

Narrating Authority Through Storytelling in Entrepreneurial Podcast Interviews

Author(s): Gelareh Farhadian¹, Faezeh S. Aarabi², Dave Keighron³

1 Fairleigh Dickson University, Vancouver, Canada.

2 University of Tehran, Tehran, Iran.

3 University Canada West, Vancouver, Canada.

Corresponding Author: g.farhadian@fdu.edu

Received: July, 2026, Published: Sep, 2026

ARTICLE INFO

Keywords:

Borrowed authority; computational discourse analysis; discursive leadership; entrepreneurial leadership; podcast interviews; reproducibility

© 2026 by the Author(s). This open-access article is distributed under a Creative Commons Attribution (CC-BY) 4.0 license, making research freely available to the public and supporting a greater global exchange of knowledge and human experiments.



ABSTRACT

How do entrepreneurs claim the authority to lead an enterprise that was never theirs alone? Although the discursive construction of leadership has attracted growing interest, most existing work is qualitative and case-limited, and leadership authority has rarely been examined as a discursive phenomenon at scale. Drawing on discursive, relational, and narrative leadership theory, this study analyzes 59 entrepreneurial podcast interviews from the Innovation Fuel series, recorded between 2020 and 2026, combining lexicon-based authority detection, sentence-embedding, topic modeling, and clustering. Under the stated lexical rules, soft classifications account for 89.8% of transcripts, and individually voiced and collectively voiced authority language routinely co-occur within the same narratives. Borrowed authority is used as a provisional interpretive label for legitimacy claimed individually yet drawn from teams and networks; it is offered as a testable interpretation rather than a separately validated construct. Recorded guest gender is neither associated with the episode-level collective ratio ($p = .205$) nor recoverable from the analyzed discourse out-of-sample (cross-validated AUC of about 0.59), and gender labels are interspersed in the embedding space. Because complete transcripts include host speech, these findings cannot be attributed to guest speech alone. Crucially, each apparent result is submitted to a matched null test: the soft-authority prevalence and the gender null survive, whereas apparent cluster and topic authority bands are shown to be aggregation artifacts over poorly separated groups. Distinguishing patterns that survive such tests from those that do not, the study models a transparent, self-auditing approach to computational discourse analysis and contributes a reproducible workflow for leadership communication research on large corpora of naturally occurring talk.

1.0 Introduction

Podcasts, interviews and digital communication are increasingly used platforms for entrepreneurial leaders to communicate their message. In situations like that, it is not just the entrepreneurs who speak about business activities; they also establish leadership authority through storytelling. The podcast interviews are, therefore, a useful context for the study of leaders' roles in conveying experience, structuring organizational development, and positioning themselves in relation to teams, communities, and the general audience.

This study is guided by two bodies of theory. Discursive and relational leadership theory views leadership as a process rather than as a trait and is developed through communicative and social practice (Fairhurst, 2007; Fairhurst & Grant, 2010; Uhl-Bien, 2006; Drath et al., 2008), which is particularly relevant for entrepreneurial discourse that focuses on cooperative action, shared goals, and group identity. Narrative leadership research adds that identity and legitimacy are constructed through storytelling (Gabriel, 2000; Shamir et

al., 2005; Auvinen et al., 2013). The following review elaborates on these strands and the computational techniques applied to their study.

Recent developments in computational social science offer more opportunities to investigate discourse at a larger scale. These approaches include topic modeling and semantic embeddings (Blei et al., 2003; DiMaggio et al., 2013; Devlin et al., 2019) that can be used to systematically study linguistic and narrative patterns in large text corpora. The methods enable the researchers to analyze themes, language patterns, and communication structures in stories of leadership.

Despite these developments, there are still a number of gaps. Although there is a significant body of literature on leadership discourse, there are few computational studies of leadership authority as a discursive phenomenon, and few small qualitative samples. Also, the emerging media storytelling (podcasts) has not been used extensively in regard to entrepreneurship. The paper fills these gaps by examining how leadership power is constructed in entrepreneurial podcast interviews. The study relies on a corpus of 59 transcripts to explore patterns of entrepreneurial storytelling through a combination of lexicon-based authority detection, semantic embedding, clustering and topic modelling.

2.0 Literature Review

The study is mainly conceptually grounded in discursive leadership theory. Instead of seeing leadership as an individual trait, it views leadership as a result of communicative processes that produce influence, meaning, and coordination (Fairhurst 2007), and argues that the focus should be on interaction, framing, and the negotiation of meaning (Fairhurst and Grant 2010). This is supported by the organizational discourse scholarship, which demonstrates that discourse is not just descriptive but also helps to shape organizational realities (Grant et al., 2004; Jian et al., 2008).

In this perspective, power is manifested in the ways speakers use language to claim leadership, position themselves as experts and as legitimate, grant themselves agency, and frame problems. Empirical studies demonstrate the ways in which leadership and workplace culture are manifested in everyday discourse (Holmes et al., 2007), and how organizational power and decision-making are discursively negotiated (Kwon et al., 2009). Recent research highlights the importance of examining leaders' actual textual and audiovisual communication in order to grasp the leadership process (Liu et al., 2023). This is reflected in the rhetorical framing, stance and positioning of narrative in entrepreneurial podcast interviews, where speakers assert their agency and expertise.

2.1 Relational Leadership Theory

Relational leadership theory focuses on the social relationships out of which leadership arises, while discursive leadership focuses on the communicative process. Uhl-Bien (2006) views leadership as a social process of influence that coordinates action and creates change through interaction; Cunliffe and Eriksen (2011) see leadership as a responsive and relational practice of dialogue and mutual accountability; and Drath et al. (2008) reimagine leadership as the production of direction, alignment, and commitment instead of fixed leader/follower roles.

This relational perspective sees power not as something that can be owned by individuals but as a social process produced by interaction. Entrepreneurial leaders frequently rely on partners, teams, investors, and the community as a way to validate themselves, and stories of community goals and shared accomplishments serve as a way to gain credibility. Leadership identity research reveals that individual leadership identities exist alongside collective leadership discourses (Empson et al., 2023), and leadership-as-practice research highlights the collaborative nature of leadership (Raelin, 2023). This means that entrepreneurial authority can be expressed collectively – we, our team, our community –

and that it can transform personal achievement into relational legitimacy for discourse analysis.

Narrative leadership research explains the centrality of storytelling in the construction of leadership authority. Narratives structure experience in terms of time, cause, and value, allowing speakers to tell coherent stories about themselves and the significance of their actions. Organizational stories are seen as carrying meaning, identity and emotion (Gabriel, 2000), a life-story approach is proposed to connect a leader's biography with identity (Shamir et al., 2005) and managerial storytelling is seen to build trust and shape interpretation (Auvinen et al., 2013).

Storytelling is particularly relevant in entrepreneurial settings where founders establish legitimacy through origin, struggle, and success stories. Recent research reveals the use of storytelling for identity work, authenticity, and impression management in leadership transitions (Felix et al., 2023), as well as public success stories for the impression of legitimacy and competence (Srivastava et al., 2023). Podcasts are therefore a genre well suited to the study of recurring narrative patterns, such as creation stories, failure-and-recovery narratives, and purpose-driven mission narratives, as these often form the basis of entrepreneurial authority.

The theories make different contributions to the empirical design. Discursive leadership theory identifies language as the site where authority claims are made; relational leadership theory focuses on whether agency and legitimacy are attributed to an individual or a collective and on whether the tone is assertive or relational; and narrative leadership theory situates those claims within temporally organized public stories. The computational measures operationalize parts of this framework, but no single lexical category is treated as equivalent to a theoretical construct.

Computational discourse analysis provides a methodological approach for analyzing leadership accounts in large text collections. The researcher can be given the opportunity to discover patterns, collocations, and discourse frames with the aid of corpus-based methods and a vast corpus of texts.

The methodological literature of recent years emphasizes the use of computational tools and theoretically informed discourse analysis. Some researchers, such as Chen et al. (2023), believe that topic modeling should be linked to research design and theoretical interpretation, and others like Matheson (2023) and Piper (2023) warn that computational analysis should not ignore the social context of discourse but should instead be mindful of it. To this end, machine learning techniques are now being employed to analyze leaders' language at scale (Liu et al., 2023). The developments are suitable for the analysis of entrepreneurial podcast interviews because of the high volume of narrative data that can be systematically analyzed using computational approaches while remaining grounded in discourse theory.

2.2 Research Gap

Although there is increased focus on the discursive construction of leadership, much of the research conducted so far is qualitative and relatively limited in terms of the number of cases. Research on entrepreneurship has explored issues of legitimacy and identity discourses, but leadership authority has not been studied specifically as a discursive phenomenon. Computational methods offer the possibility of scaling leadership discourse analysis, but they might lose track of language and interaction without a connection to discourse theory.

Podcast interviews can be a good way to bridge this gap. They are public, communicative events in which entrepreneurs seek to establish leadership authority before interviewers and socially oriented audiences. The present study draws on discursive, relational, and narrative leadership theory, as well as computational discourse analysis, to explore the construction of leadership authority in entrepreneurial storytelling, a new genre of

discourse, and to demonstrate that computational methods can be used to analyze large amounts of data without neglecting language, interaction, and narrative meaning.

3.0 Methodology

3.1 Research Design

This study uses a computational discourse analysis approach to investigate the construction of leadership authority in entrepreneurial storytelling. The analysis integrates lexicon-based authority indicators with semantic embeddings, topic modelling, and clustering, allowing for systematic analysis of leadership discourse in a large corpus while remaining rooted in discursive and relational leadership theory.

3.2 Podcast Corpus

The empirical information is from the Innovation Fuel podcast, a University Canada West business and entrepreneurship series. Its episodes consist of semi-structured interviews with entrepreneurs, executives and industry professionals, where they describe their leadership experiences, strategies and challenges, resulting in extensive naturally occurring leadership discourse in conversational, extended narratives. The series did not use a uniform research-interview schedule; archived host prompts recurrently addressed entrepreneurial origins, the business problem or offering, leadership and team building, partnerships, customers and markets, financing and scaling, technology and innovation, setbacks, community impact, gender or structural barriers where relevant, and closing advice.

There are 59 interview transcripts in the corpus, one for each episode. In both, a guest is interviewed about their journey as an entrepreneur, their experiences in leadership, and their knowledge of the industry. The conversational style allows the participants to express their intentions, means of action and shared experiences, and the transcripts are thus appropriate for analysing the construction of leadership authority in the form of a narrative.

The 59-row analysis dataset contains complete episode transcripts, including guest speech, interviewer prompts, introductions, and other conversational scaffolding. Speaker turns were not consistently separable in the source version used for the reported 59-transcript results. The measures must therefore be interpreted as episode-level discourse labeled by guest characteristics, not as pure measures of guest speech. This feature is especially relevant to collective pronouns and repeated introductory language and is treated as a limitation rather than silently attributed to the guest.

The corpus contains 304,202 words, with an average of approximately 5,156 words per transcript. It covers topics such as technology, consulting, healthcare, finance, and digital marketing, and features a diverse group of guests, mostly from North America but with international experience.

For the authority measures, complete transcripts were converted to plain text, lowercased, and searched directly using case-insensitive whole-word lexical matching. The authority analysis did not lemmatize the text or remove stop words.

3.3 Transcript Preprocessing

All transcripts had been preprocessed with standard pretexts in the computational text analysis before analysis so that there is uniformity in documents. To improve the quality of linguistic processing in the downstream stage, text pre-processing helps to reduce noise (Jurafsky, and Martin, 2023).

Preprocessing differed by analytic component. Authority counts used lowercased complete transcript strings and whole-word lexical matching. LDA used CountVectorizer with English stop-word removal, max_df = 0.9, and min_df = 3.

The corpus was cleaned and used for authority detection using lexicon based and semantic embedding analysis. These kinds of preprocessing are common in the study of computational discourse where the patterns formed in the corpus are required to capture a significant linguistic composition and not formatting artifacts (DiMaggio, Nag, and Blei, 2013).

3.4 Lexicon Construction

To operationalize leadership authority, the study adopts a dictionary-based lexical approach that identifies linguistic patterns associated with various forms of authority. As leadership is constructed through language, narrative framing, and interaction (Fairhurst & Connaughton, 2014; Clifton, 2012), the analysis of the linguistic construction of leadership narratives shows how leaders frame themselves, their team, and their organization in accounts of action, responsibility, and collaboration.

The lexicon is based on two dimensions: agency framing and leadership tone. The agency dimension captures the difference between individual and collective actor language, with evidence that leadership can be expressed as self-leadership or shared leadership across teams and networks (DeRue & Ashford, 2010), through the use of pronouns (I, my, we, our) and the mention of teams, communities, and organizations.

The second dimension, leadership tone, is a distinction between assertive and relational language: assertive language emphasizes strategic action, decision-making, and performance, while relational language emphasizes mentoring, collaboration, trust, and interpersonal engagement, which are well documented in the leadership literature (Eagly & Karau, 2002).

These dimensions produce four reported categories: Individual–Assertive, Individual–Soft, Collective–Assertive, and Collective–Soft. The analysis script fixes the terms assigned to each dimension, as reproduced in Table 1. The categories were theoretically informed and include expressions present in the corpus, but the surviving project documentation does not record a formal item-generation panel, an independent item-assignment exercise, or adjudication of disputed terms. The lexicon should therefore be treated as a transparent operationalization for evaluation, not as a previously validated scale.

The dictionary consists of unigrams and exact multiword expressions, including “my vision,” “working together,” and “deliver results.” Matching was case-insensitive and used whole-word boundaries. Each listed term contributed one match whenever it occurred; multiword expressions were counted as their exact sequences.

For each transcript, the script counts matches in the four lists. The collective ratio is $\text{collective_count} / (\text{collective_count} + \text{individual_count} + 1)$, where +1 is the smoothing constant used in the original analysis. A transcript is labeled Collective when this ratio is greater than .55 and Individual otherwise. It is labeled Assertive when assertive_count is greater than relational_count and Soft otherwise; ties are therefore classified as Soft. The cross of these two rules yields the four categories. The continuous counts and ratio remain available alongside the categorical label.

The theoretically motivated lexical categories combined with the computational text analysis provide a transparent and reproducible means of analysing the expression of leadership authority in entrepreneurial discourses. Terms were not context-disambiguated, and a term could contribute to more than one dimension. The analysis did not resolve pronoun referents or distinguish guest from host speech.

Table 1 reproduces the complete operational authority lexicon used in the 59-transcript analysis. Publishing every term makes category assignments inspectable and allows the measure to be reproduced or challenged.

Table 1. Operational Authority Lexicon

Dimension	Terms and expressions	Operational role
Individual agency (pronouns and possessives)	i; me; my; mine; myself; my company; my idea; my vision; my goal; my strategy	First-person self-reference and explicit self-owned entities or intentions.
Individual agency (action terms)	found; founded; build; built; create; created; launch; launched; start; started; establish; established; lead; led; drive; drove; decide; decided; manage; managed; direct; directed; develop; developed; design; designed	Action vocabulary counted in the individual-agency list; several terms also occur in the assertive list.
Collective agency (pronouns and groups)	we; our; us; together; team; group; collective; community; network; ecosystem; organization; organizations; company; companies; people; members; stakeholders	Collective pronouns and references to groups or organizational actors.
Collective agency (coordination terms)	collaborate; collaboration; partnership; partner; partners; shared; joint; mutual; collectively; togetherness; community building; working together; support each other	Shared-action and coordination vocabulary counted as collective agency.
Assertive tone (action terms)	build; scale; grow; drive; lead; achieve; deliver; execute; expand; accelerate; compete; win; dominate; transform; innovate; disrupt	Action, competition, and execution vocabulary.
Assertive tone (strategy and performance terms)	strategic; strategy; vision; target; results; performance; success; growth; impact; launch; create value; market leadership; optimize; implement; deliver results	Strategy, outcomes, growth, and performance vocabulary.
Relational tone (support terms)	support; mentor; mentoring; coach; coaching; empower; empowering; listen; listening; care; caring; trust; trusted; relationship; relationships; connect; connection; connected; help; helping; guide; guiding; encourage; encouraging; inspire; inspiring	Support, development, listening, trust, and relationship vocabulary.
Relational tone (collaboration and communication terms)	collaborate; collaboration; together; community; engage; engagement; understand; understanding; communication; communicate	Collaboration, engagement, understanding, and communication vocabulary.

Note. Terms are reproduced from the analysis script. Matching was case-insensitive and used whole-word boundaries; listed multiword expressions were matched as exact sequences. Overlapping membership was permitted, and no contextual word-sense or pronoun-referent disambiguation was applied.

3.5 Statistical Analysis

Patterns of authority language were analyzed across the corpus by a set of descriptive and inferential tests. The key authority indicators, such as collective ratio and relational-authority measures, were first summarized using descriptive statistics.

Independent-samples t-tests were then used to determine if there were differences between male and female speakers in collective- versus relational-authority language.

The main comparisons were made using bootstrap confidence intervals, with the data being resampled many times to get a nonparametric estimate of the stability of the estimates.

These statistical processes enabled a systematic analysis of the patterns of power language, and facilitated the discourse-based analyses that were carried out at later stages of the research.

3.6 Semantic Embedding Analysis

Each transcript was represented as a vector using a Sentence-BERT transformer model to examine the semantic connections within the collection of leadership stories. Transformer-based embeddings learn contextual meaning in writing by representing documents in a

high-dimensional semantic space, where narrative comparisons are based on semantic meaning rather than word frequency alone.

The document embeddings were then reduced to 2D using Uniform Manifold Approximation and Projection (UMAP), which reveals patterns of similarity among the transcripts while preserving their underlying semantic relationships.

In this 2D space, transcripts that are closer to each other have similar discourse structures, and a representation is created that allows exploration of similarities and patterns of variation in the leadership discourse in the corpus.

3.7 Cluster Analysis

K-means clustering was used to the semantic embeddings of the transcripts to determine underlying narrative structures in the corpus. Clustering techniques group documents by their similarity in semantic representation, without predefined categories, allowing patterns of discourse to be observed.

The elbow method was used to determine the number of clusters, and a four-cluster solution was chosen as the most interpretable representation of the corpus's semantic structure, as there was no significant decrease in within-cluster variance beyond four clusters.

The clusters indicate groups of transcripts that are similar in the semantic embedding space. These patterns were seen as separate leadership narrative patterns, and they were considered to be a sign of different ways of constructing and sharing leadership by the speakers through entrepreneurial storytelling.

3.8 Topic Modeling

The processed transcripts were then analysed using Latent Dirichlet Allocation (LDA) to find common themes, and to infer latent topics based on patterns of word co-occurrence across the documents.

The LDA model estimated latent topics, where each topic is a distribution over high-probability co-occurring keywords, and each transcript was represented as a bag-of-words input to the LDA model, yielding a distribution over topics.

The number of topics was selected for interpretability and stability. These topics complement the semantic and authority analyses, and represent important aspects of the discourse on entrepreneurial leadership, including innovation, professional development, community involvement, and entrepreneurial development.

3.9 Robustness and Acceptance Checks

Multiple reliability tests were performed. Lexicon-coverage analysis initially showed that authority indicators were present in all the transcripts. Second, correlation analysis was used to assess the relationships between the lexical authority measures and the semantic embedding dimensions.

Third, topic-model stability was evaluated by running the model multiple times; fourth, cluster validity was evaluated using the silhouette score; and finally, bootstrapping was used to evaluate the stability of the observed patterns in the main statistical comparisons.

Overall, these processes strengthen computational discourse analysis, as the result is not dependent on a single analytical approach.

3.10 Reproducibility and Computational Environment

The analysis was implemented in Python in Google Colab. The preserved pipeline uses pandas and NumPy for data handling, Python regular expressions for lexicon matching, scikit-learn for TF-IDF, LDA, clustering, and prediction, sentence-transformers for all-MiniLM-L6-v2 embeddings, umap-learn for projection, and matplotlib and seaborn for visualization.

The full analysis pipeline was implemented in a Google Colab notebook, and the code to build the corpus, extract features, and model it was made publicly available at https://github.com/GelarehGF/Leadership_Through_Storytelling to ensure reproducibility.

Recorded parameters in the preserved pipeline are: Sentence-BERT model all-MiniLM-L6-v2; UMAP random_state = 42; K-means n_clusters = 4, random_state = 42, and n_init = "auto"; CountVectorizer stop_words = "english", max_df = 0.9, and min_df = 3; and LDA n_components = 5 and random_state = 42.

3.11 Generative AI Statement

Only language editing and documentation support tools were based on generative AI. All data analysis, model implementation and results interpretation were done by the authors.

4.0 Results

The Results proceed from description to increasingly exploratory analysis. First, the corpus and four authority-category frequencies are reported. Second, episode-level associations with recorded guest gender, industry, and transcript length are examined. Third, semantic embeddings, clusters, and LDA topics are presented as descriptive maps. Finally, validation and matched robustness checks identify which patterns are retained and which are treated as artifacts. Interpretive claims about borrowed authority are reserved for the Discussion.

4.1 Corpus Description

The sample includes 59 podcast interview transcripts of the Innovation Fuel series which were recorded in the period of 2020-2026. All the episodes are between 30 and 40 minutes long, and they yield a corpus of about 304,000 words, with an average transcript length of about 5,156 words.

The speakers are representatives of various entrepreneurial roles across different industries. The gender proportion is nearly equal: 30 female orators (50.8%) and 29 male orators (49.2%). Most of the participants are in Canada, with the rest commenting on the United States, Brazil and other parts of the world. The demographic properties of the corpus are shown in Table 2.

Table 2. Demographic Characteristics of the Corpus

Variable	Category	Count	Percentage
Gender	Female	30	50.8%
	Male	29	49.2%
Leadership Role	Founder / Co-Founder	35	59.3%
	Executive (CEO / Director)	7	11.9%
	Expert / Advisor	7	11.9%
	Other	10	16.9%
Industry Group	Other	31	52.5%
	Tech	19	32.2%
	Services	5	8.5%
	Real Estate	2	3.4%
	Finance	2	3.4%
Geography	Canada	51	86.4%
	USA	4	6.8%
	Brazil / Canada	2	3.4%
	Canada / Poland	1	1.7%
	Sweden	1	1.7%

Note. Counts and within-variable percentages (N = 59 transcripts) for speaker gender, leadership role, industry group, and geographic location. The table establishes the near-even gender split and the predominance of founders and Canada-based speakers that frame the interpretation of subsequent analyses.

4.2 Authority Narrative Distribution

Soft-authority narratives predominate in the corpus. The most frequent style is Individual-Soft (28 transcripts, 47.5%), followed by Collective-Soft (25 transcripts, 42.4%). The

Individual-Assertive style is rarely used (5 transcripts, 8.5%), as is the Collective-Assertive (1 transcript, 1.7%). The discourse of leadership in the corpus thus reveals a much more supportive and relational style of leadership rather than directive or assertive leadership. The distribution is summarized in Table 3. This distribution is an output of the operational classifier and is not direct evidence that 89.8% of guests enact relational leadership.

Table 3. Distribution of Authority Narrative Styles

Authority Narrative Style	Count	Percentage
Individual-Soft	28	47.46%
Collective-Soft	25	42.37%
Individual-Assertive	5	8.47%
Collective-Assertive	1	1.69%

Note. Frequency and percentage of transcripts assigned to each of the four authority styles, where a transcript's style reflects its dominant position on the agency (Individual/Collective) and tone (Soft/Assertive) dimensions. The corpus concentrates in the two soft-authority styles (89.8% combined), and assertive framing is rare.

4.3 Gender and Authority Narratives

To analyze the differences between authority narratives between different speakers, the gender difference in terms of authority styles was made. Figure 1 shows the distribution of types of authority and is a heatmap of the type of gender and authority narratives.

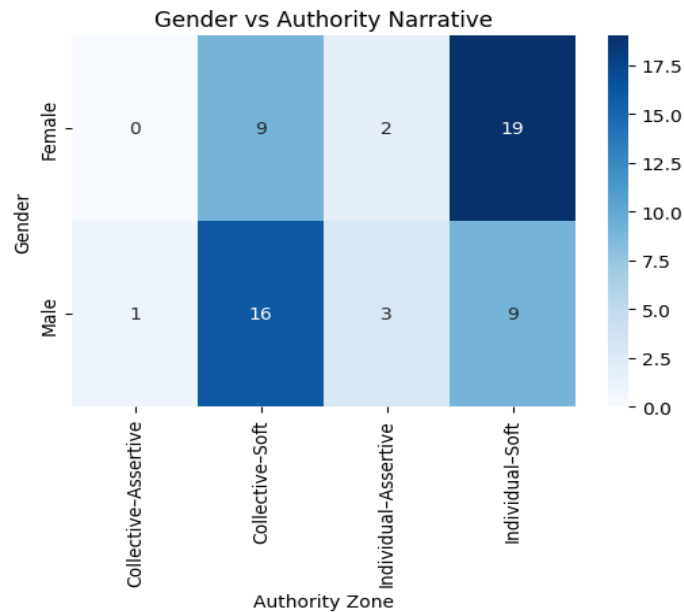


Fig 1. Gender vs. Authority Narrative

Note. Heatmap of transcript counts cross-classifying speaker gender (rows) against authority narrative style (columns), with cell shading proportional to frequency. The figure shows how soft-authority styles dominate within both gender groups and where the modest descriptive differences between male and female speakers occur.

Table 4 summarizes the cross-tabulation of gender and authority styles. Individual-Soft narratives (19 transcripts) were the most common among female speakers, then there was Collective-Soft narratives (9 transcripts). Collective-Soft narratives were more common among male speakers (16 transcripts), whereas the Individual-Soft narratives were found in 9 transcripts. Few aggressive authority narratives were found in both groups.

Table 4. Gender × Authority Narrative Styles

Gender	Collective-Soft	Individual-Soft	Individual-Assertive
Female	9	19	2
Male	16	9	3

Note. Cross-tabulation of transcript counts by speaker gender and authority style. The Collective-Assertive style is omitted because it did not co-occur across both gender groups. The table provides the descriptive basis for the non-significant gender comparison reported below.

The distribution of collective ratio in transcripts is shown in a boxplot of collective authority scores by gender (Figure 2). The distribution of the collective authority scores appears to be slightly higher for male speakers, however the median score is very similar.

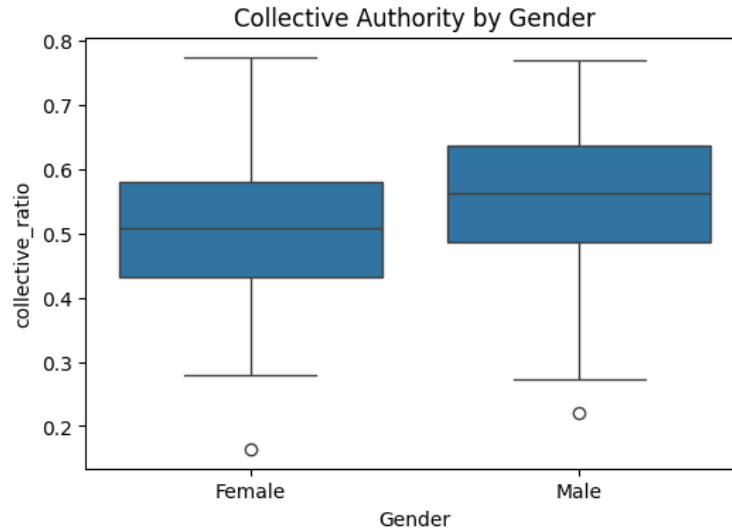


Fig 2. Collective Authority by Gender

Note. Boxplots of the collective-authority ratio (proportion of collective-agency lexical matches per transcript) for female and male speakers, showing medians, interquartile ranges, and outliers. The substantial overlap of the two distributions predicts the non-significant independent-samples t-test ($p \approx 0.205$).

An independent-sample t-test was used to determine if these differences were statistically significant, using the collective_ratio variable. The results showed that the difference between the male and female speakers is not statistically significant ($p \approx 0.205$).

The study results show that descriptive differences in narratives about authority and gender are not significant predictors of the language of collective authority in the corpus.

4.4 Does Gender Generalize? An Out-of-Sample Test

The comparisons above are in-sample and test each theme in isolation, where small uncorrected differences are most likely to occur by chance. To ask the stronger question of whether speaker gender is recoverable from the leadership self-presentation as a whole, we trained a regularized logistic classifier to predict gender from the discourse and tested it out-of-sample with repeated stratified five-fold cross-validation, testing for significance against a null model formed by permuting the gender labels. We used two feature representations: the eight theme markers and a term-frequency-inverse-document-frequency (TF-IDF) vocabulary. The classifier did not discriminate between women's and men's self-descriptions in either case (Table 5): the cross-validated ROC-AUC was 0.59 for the theme markers (permutation $p = .17$) and 0.59 for the TF-IDF vocabulary (permutation $p = .23$). Because complete transcripts include host speech, this result is episode-level non-prediction in this corpus and does not establish that guest gender is irretrievable from guest speech. The cross-validated ROC-AUC and permutation p-values are shown for each feature.

Table 5. Out-of-Sample Prediction of Guest Gender from Episode-Level Discourse

Feature representation	Features	CV ROC-AUC	Permutation p
Eight theme markers	8	0.59	.17
TF-IDF vocabulary	765	0.59	.23

Note. Repeated stratified five-fold cross-validated ROC-AUC for a regularized logistic classifier, with significance assessed against a label-permutation null. Neither representation predicts the recorded gender of the episode's guest above chance in this corpus. The analyzed complete transcripts include host speech.

This out-of-sample null does not imply that there are no individual themes that vary by gender. Women are more likely to use developmental and mentorship language in their self-descriptions than men (45% vs 19%, odds ratio 3.5, 95% CI [1.1, 10.2]; risk difference 26 percentage points) and the naming of gender or structural barriers is similarly more common in women than in men (21% vs 3%) although is limited by only one male speaker and is underpowered. These are reported as local differences of emphasis, and with two caveats put forward explicitly. First, neither difference is significant across all eight themes, even when using the Benjamini–Hochberg correction (mentorship $q = .21$); the two effects are only significant in a form of confirmatory test, but the hypotheses were selected after examining the exploratory results. Second, the theme markers are subject to the reliability caveat mentioned in the Robustness Checks. The straightforward reading is that gender influences, if at all, only the emphasis of one communal theme; it does not create a discursive signature that a model can recover.

Combined with the non-significant collective-ratio comparison ($p = .205$) and the visual interspersion of gender labels in the embedding projection, these analyses do not provide evidence of a recoverable gender-wide discourse signature in this corpus. They neither confirm distinctly gendered leadership-communication styles nor refute role-congruity theory (Eagly & Karau, 2002), because the sample is small and self-selected, the analyses are partly exploratory, the episodes share a common interview genre, and the complete transcripts include host language.

4.5 Authorities Across Industries

Authority-category counts were compared descriptively across the five archived broad industry groups: Technology, Services, Finance, Real Estate, and Other. Consumer cases are included in the heterogeneous Other group, consistent with the 59-case margins in Table 2. No inferential test of industry differences was conducted, and several groups contain too few transcripts for stable comparison.

Figure 3 displays category counts by archived broad industry group. Soft classifications appear in every group, while assertive classifications are sparse. Given the small cells, the broad residual category, and the absence of a preserved grouping rule, the figure is used only to describe this sample.

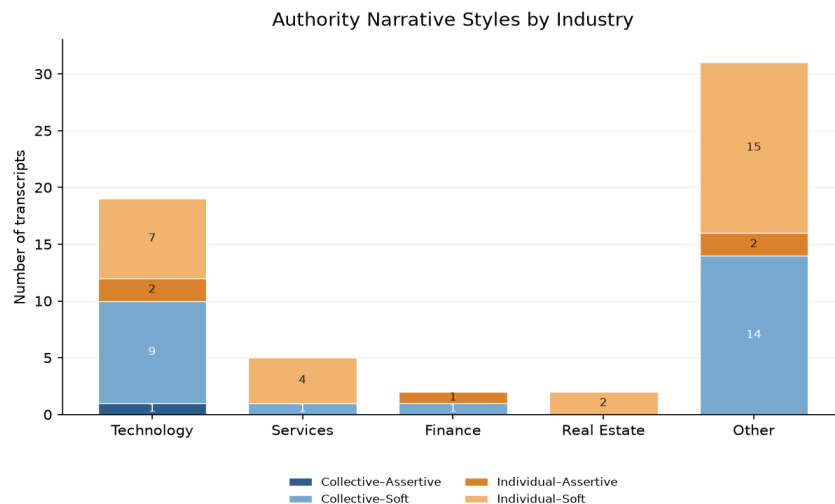


Fig 3. Authority Narrative Styles by Industry

Note. Stacked bars show the number of transcripts in each authority category within the five archived broad industry groups ($N = 59$). Consumer cases are included in Other. The cells are descriptive; the project files do not preserve the rule used to assign detailed industry labels to these broad groups.

Individual-Soft cases appear in several groups, while assertive cases are sparse. Because Finance and Real Estate contain only two transcripts each, Services contains five, and Other contains 31, these counts do not support stable conclusions about industry-specific storytelling. Table 6 cross-tabulates industry group and authority category.

Table 6. Industry × Authority Narrative Styles

Industry	Collective-Assertive	Collective-Soft	Individual-Assertive	Individual-Soft	Total
Tech	1	9	2	7	19
Consumer (included in Other)	—	—	—	—	—
Services	0	1	0	4	5
Finance	0	1	1	0	2
Real Estate	0	0	0	2	2
Other	0	14	2	15	31
Total	1	25	5	28	59

Note. Counts by archived broad industry group and authority category. Consumer cases are included in Other; dashes indicate that they are not separately tabulated. The row and column margins reconcile to $N = 59$. Because the project files do not preserve the detailed-to-broad grouping rule, the table is descriptive only.

4.6 Narrative Length and Authority

Transcript length and authority style were compared by looking at the number of words and the collective authority ratio. This relationship is shown in a scatter plot in Figure 4.

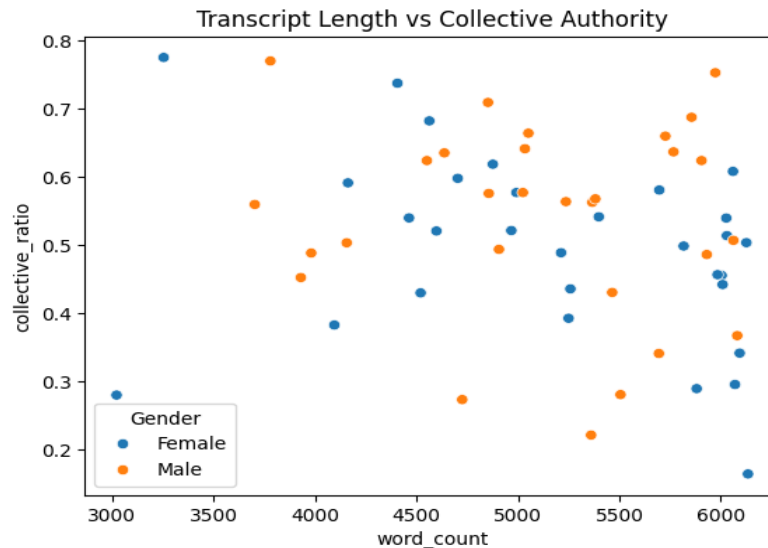


Fig 4. Transcript Length vs. Collective Authority

Note. Scatter plot of each transcript's word count (x-axis) against its collective-authority ratio (y-axis), one point per transcript. The dispersed, trendless cloud indicates that the degree of collective framing is largely independent of how long a speaker talks.

No clear linear pattern can be seen in the points: transcript length is not systematically related to the use of collective-authority language; that is, length of a narrative does not predict the strength of collective leadership framing.

4.7 Semantic Structure of Leadership Stories

To investigate the relationship between leadership-related narratives in terms of semantics, transcript embeddings were created with a Sentence-BERT model and mapped to two dimensions with UMAP. The outcome is a semantic structure of the corpus (Figure 5).

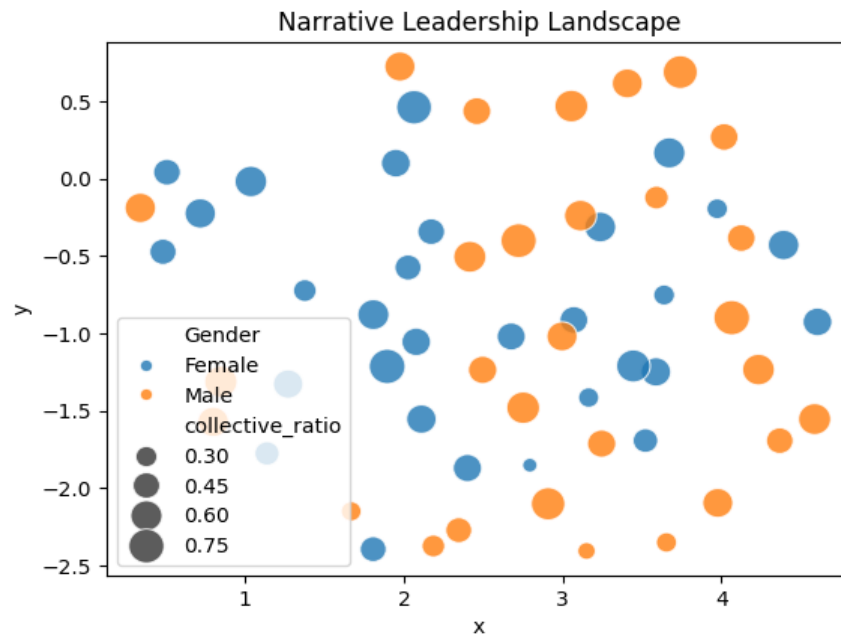


Fig 5. Narrative Leadership Landscape

Note. Two-dimensional UMAP projection of Sentence-BERT transcript embeddings, where local proximity represents similarity in the model's high-dimensional representation. Point color encodes recorded guest gender and point size encodes the collective ratio. The projection is exploratory and does not by itself test whether narrative content is more influential than gender.

Each point represents one transcript; color indicates recorded guest gender and size indicates the collective ratio. Episodes featuring female and male guests appear interspersed in the two-dimensional projection. This visual pattern is consistent with the non-predictive classifiers, but UMAP is a nonlinear exploratory projection and cannot establish that narrative content has a stronger causal influence than guest characteristics.

4.8 Narrative Clusters

K-means partitioned the transcript embeddings into four exploratory clusters. Table 7 reports the mean collective ratio within each partition. Because cluster separation is weak (silhouette = 0.058), these means are descriptive summaries rather than evidence of four distinct narrative types.

Table 7. Mean Collective Authority by Narrative Cluster

Cluster	Mean Collective Ratio
0	0.46
1	0.53
2	0.53
3	0.50

Note. Mean collective ratio for each of four K-means partitions derived from transcript embeddings. The observed range (0.46–0.53) falls within the aggregation envelope reported in Table 9 and is not interpreted as a substantive cluster effect.

The collective-authority language is concentrated in Clusters 1 and 2, indicating more collective framings of narratives, with Cluster 0 having the lowest collective-authority ratio and Cluster 3 being in between. This narrow cluster band is an aggregation artifact, however, and not a substantive concentration, and the cluster differences are not interpreted as findings, as the Robustness Checks below demonstrate.

4.9 Topic Modeling Results

LDA produced five descriptive word-co-occurrence topics. Their prominent terms suggested technology and innovation; coaching and professional development; product

development and design; international and gender-related experience; and a remaining mixed business theme. These labels are interpretive summaries of high-probability words, not themes supplied by a common interview schedule. Mean collective ratios by assigned topic are reported in Table 8.

Table 8 Mean Collective Authority Ratio by Topic

Topic	Mean Collective Ratio
0	0.539
1	0.464
2	0.588
3	0.495
4	0.525

Note. Mean collective ratio for each assigned LDA topic. The observed range (0.46–0.59) falls within the aggregation envelope in Table 9 and is not interpreted as a topic effect.

The collective-authority orientation is highest in Topic 2, suggesting that language is more collaborative and communal, and lowest in Topic 1, suggesting somewhat more individual framing. As with the clusters, the Robustness Checks below show this topic band to be an aggregation artifact.

4.10 Authority × Topic × Cluster Interaction

Figure 6 integrates each transcript's topic assignment, exploratory cluster membership, and authority category in a single visualization. This is a descriptive map rather than a statistical interaction analysis.

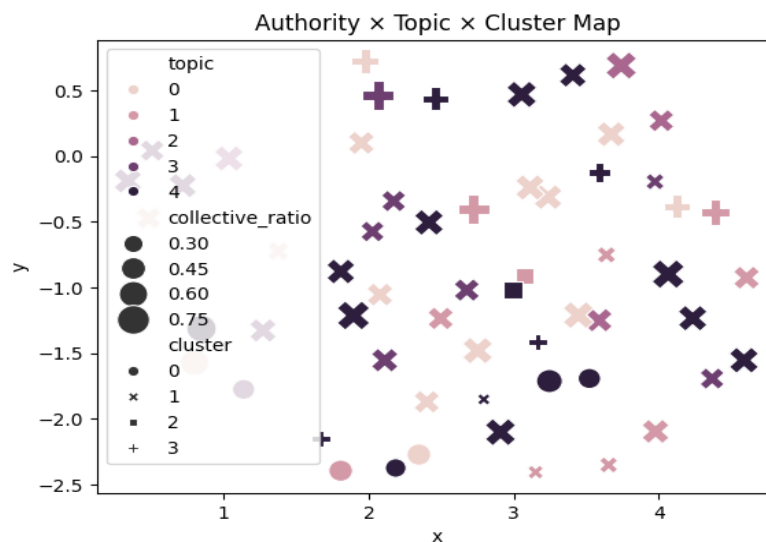


Fig 6. Authority × Topic × Cluster Map

Note. Integrated visualization positioning transcripts in the semantic embedding space while jointly encoding their LDA topic assignment, K-means cluster membership, and authority style. The figure shows that authority styles spread across clusters and topics rather than concentrating in any single thematic region.

Authority categories appear across multiple topics and clusters rather than concentrating in one region of the visualization. Given the weak cluster separation and aggregation results, the map does not support a claim that topic and cluster membership jointly affect authority style; it shows only that the classifications are distributed across the exploratory semantic representation.

4.11 Validation Results

The computational measures were evaluated in several ways. Coverage analysis found an average of 54.9 authority-lexicon matches per transcript. Coverage shows that the dictionary fires throughout the corpus; it does not demonstrate that the matches validly represent authority.

A Spearman correlation between the collective ratio and one semantic-embedding dimension was moderate and positive ($r = 0.34$, $p = .008$). This association shows that the lexical measure covaries with one coordinate of the representation, but an embedding coordinate has no fixed substantive meaning and the result is not treated as construct validation.

Topic-model consistency was evaluated across repeated runs. Similar perplexity values indicate numerical stability under the tested specifications, although stability does not establish that the inferred topics are uniquely correct or theoretically valid.

Cluster validity was assessed with the silhouette coefficient. The score of 0.058 indicates very weak separation, so the four-cluster solution is retained only as an exploratory visualization.

A prior hand-verification exercise reported Cohen's $\kappa = 0.91$. However, the surviving project records do not document the number of sampled units and constructs, an independent second coder, blinding, or an adjudication procedure. The value is therefore reported as preliminary provenance rather than sufficient evidence of reliability. The available completed coding sheet contains 20 sentence-level units and a different claim/granting codebook, so it cannot retrospectively validate the transcript-level four-way authority classifier.

4.12 Robustness Checks: The Difference Between Findings and Artifacts

Several descriptive patterns can be generated mechanically by a computational pipeline, so each was compared with a matched null or robustness expectation. The range of mean collective ratios for the four clusters (0.46–0.53; Table 7) and five topics (0.46–0.59; Table 8) is consistent with aggregation alone. With a transcript-level standard deviation of approximately 0.16, group means for groups of roughly 12–15 transcripts are expected to fall about ± 0.08 to ± 0.09 around the grand mean. All observed cluster and topic means fall inside that envelope (Table 9), and cluster separation is weak (silhouette = 0.058). The narrow bands are therefore not interpreted as substantive findings.

Table 9. Aggregation Check on Cluster and Topic Authority Bands

Grouping	Observed mean range	Aggregation envelope (mean $\pm 2 \cdot SD/\sqrt{n}$)	Silhouette
Four K-means clusters	0.46–0.53	0.42–0.59	0.058
Five LDA topics	0.46–0.59	0.41–0.60	—

Note. For each grouping, the observed range of group mean collective-authority ratios is compared with the range expected from aggregation alone (grand mean $\pm 2 \cdot SD/\sqrt{n}$, with corpus SD = 0.16 and group sizes of 12–15). Both observed ranges fall inside the aggregation envelope, and the clusters' silhouette coefficient indicates little separation.

The dispersion of the collective-authority ratio itself is similar to a token base rate, and not a spread of individual authority styles. The standard deviation of the distribution of transcript ratios is not smaller than the standard deviation of the distribution of the same transcript ratios, redrawn as Bernoulli trials at the collective rate of the corpus; the observed variation is thus what word-frequency variation alone would produce. The ratio of one transcript should therefore be interpreted in the context of a base rate for the corpus, rather than as an individual authority profile.

The lexicon's reliability and construct validity remain limited. The preliminary κ value lacks sufficient design documentation, while the available sentence-level coding exercise uses a different unit and codebook. In addition, the dictionary permits overlapping terms, does not resolve pronoun referents or word senses, and was applied to complete episodes containing host speech. The four-way classification is therefore provisional pending a documented, independently coded validation sample using the same unit, definitions, and decision rules.

Three episode-level patterns remain after these specific checks: 89.8% of transcripts receive a Soft classification under the stated rule; the collective-ratio comparison by recorded guest gender is non-significant ($p = .205$) and the gender labels are visually interspersed in the embedding projection; and no clear relation between transcript length and collective ratio appears in the scatterplot. Their survival does not remove the measurement limitations above. In particular, a non-significant result is evidence of non-detection in this sample, not proof of equality or absence of gendered discourse.

5.0 Discussions

Listen to enough founders and a paradox surfaces: they narrate triumphs in the first person yet cannot finish the story without a team, a network, or a community carrying it. Across the corpus, entrepreneurial leaders claim authority individually but ground it in the collective, soft, relational framing dominates, and individually and collectively voiced authority language routinely co-occur within the same narratives. We use borrowed authority as a provisional interpretive label for legitimacy voiced individually yet drawn from teams, networks, and shared endeavor. Reading the results through this lens unifies them. The dominance of soft-authority classifications and the bounded gender null may be related, but the analyses do not establish a single causal discursive condition. The sections that follow develop this argument and situate it against prior work.

The findings substantiate the view that leadership authority is constructed in discourse rather than conferred by formal position. Soft-authority styles, Individual-Soft and Collective-Soft, dominate the corpus, expressing leadership through supportiveness, relationship, and narration rather than directive language. Collaboration, shared responsibility, and collective action recur across transcripts, indicating that entrepreneurial leaders build their accounts around the language of teams, communities, and networks rather than personal power.

This bias toward soft authority permits comparison with earlier discourse-analytic work. Holmes et al. (2007) found that effective workplace leadership was frequently performed through relational and facilitative language rather than directive assertion. The convergence is notable because the settings differ: Holmes et al. studied talk among colleagues sharing an organizational affiliation, whereas the podcast interview addresses an interviewer and an unseen audience. Relational framing of authority thus appears not merely as in-group workplace politeness but as a general resource speakers draw on when seeking legitimacy through talk. One explanation is that the reflective, conversational interview format encourages collaborative self-explanation, and that assertive self-promotion may carry reputational cost before a large audience. The dominance of soft styles is therefore better read as a discursive requirement of the genre, which privileges legitimacy claims grounded in relationship and shared meaning (Grant et al., 2004).

This aligns with relational leadership theory, which holds that leadership is a social process emerging from interaction rather than the action of an individual (Uhl-Bien, 2006). On this view, leadership authority derives from the relationships that guide, align, and commit actors in an organizational or entrepreneurial setting.

The relational reading is developed, rather than confirmed, by the wider results. That collective framing pervades even individually oriented narratives is compatible with the borrowed-authority interpretation, and it echoes Empson et al.'s (2023) observation that

individual leadership claims and collective leadership stories are presented together rather than as alternatives, and Raelin's (2023) leadership-as-practice account, which treats collaborative activity as the common ground of leadership. The present corpus supplies an episode-level lexical pattern consistent with these claims, but it does not validate them. A structural reading is that the entrepreneurial enterprise is a material joint effort of founders, teams, investors, and customers, so a speaker who emphasizes personal agency must also invoke the collective to make the account credible. This remains a possible structural interpretation rather than a tested conclusion.

The findings also underscore the role of storytelling in constructing entrepreneurial leadership authority. The topic-modeling, semantic-landscape, and clustering results show that leaders organize their experiences into recurring narrative patterns, so that authority is expressed less through claims of intellect or decisiveness than through narratives that contextualize challenges and success. This is consistent with work on narrative leadership emphasizing the role of stories in building leadership identity and legitimacy (Gabriel, 2000; Shamir et al., 2005; Auvinen et al., 2013): entrepreneurial authority in podcast interviews is mediated by narrative, through which leaders claim legitimacy by recounting experience, relationships, and strategic decisions.

The gender analyses should be interpreted as bounded null results. The collective-ratio comparison is non-significant ($p = .205$), the classifiers do not predict recorded guest gender out-of-sample ($AUC \approx 0.59$; Table 5), and gender labels are interspersed in the UMAP projection. These results do not show that women and men communicate identically, nor do they broadly challenge role-congruity theory (Eagly & Karau, 2002). The sample contains only 59 episodes from a single series; the guests are self-selected public participants; the study may lack power to detect modest effects; some theme comparisons were selected after exploration, and the analyzed transcripts contain shared host language. A common interview frame may reduce detectable differences, but this remains an alternative explanation rather than a tested boundary condition.

Methodologically, the study joins computational discourse analysis to leadership communication. While prior analyses of leadership discourse have largely relied on qualitative tools or small-case data, this analysis shows how computational methods can scale to a larger body of naturally occurring communication. The framework combines lexicon-based authority tagging, transformer-based semantic embedding, topic modeling, and clustering to capture linguistic, thematic, and semantic facets of leadership discourse, following the move in computational social science toward large-scale study of cultural and communicative patterns (DiMaggio et al., 2013). Applied at this scale, the approach identifies a pattern that motivates the provisional borrowed-authority interpretation, but the present lexicon and complete-episode data do not validate it as the default grammar of entrepreneurial leadership talk.

6.0 Conclusion and Future Works

This study examined how authority is represented in 59 entrepreneurial podcast transcripts through lexical, semantic, and thematic analyses. Under the stated classifier, 89.8% of transcripts are relational-dominant; the episode-level collective ratio does not differ significantly by recorded guest gender; and guest gender is not predicted out-of-sample. Apparent cluster and topic authority bands do not survive aggregation checks, and weak cluster separation limits claims about distinct narrative types.

The study's conceptual contribution is the provisional idea of borrowed authority: individual leadership legitimacy narrated with linguistic resources drawn from teams, organizations, networks, customers, and communities. The evidence supports this as a testable interpretation rather than a validated construct. Its methodological contribution is equally important: transparent term lists, explicit decision rules, and null checks expose where computational discourse findings are robust and where they depend on measurement and

genre. Stronger tests require guest-only speech, multiple interview settings, contextual coding of claim-granting sequences, and independently documented reliability.

Several limitations qualify the conclusions. First, the corpus comes from one English-language podcast series, limiting generalization across genres, cultures, and languages. Second, the analyzed 59-transcript dataset contains host speech, introductions, and prompts; collective vocabulary and non-prediction of guest gender may partly reflect this shared conversational scaffolding. Third, the lexicon permits overlapping terms and does not resolve context, word sense, or pronoun referents. Its reliability has not been established through a documented independent multi-coder study using the same transcript-level definitions, so the four-way classification is provisional. Fourth, the sample is small for gender inference, statistical power is limited, and the mentorship and barrier comparisons were selected after exploration and do not survive correction across all eight themes. Fifth, topic labels and clusters are exploratory, with weak separation between clusters.

Future research should separate host and guest turns, compare multiple podcast series and non-podcast genres, validate the codebook with independent coders, test borrowed authority through contextual claim-granting sequences, and examine audio and visual cues that convey authority through tone, timing, and gesture.

References

- [1] Auvinen, T., Aaltio, I., & Blomqvist, K. (2013). Constructing leadership by storytelling: The meaning of trust and narratives. *Leadership & Organization Development Journal*, 34(6), 496–514. <https://doi.org/10.1108/LODJ-10-2011-0102>
- [2] Bird, S., Klein, E., & Loper, E. (2009). *Natural language processing with Python*. O'Reilly Media.
- [3] Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003). Latent Dirichlet allocation. *Journal of Machine Learning Research*, 3, 993–1022.
- [4] Chen, Y., Peng, Z., Kim, S. H., & Choi, C. W. (2023). What we can do and cannot do with topic modeling: A systematic review. *Communication Methods and Measures*, 17(2), 111–130. <https://doi.org/10.1080/19312458.2023.2167965>
- [5] Clifton, J. (2012). A discursive approach to leadership: Doing assessments and managing organizational meanings. *Journal of Business Communication*, 49(2), 148–168. <https://doi.org/10.1177/0021943612437762>
- [6] Cunliffe, A. L., & Eriksen, M. (2011). Relational leadership. *Human Relations*, 64(11), 1425–1449. <https://doi.org/10.1177/0018726711418388>
- [7] DeRue, D. S., & Ashford, S. J. (2010). Who will lead and who will follow? A social process of leadership identity construction in organizations. *Academy of Management Review*, 35(4), 627–647.
- [8] Devlin, J., Chang, M. W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. *arXiv*. <https://doi.org/10.48550/arXiv.1810.04805>
- [9] DiMaggio, P., Nag, M., & Blei, D. M. (2013). Exploiting affinities between topic modeling and the sociological perspective on culture: Application to newspaper coverage of U.S. government arts funding. *Poetics*, 41(6), 570–606. <https://doi.org/10.1016/j.poetic.2013.08.004>
- [10] Drath, W. H., McCauley, C. D., Palus, C. J., Van Velsor, E., O'Connor, P. M. G., & McGuire, J. B. (2008). Direction, alignment, commitment: Toward a more integrative ontology of leadership. *The Leadership Quarterly*, 19(6), 635–653. <https://doi.org/10.1016/j.leaqua.2008.09.003>
- [11] Eagly, A. H., & Karau, S. J. (2002). Role congruity theory of prejudice toward female leaders. *Psychological Review*, 109(3), 573–598. <https://doi.org/10.1037/0033-295X.109.3.573>
- [12] Empson, L., Langley, A., & Sergi, V. (2023). When everyone and no one is a leader: Constructing individual leadership identities while sustaining an organizational narrative of collective leadership. *Organization Studies*, 44(2), 201–227. <https://doi.org/10.1177/01708406221135225>

- [13] Fairhurst, G. T. (2007). Discursive leadership: In conversation with leadership psychology. Sage. <https://doi.org/10.4135/9781452231051>
- [14] Fairhurst, G. T., & Connaughton, S. L. (2014). Leadership: A communicative perspective. *Leadership*, 10(1), 7–35. <https://doi.org/10.1177/1742715013509396>
- [15] Fairhurst, G. T., & Grant, D. (2010). The social construction of leadership: A sailing guide. *Management Communication Quarterly*, 24(2), 171–210. <https://doi.org/10.1177/0893318909359697>
- [16] Felix, B., dos Santos, R., & Teixeira, A. (2023). Tales of me: Storytelling identity work, authenticity, and impression management during new CEOs' work role transitions. *Frontiers in Psychology*, 14, Article 1246887. <https://doi.org/10.3389/fpsyg.2023.1246887>
- [17] Gabriel, Y. (2000). *Storytelling in organizations: Facts, fictions, and fantasies*. Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780198290957.001.0001>
- [18] Grant, D., Hardy, C., Oswick, C., & Putnam, L. L. (2004). Introduction: Organizational discourse: Exploring the field. In D. Grant, C. Hardy, C. Oswick, & L. L. Putnam (Eds.), *The Sage handbook of organizational discourse* (pp. 1–36). Sage. <https://doi.org/10.4135/9781848608122.n1>
- [19] Holmes, J., Schnurr, S., & Marra, M. (2007). Leadership and communication: Discursive evidence of a workplace culture change. *Discourse & Communication*, 1(4), 433–451. <https://doi.org/10.1177/1750481307082207>
- [20] Jacobs, G. (2024). Entrepreneurial podcasting: A linguistic ethnographic perspective. *International Journal of Business Communication*. Advance online publication. <https://doi.org/10.1177/23294884241255900>
- [21] Jian, G., Schmisser, A. M., & Fairhurst, G. T. (2008). Organizational discourse and communication: The progeny of Proteus. *Discourse & Communication*, 2(3), 299–320. <https://doi.org/10.1177/1750481308091912>
- [22] Jurafsky, D., & Martin, J. H. (2023). *Speech and language processing* (3rd ed. draft). <https://web.stanford.edu/~jurafsky/slp3/>
- [23] Kwon, W., Clarke, I., & Wodak, R. (2009). Organizational decision-making, discourse, and power: Integrating across contexts and scales. *Discourse & Communication*, 3(3), 273–302. <https://doi.org/10.1177/1750481309337208>
- [24] Liu, E. H., Chambers, C. R., & Moore, C. (2023). Fifty years of research on leader communication: What we know and where we are going. *The Leadership Quarterly*, 34(6), Article 101734. <https://doi.org/10.1016/j.leaqua.2023.101734>
- [25] Manning, C. D., Raghavan, P., & Schütze, H. (2008). *Introduction to information retrieval*. Cambridge University Press.
- [26] Matheson, D. (2023). Discourse analysis after the computational turn: A mixed bag. *Communication Research and Practice*, 9(1), 3–15. <https://doi.org/10.1080/22041451.2023.2190531>
- [27] Piper, A. (2023). Computational narrative understanding: A big picture analysis. In *Proceedings of the Big Picture Workshop* (pp. 28–39). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2023.bigpicture-1.3>
- [28] Raelin, J. A. (2023). Leadership-as-practice: Its past history, present emergence, and future potential. *Academy of Management Collections*, 2(2), 19–30. <https://doi.org/10.5465/amc.2021.0005>
- [29] Shamir, B., Dayan-Horesh, H., & Adler, D. (2005). Leading by biography: Towards a life-story approach to the study of leadership. *Leadership*, 1(1), 13–29. <https://doi.org/10.1177/1742715005049348>
- [30] Srivastava, S., Oberoi, S., & Gupta, V. K. (2023). The story and the storyteller: Strategic storytelling that gets human attention for entrepreneurs. *Business Horizons*, 66(3), 347–358. <https://doi.org/10.1016/j.bushor.2023.02.003>
- [31] Uhl-Bien, M. (2006). Relational leadership theory: Exploring the social processes of leadership and organizing. *The Leadership Quarterly*, 17(6), 654–676. <https://doi.org/10.1016/j.leaqua.2006.10.007>
- [32] van Dijk, T. A. (2007). *Discourse & Communication: A new journal to bridge two fields*. *Discourse & Communication*, 1(1), 5–7. <https://doi.org/10.1177/1750481306072182>