

The Critical Role of Unstructured Data in Modern Marketing Research

Reza (Pouria) Khansari

Ivey Business School at Western University

Kersi Antia

Ivey Business School at Western University

Abhishek Borah

Professor of Marketing, INSEAD

Khaled Boughanmi

Columbia University

Ishita Chakraborty

UW Madison

Cite as:

Khansari Reza (Pouria), Antia Kersi, Borah Abhishek, Boughanmi Khaled, Chakraborty Ishita (2026), The Critical Role of Unstructured Data in Modern Marketing Research. *Proceedings of the European Marketing Academy*, 55th, (133097)

Paper from the 55th Annual EMAC Conference, Bath, United Kingdom, June 2-5, 2026



The Critical Role of Unstructured Data in Modern Marketing Research

Session Chair(s):

Pouria Khansari, Ivey Business School at Western University

Kersi D. Antia, Ivey Business School at Western University

Brands and Bands: Do Brand Names in Songs Help or Hinder Song Popularity?

Abhishek Borah, INSEAD Business School (presenting author)

Brendon Rhodes, INSEAD Business School

Raoul V. Kübler, ESSEC Business School

Giovanni Luca Cascio Rizzo, University of Southern California

Sourindra Banerjee, University of Leeds

Modeling Dynamic Consumer Preferences from Few-shot Data: A Meta-Learning Approach

Mingzhang Yin, University of Florida

Khaled Boughanmi, Cornell University (presenting author)

Anirban Mukherjee, Avyayam Holdings

From Reviews to Responses: Bridging Pre- and Post-Purchase Consumers through AI-Enhanced QA with RAG

Xingjian You, University of Wisconsin-Madison

Ishita Chakraborty, University of Wisconsin-Madison (presenting author)

Neeraj Arora, University of Wisconsin-Madison

Customer-Employee Dyads in Complaint Handling Interactions

Pouria Khansari, Western University (presenting author)

Hoorsana Damavandi, University of Tennessee

Kersi D. Antia, Western University

Declaration:

Each presenter has agreed to register for the conference and to present the paper, if the proposal is accepted; none of the papers has been submitted to other conference tracks, and none has previously been presented at EMAC.

Abstract

The proliferation of unstructured data—whether questions and answers, reviews, music, or speech—has given marketing researchers unprecedented access to subtleties of human expression that were previously impossible to study at scale. Recent advances in natural language processing and deep learning now allow scholars to extract meaningful insights from words, voice tonality, and music acoustic features. This special session comprises four studies that illustrate how modern computational tools uncover latent meaning from unstructured data: from brand mentions in song lyrics to evolving user preferences on streaming platforms, AI-augmented customer Q&A systems, and the dynamics of speech in customer-employee service interactions. Together, these papers demonstrate how unstructured data may be leveraged to advance empirical and theoretical frontiers of marketing research.

The first paper by Borah, Rhodes, V. Kübler, Cascio Rizzo, and Banerjee investigates whether brand mentions in song lyrics—known as *Brand Name Dropping (BND)*—, especially luxury BND, help or hinder a song’s popularity and why. The study answers these questions by utilizing a multimethod approach combining large-scale field data on 2,966 songs from the Billboard Hot 100 and Spotify with controlled experimentation. The second paper presented in this session, coauthored by Yin, Boughanmi, and Mukherjee, introduces a Transformer-based recommendation system that learns user preferences dynamically from limited interactions during consumption. Incorporating acoustic features of songs, their meta-learning framework demonstrates superior predictive performance compared with other state-of-the-art models. In the third presentation, You, Chakraborty, and Arora differentiate user-generated content (UGC) on online platforms that is generated pre- (Question Answering (QA)) or post-purchase and leverage this distinction to propose a novel generative QA model. Drawing on Amazon QAs and reviews, it develops a retrieval-augmented-generation framework with answerability detection and human-centered evaluation to assess the quality and trustworthiness of AI-generated responses. In the fourth and final presentation, Khansari, Damavandi, and Antia delve into the three understudied aspects of speech-based communication in customer complaint handling interactions. Analyzing more than 30,000 conversational turns from 1069 recorded service calls, their study models the dyadic exchange and trajectory of verbal and vocal empathy between customers and employees and links them to immediate (satisfaction) and longer-term (churn) outcomes.

By offering insights based on fresh perspectives and innovative methodologies, our proposed session is likely to be of significant interest to academics interested in digital and service marketing, as well as those interested in how AI and unstructured data are reshaping marketing theory and practice.

Research Presentation 1

Title: Brands and Bands: Do Brand Names in Songs Help or Hinder Song Popularity?

Authors: Abhishek Borah (Associate Professor of Marketing, INSEAD Business School), Brendon Rhodes (PhD Student in Marketing, INSEAD Business School), Raoul V. Kübler (Associate Professor of Marketing, ESSEC Business School), Giovanni Luca Cascio Rizzo (Assistant Professor of Marketing, University of Southern California), Sourindra Banerjee (Associate Professor of Marketing, University of Leeds)

Brand mentions in song lyrics—known as *Brand Name Dropping (BND)*—have become increasingly prevalent in contemporary music, from The Beatles’ “he shoot Coca-Cola” to Beyoncé’s “Givenchy dress.” Yet, while BND has become integral to music’s cultural and commercial vocabulary, its impact on *song performance* remains surprisingly understudied. Prior research has examined product placement effects on *brands*—such as awareness, attitude, and equity—but not on the *songs* that feature them. This paper addresses that gap by investigating whether BND, especially luxury BND, helps or hinders a song’s popularity and why. Grounded in identity-based consumption theory and the literature on authenticity and expressive media, the study posits that brand inclusion in music operates differently from other media contexts. Whereas brand placement in films or TV can enhance realism and immersion, in music—an art form prized for emotional honesty and social resonance—brand references may disrupt authenticity, alienate audiences, and lower engagement. Because luxury brands evoke exclusivity and status signaling, they may further erode relatability among listeners who perceive music as a democratizing and inclusive form of self-expression. Accordingly, the paper proposes three hypotheses: (H1) BND in lyrics is negatively related to song popularity; (H2) luxury brand mentions are more detrimental than non-luxury mentions; and (H3) this relationship is mediated by listeners’ reduced relatability to the song.

To test these predictions, the authors adopt a mixed-method approach combining large-scale field data with controlled experimentation. Study 1 analyzes 2,966 songs by 1,518 artists that

appeared on the Billboard Hot 100 from 2017 to 2021. Song lyrics were manually coded for brand mentions (luxury vs. non-luxury) using a database of over one million brand names and matched with Billboard rankings and Spotify audio features (danceability, energy, valence, etc.). The model includes artist and genre fixed effects, linguistic variables (LIWC), thematic features (LDA topics), and artist popularity, with endogeneity addressed via a Gaussian Copula approach and clustered standard errors. The results reveal that total brand mentions do not significantly affect popularity, but luxury brand mentions *decrease* it sharply: each luxury name-drop corresponds to roughly six positions lower on the Billboard chart ($p < .001$). Non-luxury brands show no negative effect. Moreover, the detrimental impact of luxury BND is stronger in identity-driven genres such as hip-hop, rap, and experimental music, and weaker in rock or dance-oriented tracks. Study 2 strengthens causal inference through a randomized experiment ($N = 301$, Prolific) where participants evaluated identical song lyrics containing a luxury brand (“Aurastar” high-end car), a non-luxury brand (“Aurastar” mid-range car), or no brand. Participants rated song liking, relatability, and perceptions of artist showmanship and commercial intent. Results mirrored the field findings: songs mentioning luxury brands were liked less ($M = 3.24$) and felt less relatable ($M = 5.32$) than those with non-luxury ($M = 3.68$; $M = 4.51$) or no brands ($M = 3.73$; $M = 3.95$). Mediation analysis (20,000 Monte Carlo simulations) confirmed that relatability fully mediated the negative relationship between luxury BND and song liking (indirect effect = $-.53$, 95% CI $[-.84, -.25]$), while alternative explanations were unsupported. Together, these findings challenge conventional assumptions about brand placement. In contrast to film or television—where brand integration typically enhances audience engagement—music’s emotive and self-expressive nature makes it uniquely sensitive to perceived inauthenticity. Theoretically, this study extends branding and media persuasion literature by identifying *relatability* as a key psychological mechanism linking symbolic content and cultural success, and by showing that luxury signaling can backfire in identity-based media. Managerially, the results caution both artists and marketers: luxury name-dropping can reduce song appeal, particularly when it disrupts listeners’ emotional connection. Rather than relying on overt brand mentions, marketers should pursue authentic, value-aligned collaborations that enhance, rather than undermine, artistic integrity. Overall, the study illuminates how branding operates as both a cultural and commercial signal—and how, in music, luxury associations may resonate discordantly with the universal language of song.

Research Presentation 2

Title: Modeling Dynamic Consumer Preferences from Few-shot Data: A Meta-Learning Approach

Authors: Mingzhang Yin (Assistant Professor of Marketing, University of Florida), Khaled

Boughanmi (Assistant Professor of Marketing, Cornell University), Anirban Mukherjee

Digital media platforms increasingly operate in environments characterized by rapidly changing consumer tastes (Brost, Mehrotra, and Jehan 2019) and data privacy regulations that limit the use of third-party information (Voigt and Von Dem Bussche 2017). As a result, platforms must learn user preferences quickly and efficiently using only a few observed interactions within a session.

Traditional personalization systems rely heavily on large amounts of historical user-level data or population-level averages, which limits their ability to provide meaningful recommendations for new users or new consumption sessions. This challenge is particularly pronounced in digital streaming platforms, where customer preferences evolve continuously in response to content exposure. Thus, personalizing recommendations based on static or slowly adapting methods can result in mismatches that lead to disengagement and reduced retention.

In this research, we introduce Meta-Temporal Processes (MetaTP), a novel meta-learning framework designed to infer dynamic, individual-specific preference trajectories using few-shot interactions (Finn, Abbeel, and Levine 2017). MetaTP is specifically trained to generalize from a distribution of sequential consumption tasks, enabling it to quickly adapt to new customers or sessions with minimal data. The framework employs a meta-learning paradigm in which individual-level interactions are partitioned into context and target sets, allowing the model to simulate real-world personalization scenarios during training. This structure enables rapid inference for each new customer or session by learning how to learn preferences, rather than learning a single fixed preference function.

Our modeling approach is built on a Transformer-based encoder that capitalizes on self-attention mechanisms to capture the temporal and contextual dependencies inherent in sequential decision-making and media consumption (Vaswani et al. 2017). This encoder extracts rich representation vectors that summarize individual behavior in real time. These representations are decoded through a structural and probabilistic decoder that maintains interpretability by parameterizing evolving latent preferences and their influence on choice. By combining flexibility in representation with structure in interpretation, MetaTP bridges

the gap between deep learning models that offer predictive accuracy and econometric models that offer interpretability.

We illustrate our proposed model on music listening sessions on a major digital streaming platform. In this empirical context, each session comprises a sequence of songs consumed by a user, where the platform must decide which track to recommend next. We incorporate acoustic features of songs, which include rhythm, energy, valence, instrumentation, and timbre, to learn how users dynamically update their preferences based on prior content exposure (Boughanmi and Ansari 2021). This enables MetaTP to uncover evolving musical tastes not only across users but also within individual sessions, demonstrating its ability to rapidly infer preference.

We benchmark the proposed approach against leading static and adaptive models, including hierarchical Bayesian choice models, recurrent neural networks, and other state-of-the-art models. Across all benchmarks, our model demonstrates superior predictive performance, higher data efficiency, and a dramatic reduction in the number of interactions required to make accurate recommendations. Unlike traditional models that must relearn parameters for each new user, MetaTP leverages meta-learned priors to instantiate a personalized preference function instantly, updating it dynamically with each new interaction.

Research Presentation 3

Title: From Reviews to Responses: Bridging Pre- and Post-Purchase Consumers through AI-Enhanced QA with RAG

Authors: Xingjian You (PhD Student of Marketing, University of Wisconsin-Madison), Ishita Chakraborty (Assistant Professor of Marketing, University of Wisconsin-Madison), Neeraj Arora (Assistant Professor of Marketing, University of Wisconsin-Madison)

Question Answering (QA) on customer-facing platforms (e-commerce, travel, education, and brand websites) often suffers from delayed, low-quality responses and limited user participation. Another prevalent form of UGC that prospective buyers rely on before making a purchase is reviews or feedback from past customers, which describe their experiences with different aspects of the product. Although both QAs and reviews inform customer decision-making, they are structurally different: QAs typically occur pre-purchase and address specific questions, whereas reviews are written post-purchase and tend to be broader in scope. Platforms explicitly separate these two content types to facilitate distinct stages of the purchase cycle. While customer reviews are abundant, their unstructured nature limits direct

use to answer specific questions. Recent developments in Large Language Models (LLMs), such as GPT-4 and Gemini, present promising opportunities to automate question-answering systems on different platforms, significantly reducing response times. These models can produce human-like responses and have demonstrated strong performance in general-purpose text generation and reasoning tasks. However, Large Language Models (LLMs) alone lack product-specific details and subjective insights to provide high-quality responses. Against this backdrop of QA, reviews, and LLMs, in this paper we seek to answer three research questions.

- How can one form of UGC be used to enrich another with the help of LLMs? Specifically, how might the generative strengths of LLMs, when combined with grounded, experience-rich post-purchase reviews, enhance the QA system on platforms?
- Can this hybrid LLM-enabled approach deliver high-quality answers across question types, particularly for subjective vs. objective questions?
- How should we evaluate LLM-generated answers to customer queries? Are lexical overlap metrics sufficient, or is consumer-centered human evaluation essential? How do human and automatic assessments differ, and what insights arise from their divergence?

To answer our research questions, we propose a framework in which the LLM bridges pre-purchase and post-purchase customers, drawing on both internal knowledge of the LLM and user-generated content to create informed responses. The central piece of our framework is Retrieval-Augmented Generation (RAG), which is a hybrid approach that augments the output of a language model with information retrieved from an external knowledge source. In this paper, this external knowledge source is consumer reviews. By injecting important contextual knowledge, RAG enables LLMs to dynamically incorporate the most relevant, up-to-date review content at inference time. This makes it particularly suitable for enhancing customer-facing QAs with transient, large volumes of review data.

Although RAG improves responses by incorporating contextual information, a persistent challenge lies in retrieval precision. While several platforms typically have a large pool of reviews, much of this content does not directly address the specific customer question. Standard RAG implementations often rely solely on the lexical similarity between questions and reviews, which can fail to capture answer relevance. To address this limitation, we propose a question–review matching module in our framework that re-ranks the reviews

based on how similar the content is to the question, and whether the review talks about the same topic as the question (e.g. fit, usability, customer service). In other words, if the question is asking about how well a product fits, the module gives higher priority to reviews that also discuss fit, rather than unrelated aspects.

Finally, while the RAG and matching modules improve contextual grounding and relevance, not all questions are answerable, even when extensive information is available. Attempting to answer such unanswerable questions can lead to hallucinated, irrelevant, or low-quality responses. This behavior can undermine user trust. To mitigate this natural behavior of LLMs, we introduce an answerability classifier in our framework that evaluates whether a given question can be reliably answered using the available knowledge base. Its sole purpose is to determine the presence or absence of sufficient supporting information before any response is generated.

We empirically validate this framework using a data set of 500 Amazon questions, 2000+ human responses, and 14,000 review sentences. Each question is typically answered by multiple individuals, thus resulting in more answers than questions. For each question, we generate an answer using a variety of models (LLM-only, LLM+RAG, LLM+RAG+ Matching, etc.) and benchmark the quality of these answers against human answers. To evaluate the quality of the generated answers, we adopt a dual approach: automatic lexical similarity metrics and human evaluation.

We begin our analysis with automatic evaluation using a widely used metric (ROGUE-L) that measures the longest common subsequence between generated and reference answers to capture lexical overlap. Although this metric captures the similarity at the word level, it overlooks deeper aspects of the answer quality. To complement this analysis, we perform a theory-driven human evaluation that focuses on five key dimensions: clarity, relevance, informativeness, trustworthiness, and helpfulness. We recruited 582 human annotators from Prolific to evaluate both model-generated and human-written answers, allowing us to benchmark model-based (LLM-only, LLM+RAG, LLM+ RAG+ Matching, etc.) answers against human answers.

The automatic and human evaluation results highlight the value of our framework.

Incorporating reviews through RAG increases lexical similarity by 30%—from .105 (LLM-only) to .143 ($p < .001$). Adding the answerability classifier yields a further 10% improvement, but restricts the system to 62.6% of questions. Matching question type to reviews slightly raises the answer rate to 64.4%, with notable gains in review-sparse categories such as customer service.

Human evaluations reveal deeper insights. Our full model achieves an average answer-quality score of 4.14 on a 5-point scale, significantly exceeding the mean human score of 3.50 ($p < .001$) and approaching the best-human score of 4.34 ($p = .007$). The model meets or exceeds 72% of human answers and closely matches top-human performance on relevance, clarity, and informativeness—for instance, scoring 4.32 vs. 4.51 ($p = .028$) on relevance. Interestingly, when raters are unaware of the source, model responses are judged more trustworthy than average human ones.

Our model performs better on subjective than objective questions, contrary to prior findings that language models struggle in subjective domains. This advantage arises because model answers blend the precision of LLMs with the “word-of-mouth” tone of reviews, making them particularly suitable for subjective queries. In some cases, model-based responses surpass human ones, yet lexical metrics penalize such deviations from benchmark answers regardless of whether the divergence reflects improvement or error. This exposes a limitation of lexical metrics, which conflate difference with inferiority.

Overall, our results underscore the importance of human-centered evaluation: similarity to benchmark answers should not be treated as the gold standard, especially in subjective or interpretive domains where higher-quality model responses may appropriately diverge from human baselines.

Figure 1: Proposed Framework

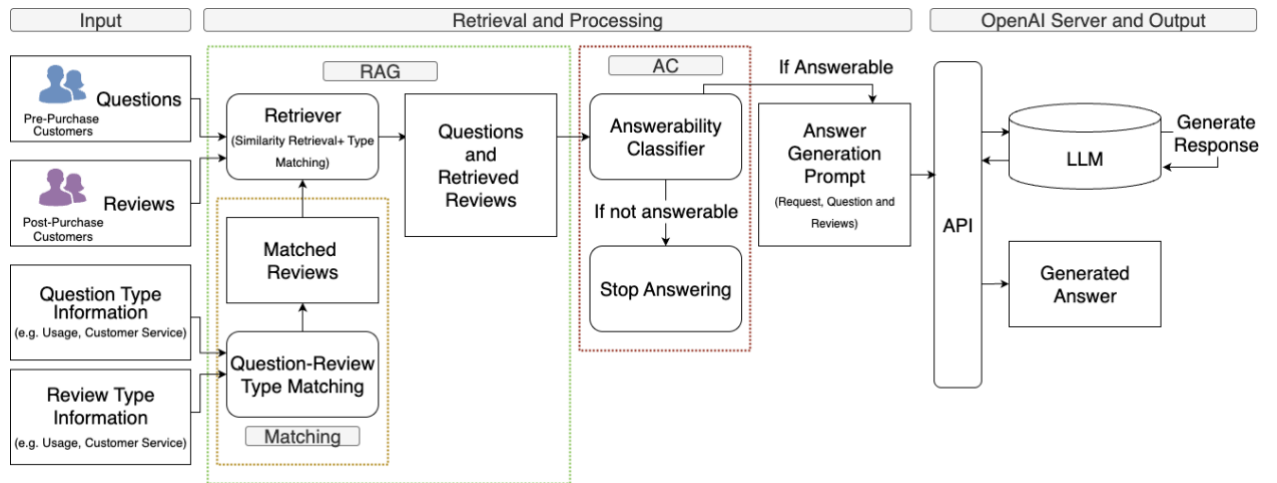


Table 1: Human Evaluation Results by Answer Source

Answer Source	Average Score	Clarity	Informa tivity	Rele vance	Trust	Helpful ness	% Over Human
Best Human Answer	4.34 (.05)	4.41 (.06)	4.29 (.06)	4.51 (.05)	4.11 (.06)	4.38 (.06)	-
Human Answer	3.50 (.04)	3.50 (.04)	3.43 (.04)	3.71 (.04)	3.37 (.04)	3.46 (.04)	-
Worst Human Answer	2.42 (.07)	2.38 (.08)	2.35 (.08)	2.67 (.08)	2.42 (.07)	2.27 (.08)	-
LLM Only	3.50 (.08)	3.61 (.08)	3.58 (.09)	3.46 (.09)	3.54 (.07)	3.29 (.10)	53%
LLM+RAG	4.09 (.06)	4.10 (.07)	4.12 (.07)	4.23 (.06)	3.86 (.06)	4.13 (.07)	72%
LLM+RAG+AC	4.11 (.06)	4.16 (.07)	4.17 (.06)	4.36 (.06)	3.80 (.07)	4.06 (.08)	69%
LLM+RAG+AC+Weighted Match (Full Model)	4.14 (.06)	4.17 (.07)	4.17 (.06)	4.32 (.07)	3.84 (.07)	4.19 (.07)	72%

Notes. Values are reported as mean (standard error). For methods incorporating AC, the comparison is based only on the portion of the question they answered.

Research Presentation 4

Title: Customer-Employee Dyads in Complaint Handling Interactions

Authors: Pouria Khansari (PhD Student of Marketing, Western University), Hoorsana Damavandi (Assistant Professor of Marketing, University of Tennessee), Kersi D. Antia (Professor of Marketing, Western University)

In 2024, businesses worldwide incurred \$3.7 trillion in losses attributable to poor customer experiences (Qualtrics 2024). Nearly half of customers report cutting their spending in response to a single negative encounter, and one in five terminate their relationship with the company (PWC 2018). Importantly, 70% of customers report that it is the firms' employees who shape their experience (*ibid.*), and who comprise the cornerstone of firms' customer complaint handling efforts. At the heart of these interactions, employees' empathy—feeling what customers feel and reacting to it—is widely acknowledged as the critical driver of customer satisfaction and loyalty (Davidow 2003; Duan and Hill 1996; Smith and Bolton 1998; Smith, Bolton, and Wagner 1999).

However, like every social interaction, service encounters are co-produced by both the customer and the employee and are inherently dynamic. It is therefore logical to expect that the outcomes of these interactions are shaped not only by the level of employee empathy but also by how customers *reciprocate* employee empathy and by how empathy *rises or falls* throughout the call. Additionally, individuals do not communicate words in isolation; spoken language is expressed through voice, which can shape the meaning of words. Thus, both *what*

is said (verbal communication) and *how it is said* (vocal communication) play critical roles in driving outcomes in interactions. This research builds on these intuitively appealing yet hitherto unaddressed ideas to answer the following key research questions:

First, *how does the interplay between a customer's and an employee's empathy individually and jointly shape customer satisfaction and their exit behavior?* Although employees' capabilities in complaint handling have been examined, the customers' role has received less attention. Because service encounters are inherently interactive, a satisfactory outcome depends on the coordinated actions of both parties (Ma and Dubé 2011; Solomon et al. 1985; Wieseke, Geigenmüller, and Kraus 2012). Building on reciprocity theory, I propose that customers and employees reciprocate empathy, since people feel obliged to return favors and acts of kindness in roughly equivalent form; furthermore, I posit satisfactory complaint resolution to be a function of both parties' joint behavior (Gouldner 1960).

Second, *how do the evolving empathy trajectories of both parties influence customer satisfaction and exit behavior?* Customer–employee conversations unfold turn by turn, with each exchange shaping the outcome (Goffman, 1981; Schegloff, 1999). Empathy is dynamic: an employee who begins highly empathetic may lose patience when faced with an irate customer, while a customer may soften in response to an employee's empathy. I conceptualize empathy trajectories to capture these patterns. According to peak–end theory (Kahneman et al. 1993), customers' retrospective evaluations of the call are not formed by averaging every moment; the most intense (peak) and final moments carry greater weight. Thus, trajectories that rise toward the end or include pronounced peaks should yield higher satisfaction and lower exit likelihood. Accordingly, I hypothesize that rising empathy amplifies the positive impact of empathy, whereas declining empathy diminishes it (Johnson, Tellis, and Macinnis 2005; Palmatier et al. 2013).

Third, *do vocal empathy cues affect customer satisfaction and exit behavior?* I emphasize that communication involves not only *what* is said, but also *how* it is said, the latter conveying more than one-third of our emotions (Gabbott and Hogg 2000; Mehrabian and Ferris 1967). Employees or customers might speak, for example, faster or slower (*articulation rate*), softer or louder (*loudness*), and/or use a higher or lower tone (*pitch*). These variations in vocalization communicate the speakers' emotional state and consequently their empathy as well (Cascio Rizzo and Berger 2023; Juslin and Laukka 2003; Xiao et al. 2014). Building on prior research that stresses the value of simultaneous verbal and non-verbal communication (such as body language or facial gestures) (Marinova et al., 2018; Packard et al., 2023;

Packard & Berger, 2023), we assess, for the first time to the best of my knowledge, multimodal empathy as communicated by both vocal and verbal means.

To address these questions, we partnered with a Toronto-based multi-site self-storage rental company to analyze a unique dataset of 1069 audio recordings of complaint calls from customers who came onboard between November 2021 and August 2022, along with their contract details. The key independent variables are employees' and customers' evolving empathy trajectories, reflecting how empathy unfolds, intensifies, or diminishes throughout the customer–employee interaction. To measure empathy, we apply advanced speech analysis techniques in two steps. First, we convert the audio to text (transcription) and parse the text into more than 30,000 conversational turns involving the customer calling in and the employee fielding the call. Then, we apply large language models (LLMs) and embedding methods, numerical representations of sentences' semantic meaning, to quantify the empathy expressed through language (Chakraborty et al. 2025). We validate the text analysis by running an online experiment in which participants rate the perceived empathy of sentences from the previous step. In the second step, we use Audio Analysis techniques to extract vocal features, including pitch, loudness, and articulation rate, as the most relevant predictors of empathy to develop a formative measure of vocal empathy.

As we model the complaint-handling encounters as *dyadic*, we expect our findings to reveal that effective encounters depend not only on how empathetic employees are, but also on how customers reciprocate. Interactions in which customers express empathy are anticipated to intensify the positive impact of employee empathy and yield higher satisfaction. Furthermore, as conversation unfolds, we predict that encounters featuring increasing empathy *trajectories* are likely to foster greater satisfaction than those with equivalent average empathy but a downward trend. We also expect that *vocal empathy* reinforces the positive impact of verbal empathy. Finally, we anticipate that positive experiences in complaint-handling encounters will shape customer satisfaction as well as churn behavior, as customers who feel understood and respected are more likely to remain longer with the firm.

References

- Boughanmi, Khaled, and Asim Ansari. 2021. "Dynamics of Musical Success: A Machine Learning Approach for Multimedia Data Fusion." *Journal of Marketing Research* 58(6):1034–57. doi:10.1177/00222437211016495.
- Brost, Brian, Rishabh Mehrotra, and Tristan Jehan. 2019. "The Music Streaming Sessions Dataset." Pp. 2594–2600 in *The World Wide Web Conference, WWW '19*. New York, NY, USA: Association for Computing Machinery.

- Cascio Rizzo, Giovanni Luca, and Jonah A. Berger. 2023. "The Power of Speaking Slower."
- Chakraborty, Ishita, Khai Chiong, Howard Dover, and K. Sudhir. 2025. "Can AI and AI-Hybrids Detect Persuasion Skills? Salesforce Hiring with Conversational Video Interviews." *Marketing Science* 44(1):30–53. doi:10.1287/mksc.2023.0149.
- Davidow, Moshe. 2003. "Organizational Responses to Customer Complaints: What Works and What Doesn't." *Journal of Service Research* 5(3):225–50. doi:10.1177/1094670502238917.
- Duan, Changming, and Clara E. Hill. 1996. "The Current State of Empathy Research." *Journal of Counseling Psychology* 43(3):261–74. doi:10.1037/0022-0167.43.3.261.
- Finn, Chelsea, Pieter Abbeel, and Sergey Levine. 2017. "Model-Agnostic Meta-Learning for Fast Adaptation of Deep Networks." Pp. 1126–35 in *Proceedings of the 34th International Conference on Machine Learning*. PMLR.
- Gabbott, Mark, and Gillian Hogg. 2000. "An Empirical Investigation of the Impact of Non-verbal Communication on Service Evaluation." *European Journal of Marketing* 34(3/4):384–98. doi:10.1108/03090560010311911.
- Goffman, Erving. 1981. *Forms of Talk*. University of Pennsylvania Press.
- Gouldner, Alvin W. 1960. "The Norm of Reciprocity: A Preliminary Statement." *American Sociological Review* 25(2):161–78. doi:10.2307/2092623.
- Johnson, Joseph, Gerard J. Tellis, and Deborah J. Macinnis. 2005. "Losers, Winners, and Biased Trades." *Journal of Consumer Research* 32(2):324–29. doi:10.1086/432241.
- Juslin, Patrik N., and Petri Laukka. 2003. "Communication of Emotions in Vocal Expression and Music Performance: Different Channels, Same Code?" *Psychological Bulletin* 129(5):770–814. doi:10.1037/0033-2909.129.5.770.
- Kahneman, Daniel, Barbara L. Fredrickson, Charles A. Schreiber, and Donald A. Redelmeier. 1993. "When More Pain Is Preferred to Less: Adding a Better End." *Psychological Science* 4(6):401–5. doi:10.1111/j.1467-9280.1993.tb00589.x.
- Ma, Zhenfeng, and Laurette Dubé. 2011. "Process and Outcome Interdependency in Frontline Service Encounters." *Journal of Marketing* 75(3):83–98. doi:10.1509/jmkg.75.3.83.
- Marinova, Detelina, Sunil K. Singh, and Jagdip Singh. 2018. "Frontline Problem-Solving Effectiveness: A Dynamic Analysis of Verbal and Nonverbal Cues." *Journal of Marketing Research* 55(2):178–92. doi:10.1509/jmr.15.0243.
- Mehrabian, Albert, and Susan R. Ferris. 1967. "Inference of Attitudes from Nonverbal Communication in Two Channels." *Journal of Consulting Psychology* 31(3):248–52. doi:10.1037/h0024648.
- Packard, Grant, and Jonah Berger. 2023. "The Emergence and Evolution of Consumer Language Research." *Journal of Consumer Research* ucad013. doi:10.1093/jcr/ucad013.

- Packard, Grant, Yang Li, and Jonah Berger. 2023. "When Language Matters." *Journal of Consumer Research* ucad080. doi:10.1093/jcr/ucad080.
- Palmatier, Robert W., Mark B. Houston, Rajiv P. Dant, and Dhruv Grewal. 2013. "Relationship Velocity: Toward a Theory of Relationship Dynamics." *Journal of Marketing* 77(1):13–30. doi:10.1509/jm.11.0219.
- PWC. 2018. "Experience Is Everything: Here's How to Get It Right."
- Qualtrics. 2024. "Bad Customer Service." <https://www.qualtrics.com/news/bad-customer-service-threatens-3-7-trillion-annually-as-frontline-workers-reach-a-breaking-point/>.
- SCHEGLOFF, EMANUEL A. 1999. "Discourse, Pragmatics, Conversation, Analysis." *Discourse Studies* 1(4):405–35. doi:10.1177/1461445699001004002.
- Smith, Amy K., and Ruth N. Bolton. 1998. "An Experimental Investigation of Customer Reactions to Service Failure and Recovery Encounters: Paradox or Peril?" *Journal of Service Research* 1(1):65–81. doi:10.1177/109467059800100106.
- Smith, Amy K., Ruth N. Bolton, and Janet Wagner. 1999. "A Model of Customer Satisfaction with Service Encounters Involving Failure and Recovery." *Journal of Marketing Research* 36(3):356–72. doi:10.1177/002224379903600305.
- Solomon, Michael R., Carol Surprenant, John A. Czepiel, and Evelyn G. Gutman. 1985. "A Role Theory Perspective on Dyadic Interactions: The Service Encounter." *Journal of Marketing* 49(1):99–111. doi:10.1177/002224298504900110.
- Vaswani, Ashish, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Ł. ukasz Kaiser, and Illia Polosukhin. 2017. "Attention Is All You Need." in *Advances in Neural Information Processing Systems*. Vol. 30. Curran Associates, Inc.
- Voigt, Paul, and Axel Von Dem Bussche. 2017. *The EU General Data Protection Regulation (GDPR)*. Cham: Springer International Publishing.
- Wieseke, Jan, Anja Geigenmüller, and Florian Kraus. 2012. "On the Role of Empathy in Customer-Employee Interactions." *Journal of Service Research* 15(3):316–31. doi:10.1177/1094670512439743.
- Xiao, Bo, Daniel Bone, Maarten Van Segbroeck, Zac E. Imel, David C. Atkins, Panayiotis G. Georgiou, and Shrikanth S. Narayanan. 2014. "Modeling Therapist Empathy through Prosody in Drug Addiction Counseling." Pp. 213–17 in.