

The company we keep: Prosodic similarity within and across organizations

Oliver Niebuhr¹

¹*Institute for Mechanical and Electrical Engineering, CIE, University of Southern Denmark
Alsion 2, DK-6400 Sonderborg, Denmark*

This paper is a revised version of the author's submitted manuscript that was peer-reviewed by at least two anonymous reviewers.

Abstract

Do speakers from different organizations sound systematically different? Using the COMPASS corpus of more than 500 L2 English business presentations from 18 companies across 10 industry sectors, we analyzed 28 acoustic-prosodic features of speaking style. Linear discriminant analyses and MANOVA reveal clear non-random clustering of prosodic patterns. Both company membership and industry-sector influence speaking style, but company-level effects are consistently stronger. Tempo, vocal effort, formants, and pausing are the most influential company predictors, particularly in organizations where employees interact frequently in stable teams or social-service contexts, while pitch variation becomes more relevant at the sector level, particularly in communication-intensive professions.

Keywords: *prosody, speaking style, workplace, communication*

1. Introduction

Many people share the same impression, e.g., when calling a company's customer-service line. You speak to one agent, then the call is transferred, and another person answers. Yet the voice sounds surprisingly similar, i.e. the tempo, the voice quality, the rhythm of the sentences etc. About two years ago, a Reddit user [1] stated: "Everytime I call customer service for literally anything ! Like bell, koodo, FedEx, ikea, etc ... I fee[l] like it's literally the exact same woman I'm speaking with." The Reddit post then culminates in the question: "Why does every customer services person sound like the exact same person?" Whether this impression is accurate is rarely addressed. But it raises an interesting question: Do groups of people who work together develop similar speaking styles? That is, does

a whole company speak 'with one voice'?

In addition to the anecdotal evidence above, research in sociolinguistics and phonetics suggests that this idea is not implausible. Speakers constantly adapt aspects of their speech to their communicative environment. One well-known framework describing such adjustments is Communication Accommodation Theory. It states that speakers tend to converge toward the speech patterns of their interlocutors in order to manage social relations and communicative goals [2]. Experimental studies have repeatedly demonstrated such phonetic convergence, showing that speakers often align features such as speaking rate, vowel quality, fundamental frequency (f0) level, consonant articulation, rhythm, smiling, etc. during interaction [3,4,21].

If people interact repeatedly, these momentary adjustments can accumulate. Over time, shared speaking habits may emerge. Sociolinguistic research thus treats speech styles as social practice. In the Communities-of-Practice framework, groups that collaborate regularly and pursue shared goals gradually develop characteristic communicative norms [5]. These patterns may become socially recognizable and acquire meaning. Processes of enregisterment describe how certain ways of speaking come to signal particular identities, roles, or contexts [6,7]. From this perspective, speaking styles are not only individual traits. Rather, they reflect the environments in which speakers work and interact.

Modern organizations provide a setting where such processes might occur. Companies are dense communicative networks [8]. Employees interact in meetings, presentations, negotiations, and customer conversations. In these situations, expectations about how people should speak may develop over time [9]. For example, organizations often value clarity, confidence, or persuasive delivery in professional communication. Research on corporate culture has long emphasized that shared norms influence behavior inside organizations [10]. Studies of workplace interaction similarly show that communication practices are shaped by institutional roles and, in addition, hierarchies [11,12,13]. Many companies also invest in presentation

Corresponding author: olni@sdu.dk

Copyright: ©2026 Oliver Niebuhr. This is an open-access article distributed under the terms of the Creative Commons Attribution 3.0 Un-ported License, which permits unrestricted use, distribution, and re-production in any medium, provided the original author and source are credited.

training to improve the skills of their employees.

On the other hand, companies are rarely homogeneous environments. Employees belong to different departments, perform different tasks, and interact with different audiences. Engineers, managers, consultants, and customer service agents often communicate in very settings—and even with different levels of vocal skill and charisma [14]. These differences may prevent a unified speaking style from emerging at the organizational level. In other words, companies provide conditions that might encourage similar speech—but they also contain strong sources of variation.

A second possibility is therefore that prosodic similarities emerge at a different level: not within single organizations, but across professional or industry contexts [15]. Rather than companies themselves speaking “with one voice” (if such a phenomenon exists at all), industry sectors or professions may do so. Many industry sectors and professions involve similar communicative tasks as well as communication frequencies and occurrences. Teachers explain complex ideas. Consultants present recommendations. Sales professionals persuade customers. Engineers describe technical solutions. These recurring activities place specific demands on delivery.

Research on professional discourse has shown that communication practices are closely tied to the tasks and identities associated with particular forms of work [12,14]. Linguistically, such recurring contexts often give rise to professional registers, i.e., systematic patterns of language use associated with particular activities or domains [15,16,17].

Evidence from voice research supports this view. Studies on occupational voice use show that different professions impose distinct vocal demands [18]. Sociophonetic research also demonstrates that phonetic variation often carries social meaning and contributes to how listeners perceive speaker identity [19,20,22].

In summary, there are a lot of indirect indications that “the company we keep”... shapes the way speak. However, much sociophonetic research has focused on regional dialects or social categories such as age, gender, or class. Comparatively little work (if any so far) has examined whether speakers in the same company—or in the same industry—develop measurable similarities in their prosodic behavior. One reason for this gap is practical: it is difficult to obtain large datasets of real-world professional speech that allow systematic comparison across organizations.

The present study addresses this gap using a unique corpus of speech recordings collected in professional public-speaking trainings. This allows to examine, for the first time, whether prosodic similarity emerges primarily within companies and/or whether broader professional or industry-related patterns structure the prosodic space of professional speech. In particular, the present study addresses the following questions:

- Is the interindividual variation in the dataset randomly distributed across speakers, or do clusters of similar prosodic patterns emerge?
- If such clusters exist, can they be associated with specific companies or broader industry sectors?

- Which factor provides the stronger grouping principle for prosodic similarity: membership in a particular company or participation in a shared professional or industry context? In other words, do speakers primarily sound like their colleagues within the same organization, or like professionals who perform similar communicative tasks or work within the same industry sector across different organizations?

2. Method

2.1 Speech material: COMPASS corpus

The analyses of speech prosody conducted in the present study were based on the COMPASS corpus (Corporate and Organizational Measurement of Prosody in Applied Speaking Situations.) The corpus contains speech recordings from authentic public-speaking scenarios. More specifically, the recordings come from workshop-based communication trainings run or equipped by AllGoodSpeakers ApS, a Danish speech-technology startup that develops science grounded, voice-centered communication training formats, supplemented by dedicated hardware and software tools.

The trainings were typically carried out in L2 English with groups of 8 to 12 participants and focused on improving professional speaking performance in business-related presentation contexts. Only recordings from participants who had explicitly consented to the storage and anonymized scientific use of their audio data were included in COM-PASS. Company names are not disclosed in the present study for reasons of confidentiality. Instead, the dataset is described in terms of speaker numbers, recording durations, and broad industry sectors.

In total, the COMPASS corpus comprises recordings from 540 speakers affiliated with 18 different companies. The total amount of analyzed speech material is 18 h 8 min 28 s. The gender distribution of the speaker sample is 63 % male (N=340) and 37 % female (N=200). The represented industry sectors cover a broad range of professional contexts, including international mail order, IT and software development, business consultancy, accounting, auditing, tourism, banking and finance, recruitment, humanitarian aid, engineering, health organizations, and the automotive sector.

In terms of speaker distribution, the largest group in the corpus comes from business consultancy (178 speakers), followed by accounting and auditing (92), IT and software development (69), and banking and finance (64). Smaller but still clearly represented groups include recruitment (34), automotive (23), health organizations (23), and engineering (19). Additional speakers come from international mail order (13), humanitarian NGOs (13), and tourism (12).

This diversity is important for the present study, as it allows us to compare potential prosodic grouping effects at the level of both individual companies and broader industry sectors. Table 1 summarizes the composition of the corpus across the 18 company-specific subsamples.

Table 1. Breakdown of the COMPASS corpus.

ID	SECTOR	SAMPLE	TIME
C1	International mail order	13	34:26
C2	IT / software development	15	20:40
C3	IT / software development	54	93:15
C4	Business consultancy	29	61:13
C5	Accounting and auditing	11	18:11
C6	Tourism	12	28:25
C7	Banking and finance	28	59:52
C8	Business consultancy	65	101:38
C9	Accounting and auditing	81	147:58
C10	Recruitment	34	77:57
C11	NGO / humanitarian aid	13	31:48
C12	Engineering	19	56:10
C13	Business consultancy	30	62:52
C14	Health organization	23	32:52
C15	Banking and finance	36	81:14
C16	Engineering	23	54:23
C17	Business consultancy	44	104:27
C18	Business consultancy	10	21:07
TOTAL	—	540	1088:28

2.2 Speech-production scenarios

The speech material consists of short professional presentations produced during the communication trainings described above. In these workshops, participants were asked to present a topic from their everyday work context. Typical presentation scenarios included idea pitches, short project presentations, or explanations of products, services, or strategic initiatives. The tasks therefore reflected communicative situations commonly encountered in business environments. The presentations were delivered in front of the other workshop participants and the trainer, either in a shared physical training space or within an online video-conferencing setting. The recordings can thus be classified as instances of semi-spontaneous speech.

The recordings themselves were made either in typical workshop environments such as meeting rooms or, in the case of online sessions, through standard video-conferencing setups using hi-fi microphones and no additional noise-cancellation and audio-compression software settings. Before acoustic feature extraction, all recordings were linearly amplitude-normalized to the maximum signal amplitude. This procedure removed differences in recording gain and overall signal level from the analysis. Recording conditions were never-

theless not fully standardized with respect to room acoustics, microphone type, and mouth-to-microphone distance. Such differences may mainly affect voice-quality and partly formant-related spectral measures, whereas temporal and f0-based parameters are comparatively robust. This is returned to in the Discussion. Importantly, both presentation settings (face-to-face and video-mediated) reflect common communication practices in contemporary professional life. So, the corpus captures a range of authentic speaking situations representative of modern workplace communication.

2.3 Acoustic analysis

All recordings were analyzed using a custom Python-based feature extraction pipeline, kindly provided by AllGoodSpeakers Aps. The selected set of 28 features was designed to capture central acoustic parameters that have been shown in previous studies to be particularly informative for the differentiation and classification of speaker-related traits such as trustworthiness, competence, enthusiasm, attractiveness, and passion, as well as for the quantification of presentation and vocal performance skills [28,29,30].

The pipeline first applied a band-pass filter (80–8000 Hz) to retain the speech-relevant frequency range. From the filtered signal, several acoustic representations were derived, including the f0 contour, the intensity (RMS) contour, formant (F1–F3) trajectories, and phrasing based on inter-pausal units (IPUs, silent-pause threshold 200 ms).

From these representations, a set of global prosodic features was computed for each recording. F0-related measures comprised mean, and standard deviation of f0, the ratio between f0 range and f0 standard deviation ($f0_{range}/f0_{std}$) as an indicator of pitch distribution structure, as well as the 10th and 90th percentiles of the F0 distribution as f0 minima and maxima. The f0 range used in the present study was defined as the difference between the 90th and 10th percentile (Q90–Q10) of the f0 distribution. Excluding the upper and lower 10 % of f0 observations was meant to reduce the influence of potential octave errors and other extreme outliers on f0 measurements.

Furthermore, the analysis also took into account several spectral-energy parameters like the mean levels of F1, F2, and F3, the mean RMS level and its standard deviation (in dBA); and several spectral tilt or voice quality features that are well known to vary across speaking styles and vocal effort levels: H1–H2, H1–A1, H1–A2, and H1–A3 and their corresponding standard deviations (all formant-corrected).

Additional temporal measures were obtained using the Praat-based syllable nuclei detection script developed by de Jong and Wempe [23]. From this script, three measures were directly included in the analysis: speech rate, articulation rate, and average syllable duration (ASD). These measures complement the timing-related parameters derived from the Python pipeline. Moreover, from the combination of both analysis pipelines, we also included in our set of 28 measures the frequency of pauses expressed as pauses per second, and the mean

pause duration, computed as the ratio between total phonation time and the number of detected pauses.

To reduce the influence of physiological differences between male and female speakers, several acoustic measures were additionally gender-normalized. Since male speakers constituted the larger group in the corpus, acoustic measures from female speakers were scaled toward the male reference range. Following empirical estimates of typical gender differences in speech acoustics, female F0 values were multiplied by 0.57 [24,25], and formant frequencies were multiplied by 0.85, following [26,27].

2.4 Inferential statistics

In order to test whether speakers from different companies exhibit distinguishable prosodic patterns (RQ1), we performed a discriminant analysis (LDA) with company membership (1-18, see Tab.1) as the grouping variable and the 28 extracted acoustic features as predictors. To examine whether prosodic similarities might instead align with professional contexts (RQ2), we conducted a second LDA using the same 28 acoustic predictors but industry sector (1-10, see Tab.1) as the new grouping variable. While discriminant analyses focus on classification performance, the third research question (RQ3) concerns the existence of statistically reliable multivariate differences between groups. To address this question, we conducted multivariate analyses of variance (MANOVA) with the 28 acoustic features as dependent variables and either company or industry sector as the grouping factor.

3. Results

3.1 Company-level prosodic differentiation

The discriminant analysis conducted at the level of individual companies revealed clear multivariate differences in the prosodic profiles of speakers from different companies. Overall discrimination across the 18 companies was highly significant (Wilks' $L = .098$, $X^2(442) = 1045.37$, $p < .001$), indicating that the combined 28-item acoustic feature set reliably distinguishes between organizational groups.

In the LDA model itself, the first six discriminant functions contributed significantly to group separation. The first function alone already accounted for 28.5 % of the discriminant variance, followed by the second (13.6 %) and third function (11.4 %). All statistically significant six functions jointly accounted for 79.2 %. In terms of classification performance, on average 39.7 % of the 540 speakers in the COMPASS corpus were correctly assigned to their original companies solely based on their ways of speaking (23.6 % under leave-one-out cross-validation). This level of accuracy exceeds the chance level for 18 equally likely groups (5.6 %) by factors of 4-8 indicating clear company prosodies. Classification accuracy varied strongly across companies, though. Several companies showed very high classification rates, with correct identification reaching between roughly 60 % and 90 % of the cases/speakers. By contrast, other companies were much more difficult

to distinguish, with classification rates falling well below 30 %.

Inspection of the structure matrix shows that the strongest contributors to group differentiation were primarily tempo- and pause-related variables, including mean speaking rate, pause-per-s rate, and mean articulation rate, as well as several spectral and resonance-related parameters such as the mean F2 and F3 levels, and the standard deviations of H1-A3 and RMS (dBA). By contrast, f0 measures contributed little to the discrimination between companies.

Sector-level prosodic differentiation

Overall discrimination across the 10 industry sectors was statistically significant as well (Wilks' $L = .314$, $X^2(234) = 525.53$, $p < .001$), i.e. prosodic patterns vary systematically across professional domains. However, the strength of this discrimination was clearly weaker than for the company-level analysis (Wilks' $L = .098$). This difference is notable as the industry model involved substantially fewer categories (10 sectors instead of 18 companies), which would normally facilitate classification rather than reduce it. The first four discriminant functions of the LDA model contributed significantly to sector separation. The first function accounted for 28.9 % of the discriminant variance, followed by the second (20.5 %), third (16.0 %), and fourth function (10.9 %). Together, they explained 76.4 % of the total discriminatory variance. Classification performance was also lower than at the company level. 35.7 % of the originally grouped cases were correctly classified (or only 23.0 % were correctly classified under leave-one-out cross-validation.) Although both percentages are clearly above chance level (10 %), the classification accuracy is nonetheless considerably lower than for the company-level analysis.

The best-performing sectors reached classification accuracies of 60–70 %, whereas many sectors were much harder to distinguish, with rates typically falling between 25 % and 35 %. Compared to the company-level analysis (i.e. about 30–90 %), the industry-level model shows a narrower range of classification performances, i.e. no sectors approached near-perfect discrimination, but also none were close to chance-level performance.

Inspection of the structure matrix indicates that the same general parameter classes contributed most strongly to discrimination as in the company-level model. The strongest contributors again came primarily from tempo- and pause-related variables, including speech rate, articulation rate, pause-per-s rate, and mean pause duration, together with several spectral and resonance-related parameters.

However, unlike in the company-level analysis, some pitch-related measures showed a more noticeable contribution to group separation. In particular, standard deviation of f0 and the ratio between f0 range and f0 standard deviation had relevant loadings in the first and third discriminant functions (.379 and .229). So, while pitch-related parameters still played a secondary role compared with temporal and spectral features, this pattern suggests that especially pitch variability contributes somewhat more to differentiation across industry sectors than across individual companies.

3.2 MANOVAs at company and sector level

A multivariate analysis of variance confirmed that the acoustic feature space also differed significantly across the 10 industry sectors (Wilks' $L = .341$, $F(216, 3775.61) = 2.372$, $p < .001$, partial $N^2 = .113$). This result supports the sector-level discriminant analysis and shows that the observed differences between sectors reflect a statistically reliable multivariate separation in the prosodic feature space rather than classification performance alone. At the same time, the sector-level MANOVA was somewhat weaker than its company-level counterpart (Wilks' $L = .118$, $F(408, 6127.20) = 2.494$, $p < .001$, partial $N^2 = .118$).

4. Discussion

The present study investigated whether speakers resemble “the company they keep” – not only socially, but also acoustically. The first research question asked whether prosodic variation in the dataset is randomly distributed across speakers or whether clusters of similar speaking styles emerge. The answer is clear: variation is not random. Both discriminant analyses and MANOVAs showed reliable clustering in the 28-dimensional prosodic feature space.

The second research question asked whether these clusters can be associated with companies or industry sectors. The results suggest that both levels matter. Significant discrimination was observed at company and sector level, indicating that professional speaking styles reflect both broader communicative tasks and the organizational environments in which speakers work.

The third research question asked which level provides the stronger grouping principle. Here, the results were also clear: company membership was more predictive than industry sector. Classification accuracy, effect sizes, and discriminant structure were consistently more pronounced for companies. In short, speakers resembled the “company they keep” more strongly than the industry or profession in which they operate.

The variation in company-level classification performance supports this interpretation. The most reliably discriminated companies were in international mail order, accounting/auditing, healthcare, tourism, and banking/finance, i.e. environments that often involve relatively stable teams or offices and frequent internal interaction. Such conditions favor prosodic accommodation and alignment, through which locally shared speaking habits may emerge [2,5,6,7]. By contrast, less discriminable companies belonged mainly to business consulting, engineering, and software development, where mobile, project-based work and changing teams may reduce sustained within-company convergence [31,32,33]. At the sector level, classification performance appeared to track communicative intensity. The most clearly discriminated sectors, i.e. international mail order, tourism, healthcare, and humanitarian aid, all involve regular spoken interaction with customers, patients, or the public. Sectors with less intensive everyday speaking demands showed weaker separation. This pattern, including the larger role of f_0 in the sector

analyses, is consistent with evidence that communication-intensive professions use and develop more distinctive prosodic patterns [18,34,35].

Although the results suggest that organizational affiliation shapes prosodic similarity more strongly than industry sector, this conclusion remains provisional because company, sector, and institutional role are only partly separable in the present sample and role annotations are unavailable. Variation in recording conditions probably weakened rather than created the observed clustering. Larger, sector-balanced datasets with explicit role labels are needed to disentangle these influences. To conclude, the results provide first empirical support for the idea that professional speaking styles partly emerge from organizational communication environments. Industry contexts and communicative demands clearly matter, but everyday interaction within organizations also shapes how professionals sound. This has implications for communication training and strategic voice branding: change may spread through key individuals, but sustainable effects may require group-level interventions.

5. Acknowledgments

The author thanks his two reviewers and AllGoodSpeakers ApS for the kind support and herewith discloses that he is also the CEO of that company, please see <https://oliver-niebuhr.com/conflict-of-interest.html> for a COI statement.

References

- [1] Reddit: “Why does every customer service person sound the same?”, Available: https://www.reddit.com/r/NoStupidQuestions/comments/1aqr8w/why_does_every_customer_services_person_sound/
- [2] H. Giles, J. Coupland, and N. Coupland (eds.): *Contexts of Accommodation: Developments in Applied Socio-linguistics*. Cambridge: CUP, 1991.
- [3] J.S. Pardo: “On phonetic convergence during conversational interaction,” *J. of the Acoustical Society of America*, 119, 4, pp. 2382–2393, 2006.
- [4] R. Levitan and J. B. Hirschberg, “Measuring acoustic-prosodic entrainment with respect to multiple levels and dimensions,” in *Proc. Conference on Empirical Methods in NLP*, Edinburgh, UK, pp. 1–5, 2011.
- [5] P. Eckert and K. Brown: “Communities of practice,” in *Concise Encycl. of Pragmatics*, pp. 109–112, 2006.
- [6] A. Agha: “Enregisterment,” in *Oxford Research Encyclopedia of Linguistics*, 2024.
- [7] M. Elmentaler and O. Niebuhr: “Platt –Missingsch –Petuh: Enregisterment zwischen Hamburg und Flensburg,” in L. Anderwald and J. Hoekstra (eds.), *Enregisterment*. Frankfurt: Peter Lang, 2017.



- [8] J. Kush, L. Argote, and B. Aven: "The effects of communication networks on shared social identity and group performance," *Small Group Research*, 56, 6, pp. 899–949, 2025.
- [9] G. Hofstede: *Culture's Consequences: Comparing Values, Behaviors, Institutions and Organizations Across Nations*. Thousand Oaks: Sage, 2001.
- [10] E. H. Schein: *Organizational Culture and Leadership*. San Francisco: Jossey-Bass, 2010.
- [11] S. Sarangi and C. Roberts: *Talk, Work and Institutional Order: Discourse in Medical, Mediation and Management Settings*. Berlin: M. de Gruyter, 1999.
- [12] K. C. C. Kong: *Professional Discourse*. Cambridge: Cambridge University Press, 2014.
- [13] S. W. Gregory Jr. and S. Webster: "A nonverbal signal in voices of interview partners effectively predicts communication accommodation and social status perceptions," *Journal of Personality and Social Psychology*, 70, 6, pp. 1231–1240, 1996.
- [14] O. Niebuhr: "Do engineers need to become more persuasive? About the why and how of vocal-charisma training," *IEEE Eng. Management Rev.*, pp. 1-9, 2025.
- [15] G. O'Grady: "Is there a role for prosody within register studies: and if so what and how?," *Language, Context and Text*, 2, 1, pp. 59–90, 2020.
- [16] D. Biber and S. Conrad: *Register, Genre, and Style*. Cambridge: Cambridge University Press, 2019.
- [17] P. A. Barbosa, S. Madureira, and P. B. de Mareüil: "Cross-linguistic distinctions between professional and non-professional speaking styles," in *Proc. Inter-speech*, Stockholm, Sweden, pp. 3921–3925, 2017.
- [18] I. R. Titze, J. Lemke, and D. Montequin: "Population in the U.S. workforce who rely on voice as a primary tool of trade: A preliminary report," *Journal of Voice*, 11, 3, pp. 254–259, 1997.
- [19] M. Armstrong, M. Breen, S. Gooden, E. Levon, and K. M. Yu: "Sociolectal and dialectal variation in prosody," *Language and Speech*, 65, 4, pp. 783–790, 2022.
- [20] D. Crystal: "Prosodic and paralinguistic correlates of social categories," in *Social Anthropology and Language*, pp. 185–206. London: Routledge, 2013.
- [21] J. Michalsky, H. Schoormann, and O. Niebuhr, "Conversational quality is affected by and reflected in prosodic entrainment," in *Proc. 9th International Conference on Speech Prosody*, Poznań, Poland, pp. 389–392, 2018.
- [22] R. M. Olsen: *It's All in How You Say It: Prosodic Cues to Social Identity and Emotion*. PhD dissertation, University of Georgia, USA, 2022.
- [23] N. H. de Jong and T. Wempe: "Praat script to detect syllable nuclei and measure speech rate automatically," *Beh. Res. Methods*, 41, pp. 385–390, 2009.
- [24] G. Demenko, B. Möbius, and B. Andreeva, "Analysis of pitch profiles in Germanic and Slavic languages," in *Proc. of Forum Acusticum 2014*, pp. 7–12, 2014.
- [25] E. Pépiot, "Male and female speech: A study of mean F0, F0 range, phonation type and speech rate in French and American English speakers," in *Proc. Speech Prosody*, Dublin, Ireland, pp. 305–309, 2014.
- [26] M. Lincoln, S. Cox, and S. Ringland: "A fast method of speaker normalisation using formant estimation," in *Proc. ICSLP*, Rhodes, Greece, pp. 1-4, 1997.
- [27] R. E. Turner, T. C. Walters, J. J. Monaghan, and R. D. Patterson: "A statistical, formant-pattern model for segregating vowel type and vocal-tract length in developmental formant data," *J. of the Ac. Soc. of America*, 125, 4, pp. 2374–2386, 2009.
- [28] Ī. Valls-Ratés, O. Niebuhr, and P. Prieto: "Encouraging participant embodiment during VR-assisted public speaking training improves persuasiveness and charisma and reduces anxiety in secondary school students," *Frontiers in VR*, 4, 1074062, 2023.
- [29] O. Niebuhr and R. Skarnitzl, "Measuring a speaker's acoustic correlates of pitch – but which?" in *Proc. of the 19th Int. Congr. of Phonetic Sciences*, pp. 1-5, 2019.
- [30] O. Niebuhr and N. Taghva: "Prosodic interfaces: Inter-disciplinary perspectives on sound patterns and human interaction," in M. Oliveira Jr. (ed.), *Prosodic Interfaces: Interdisciplinary Perspectives on Sound Patterns and Human Interaction*. Berlin: De Gruyter, pp. 181–221, 2025.
- [31] A. Chia: "Distilling the essence of the McKinsey way: The problem-solving cycle," *Management Teaching Review*, 4, 4, pp. 355–370, 2019.
- [32] R. G. Lazar and A. D. Kellner: "Personnel and organization development in an R&D matrix-overlay operation," *IEEE Trans. on Eng. Man.*, 2, pp. 78–82, 1964.
- [33] A. Buch and V. Andersen: "Team and project work in engineering practices," *Nordic Journal of Working Life Studies*, 5, 3a, pp. 27–46, 2015.
- [34] D. Tannen: *Conversational Style: Analyzing Talk Among Friends*. Oxford: OUP, 2005.
- [35] A. Buxó-Lugo, J. C. Toscano, and D. G. Watson: "Effects of participant engagement on prosodic prominence," *Discourse Processes*, 55, 3, pp. 305–323, 2018.

