: A perspective on intuitive judgment and choice. *Nobel Prize Lecture*, 8(1), 351–401.
- Kaplan, D. M., Rentscher, K. E., Lim, M., Reyes, R., Keating, D., Romero, J., Shah, A., Smith, A. D., York, K. A., Milek, A., Tackman, A. M., & Mehl, M. R. (2020). Best practices for Electronically Activated Recorder (EAR) research: A practical guide to coding and processing EAR data. *Behavior Research Methods*, 52(4), 1538–1551. <https://doi.org/10.3758/s13428-019-01333-y>
- Kargupta, H., Datta, S., Wang, Q., & Sivakumar, K. (2003, November). On the privacy preserving properties of random data perturbation techniques. In *Third IEEE International Conference on Data Mining* (pp. 99–106). IEEE.
- Kiritchenko, S., & Mohammad, S. (2017, July). Best-worst scaling more reliable than rating scales: A case study on sentiment intensity annotation. In *Proceedings of the 55th annual meeting of the Association for Computational Linguistics* (Vol. 2: Short Papers; pp. 465–470). Association for Computational Linguistics.
- Kleinberg, B. (2023). *Textwasb* [Software].
- Kleinberg, J., Ludwig, J., Mullainathan, S., & Obermeyer, Z. (2015). Prediction policy problems. *American Economic Review*, 105(5), 491–495.
- Knight, A. (2023). *zoomGroupStats* (R package).
- Koenecke, A., Nam, A., Lake, E., Nudell, J., Quartey, M., Mengesha, Z., Touns, C., Rickford, J. R., Jurafsky, D., & Goel, S. (2020). Racial disparities in automated speech recognition. *Proceedings of the National Academy of Sciences, USA*, 117(14), 7684–7689. <https://doi.org/10.1073/pnas.1915768117>
- Kordzadeh, N., & Ghasemaghaei, M. (2022). Algorithmic bias: Review, synthesis, and future research directions. *European Journal of Information Systems*, 31(3), 388–409.
- Koshiyama, A., Kazim, E., & Treleaven, P. (2022). Algorithm auditing: Managing the legal, ethical, and technological risks of artificial intelligence, machine learning, and associated algorithms. *Computer*, 55(4), 40–50.
- Kross, E., Verduyn, P., Boyer, M., Drake, B., Gainsburg, I., Vickers, B., Ybarra, O., & Jonides, J. (2019). Does counting emotion words on online social networks provide a window into people's subjective experience of emotion? A case study on Facebook. *Emotion*, 19(1), 97–107. <https://doi.org/10.1037/emo0000416>
- LaForme, R. (2018, November 16). The best automatic transcription tools for journalists. *Poynter*. <https://www.poynter.org/tech-tools/2017/the-best-automatic-transcription-tools-for-journalists/>
- Landauer, T. K., & Dumais, S. T. (1997). A solution to Plato's problem: The latent semantic analysis theory of acquisition, induction, and representation of knowledge. *Psychological Review*, 104(2), 211–240.
- Lapakko, D. (1997). Three cheers for language: A closer examination of a widely cited study of nonverbal communication. *Communication Education*, 46(1), 63–67.
- Levinson, S. C. (2016). Turn-taking in human communication—origins and implications for language processing. *Trends in Cognitive Sciences*, 20(1), 6–14.
- Lewis, M., Yarats, D., Dauphin, Y. N., Parikh, D., & Batra, D. (2017). *Deal or no deal? End-to-end learning for negotiation dialogues*. arXiv. <https://doi.org/10.48550/arXiv.1706.05125>
- Li, H. Z., Krysko, M., Desroches, N. G., & DEagle, G. (2004). Reconceptualizing interruptions in physician-patient interviews: Cooperative and intrusive. *Communication and Medicine*, 1(2), 145–157.

- Li, Y., Packard, G., & Berger, J. (2020). Dynamically solving the self-presenter's paradox: When customer care should be warm vs. competent. In J. Argo, T. M. Lowrey, and H. J. Schau (Eds.), *NA - Advances in consumer research: Vol. 48* (pp. 981–986). Association for Consumer Research.
- Lieberman, M., & Cieri, C. (1998). The creation, distribution and use of linguistic data: The case of the linguistic data consortium. In *Proceedings of the 1st international conference on language resources and evaluation (LREC)* (pp. 159–164). European Language Resources Association.
- Lipton, Z. C. (2018). The mythos of model interpretability: In machine learning, the concept of interpretability is both important and slippery. *Queue*, 16(3), 31–57.
- Loughran, T., & McDonald, B. (2016). Textual analysis in accounting and finance: A survey. *Journal of Accounting Research*, 54(4), 1187–1230.
- Madsen, A., Reddy, S., & Chandar, S. (2021). *Post-hoc interpretability for neural nlp: A survey*. arXiv. <https://doi.org/10.48550/arXiv.2108.04840>
- Manning, C. D., Clark, K., Hewitt, J., Khandelwal, U., & Levy, O. (2020). Emergent linguistic structure in artificial neural networks trained by self-supervision. *Proceedings of the National Academy of Sciences, USA*, 117(48), 30046–30054.
- Manning, C. D., Surdeanu, M., Bauer, J., Finkel, J. R., Bethard, S., & McClosky, D. (2014, June). The Stanford CoreNLP natural language processing toolkit. In *Proceedings of 52nd annual meeting of the Association for Computational Linguistics: System demonstrations* (pp. 55–60). Association for Computational Linguistics.
- McCabe, R., & Healey, P. G. (2018). Miscommunication in doctor–patient communication. *Topics in Cognitive Science*, 10(2), 409–424.
- Mehl, M. R. (2017). The electronically activated recorder (EAR): A method for the naturalistic observation of daily social behavior. *Current Directions in Psychological Science*, 26(2), 184–190.
- Mehl, M. R., & Pennebaker, J. W. (2003). The sounds of social life: A psychometric analysis of students' daily social environments and natural conversations. *Journal of Personality and Social Psychology*, 84, 857–870.
- Mehl, M. R., Pennebaker, J. W., Crow, D. M., Dabbs, J., & Price, J. H. (2001). The Electronically Activated Recorder (EAR): A device for sampling naturalistic daily activities and conversations. *Behavior Research Methods, Instruments, & Computers*, 33(4), 517–523.
- Mehl, M. R., Vazire, S., Holleran, S. E., & Clark, C. S. (2010). Eavesdropping on happiness: Well-being is related to having less small talk and more substantive conversations. *Psychological Science*, 21(4), 539–541.
- Meier, T., Boyd, R. L., Mehl, M. R., Milek, A., Pennebaker, J. W., Martin, M., Wolf, M., & Horn, A. B. (2021). (Not) lost in translation: Psychological adaptation occurs during speech translation. *Social Psychological and Personality Science*, 12(1), 131–142. <https://doi.org/10.1177/1948550619899258>
- Mendelberg, T., & Karpowitz, C. F. (2016). Power, gender, and group discussion. *Political Psychology*, 37, 23–60.
- Mendels, O., Peled, C., Levy, N. V., Rosenthal, T., & Lahiani, L. (2018). *Microsoft Presidio: Context aware, pluggable and customizable PII anonymization service for text and images*.
- Meredith, J., & Stokoe, E. (2014). Repair: Comparing Facebook 'chat' with spoken interaction. *Discourse & Communication*, 8(2), 181–207.
- Meyer, M. N. (2018). Practical tips for ethical data sharing. *Advances in Methods and Practices in Psychological Science*, 1(1), 131–144.
- Mikolov, T., Sutskever, I., Chen, K., Corrado, G. S., & Dean, J. (2013). Distributed representations of words and phrases and their compositionality. In M. I. Jordan, Y. LeCun, & S. A. Solla (Eds.), *Advances in neural information processing systems* (pp. 3111–3119). MIT Press.
- Miller, A. H., Feng, W., Fisch, A., Lu, J., Batra, D., Bordes, A., Parikh, D., & Weston, J. (2017). *Parlai: A dialog research software platform*. arXiv. <https://doi.org/10.48550/arXiv.1705.06476>
- Misyak, J. B., Melkonian, T., Zeitoun, H., & Chater, N. (2014). Unwritten rules: Virtual bargaining underpins social interaction, culture, and society. *Trends in Cognitive Sciences*, 18(10), 512–519.
- Mitchell, M., Wu, S., Zaldivar, A., Barnes, P., Vasserman, L., Hutchinson, B., Spitzer, E., Raji, I. D., & Gebru, T. (2019, January). Model cards for model reporting. In *Proceedings of the conference on fairness, accountability, and transparency* (pp. 220–229). Association for Computing Machinery.
- Molnar, A. (2019). SMARTRIQS: A simple method allowing real-time respondent interaction in qualtrics surveys. *Journal of Behavioral and Experimental Finance*, 22, 161–169.
- Moore, B. (2018, August 22). The best transcription services. *PCMag*. <https://www.pcmag.com/picks/the-best-transcription-services>
- Moore, D. A. (2016). Preregister if you want to. *American Psychologist*, 71(3), 238–239.
- Mullainathan, S., & Spiess, J. (2017). Machine learning: An applied econometric approach. *Journal of Economic Perspectives*, 31(2), 87–106.
- National Academies of Sciences, Engineering, and Medicine. (2018). *Open science by design: Realizing a vision for 21st century research*. National Academies Press.
- Nelson, L., Simmons, J., & Simonsohn, U. (2018). Psychology's renaissance. *Annual Review of Psychology*, 69, 511–534.
- Nguyen, M., Gruber, J., Fuchs, J., Marler, W., Hunsaker, A., & Hargittai, E. (2020). Changes in digital communication during the COVID-19 global pandemic: Implications for digital inequality and future research. *Social Media+ Society*, 6(3). <https://doi.org/10.1177/2056305120948255>
- Nguyen, V. A., Boyd-Graber, J., Resnik, P., Cai, D. A., Midberry, J. E., & Wang, Y. (2014). Modeling topic control to detect influence in conversations using nonparametric topic models. *Machine Learning*, 95(3), 381–421.
- Niederhoffer, K. G., & Pennebaker, J. W. (2002). Linguistic style matching in social interaction. *Journal of Language and Social Psychology*, 21(4), 337–360.
- Nisbett, R. E., & Wilson, T. D. (1977). Telling more than we can know: Verbal reports on mental processes. *Psychological Review*, 84(3), 231–259.
- Oba, D., & Berger, J. (in press). How communication mediums shape the message. *Journal of Consumer Psychology*.
- OpenAI. (2022). *GPT-3.5 architecture* [Computer software]. <https://openai.com>

- Park, T. J., Kanda, N., Dimitriadis, D., Han, K. J., Watanabe, S., & Narayanan, S. (2022). A review of speaker diarization: Recent advances with deep learning. *Computer Speech & Language*, 72, Article 101317. <https://doi.org/10.1016/j.csl.2021.101317>
- Passonneau, R. J., & Litman, D. (1993). Intention-based segmentation: Human reliability and correlation with linguistic cues. In *31st Annual Meeting of the Association for Computational Linguistics* (pp. 148–155). Association for Computational Linguistics.
- Pennebaker, J. W., Mehl, M. R., & Niederhoffer, K. G. (2003). Psychological aspects of natural language use: Our words, our selves. *Annual Review of Psychology*, 54(1), 547–577.
- Pickering, M. J., & Garrod, S. (2004). Toward a mechanistic psychology of dialogue. *Behavioral and Brain Sciences*, 27(2), 169–190.
- Poldrack, R. A., Huckins, G., & Varoquaux, G. (2020). Establishment of best practices for evidence for prediction: A review. *JAMA Psychiatry*, 77(5), 534–540.
- Pomerantz, A. (1990). Conversation analytic claims. *Communications Monographs*, 57(3), 231–235.
- Purver, M. (2011). Topic segmentation. In *Spoken language understanding: Systems for extracting semantic information from speech* (pp. 291–317). Wiley & Sons, Ltd.
- Quoidbach, J., Taquet, M., Desseilles, M., de Montjoye, Y. A., & Gross, J. J. (2019). Happiness and social behavior. *Psychological Science*, 30(8), 1111–1122.
- Rainie, H., & Wellman, B. (2012). *Networked: The new social operating system* (Vol. 10). MIT Press.
- Ranganath, R., Jurafsky, D., & McFarland, D. (2009, August). It's not you, it's me: Detecting flirting and its misperception in speed-dates. In *Proceedings of the 2009 conference on empirical methods in natural language processing* (pp. 334–342). Association for Computational Linguistics.
- Reece, A., Cooney, G., Bull, P., Chung, C., Dawson, B., Fitzpatrick, C., Glazer, T., Knox, D., Liebscher, A., & Marin, S. (2022). *Advancing an interdisciplinary science of conversation: Insights from a large multimodal corpus of human speech*. arXiv. <https://doi.org/10.48550/arXiv.2203.00674>
- Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). *Model-agnostic interpretability of machine learning*. arXiv. <https://doi.org/10.48550/arXiv.1606.05386>
- Robbins, M. L. (2017). Practical suggestions for legal and ethical concerns with social environment sampling methods. *Social Psychological and Personality Science*, 8(5), 573–580.
- Robbins, M. L., Mehl, M. R., Holleran, S. E., & Kastle, S. (2011). Naturalistically observed sighing and depression in rheumatoid arthritis patients: A preliminary study. *Health Psychology*, 30(1), 129–133.
- Roberts, M. E., Stewart, B. M., & Tingley, D. (2019). Stm: An R package for structural topic models. *Journal of Statistical Software*, 91(1), 1–40.
- Rocklage, M. D., He, S., Rucker, D. D., & Nordgren, L. F. (2022). Beyond sentiment: The value and measurement of consumer certainty in language. *Journal of Marketing Research*. Advance online publication. <https://doi.org/10.1177/00222437221134802>
- Rogers, T., Ten Brinke, L., & Carney, D. R. (2016). Unacquainted callers can predict which citizens will vote over and above citizens' stated self-predictions. *Proceedings of the National Academy of Sciences, USA*, 113(23), 6449–6453.
- Romero, D. M., Uzzi, B., & Kleinberg, J. (2016, April). Social networks under stress. In *Proceedings of the 25th International Conference on World Wide Web* (pp. 9–20). Association for Computing Machinery.
- Rossignac-Milon, M., Bolger, N., Zee, K. S., Boothby, E. J., & Higgins, E. T. (2021). Merged minds: Generalized shared reality in dyadic relationships. *Journal of Personality and Social Psychology*, 120(4), 882–911.
- Rubinstein, I. S., & Hartzog, W. (2016). Anonymization and risk. *Washington Law Review*, 91, Article 703. <https://digitalcommons.law.uw.edu/wlr/vol91/iss2/18>
- Rudin, C. (2019). Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nature Machine Intelligence*, 1(5), 206–215.
- Sacks, H., Schegloff, E. A., & Jefferson, G. (1974). A simplest systematics for the organization of turn-taking for conversation. *Language*, 50, 696–735.
- Sagi, E., & Dehghani, M. (2014). Measuring moral rhetoric in text. *Social Science Computer Review*, 32(2), 132–144.
- Samarati, P., & Sweeney, L. (1998). *Protecting privacy when disclosing information: K-anonymity and its enforcement through generalization and suppression* (Technical Report SRI-CSL-98-04). SRI International Computer Science Laboratory.
- Schegloff, E. A. (1968). Sequencing in conversational openings. *American Anthropologist*, 70, 1075–1095.
- Schegloff, E. A. (2007). *Sequence organization in interaction: A primer in conversation analysis I* (Vol. 1). Cambridge University Press.
- Schegloff, E. A., & Sacks, H. (1973). Opening up closings. *Semiotica*, 8, 289–327.
- Schweinsberg, M., Feldman, M., Staub, N., van den Akker, O. R., van Aert, R. C., Van Assen, M. A., Liu, Y., Althoff, T., Heer, J., Kale, A., Mohamed, Z., Amireh, H., Venkatesh Prasad, V., Bernstein, A., Robinson, E., Snellman, K., Amy Sommer, S., Otner, S. M. G., Robinson, D., Madan, N., . . . Uhlmann, E. L. (2021). Same data, different conclusions: Radical dispersion in empirical results when independent analysts operationalize and test the same hypothesis. *Organizational Behavior and Human Decision Processes*, 165, 228–249.
- Searle, J. R. (1965). What is a speech act. In R. J. Stainton (Ed.), *Perspectives in the philosophy of language: A concise anthology, 2000* (pp. 253–268). Broadview Press.
- Shi, Y., Yeomans, M., Truong, M., & Fast, N. (2023). *What you say in the conversation affects the flow: Modelling conversational flow using NLP methods* [Working paper].
- Simons, D. J., Shoda, Y., & Lindsay, D. S. (2017). Constraints on generality (COG): A proposed addition to all empirical papers. *Perspectives on Psychological Science*, 12(6), 1123–1128.
- Simonsohn, U., Simmons, J. P., & Nelson, L. D. (2020). Specification curve analysis. *Nature Human Behaviour*, 4(11), 1208–1214.
- Sommers, S. R. (2006). On racial diversity and group decision making: Identifying multiple effects of racial composition

- on jury deliberations. *Journal of Personality and Social Psychology*, 90(4), 597–612.
- Srivastava, S. B., Goldberg, A., Manian, V. G., & Potts, C. (2018). Enculturation trajectories: Language, cultural adaptation, and individual outcomes in organizations. *Management Science*, 64(3), 1348–1364.
- Stillwell, D. J., & Kosinski, M. (2004). MyPersonality project: Example of successful utilization of online social networks for large-scale social research. *American Psychologist*, 59(2), 93–104.
- Stivers, T., Enfield, N. J., & Levinson, S. C. (2010). Question-response sequences in conversation across ten languages: An introduction. *Journal of Pragmatics*, 42, 2615–2619.
- Stivers, T., & Sidnell, J. (Eds.). (2012). *The handbook of conversation analysis*. John Wiley & Sons.
- Stokoe, E. (2010). 'I'm not gonna hit a lady': Conversation analysis, membership categorization and men's denials of violence towards women. *Discourse & Society*, 21(1), 59–82.
- Stokoe, E. (2021). The sense of a conversational ending. *Nature*, 593(7859), 347–349.
- Stolcke, A., Ries, K., Coccaro, N., Shriberg, E., Bates, R., Jurafsky, D., Taylor, P., Martin, R., Van Ess-Dykema, C., & Meteer, M. (2000). Dialogue act modeling for automatic tagging and recognition of conversational speech. *Computational Linguistics*, 26(3), 339–373.
- Sun, J., Harris, K., & Vazire, S. (2020). Is well-being associated with the quantity and quality of social interactions? *Journal of Personality and Social Psychology*, 119(6), 1478–1496. <https://doi.org/10.1037/pspp0000272>
- Sun, J., Schwartz, H. A., Son, Y., Kern, M. L., & Vazire, S. (2020). The language of well-being: Tracking fluctuations in emotion experience through everyday speech. *Journal of Personality and Social Psychology*, 118(2), 364–387.
- Swaab, R. I., Lount, R. B., Jr., Chung, S., & Brett, J. M. (2021). Setting the stage for negotiations: How superordinate goal dialogues promote trust and joint gain in negotiations between teams. *Organizational Behavior and Human Decision Processes*, 167, 157–169.
- Sweeney, L. (2002). k-anonymity: A model for protecting privacy. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 10(5), 557–570.
- Takanobu, R., Huang, M., Zhao, Z., Li, F. L., Chen, H., Zhu, X., & Nie, L. (2018, July). A weakly supervised method for topic segmentation and labeling in goal-oriented dialogues via reinforcement learning. In *Proceedings of the Twenty-Seventh International Joint Conference on Artificial Intelligence* (pp. 4403–4410). AAAI Press. <https://doi.org/10.24963/ijcai.2018/612>
- Tan, C., Niculae, V., Danescu-Niculescu-Mizil, C., & Lee, L. (2016, April). Winning arguments: Interaction dynamics and persuasion strategies in good-faith online discussions. In *Proceedings of the 25th International Conference on World Wide Web* (pp. 613–624). Association for Computational Linguistics.
- Tausczik, Y. R., & Pennebaker, J. W. (2010). The psychological meaning of words: LIWC and computerized text analysis methods. *Journal of Language and Social Psychology*, 29(1), 24–54.
- Templeton, E. M., Chang, L. J., Reynolds, E. A., Cone LeBeaumont, M. D., & Wheatley, T. (2022). Fast response times signal social connection in conversation. *Proceedings of the National Academy of Sciences*, 119(4), e2116915119.
- Templeton, E. M., Chang, L. J., Reynolds, E. A., Cone LeBeaumont, M. D., & Wheatley, T. (2023). Long gaps between turns are awkward for strangers but not for friends. *Philosophical Transactions of the Royal Society B*, 378(1875), 20210471.
- Traeger, M. L., Strohkorb Sebo, S., Jung, M., Scassellati, B., & Christakis, N. A. (2020). Vulnerable robots positively shape human conversational dynamics in a human–robot team. *Proceedings of the National Academy of Sciences, USA*, 117(12), 6370–6375.
- Turmunkh, U., Van den Assem, M. J., & Van Dolder, D. (2019). Malleable lies: Communication and cooperation in a high stakes TV game show. *Management Science*, 65(10), 4795–4812.
- van Werven, R., Bouwmeester, O., & Cornelissen, J. P. (2019). Pitching a business idea to investors: How new venture founders use micro-level rhetoric to achieve narrative plausibility and resonance. *International Small Business Journal*, 37(3), 193–214.
- Voigt, R., Camp, N. P., Prabhakaran, V., Hamilton, W. L., Hetey, R. C., Griffiths, C. M., Jurgens, D., Jurafsky, D., & Eberhardt, J. L. (2017). Language from police body camera footage shows racial disparities in officer respect. *Proceedings of the National Academy of Sciences, USA*, 114(25), 6521–6526. <https://doi.org/10.1073/pnas.1702413114>
- Wang, X., Yang, D., Wen, M., Koedinger, K., & Rosé, C. P. (2015). *Investigating how student's cognitive behavior in MOOC discussion forums affect learning gains*. International Educational Data Mining Society.
- Warriner, A. B., Kuperman, V., & Brysbaert, M. (2013). Norms of valence, arousal, and dominance for 13,915 English lemmas. *Behavior Research Methods*, 45(4), 1191–1207.
- Weingart, L., Smith, P., & Olekalns, M. (2004). Quantitative coding of negotiation behavior. *International Negotiation*, 9(3), 441–456.
- Weiss, K., Khoshgoftaar, T. M., & Wang, D. (2016). A survey of transfer learning. *Journal of Big Data*, 3(1), 1–40.
- Wickham, H., Averick, M., Bryan, J., Chang, W., McGowan, L. D. A., François, R., Grolemond, G., Hayes, A., Henry, L., Hester, J., Kuhn, M., Pedersen, T. L., Miller, E., Bache, S. M., Müller, K., Ooms, J., Robinson, D., Seidel, D. P., Spinu, V., . . . Yutani, H. (2019). Welcome to the Tidyverse. *Journal of Open Source Software*, 4(43), Article 1686. <https://doi.org/10.21105/joss.01686>
- Word, C. O., Zanna, M. P., & Cooper, J. (1974). The nonverbal mediation of self-fulfilling prophecies in interracial interaction. *Journal of Experimental Social Psychology*, 10(2), 109–120.
- Xing, L., & Carenini, G. (2021, July). Improving unsupervised dialogue topic segmentation with utterance-pair coherence scoring. In *Proceedings of the 22nd annual meeting of the Special Interest Group on discourse and dialogue* (pp. 167–177). Association for Computational Linguistics.
- Xu, Y., Zhao, H., & Zhang, Z. (2021). Topic-aware multi-turn dialogue modeling. *Proceedings of the AAAI Conference on Artificial Intelligence*, 35(16), 14176–14184.

- Yeomans, M. (2021). A concrete example of construct construction in natural language. *Organizational Behavior and Human Decision Processes*, 162, 81–94.
- Yeomans, M., & Brooks, A. W. (2023). *What do you want to talk about? Topic preference detection in conversation*. Working paper.
- Yeomans, M., Brooks, A. W., Huang, K., Minson, J., & Gino, F. (2019). It helps to ask: The cumulative benefits of asking follow-up questions. *Journal of Personality and Social Psychology*, 117(6), 1139–1144.
- Yeomans, M., Kantor, A., & Tingley, D. (2018). The politeness package: Detecting politeness in natural language. *R Journal*, 10(2), 489–502.
- Yeomans, M., Minson, J., Collins, H., Chen, F., & Gino, F. (2020). Conversational receptiveness: Improving engagement with opposing views. *Organizational Behavior and Human Decision Processes*, 160, 131–148.
- Yeomans, M., Schweitzer, M. E., & Brooks, A. W. (2022). The conversational circumplex: Identifying, prioritizing, and pursuing informational and relational motives in conversation. *Current Opinion in Psychology*, 44, 293–302. <https://doi.org/10.1016/j.copsyc.2021.10.001>
- Yeomans, M., Shah, A., Mullainathan, S., & Kleinberg, J. (2019). Making sense of recommendations. *Journal of Behavioral Decision Making*, 32(4), 403–414.
- Yeomans, M., Stewart, B. M., Mavon, K., Kindel, A., Tingley, D., & Reich, J. (2018). The civic mission of MOOCs: Engagement across political differences in online forums. *International Journal of Artificial Intelligence in Education*, 28(4), 553–589.
- Zeileis, A., Köll, S., & Graham, N. (2020). Various versatile variances: An object-oriented implementation of clustered covariances in R. *Journal of Statistical Software*, 95(1), 1–36.
- Zhang, J., Chang, J., Danescu-Niculescu-Mizil, C., Dixon, L., Hua, Y., Taraborelli, D., & Thain, N. (2018, July). Conversations gone awry: Detecting early signs of conversational failure. In *Proceedings of the 56th annual meeting of the Association for Computational Linguistics* (Vol. 1: Long Papers, pp. 1350–1361). Association for Computational Linguistics.
- Zhang, J., Mullainathan, S., & Danescu-Niculescu-Mizil, C. (2020). Quantifying the causal effects of conversational tendencies. *Proceedings of the ACM on Human-Computer Interaction*, 4(CSCW2), 1–24.
- Zhang, J., Spirling, A., & Danescu-Niculescu-Mizil, C. (2017, September). Asking too much? The rhetorical role of questions in political discourse. In *Proceedings of the 2017 conference on empirical methods in natural language processing* (pp. 1558–1572). Association for Computational Linguistics.
- Zheng, W., Yan, L., Gou, C., Zhang, Z. C., Zhang, J. J., Hu, M., & Wang, F. Y. (2021). Pay attention to doctor-patient dialogues: Multi-modal knowledge graph attention image-text embedding for COVID-19 diagnosis. *Information Fusion*, 75, 168–185. <https://doi.org/10.1016/j.inffus.2021.05.015>