

Extracting Narrative Structure from Online Labor Discourse: A Prompt-Based Operationalization of the Actantial Model

Yuntong Ji^{1*,†}

¹Northeastern University, Department of Sociology and Anthropology, Boston, MA, 02115, USA

Abstract

This paper introduces a prompt-based approach to extracting narrative structure from online labor discourse using large language models. Building on Greimas's Actantial Model, it develops a structured prompt that constrains LLM outputs to a closed set of narrative roles (subject, opponent, helper, and object). Applying this method to over 26,000 posts and comments from Reddit labor communities (r/antiwork and r/WorkReform) surrounding a major reputational crisis, the analysis identifies tendencies of narrative transformation. The paper illustrates the potential of structured prompting for scalable narrative analysis in computational social science, while highlighting the importance of prompt-model co-optimization.

Keywords

Narrative analysis, Large language models, Prompt, Online labor discourse, Actantial model, Social movements

1. Introduction

Recent advances in large language models (LLMs) provide new opportunities for social scientists to perform rich text interpretation beyond topics, sentiment, or lexical change. Building on the work of Elfes (2024), this paper proposes a prompt-based approach to analyzing online labor discourse. I introduce a structured LLM prompt that operationalizes Greimas's Actantial Model to extract narrative roles from text in a consistent and machine-readable format. Rather than relying on open-ended generation, the prompt constrains the model to a closed ontology of narrative roles (e.g., subject, opponent, helper, and object) and produces standardized JSON outputs. I apply this prompt to discourse from online labor communities on Reddit, focusing on the r/antiwork subreddit and its splinter community, r/WorkReform. Online labor communities provide a particularly revealing empirical setting: they are decentralized, lack formal leadership, and depend almost entirely on ongoing discourse to sustain collective identity and interpret workplace dissatisfaction. As a result, shifts in narrative structure are especially consequential for their cohesion and capacity for collective action.

In January 2022, r/antiwork experienced a major reputational crisis following a widely criticized Fox News interview given by one of its moderators. The interview triggered intense internal conflict, the temporary privatization of the subreddit, and the exit of thousands of users who went on to form r/WorkReform. Despite widespread expectations of collapse, r/antiwork continued to operate in the months and years following the incident. This event offers a valuable opportunity to examine how narrative structure changes under conditions of public legitimacy loss. Applying the actantial prompt uniformly to over 26,000 posts and comments across multiple crisis periods, this paper try to show that narrative extraction reveals coordinated patterns of narrative transformation. Specifically, the results indicate shifts from individualized antagonism toward structural attribution, from individual to collective protagonists, and from exit-oriented goals toward collective action and systemic change. These patterns suggest that reputational crises have the possibility to trigger systematic narrative repair in decentralized online movements.

[PromptEng] Third International Workshop on Prompt Engineering for Pre-Trained Language Models co-located with the ACM WebConf, April.13, 2026, Dubai, United Arab Emirates

✉ ji.yunt@northeastern.edu (Y. Ji)

🌐 <https://cssh.northeastern.edu/student/yuntong-ji/> (Y. Ji)

🆔 0009-0000-7424-7629 (Y. Ji)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

Methodologically, this study demonstrates how narrative theory can be embedded in a constrained prompt to enable scalable empirical analysis of meaning-making processes. It illustrates how online labor communities reorganize narratives of responsibility, identity, and action following crisis. Together, these contributions highlight the potential of structured prompting as a generalizable method for studying collective narratives in digital environments.

2. Related Work

In social movement research, the “framing perspective” developed by Benford and Snow [1] has become one of the central analytical paradigms. It argues that social problems do not automatically generate collective action; rather, mobilization becomes possible when actors actively interpret and transform personal experiences and diffuse discontent into coherent and actionable meaning structures. This process is inherently interpretive. Within this perspective, collective action frames are defined as “action-oriented sets of beliefs and meanings” that selectively punctuate, simplify, and rearticulate complex realities in ways that identify problems, attribute responsibility, and provide impetus for mobilization.

Framing perspective typically involves three core framing tasks: diagnostic framing, which identifies problems and attributes blame by asking who or what is responsible; prognostic framing, which articulates proposed solutions, strategies, and lines of action; and motivational framing, which offers “vocabularies of motive” that encourage participation. Building on this, Klandermans’s distinction between consensus mobilization and action mobilization provides a continuous pathway linking meaning construction to behavioral engagement [2]. Simply put, consensus mobilization relies on the construction and diffusion of agreement, whereas action mobilization transforms shared meanings into action. The framing perspective becomes particularly important in the digital era, as online discussions constitute a decentralized and ongoing process of meaning production. Through narrating personal experiences, participants collectively reframe individual stories into publicly recognizable interpretations. In this sense, framing theory offers a suitable foundation for analyzing online labor discourse, revealing how emotions, narratives, and interpretive practices jointly constitute mechanisms of mobilization and reshape the potential for collective action.

Hirschman’s [3] theory of exit, voice, and loyalty (EVL) provides another conceptual lens for understanding narrative transformation. Hirschman argues that when facing organizational dissatisfaction or perceived injustice, individuals can adopt three responses: exit (leaving the situation), voice (protesting to change it), or loyalty (remaining despite dissatisfaction). Although originally developed for consumer and organizational contexts, subsequent scholars have extended EVL to social movements and collective action [4, 5, 6]. Exit and voice represent two distinct strategies with different implications for collective mobilization: exit is an individual strategy requiring minimal coordination, while voice is a collective strategy that requires organization, solidarity, and sustained engagement. Movements emphasizing exit narratives risk demobilization because individual departure prevents structural transformation, while movements emphasizing voice facilitate mobilization by framing problems as requiring coordinated action.

To operationalize narrative transformation more systematically, I further employ Greimas’s actantial model, which identifies narrative roles across diverse texts. Greimas argues that narratives can be analyzed through six functional positions: Subject, Object, Sender, Receiver, Helper, and Opponent [7]. This framework offers analytical advantages for studying social movement narratives because it moves beyond content to examine relational structure: who acts, against whom, with whose support, and toward what ends. It also allows consistent comparison across time and communities.

Methodologically, this study builds on two methods: semantic analysis using word embeddings and narrative role extraction using large language models. Word embeddings such as Word2Vec [8] have become foundational tools for detecting semantic shifts over time. Hamilton et al. [9] showed that local neighborhood measures and global distance measures capture different dimensions of semantic change, highlighting the need for researchers to select metrics appropriate to their theoretical aims.

Embedding-based approaches have already been applied to studies of gender stereotypes [10] and political discourse [11]. Extending this line of work, my study applies semantic drift analysis to trace how labor movement vocabulary evolves following a reputational crisis. Results from the semantic drift analysis can also provide a basis for selecting example concepts used in the Actantial Model.

Finally, recent work shows that instruction-tuned LLMs can perform zero-shot extraction of narrative elements well without labeled data. Elfes [12], for instance, developed an LLM-based framework for recovering Greimas’s Actantial Model from news articles. Applying Llama-3-8B to more than 5,000 articles on the Israel–Palestine conflict, Elfes demonstrated that narrative-structured embeddings could successfully differentiate texts with similar topics but distinct narrative framings. My study extends this innovation to labor movement discourse by using LLMs to extract actantial roles across thousands of Reddit posts.

3. Data and Methodology

3.1. Data and Pipeline

I collected posts and comments from r/antiwork and r/WorkReform using the Python Reddit API across August 26, 2021–November 1, 2025. All Posts in the two online community were extracted, and top-level comments were also included to capture community dialogue; comments outnumbered posts by roughly 29:1. The final dataset contains 26,104 texts (866 posts and 25,238 comments). To analyze narrative changes around the Fox News crisis, I divided the corpus into four periods: (1) a pre-crisis baseline (P1: Aug 26, 2021–Jan 25, 2022); (2) post-crisis discourse within r/antiwork (P2A: Jan 26–Jun 30, 2022); (3) parallel post-crisis discourse in the newly formed r/WorkReform (P2B: Jan 26–Jun 30, 2022); and (4) a long-term stabilization period nearly three years later (P3: Nov 1, 2024–Nov 1, 2025). Table 1 reports the distribution of texts across periods; despite variation in period length, P1 and P3 include similar volumes (7,262 and 7,214 texts respectively), while P2A and P2B capture the intensive crisis response phase. The temporal gap between P2 and P3 allows narrative changes to be distinguished from short-term reactions.

Table 1
Data Collection Overview

Period	Label	Timeframe	Subreddit	Posts	Comments	Total
P1	Pre-Crisis Baseline	Aug 2021 – Jan 2022	r/antiwork	242	7,020	7,262
P2A	Post-Crisis (Original)	Jan – Jun 2022	r/antiwork	194	6,006	6,200
P2B	Post-Crisis (Splinter)	Jan – Jun 2022	r/WorkReform	181	5,247	5,428
P3	Long-term	Nov 2024 – Nov 2025	r/antiwork	249	6,965	7,214
Total				866	25,238	26,104

I applied a preprocessing pipeline to ensure text quality. First, I filtered out posts under 50 characters, non-English texts, and exact duplicates. I then normalized the text by stripping Reddit-specific markup, converting to lowercase for embedding and emotion models, standardizing punctuation, and replacing emojis or special characters with simple descriptive tokens. For Word2Vec, I tokenized using spaCy, retained stopwords to preserve contextual structure, and applied lemmatization to reduce morphological noise. In contrast, original text was kept intact for LLM-based actant extraction.

3.2. Semantic Analysis

For the semantic analysis, I employ Word2Vec embeddings to capture distribution semantics. I trained period-specific Word2Vec models using the Skip-gram. Skip-gram is preferable for this application, as it better captures meanings of less frequent terms. Table 2 summarizes the model parameters.

Rather than analyzing the entire vocabulary, I focused on 34 theoretically motivated terms grouped into five conceptual categories (Table 3). The selection followed several criteria. First, I applied frequency-

Table 2
Word2Vec Model Parameters

Parameter	Value
Architecture	Skip-gram
Vector dimensionality	100
Context window	5
Minimum word frequency	3 occurrences
Training epochs	10

based filtering to ensure that chosen terms appeared often enough across periods to generate stable embeddings, while excluding rare or overly specialized vocabulary. Second, I selected terms on the basis of theoretical relevance. This included core labor concepts as well as terms tied to the frameworks guiding this study; for example, *quit* and *union* reflect Hirschman’s [3] exit–voice distinction, while authority and system terms correspond to expected framing shifts. Finally, I ensured that the selected terms captured the dimensions anticipated by the actantial model, including protagonist identity, antagonist identity, and strategic methods.

Table 3
Term Selection and Categorization

Category	Terms	Theoretical Rationale
Core Concepts (7)	work, job, labor, employee, people, worker, live	Fundamental labor vocabulary
Authority (5)	boss, manager, employer, company, business	Antagonist framing; tests depersonalization
Economic Terms (8)	money, pay, wage, salary, raise, hour, week, office	Material conditions; individual vs. structural framing
Action Terms (6)	quit, leave, fire, refuse, organize, strike	Strategic terms; tests exit-voice
System Terms (8)	capitalism, exploit, system, law, government, union	Structural critique; system-level discourse

3.3. Narrative Role Extraction

For the structural role analysis, I apply Greimas’s actantial model [7] to extract key narrative roles. The analysis focuses on four actants most relevant to labor movement narratives: the **Subject** (the protagonist or actor), the **Opponent** (the entity responsible for or causing the problem), the **Helper** (supporting forces or resources), and the **Object** (the desired outcome or goal).

Model Configuration. I used a large language model (LLM) in a zero-shot setting to extract actantial structure from each text. Specifically, I employed **Qwen2.5-14B-Instruct** [13], an open-source instruction-tuned LLM developed by Alibaba (Table 4 summarizes the model configuration). Each post was processed using a standardized, structured prompt designed to identify narrative roles and framing components. The low temperature setting (0.1) was chosen to maximize output consistency and reproducibility.

Label Selection. The selection of label categories for each actant field was guided by theoretical knowledge of the literature on social movements. Simply put, for the **subject** role, the categories `individual_worker` and `collective_workers` capture a core distinction between personal and

Table 4
LLM Configuration for Narrative Extraction

Parameter	Value
Model	Qwen2.5-14B-Instruct
Parameters	14 billion
Setting	Zero-shot
Temperature	0.1
Max tokens	300
Output format	Structured JSON

collective identity. This distinction operationalizes the shift from “I” to “we” narratives, which Benford and Snow [1] identify as a central mechanism in mobilization processes and the formation of collective action capacity. For the **opponent** role, the categories `individual_boss`, `company`, and `capitalism_system` represent increasing levels of attribution in diagnostic framing. These distinctions move from interpersonal conflict to organizational responsibility and, finally, to systemic critique. For the **object** role, the categories `escape_quit`, `fair_treatment`, and `collective_action` are grounded in Hirschman’s [3] exit–voice framework. Exit-oriented goals emphasize individual withdrawal from adverse conditions, whereas voice-oriented goals frame dissatisfaction as requiring collective engagement and coordinated action. In addition to theoretical grounding, I employed semantic drift analysis to provide preliminary empirical support for label selection.

Prompt Iteration. The prompt was iteratively refined to improve structured output compliance. Table 5 summarizes the progression in three versions.

Table 5
Prompt Iteration and Validity Rates (Qwen2.5-14B)

Version	Key Modification	Valid Rate
V1	Comma-separated options in JSON example	3.3%
V2	Pipe-separated options in JSON example	46.5%
V3	Explicit “choose ONE” instruction; Comma-separated option	90.4%

V1: Initial Attempt: The first prompt version presented label options as comma-separated values within the JSON structure (e.g., `"subject": "individual_worker", "collective_workers", "specific_group"`). This format proved highly ambiguous. The model frequently included multiple labels per field, resulting in only 3.3% valid outputs.

V2: Second Attempt: The second version reformatted options using pipe separators within quoted strings (e.g., `"subject": "individual_worker|collective_workers|specific_group"`) and provided a complete JSON schema including both actant and frame fields. This clarification improved validity to 46.5%.

V3: Explicit Constraint: The final version introduced a key modification by explicitly instructing the model to choose exactly one label from the allowed options. This change increased the valid output rate to 90.4%. For the final analysis, results from Versions 2 and 3 were combined. Version 2 was first applied due to its clearer instructions, and Version 3 was subsequently used to process cases where Version 2 failed but the outputs were still informative (e.g., sufficiently length).

This iterative process demonstrates that prompt engineering for structured extraction requires careful attention to both format specification and explicit constraints. The progression from 46.5% to 90.4% validity underscores the sensitivity of LLM outputs to prompt design choices (All prompt versions are included in the Appendix).

Model Comparison. To assess whether performance can generalize between models, I replicated the analysis using DeepSeek-V2.5, an alternative open-source LLM. Table 6 compares the two models using identical prompts.

Table 6
Cross-Model Comparison: Validity Rates by Prompt Version

Prompt Version	Qwen2.5-14B	DeepSeek-V2.5
V2	46.5%	14.5%
V3	90.4%	41.4%

Model comparisons revealed differences in structured output performance. With the V2 prompt, Qwen2.5-14B achieved a performance rate of 46.5%, compared to only 14.5% for DeepSeek-V2.5. Even with the improved V3 prompt, DeepSeek-V2.5 reached 41.4% successful rate, which remains less than half of Qwen2.5-14B’s 90.4%. This persistent performance gap suggests that instruction-tuned models differ in their ability to follow the same format prompt.

Importantly, the V3 prompt improved structured output compliance for both models, indicating that explicit constraints benefit structured extraction regardless of model choice. At the same time, these results show that prompts achieving high performance on one model cannot be assumed to transfer directly to another. Based on its higher performance with structured extraction, Qwen2.5-14B with the V3 prompt was selected for the main analysis.

Validation. To assess the accuracy of the LLM-based extraction, I randomly selected sample of 180 texts with 60 per period from P1, P2A, and P2B and manually recode each text. I independently coded all four actant roles for each selected text and use the final code to compare with the LLM labels. Table 7 reports the agreement rates.

Table 7
Human Validation: LLM-Human Agreement Rates

Actant Field	Matches	Total	Accuracy
Subject	122	180	67.8%
Opponent	104	180	57.8%
Helper	165	180	91.7%
Object	95	180	52.8%
Overall	486	720	67.5%

The helper role achieved the highest agreement rate (91.7%), reflecting its relatively low level of ambiguity. Texts either explicitly mention support structures or they do not, making this role easier to identify consistently. The subject field showed moderate agreement (67.8%), with the most common error occurring when the LLM classified collective narratives as individual worker accounts.

In contrast, the opponent and object fields exhibited lower agreement rates (57.8% and 52.8%, respectively), largely due to greater interpretive ambiguity. For example, distinguishing between “fair_treatment” and “work_life_balance” often requires subjective judgment regarding the narrator’s primary concern, which motivated subsequent prompt refinement. In the opponent field, the LLM also tended to assign the label “none”, but I was more likely to identify implicit systemic critiques (e.g., references to “capitalism” or “society”). These patterns suggest that the LLM adopts a more conservative interpretive strategy, while human coders are more inclined to infer latent meanings from contextual cues.¹

¹Accuracy rates remained consistent across periods.

4. Results

4.1. Semantic Transformations

Figure 1 and Figure 2 report average semantic drift scores across the five conceptual categories as well as 34 target terms, comparing Period 1 (pre-crisis baseline) with Period 2 (post-crisis). Terms associated with authority exhibit the highest average drift (0.692), exceeding all other categories. System-level terms show the second-highest drift (0.622). Action (0.585) and economic (0.570) vocabularies also demonstrate notable shifts, though to a lesser extent. In contrast, core labor concepts display the lowest drift (0.549), indicating that foundational descriptors of work maintained relatively stable meanings.

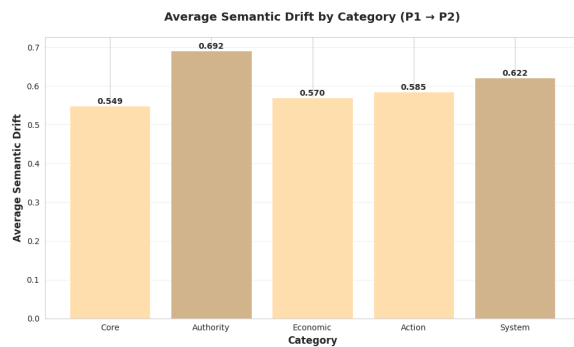


Figure 1: Average Semantic Drift by Category

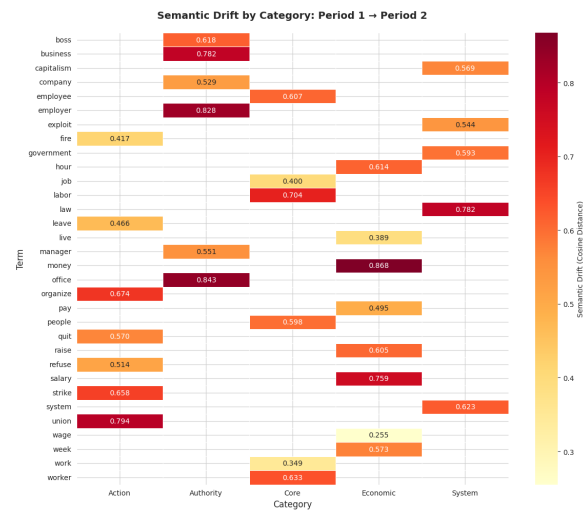


Figure 2: Semantic Drift by Category: P1 to P2

I further analyzed semantic neighbors (the five most similar words in each period) for high-drift terms to examine the qualitative nature of the shifts. This analysis reveals some possible transformation patterns. Table 8 highlights the ten terms with the greatest drift scores along with their primary shift patterns.

First, economic terms shifted from individual aspiration to redistribution discourse. Terms such as *money* (drift = 0.868) and *salary* (0.759) underwent some of substantial shifts. In Period 1, *money* is embedded in narratives of individual advancement (“earn,” “save,” “million”). By Period 3, the term is surrounded by discourses of structural redistribution (“taxpayer,” “hoard,” “pocket”). *Salary* similarly evolves from competitive wage descriptors (“competitive,” “significantly”) toward negotiation terminology (“range,” “adjustment,” “workload”).

Second, authority-related terms shifted away from interpersonal discourse. *Boss* (0.618) and *employer* (0.828) exhibited notable transformations. In P1, *boss* neighbors arouse interpersonal harm (“slap,” “berate,” “yell”). By P3, the term is embedded in organizational contexts (“colleague,” “review,” “meeting”). *Employer* transitions from hierarchical authority (“obligation,” “terminate”) to legalistic resistance (“prohibit,” “unlawful”) and ultimately to negotiated relations (“fair,” “negotiate”).

Last, collective action terms shifted from generalized sentiment to concrete actions. *Union* (0.794) moves from vague vocabulary (“pro,” “join”) to operational vocabulary (“strike,” “recognize,” “bust,” “unionize”). *Strike* (0.658) transitions to campaign-specific mobilization (“costco”).

These preliminary qualitative patterns suggest that post-crisis semantic change occurs. The semantic analysis therefore provides complementary support for the results of the structural role analysis presented in the following section.

Table 8
The Top 10 Terms and Primary Shift Pattern

Rank	Term	Category	Drift Score	Primary Shift Pattern
1	money	Economic	0.868	Saving/earning → Spending/re-distribution
2	office	Authority	0.843	Physical space → Remote work policies
3	employer	Authority	0.828	Obligations → Fairness/negotiation
4	union	Action	0.794	General support → Active organizing
5	law	Authority	0.782	Federal protection → Enforcement/illegal
6	business	Authority	0.782	Success model → Ownership/-operations
7	salary	Economic	0.759	Competitive scales → Range transparency
8	labor	Core	0.704	Relations board → Complaints/practices
9	organize	Action	0.674	Community solidarity → Collective bargaining
10	strike	Action	0.658	General strikes → Union recognition

4.2. Actantial Role Analysis

Actantial role distributions show evidence of narrative restructuring across the four periods (Figures 3–6).

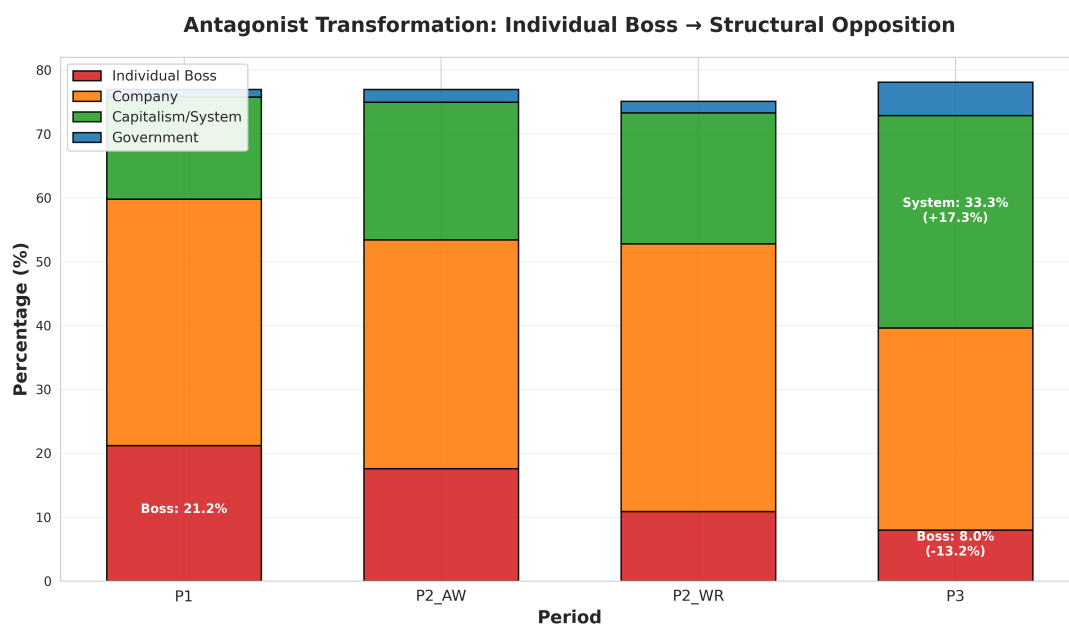


Figure 3: Antagonist Transformation

Antagonist role has the strongest transformation. In the pre-crisis baseline, individual bosses represented 21.2% of all antagonists; this declined to 17.8% in P2 in original community, 10.9% in splinter community, and 8.0% in P3 with a total reduction of 13.2 percentage points. In contrast, system-level antagonists increased from 16.0% in P1 to 22.5% (P2 in original community), 21.2% (P2 in splinter

community), and ultimately 33.3% in P3, marking a 17.3-point increase. This is the largest shift among all antagonist categories. Company-level blame decreased moderately from 38.8% (P1) to 31.8% (P3) (7.0 points). These patterns show a statistically meaningful shift from individual to structural opposition.

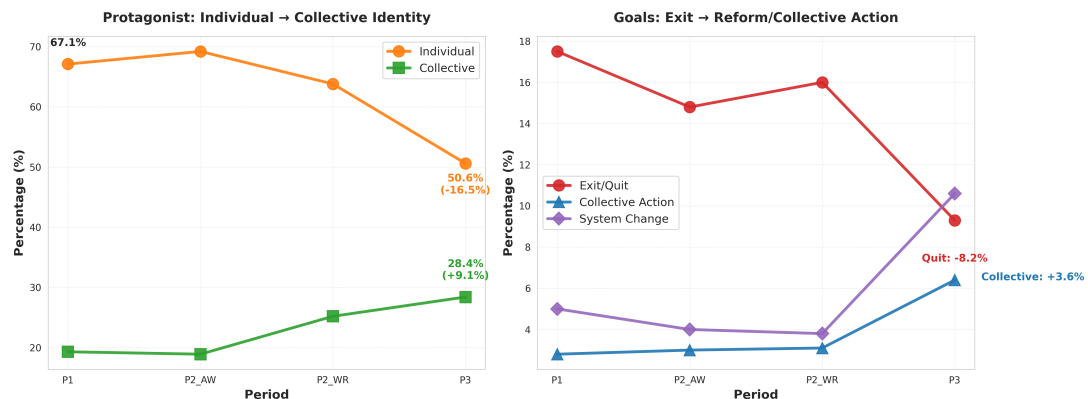


Figure 4: Protagonist and Goals Transformation

Protagonist roles show similar trend. Individual workers made up 67.1% of subjects in P1 and increased slightly to 68.5% in P2 within the original community. However, this share dropped to 64.0% in the splinter community’s P2 and declined further to 50.6% in P3, marking an overall decrease of 16.5 percentage points. Collective worker identities, meanwhile, increased from 19.3% in P1 to 18.6% in P2 within the original community, 25.3% in the splinter community’s P2, and 28.4% in P3—an overall rise of 9.1 percentage points. This pattern indicates a gradual shift from individualized narratives toward more collectivized identities over time.

Goal/object roles also show a shift away from exit and toward collective strategies. Quit/exit goals declined from 17.4% in P1 to 14.8% in P2 within the original community, 15.9% in the splinter community, and ultimately 9.2% in P3; an overall decrease of 8.2 points. In contrast, collective action goals rose from 2.9% in P1 to 3.2% in P2 in the original community, 3.3% in the splinter community, and 6.5% in P3, marking a 3.6-point increase. System-change goals followed a similar trajectory, increasing from 5.0% in P1 to 4.0% in P2 (original community), 3.8% in P2 (splinter community), and reaching 10.6% in P3; a total rise of 5.6 points. These patterns reflect a gradual pattern from individual exit pathways toward collective or systemic aims.

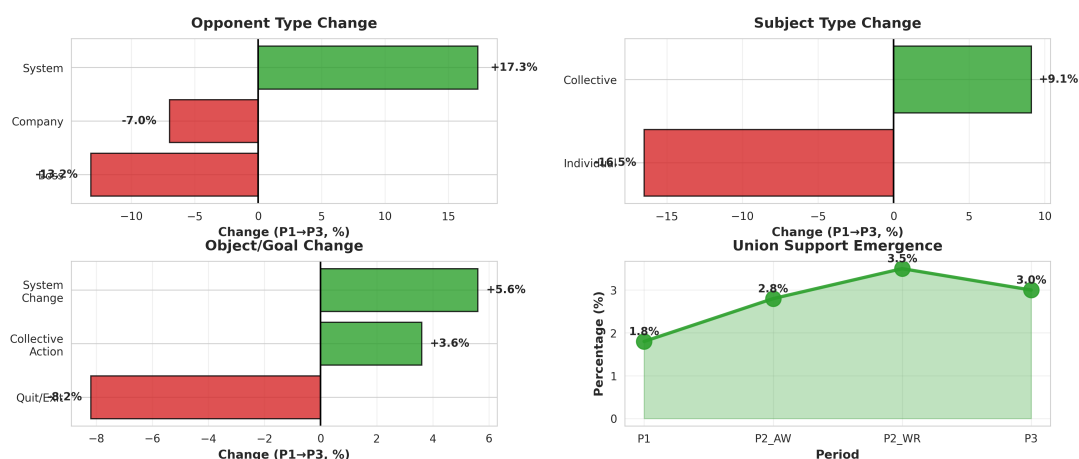


Figure 5: Narrative Role Shifts from Pre-Crisis to Long-Term Period (P1→P3)

Additionally, explicit union support grew from 1.8% in P1 to 2.8% in P2 in the original community, peaked at 3.5% in the splinter community, and settled at 3.0% in P3, with most of the increase driven by users who migrated to r/WorkReform.

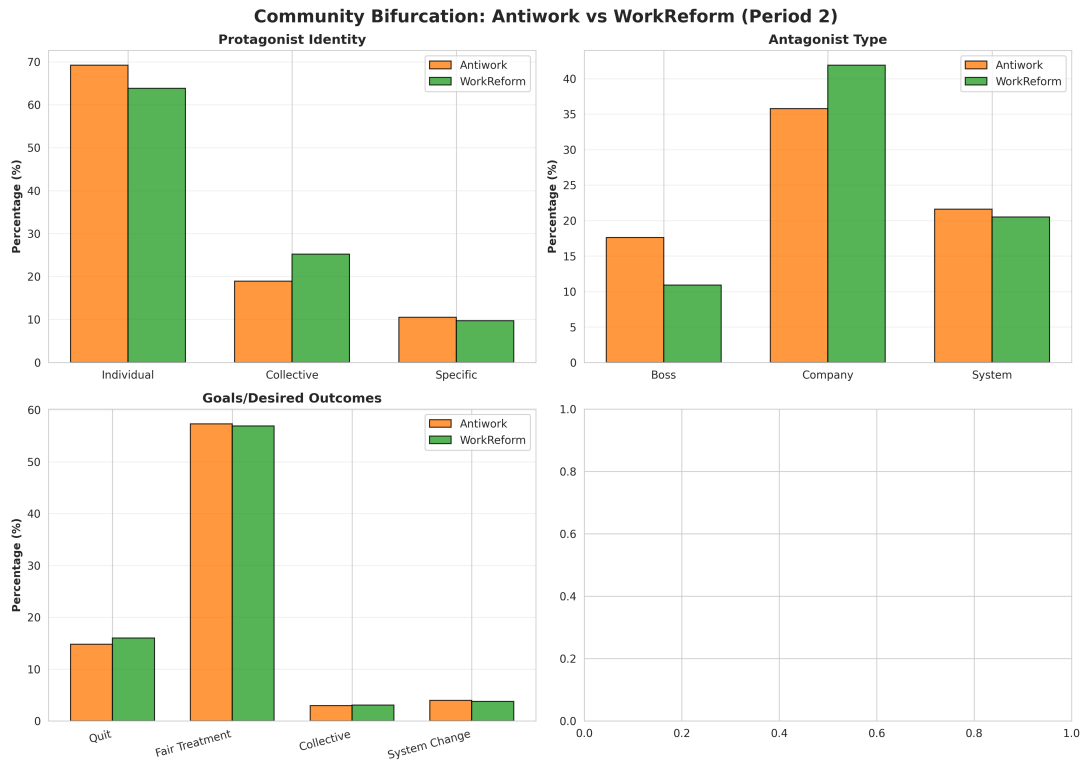


Figure 6: Antiwork vs WorkReform (P2)

Period 2 comparisons highlight bifurcation between communities. In P2 in splinter community, boss-blame (10.9%) was substantially lower than in P2 in original community (17.8%), and collective identity (25.3% vs. 18.6%) and company/system blame were higher. This indicates that users who disagreed with the original community reorganized their narratives more quickly along collective and structural lines.

The total shift shows that narrative after the crisis involved coordinated transformations across actant roles: (1) declining individualized antagonism and rising structural attribution, (2) shifts from individual to collective protagonist identities, and (3) movement from exit-oriented goals toward collective action and system change. The direction and degree of these changes provide empirical support for the hypotheses.

Statistical Significance. To assess whether the observed distributional shifts are statistically meaningful, I conducted chi-square tests of independence for each actant field across the four periods. All four actant fields showed statistically significant differences in label distributions across periods ($p < 0.05$). The statistical tests confirm that the narrative shifts from individual to structural attribution are unlikely to have arisen by chance. Table 9 reports the results.

Table 9
Chi-Square Tests for Actant Distribution Differences Across Periods

Actant Field	χ^2	p -value	Cramér's
Subject	376.59	<0.001***	0.108
Opponent	861.66	<0.001***	0.164
Helper	184.22	0.030*	0.076
Object	584.34	<0.001***	0.135

* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

5. Conclusion and Limitation

This study demonstrates how structured prompting enables the large-scale analysis of narrative structure in online labor discourse and provides empirical evidence that the r/antiwork crisis triggered shifts in narrative structure. These transformations suggest that the crisis served as a critical point that redirected how the community understood workplace injustice, how users positioned themselves as political actors, and how they imagined possible responses.

Actantial structures reveal the clearest evidence of narrative change. After the crisis, antagonists were described less as individual bosses and more as structural or institutional forces. Protagonist roles moved in the same direction: narratives centered on isolated workers gradually gave way to more collective identities. Likewise, the goals expressed in the narratives shifted away from individual exit strategies toward collective action and broader demands for institutional change. These patterns suggest that narrative change was not merely rhetorical but involved rethinking the responsibility and the possibilities for action. More importantly, the splinter community, r/WorkReform, adopted new framings more quickly. This finding highlights that crisis fragmentation may facilitate narrative innovation. Users who left the original community presented faster shifts toward collective and constructive narratives, while those who remained maintained more continuity. This dynamic suggests that fragmentation, often viewed as a sign of movement weakening, may create spaces to shape new meaning-making strategies.

Several limitations should be acknowledged. First, the actantial labels are generated through zero-shot prompting and are therefore subject to model error. Sarcasm, humor, and multi-layered narratives remain challenging for automated extraction, particularly in informal online discourse. Second, full ground-truth verification remains limited in scale. The size of the manually labeled sample included in this study is not yet sufficient to support comprehensive accuracy evaluation, which constrains the assessment of label correctness at the individual-text level. Third, differences across models further limit generalizability. As demonstrated in the model comparison, prompts that achieve high structured-output performance for one model do not necessarily transfer effectively to another. As a result, this study should be understood as exploratory in nature. Future work could extend this approach by increasing the scale of manual validation, incorporating multi-label outputs, and conducting few-shot calibration across multiple models.

Despite these limitations, the study demonstrates the potential of structured prompting as an exploratory method for large-scale narrative analysis.

Acknowledgments

Thanks to Prof. Yang for the course that introduced and inspired this line of research. Special thanks are due for the weekly guidance, patience, and insightful discussions. Thanks also to the workshop for providing this platform for newcomers to explore. Finally, heartfelt thanks go to the family for their continuous support.

Declaration on Generative AI

During the preparation of this work, the author used ChatGPT, Claude in order to: Grammar and spelling check, Paraphrase and reword, and Check Python code. After using this tool/service, the author reviewed and edited the content as needed and take full responsibility for the publication's content.

References

- [1] R. D. Benford, D. A. Snow, Framing processes and social movements: An overview and assessment, *Annual Review of Sociology* 26 (2000) 611–639. doi:10.1146/annurev.soc.26.1.611.

- [2] B. Klandermans, Consensus and action mobilization, in: *The Wiley-Blackwell Encyclopedia of Social and Political Movements*, John Wiley & Sons, Ltd, 2013. doi:10.1002/9780470674871.wbespm048.
- [3] A. O. Hirschman, *Exit, Voice, and Loyalty: Responses to Decline in Firms, Organizations, and States*, Harvard University Press, 1970.
- [4] R. Isaacs, Exit, voice and loyalty: An analytical framework for opposition agency in authoritarian regimes, in: *Political Opposition in Authoritarianism: Exit, Voice and Loyalty in Kazakhstan*, Springer International Publishing, 2022. doi:10.1007/978-3-031-06536-1_3.
- [5] P. John, K. Dowding, Spanning exit and voice: Albert hirschman's contribution to political science, in: *Research in the History of Economic Thought and Methodology*, volume 34B, Emerald Group Publishing Limited, 2016. doi:10.1108/S0743-41542016000034B007.
- [6] J. Lagerkvist, The unknown terrain of social protests in china: 'exit', 'voice', 'loyalty', and 'shadow', *Journal of Civil Society* 11 (2015) 137–153. doi:10.1080/17448689.2015.1052229.
- [7] Y. Wang, C. W. Roberts, Actantial analysis: Greimas's structural approach to the analysis of self-narratives, *Narrative Inquiry* 15 (2005) 51–74. doi:10.1075/ni.15.1.04wan.
- [8] T. Mikolov, K. Chen, G. Corrado, J. Dean, Efficient estimation of word representations in vector space, *arXiv preprint arXiv:1301.3781* (2013).
- [9] W. L. Hamilton, J. Leskovec, D. Jurafsky, Cultural shift or linguistic drift? comparing two computational measures of semantic change, in: *Proceedings of the Conference on Empirical Methods in Natural Language Processing*, 2016, pp. 2116–2121. doi:10.18653/v1/d16-1229.
- [10] N. Garg, L. Schiebinger, D. Jurafsky, J. Zou, Word embeddings quantify 100 years of gender and ethnic stereotypes, *Proceedings of the National Academy of Sciences* 115 (2018) E3635–E3644. doi:10.1073/pnas.1720347115.
- [11] E. Rodman, On political theory and large language models, *Political Theory* 52 (2024) 548–580. doi:10.1177/00905917231200826.
- [12] J. Elfes, Mapping news narratives using llms and narrative-structured text embeddings, *arXiv preprint arXiv:2409.06540* (2024). doi:10.48550/arXiv.2409.06540.
- [13] Qwen Team, Qwen2.5: A party of foundation models, <https://qwenlm.github.io/blog/qwen2.5/>, 2024. Technical Report.

A. Prompt Versions

This appendix presents three prompt versions used during iterative development.

A.1. Prompt V1

You are analyzing a Reddit post about labor and workplace issues.
Extract narrative structure using the Actantial Model (Greimas):

POST:
{text}

Respond with ONLY valid JSON (no markdown, no explanation):

```
{{"actants": {{
  "subject": "individual_worker", "collective_workers", "specific_group"
  "opponent": "individual_boss", "company", "capitalism_system",
    "government", "none"
  "helper": "coworkers", "union", "community", "government", "none"
  "object": "escape_quit", "fair_treatment", "collective_action",
    "system_change", "work_life_balance"
}},
```

Definitions:

Actants (who plays what role):

- subject: Main actor/protagonist identity
(individual "I" vs collective "we"/"workers")
- opponent: Who/what is antagonized
(specific person vs institution vs abstract system)
- helper: Support mentioned (coworkers, union, etc.; "none" if no support)
- object: Desired outcome/goal (what does subject want?)

Output ONLY the JSON object.

A.2. Prompt V2

You are analyzing a Reddit post about labor and workplace issues.

Extract narrative structure using the Actantial Model (Greimas)
and Frame Analysis:

POST:

{text}

Respond with ONLY valid JSON (no markdown, no explanation):

```
{{
  "actants": {{
    "subject": "individual_worker|collective_workers|specific_group",
    "opponent": "individual_boss|company|capitalism_system|
government|none",
    "helper": "coworkers|union|community|government|none",
    "object": "escape_quit|fair_treatment|collective_action|
system_change|work_life_balance"
  }},
  "frames": {{
    "problem_type": "unfair_pay|toxic_boss|long_hours|lack_respect|
exploitation|burnout|other",
    "attribution": "individual_fault|bad_company|systemic_issue|
capitalism",
    "action_taken": "quit|confront|organize|legal|vent_only",
    "collective_mentioned": "yes|no"
  }},
  "confidence": "high|medium|low"
}}
```

Definitions:

ACTANTS (who plays what role):

- subject: Main actor/protagonist identity
(individual "I" vs collective "we"/"workers")
- opponent: Who/what is antagonized
(specific person vs institution vs abstract system)
- helper: Support mentioned (coworkers, union, etc.; "none" if no support)
- object: Desired outcome/goal (what does subject want?)

FRAMES (how problem is understood):

- problem_type: Main complaint/grievance
- attribution: Where is blame/responsibility placed?
- action_taken: Response taken or proposed
- collective_mentioned: Any mention of unions/strikes/organizing?

Output ONLY the JSON object.

A.3. Prompt V3

You are coding a Reddit post about labor issues into a JSON schema.

POST:
{text}

For each field, **choose exactly ONE** label from the allowed options.
Do NOT output the list of options. Output only the chosen label.

Allowed labels:

- actants.subject: "individual_worker", "collective_workers"
- actants.opponent: "individual_boss", "company", "capitalism_system", "none"
- actants.helper: "none", "union", "coworkers"
- actants.object: "escape_quit", "fair_treatment", "collective_action"

- frames.problem_type: "unfair_pay", "toxic_boss", "long_hours",
"lack_respect", "exploitation", "burnout", "other"
- frames.attribution: "individual_fault", "bad_company",
"systemic_issue", "capitalism"
- frames.action_taken: "quit", "confront", "organize", "legal", "vent_only"
- frames.collective_mentioned: "yes", "no"

Output ONLY valid JSON in this format (no explanations):

```
{
  "actants": {
    "subject": "one_label_here",
    "opponent": "one_label_here",
    "helper": "one_label_here",
    "object": "one_label_here"
  },
  "frames": {
    "problem_type": "one_label_here",
    "attribution": "one_label_here",
    "action_taken": "one_label_here",
    "collective_mentioned": "one_label_here"
  },
  "confidence": "high|medium|low"
}
```