

Electoral Spectators: 2024 US Presidential Election and Online Political Discourse in China*

Muzhi Liu[†]

This version: December 2025

Abstract: Democratic elections abroad provide rare, observable signals that citizens under authoritarian rule consume as electoral spectators. I study Chinese-language social media discourse surrounding the 2024 U.S. presidential election using 14,606 posts and 379,725 comments from Weibo and Rednote. At baseline, Weibo discourse is substantially more anti-American and anti-democracy than Rednote, consistent with exposure to state-curated narratives. Yet when the contest resolves peacefully, evaluations shift: difference-in-differences estimates show Weibo converging toward Rednote's more favorable baseline, and within-user fixed effects models confirm that shifts emerge on Election Day and persist for two weeks. A key interpretive finding is that democracy evaluations are tightly bundled with U.S. sentiment, suggesting that "democracy" in this discourse functions as country-specific affect rather than separable institutional commitment. Post-level sentiment diffuses to comments, with significant within-thread alignment across outcomes. Three dynamics emerge among commenters: Trump support rises after victory, especially among mainland users; anti-U.S. and anti-democracy expression declines even as overall discussion contracts; and political engagement falls once uncertainty resolves. China sentiment remains flat—updating foreign evaluations does not require revising domestic attachments. These findings suggest that state-curated environments shape baseline priors but do not fully insulate publics from observable signals that democracies can function.

Word Count: 7,802

*I thank Shigeo Hirano, Donald Green, Eunji Kim, Yamil Velez, Patrick Liu, Ronald Yi Xiu Wu, and participants of the Kimuli Kasara Comparative Politics Seminar and American Politics Graduate Seminar at Columbia University for their helpful comments and suggestions. All errors are my own.

[†]Columbia University (email: ml4967@columbia.edu)

1 Introduction

Democratic elections are not only domestic political events; they are globally visible signals about whether democracy can govern conflict, transfer power, and produce order. For citizens living under non-electoral regimes, this visibility creates a distinctive informational situation: they can observe a high-stakes democratic contest abroad, but they do so as electoral spectators—with no vote, limited ability to verify competing narratives, and varying degrees of personal stake in the outcome. What happens to authoritarian publics' expressed beliefs when an election moves from uncertainty to resolution? Do they treat the outcome as confirmation that democracy is chaotic and unreliable, or as evidence that democratic institutions can function even under stress?

This question matters for three reasons. First, authoritarian regimes invest heavily in narratives of Western democratic dysfunction, framing foreign elections as proof that democracy produces chaos rather than order (Repnikova, 2022; Roberts, 2018). If foreign elections reliably produce cynicism among authoritarian audiences, then international democratic stress tests may inadvertently strengthen authoritarian narrative power. Second, elections provide a rare kind of observable evidence. Regardless of regime messaging, an election that resolves peacefully gives spectators an undeniable fact: institutions either hold or they do not. This creates a clear theoretical prediction about timing—the moment of resolution should be when belief updating is most likely. Third, elections are attention shocks. They generate intense engagement during uncertainty, but attention drops quickly once a winner emerges. The post-election window is therefore crucial for understanding what remains in expressed opinion after the spectacle fades.

I study these dynamics in the context of the 2024 U.S. presidential election, using Chinese-language social media discourse as a window into expressed political evaluation under authoritarian governance. China is a strategically important case: state media regularly frames U.S. elections, polarization, and institutional conflict as proof that democracy is unstable and unsuitable for China (Repnikova, 2017), while online platforms operate under extensive censorship and information management (King, Pan and Roberts, 2013; Roberts, 2018). Yet Chinese-speaking audiences also consume and debate U.S. politics at scale, forming impressions that may diverge from official framing—and that vary depending on whether users are domestic spectators or diaspora members with direct stakes in American policy.

The analysis draws on over 390,000 posts and comments from two major Chinese-language

platforms: Rednote (Xiaohongshu) and Weibo. These platforms differ sharply in user composition and information environment. Weibo is predominantly domestic and operates under direct state content moderation (King, Pan and Roberts, 2017). Rednote hosts a substantial overseas Chinese user base, including international students and immigrants for whom U.S. policy outcomes—on visas, immigration, and trade—carry personal consequences. This variation provides two kinds of analytic leverage. Cross-platform comparisons reveal baseline differences in how elections are framed across information environments. Within-platform analyses allow estimation of how expressed evaluations change at the moment of resolution, and whether such changes differ between domestic and diaspora audiences.

Methodologically, the paper combines several approaches to address the challenges of studying opinion dynamics in social media discourse. Social media data present well-known difficulties for political analysis: posts are often short and context-dependent, platforms contain substantial noise from bots, spam, and off-topic content, and user composition shifts over time as events attract different audiences (Ruths and Pfeffer, 2014; Tufekci, 2014; Grimmer and Stewart, 2013). These features make it difficult to distinguish genuine attitude change from compositional shifts, and standard text analysis methods struggle to capture nuanced political stances in brief, informal text (Hopkins and King, 2010). The paper addresses these challenges through a combination of measurement and identification strategies.

First, I employ large language models (LLMs) to classify multiple dimensions of political sentiment at scale. Unlike keyword-based or dictionary approaches that capture only surface-level features and perform poorly on short, informal text, LLM classification can identify nuanced evaluative stances—distinguishing, for instance, between pro-Trump sentiment, pro-U.S. sentiment, and pro-democracy sentiment within the same post. Second, I combine multiple identification strategies—each well-established in causal inference but rarely applied jointly to social media data—to triangulate post-election dynamics and separate attitude change from compositional artifacts. Post-level difference-in-differences models with platform-by-period interactions estimate whether sentiment shifted more on Weibo than on Rednote at the election resolution cutoff. Comment-level within-user fixed effects models isolate individual-level change by comparing the same user’s expressed sentiment before and after the election, absorbing time-invariant user characteristics and directly addressing the concern that post-election shifts might reflect different users entering the conversation rather than the same users changing their views. Event study specifications trace the precise timing of shifts, testing for pre-trends and estimating daily coefficients around the election. The combination of

cross-platform, within-user, and within-thread analyses provides a more comprehensive account of how foreign electoral events propagate through authoritarian-context social media than any single design could offer.

Four main findings emerge. First, baseline platform differences are large and concentrated in country- and regime-linked evaluation: Weibo content is substantially more anti-U.S. and anti-democracy than Rednote. Difference-in-differences estimates show that temporal change is also concentrated on Weibo, with evaluations partially converging toward Rednote's more favorable baseline after the election. Second, "democracy" talk is largely U.S. legitimacy talk—U.S. sentiment overwhelmingly predicts democracy views, suggesting institutional language functions as country-specific affect rather than separable democratic commitment. Third, sentiment diffuses from posts to comments, with stronger within-thread alignment on Weibo than Rednote. Fourth, within-user models reveal three post-election dynamics: a bandwagon effect as Trump support rises, especially among mainland users; dampening of anti-U.S. and anti-democracy expression; and attention withdrawal as political engagement declines. China sentiment remains stable, suggesting that updating foreign evaluations does not require revising domestic attachments.

These findings speak to debates about authoritarian information control and cross-border opinion formation. They suggest that authoritarian narrative environments shape baseline priors—Weibo's more hostile baseline reflects sustained exposure to state-curated framing—but do not fully insulate audiences from updating when democratic institutions visibly function. Yet the nature of this updating is more limited than simple "demonstration effects" accounts might predict: peaceful resolution dampens negative discourse rather than generating positive enthusiasm, and the largest shifts occur on the platform with the most negative baseline, representing convergence rather than uniform movement. For democracies concerned with international legitimacy, the implication is cautiously optimistic—functioning institutions can reduce hostile discourse even in unfavorable information environments—but the returns may prove shallow if they depend more on spectacle fatigue than genuine persuasion.

2 Literature Review

This paper asks how audiences embedded in authoritarian information environments interpret high-salience democratic events abroad. I engage four strands of scholarship: cross-border political communication and democratic diffusion, authoritarian narrative competition

and citizen learning, electoral spectatorship and post-outcome dynamics, and platform ecologies that shape both what spectators see and how sentiment propagates. I then synthesize these literatures to identify the gap this paper addresses.

2.1 Cross-Border Political Communication

Citizens are increasingly exposed to political information from beyond their national borders, and such exposure can shape domestic political attitudes. Kern and Hainmueller (2009) demonstrate that East Germans with access to West German television were more satisfied with the communist regime—an “opium of the masses” effect in which foreign entertainment substituted for political dissatisfaction. Conversely, Crabtree, Darmofal and Kern (2015) find that exposure to West German television increased support for democratization in regions where it provided information rather than mere entertainment. These findings suggest that the *content* of cross-border information—whether it conveys democratic competence or dysfunction—matters for its political effects.

A related literature examines “demonstration effects” in democratization: whether successful transitions in one country increase pressure for reform in others (Huntington, 1991; Brinks and Coppedge, 2006). The logic is that citizens in autocracies observe foreign democracies and update on whether the democratic model “works.” Yet the empirical record is mixed. Weyland (2014) argues that demonstration effects are often shallow—citizens respond to dramatic, visible events but draw limited lessons about institutional design. Gleditsch and Ward (2006) finds that regional democratic density predicts transitions, but the mechanism remains unclear: diffusion may operate through elite learning, popular mobilization, or international pressure rather than mass attitude change.

Less is known about how citizens in authoritarian contexts interpret democratic events *in real time*. Elections, unlike regime transitions, are scheduled and predictable, generating anticipatory attention and enabling researchers to observe attitudes before, during, and after the event. When citizens observe a foreign election, do they update on democracy as a system, on the specific candidates involved, or on the country holding the election? The distinction matters: if Chinese audiences respond to the 2024 U.S. election by revising their assessment of American institutions rather than their views on democracy *in general*, this suggests that “demonstration effects” may be more country-specific than the diffusion literature assumes.

2.2 China's Information Environments and Citizen Learning

State media in China invests substantial resources in shaping domestic narratives about Western democracy, often highlighting electoral dysfunction, political polarization, and social unrest in democratic countries as evidence that the Western model is unsuitable for China (Repnikova, 2017, 2022). Content moderation promotes information conveying state support (King, Pan and Roberts, 2017) and portrays society as harmonious (Lu et al., 2025). On the other hand, online censorship in China operates through “friction” and “flooding” rather than absolute prohibition (Roberts, 2018). Friction raises the cost of accessing foreign information; flooding saturates the information space with regime-friendly content. Neither mechanism guarantees that citizens internalize official narratives. King, Pan and Roberts (2013) show that censorship targets collective action potential rather than criticism per se, suggesting tolerance for heterogeneous opinion. Qin, Strömberg and Wu (2024) find that despite Weibo’s substantial state presence, Chinese protest participants use the platform as a focal point for spreading collective action messages.

Given the information environment in China, how do citizens respond to U.S.-related political events? Survey evidence indicates that Chinese public opinion about the United States responds to bilateral relations. Fang, Johnston and Shen (2022) document that Chinese evaluations of the United States declined during the Trump administration’s trade war but rebounded modestly under Biden, suggesting that citizens update based on observed events rather than simply absorbing state propaganda. However, survey research in China faces significant constraints: social desirability bias, restrictions on sensitive questions, and government approval processes limit what can be measured (Huang, Intawan and Nicholson, 2023).

These limitations have motivated increasing use of social media data to study political attitudes in authoritarian contexts. Social media captures organic, unsolicited expression at temporal granularity that surveys cannot match, enabling analysis of how attitudes evolve around specific events. The data are not without problems—users are unrepresentative (Mukerjee, Jaidka and Lelkes, 2022), content is shaped by platform affordances and censorship, and expressed attitudes may diverge from privately held beliefs (Ruths and Pfeffer, 2014; Tufekci, 2014). Yet for studying *expressed* political sentiment—the publicly visible discourse that constitutes the information environment for other users—social media offers unique advantages. The key question is whether state-curated information environments fully determine expressed sentiment, or whether observable foreign events can override official framing.

2.3 Electoral Spectatorship and Consequences

Foreign elections present authoritarian publics with a distinctive situation: they can observe a high-stakes democratic contest but participate only as spectators. This spectatorship generates at least three possible response dynamics.

First, *bandwagon effects* may lead observers to update toward the winning candidate. The bandwagon mechanism is well-documented in domestic contexts: individuals shift toward supporting winners after primary elections (Bartels, 1988), and post-election surveys reveal increased approval for victorious candidates (Bischoff and Egbert, 2013). The mechanism may involve conformity, inference that winners' positions are meritorious, or desire to align with prevailing sentiment (Mehrabian, 1998; Mutz, 1997). Whether bandwagon effects operate for *foreign* elections is less clear. Citizens observing elections abroad lack direct stakes and the informational environment—campaign exposure, social discussion, media saturation—that facilitates domestic preference formation. Yet globalized media increasingly exposes citizens to foreign political events in real time, potentially enabling preference formation even for distant contests. If bandwagon dynamics operate cross-nationally, we should observe increased support for Trump after his victory.

Second, *attention withdrawal* may follow electoral resolution. Elections generate engagement partly because uncertainty creates dramatic tension (Prior, 2007). Once the outcome is known, the spectacle value of following the contest diminishes. This entertainment dimension of foreign election-watching implies that political engagement should decline after resolution, even if substantive interest in policy implications persists. The dampening of election-related discussion may be asymmetric: negative, conflict-oriented commentary thrives on uncertainty and may decline more sharply than policy-focused discourse once the narrative closes.

Third, exposure to American democratic events may shift sentiment toward the United States and democratic governance in either direction. One possibility is that observing a peaceful electoral resolution provides evidence of institutional competence, leading spectators to express less negative views of the U.S. and its political system (Levi, 1997; Tyler, 2006). An alternative possibility is that engagement with U.S. politics reinforces pre-existing skepticism: if spectators interpret the election through the lens of state narratives emphasizing American dysfunction, or if the contest itself surfaces polarization and conflict, exposure may entrench rather than soften negative evaluations (Repnikova, 2022). Distinguishing these possibilities reveals whether authoritarian publics are capable of updating positively on foreign democra-

cies when institutions visibly function, or whether prior beliefs and information environments filter observable signals in ways that prevent such updating.

2.4 Platform Ecologies and Sentiment Diffusion

Social media platforms differ in user bases, algorithmic curation, and content moderation, with consequences for aggregate political expression. Weibo, China's dominant microblogging platform, operates under state content moderation and serves a predominantly domestic audience (King, Pan and Roberts, 2017). Rednote (Xiaohongshu), originally a lifestyle platform, has cultivated a substantial overseas Chinese user base and is subject to less intensive political moderation. These differences may produce systematic variation in expressed attitudes through multiple mechanisms.

Selection implies that users with different baseline attitudes sort into different platforms: internationally oriented users may prefer Rednote, while domestically focused users concentrate on Weibo (Prior, 2007; Sunstein, 2017). *Socialization* implies that platform environments shape attitudes through exposure to co-users' content and platform-specific norms (Bail et al., 2018). *Censorship* affects what can be expressed, shifting observed sentiment even if underlying attitudes are similar (Roberts, 2018). Distinguishing these mechanisms is challenging, but examining how platforms respond *differentially* to common shocks provides leverage: if users on both platforms update similarly despite different baselines, this suggests selection rather than platform-specific effects on opinion formation.

Beyond platform-level differences, social media architecture shapes how sentiment propagates within discussions. Research on emotional contagion demonstrates that exposure to others' expressed sentiment can shift one's own emotional expression (Kramer, Guillory and Hancock, 2014; Brady et al., 2017). In threaded discussions, posts may function as cues that shape subsequent comments—a form of within-thread sentiment diffusion. If alignment between post and comment sentiment varies across platforms, topics, or time periods, this reveals how effectively different frames propagate and whether platform structure moderates diffusion. Weibo's more public, broadcast-oriented architecture may produce tighter coupling between posts and replies than Rednote's more intimate, interest-based communities.

2.5 Synthesis: What We Do Not Know

These literatures converge on a set of open questions. Cross-border communication research establishes that foreign media exposure shapes attitudes but offers limited evidence on

how citizens process *discrete democratic events* in real time. Demonstration effects scholarship predicts that successful elections should improve views of democracy, but does not specify whether this operates through country-specific affect or abstract institutional evaluation. Research on authoritarian information environments shows that state narratives shape baseline attitudes but leaves unclear whether observable foreign events can override official framing. Work on bandwagon effects and attention dynamics has focused on domestic elections, with little evidence on whether these mechanisms operate for foreign contests observed from authoritarian contexts. And platform studies document that different environments produce different expressed sentiment, without establishing whether users update similarly or differently in response to common shocks.

This paper addresses these gaps by studying Chinese-language social media discourse surrounding the 2024 U.S. presidential election. The election provides a clear informational signal—peaceful resolution of a contested race—whose effects can be traced at the platform, user, and thread level. I examine whether baseline platform differences persist or converge after the election, whether sentiment diffuses from posts to comments, and whether users exhibit bandwagon effects, attention withdrawal, or nationalist counter-mobilization. By combining post-level comparisons, within-user panel models, and within-thread alignment analyses, I provide a more comprehensive account of how foreign electoral events propagate through authoritarian-context social media than prior work has offered.

2.6 Hypotheses

The 2024 U.S. election provides two conceptually distinct kinds of information to foreign observers. First, the process may signal procedural competence: whether democratic institutions can manage conflict, transfer power, and produce legitimate outcomes under stress. Second, the result may trigger outcome-conditioned updating, as spectators revise evaluations of the country and its leadership in light of who won. These informational channels generate the following predictions:

- **H1 (Platform Baseline):** Weibo discourse will exhibit more negative sentiment toward the United States and U.S. democracy than Rednote discourse at baseline.
- **H2 (Legitimacy Updating):** Evaluations of the United States and U.S. democracy will become less negative after the election resolves peacefully.
- **H3 (Bandwagon):** Expressed support for Trump will increase after his victory, with

larger effects among Mainland China users than among U.S.-based users.

- **H4 (Attention Withdrawal):** Political engagement will decline after the election, as uncertainty resolves and spectacle value diminishes.
- **H5 (Convergence):** Weibo and Rednote users will update in similar directions after the election, even if from different baselines.

3 Data and Methods

3.1 Data Collection

I collected public posts and their associated comment threads from Weibo and Rednote (Xiaohongshu) surrounding the 2024 U.S. presidential election. Data collection proceeded in real time as events unfolded, rather than retrospectively, helping mitigate concerns about post-hoc censorship of Chinese social media content. The empirical design treats the election outcome as a shared information shock for Chinese-speaking observers. The main analyses focus on a symmetric window around Election Day (November 5, 2024). I define *pre-election* as posts/comments created before November 5, *election day* as November 5, and *post-election* as posts/comments created after November 5. Content was retrieved using keyword-based queries designed to capture election-related discussion across platforms and across Chinese-speaking communities. Table 1 reports the keyword families (English translations shown; queries were executed in Chinese on both platforms, and in English where indicated). Candidate names include multiple transliterations in Chinese.

The note-level dataset consists of original posts returned by keyword queries. For each note, I retain the text, timestamp, platform identifiers, and engagement metadata (e.g., likes and comment counts, where available). The comment-level dataset consists of comments attached to collected notes; for each comment I retain the text, timestamp, and platform-provided coarse location (IP-based country/region) when available. Posts do not include user location information, so geographic heterogeneity is analyzed using commenters.

3.2 LLM Labeling

I use large language models (LLMs) to classify Chinese-language posts and comments along multiple dimensions relevant to the paper’s argument: (i) expressed evaluations of political objects (candidates, countries, and democratic institutions), (ii) hostile discourse toward

Category	Keyword
Event	U.S. presidential election
Candidates	Trump (transliteration 1) Trump (transliteration 2) Harris (transliteration) Harris (official Chinese name)
Candidates (English)	Trump Harris
Parties & Movements	Republican Party Democratic Party MAGA
Community	Chinese Trump supporters

Note: Data were collected in real-time during election days. Keywords were queried in Chinese on both platforms, with candidate names including both mainland Chinese transliterations and alternative forms common in Taiwan and overseas Chinese communities.

Table 1: Keywords Used for Data Collection

political outgroups, and (iii) informational quality and misinformation. This approach enables consistent annotation of nuanced political expression at a scale that would be prohibitively costly for fully manual coding. Throughout, the outcome variables are best understood as measures of expressed sentiment and framing in public discourse rather than privately held attitudes. In particular, *Democracy View* is coded as evaluation of democracy/democratic processes as invoked in the post/comment. In this election-centered corpus, references to “democracy” frequently function as shorthand for assessments of American democratic performance (e.g., electoral procedures, legitimacy, peaceful transfer), rather than as abstract regime preference. This operationalization matches the theoretical focus on how authoritarian publics interpret foreign democratic events.

Table 2 summarizes the outcome variables and coding rules. I classify three types of constructs: (A) attitude variables capturing directional sentiment toward political objects, (B) binary indicators for the presence of content features, and (C) ordinal/continuous measures of informational quality.¹

Attitude variables use a ternary scale (−1, 0, 1) because election discourse often consists

¹LLM labeling code has been pre-registered at AsPredicted.org. (AsPredicted #262,912). The full pre-registration document is presented in Appendix B.

Variable	Scale	Description
<i>Panel A: Attitude Variables (−1 = negative, 0 = neutral, 1 = positive)</i>		
Trump	{−1, 0, 1}	Attitude toward Donald Trump: supports/praises (1), neutral/not mentioned (0), criticizes/opposes (−1)
Harris	{−1, 0, 1}	Attitude toward Kamala Harris: supports/praises (1), neutral/not mentioned (0), criticizes/opposes (−1)
Republican	{−1, 0, 1}	Attitude toward Republican Party beyond Trump
Democrat	{−1, 0, 1}	Attitude toward Democratic Party beyond Harris
China	{−1, 0, 1}	Attitude toward China: nationalist pride/defense (1), neutral (0), critical of China (−1)
US	{−1, 0, 1}	Attitude toward United States: admires US system (1), neutral (0), critical/schadenfreude (−1)
Democracy View	{−1, 0, 1}	Attitude toward democracy: values democratic processes (1), neutral (0), skeptical/mockingly (−1)
<i>Panel B: Binary Variables (0 = absent, 1 = present)</i>		
Outgroup Hate	{0, 1}	Contains hostile, contemptuous, or derogatory language toward political outgroups
Policy Interest	{0, 1}	Mentions specific policy issues (tariffs, immigration, abortion, TikTok ban, etc.)
Political	{0, 1}	Content is politically relevant (discusses elections, candidates, parties, policies)
<i>Panel C: Ordinal Variables</i>		
Informed	{0, 1, 2}	Political sophistication: low/slogans (0), moderate awareness (1), specific policy knowledge (2)
Misinfo	[0, 1]	Factual accuracy: accurate (0), minor issues (0.25), mixed (0.5), mostly false (0.75), fabricated (1)

Table 2: LLM Classification Scheme

of brief endorsements, condemnations, or passing references. A binary scale would conflate neutrality/non-mention with opposition, while a five-point scale would demand distinctions short texts rarely support. Binary variables capture the presence/absence of salient content features (e.g., outgroup hate). The informed measure is ordinal because political sophistication admits meaningful gradations. Misinformation is treated as a graded index because factual accuracy varies from fully accurate to fully fabricated.

I use GPT-4.1 (OpenAI) to classify each post and comment. The prompt provides detailed coding instructions, scale definitions, and worked examples spanning the range of election-related discourse observed in the corpus. I validate the LLM labels against expert-coded labels for a random sample of 200 items from each dataset (Rednote posts, Rednote comments, Weibo posts, Weibo comments). I report agreement metrics (accuracy, MCC, and Kappa) in Appendix C. These validation results strongly support the use of LLM-based classification for my main analyses. All outcome variables achieve a 0.9 - 1.0 accuracy, an over 0.8 MCC, with the large majority reaching almost perfect agreement ($\kappa > 0.80$).

3.3 Note Data

I analyze 14,606 notes from Rednote ($N = 3,442$) and Weibo ($N = 11,164$) collected during the 2024 U.S. presidential election. The note-level analysis examines (i) what users choose to express (stance and tone) and (ii) what platforms appear to reward (engagement). I distinguish between differences in composition (the mix of content appearing on each platform) and differences in amplification (how engagement concentrates on particular content types).

3.3.1 Descriptive Statistics

Table 3 reports descriptive statistics by platform. Because sentiment measures use $\{-1, 0, 1\}$ coding, the descriptive percentages are informative not only about evaluation but also about salience (how often users take a stance versus merely reference an object). Three patterns motivate the subsequent modeling strategy.

First, election discourse is dominated by neutral references rather than explicit stance-taking. Roughly three-quarters to nine-tenths of notes are neutral toward each candidate. When notes do take a position, sentiment is asymmetric: pro-Trump notes outnumber anti-Trump notes on both platforms, whereas Harris content skews negative. Second, the sharpest cross-platform cleavage concerns evaluations of the United States and democracy. Weibo contains substantially more anti-US and anti-democracy content than Rednote. By contrast, China

sentiment is overwhelmingly neutral on both platforms. Third, Weibo shows modestly higher base rates of outgroup hate and misinformation. While these rates are small in absolute terms, they matter for understanding what kinds of content become visible and potentially contagious through engagement and comment dynamics.

These descriptive patterns motivate the note-level questions: Are platform gaps in US and democracy sentiment compositional or structural? Which content features co-occur with hostility and institutional skepticism? And do platforms differ in how discourse updates after the election outcome is resolved?

3.3.2 Identification and Estimation Strategies

The note-level analysis leverages cross-platform variation in expressed discourse and temporal dynamics. It is intentionally complementary to the comment-level analysis: whereas the comment-level models use within-user change to study updating around election resolution, the note-level models describe how two platforms differ in baseline discourse and how their public-facing content shifts after the same external shock. In the main text, I use three empirical approaches: (i) bivariate platform comparisons, (ii) multivariate models of sentiment correlates, and (iii) difference-in-differences estimates of post-election shifts.

Bivariate platform models. To summarize the directional platform gradient in each attitude outcome, I estimate ordered logistic regressions of each sentiment outcome on a platform indicator:

$$\Pr(Y_i \leq j) = \frac{1}{1 + \exp(-(\alpha_j - \beta \cdot \text{Weibo}_i))}, \quad (1)$$

where $Y_i \in \{-1, 0, 1\}$ is the sentiment outcome for note i and Weibo_i indicates platform membership. A positive β indicates that Weibo notes are more likely to fall in a higher (more positive) sentiment category. Because 0 can reflect non-mention, these bivariate models capture combined differences in salience and valence.

Multivariate correlates of democracy views. To assess whether “democracy” functions as an abstract institutional evaluation or as a proxy for country-specific affect, I estimate an ordered logistic regression predicting democracy views from multiple covariates:

$$\Pr(\text{Democracy}_i \leq j) = \frac{1}{1 + \exp(-(\alpha_j - \mathbf{X}'_i \boldsymbol{\beta}))}, \quad (2)$$

	Posts		Comments	
	Rednote (<i>N</i> = 3, 442)	Weibo (<i>N</i> = 11, 164)	Rednote (<i>N</i> = 138, 246)	Weibo (<i>N</i> = 241, 479)
<i>Panel A: Candidate Sentiment</i>				
Pro-Trump (%)	12.4	17.5	4.0	5.2
Anti-Trump (%)	8.1	8.3	3.3	4.4
Neutral Trump (%)	79.5	74.2	92.7	90.4
Mean Trump sentiment	0.043	0.091	0.007	0.008
Pro-Harris (%)	3.8	2.2	0.7	0.9
Anti-Harris (%)	5.9	7.8	4.1	4.1
Neutral Harris (%)	90.3	90.0	95.2	95.1
Mean Harris sentiment	−0.021	−0.056	−0.034	−0.032
<i>Panel B: Country Sentiment</i>				
Pro-China (%)	3.5	4.5	1.2	1.8
Anti-China (%)	0.3	0.4	0.6	0.7
Neutral China (%)	96.1	95.2	98.2	97.6
Mean China sentiment	0.032	0.041	0.006	0.011
Pro-US (%)	7.9	4.0	1.0	0.9
Anti-US (%)	8.9	15.1	5.1	7.2
Neutral US (%)	83.2	80.9	94.0	92.0
Mean US sentiment	−0.010	−0.111	−0.041	−0.063
<i>Panel C: Democracy Views</i>				
Pro-democracy (%)	6.5	3.4	0.8	0.7
Anti-democracy (%)	5.8	9.5	3.5	5.2
Neutral democracy (%)	87.6	87.1	95.7	94.2
Mean democracy view	0.007	−0.061	−0.027	−0.045
<i>Panel D: Discourse Quality</i>				
Political content (%)	91.2	97.6	33.9	38.8
Policy interest (%)	34.7	31.2	9.0	7.4
Mean informed score	1.09	1.17	0.24	0.22
Outgroup hate (%)	3.7	4.7	3.7	5.5
Mean misinfo score	0.036	0.055	0.022	0.024
High misinfo (%)	2.8	4.7	2.2	2.5
<i>Panel E: Commenter Location</i>				
Mainland China	—	—	71,141 (75.4%)	110,996 (95.6%)
United States	—	—	11,009 (11.7%)	1,580 (1.4%)
Canada	—	—	2,206 (2.3%)	365 (0.3%)
Hong Kong	—	—	1,028 (1.1%)	581 (0.5%)
Australia	—	—	1,029 (1.1%)	346 (0.3%)
Other	—	—	6,913 (7.3%)	1,886 (1.6%)

Note: Sentiment variables coded as −1 (negative), 0 (neutral), 1 (positive). Informed score ranges from 0 to 2. Misinfo score ranges from 0 to 1; high misinfo defined as ≥ 0.5 . Panel E reports unique commenters by IP location; “Other” includes Japan, Singapore, Europe, Southeast Asia, Taiwan, Macau, South Korea, United Kingdom, and unknown locations. Posts do not include user location information.

Table 3: Descriptive Statistics by Platform and Content Type

where $\mathbf{X}_i$ includes platform indicator, candidate sentiments (Trump, Harris), country sentiments (US, China), policy interest, informational quality, and misinformation score. If democracy views are tightly bundled with US sentiment—with the US sentiment coefficient dominating other predictors—this suggests that “democracy” in this discourse operates primarily as a label for country-specific affect rather than a separable institutional commitment.

Difference-in-differences (platform-by-post-election updating). To estimate whether platforms update differently after election resolution, I estimate:

$$Y_i = \beta_0 + \beta_1 \cdot \text{Weibo}_i + \beta_2 \cdot \text{Post}_i + \beta_3 \cdot (\text{Weibo}_i \times \text{Post}_i) + \beta_4 \cdot \text{Days}_i + \varepsilon_i, \quad (3)$$

where Post_i indicates notes posted after November 5, 2024 and Days_i is the number of days from the election (centered at zero). The coefficient β_2 captures the average discontinuous shift on Rednote, while β_3 captures the platform-specific differential. Election-day observations are excluded.

The identifying assumption is that, absent the election, sentiment trends would have been parallel across platforms. I cannot fully test this assumption, but including Days_i absorbs gradual drift that might otherwise confound post-election discontinuities. As with the comment-level analysis, these estimates capture changes in expressed sentiment and framing, which may reflect genuine updating and/or strategic self-presentation in public discourse.

3.4 Comment Data

3.4.1 Descriptive Statistics

The comment-level corpus comprises 379,725 comments: 138,246 from Rednote and 241,479 from Weibo. Comments provide a different window on public opinion than notes: they are more interactive, more immediately responsive to salient events, and more tightly linked to the discourse environments created by platform sorting and thread-level dynamics.

Panel E of Table 3 highlights sharp cross-platform differences in commenter geography. Rednote’s commenter base is comparatively diverse, with a substantial overseas Chinese presence, whereas Weibo commenters are overwhelmingly domestic. This compositional contrast matters for interpretation: Rednote enables heterogeneity analyses by political context (Mainland vs. U.S.-based users), while Weibo provides high-volume evidence for domestic Chinese discourse.

Comments exhibit lower political engagement and lower informational content than posts. Many comments are brief reactions or social signals rather than substantive contributions, producing higher neutrality rates and more compressed sentiment distributions. Despite this compression, comments skew more skeptical of the United States and democracy than posts, especially on Weibo—a pattern consistent with comments functioning as a venue for sharper, more evaluative reactions.

3.4.2 Identification and Estimation Strategies

I use four complementary strategies to examine comment discourse on each platform: (i) content diffusion within threads, (ii) within-user pre/post estimates, (iii) an event study design, and (iv) unconditional logistic regressions to distinguish dampening from mobilization. Together, these strategies distinguish between (a) thread-level alignment and contagion, (b) within-user updating around election resolution, and (c) shifts in the overall volume of positive versus negative discourse.

Within-thread alignment. To assess whether post-level sentiment predicts comment-level responses within the same thread, I estimate:

$$Y_{ijt} = \beta \cdot X_{jt} + \gamma_h + \delta_d + \varepsilon_{ijt},$$

where Y_{ijt} is the sentiment outcome for comment i on post j at time t , X_{jt} is the corresponding sentiment measure for post j , γ_h denotes hour-of-day fixed effects, and δ_d denotes day-of-week fixed effects. Standard errors are clustered at the post level.

The coefficient β captures how closely comment sentiment aligns with the originating post. Because users may self-select into congenial threads and platforms may sort users into content, these estimates are interpreted as thread-level alignment rather than a pure causal effect of exposure.

Within-user pre/post estimates (difference-in-means with user fixed effects). To quantify how the *same users* shift their expressed sentiment after the election outcome is known, I estimate:

$$Y_{it} = \alpha_i + \beta \cdot \text{Post}_t + \varepsilon_{it},$$

where Y_{it} is the user-day average sentiment for user i on day t , α_i denotes user fixed effects, and Post_t indicates days on or after November 5, 2024. Standard errors are two-way clustered

by user and date. This design isolates within-person updating and is therefore closely aligned with the paper’s interpretation of election resolution as an information shock.

Event study (day-by-day updating around election resolution). To trace the dynamics of updating, I estimate an event study at the user-day level:

$$Y_{it} = \alpha_i + \sum_{k=-K}^K \beta_k \cdot D_{it}^k + \varepsilon_{it}, \quad k \neq -1,$$

where D_{it}^k indicates that day t falls k days from Election Day (November 5, 2024, aligned to platform timestamps in China Standard Time). Day -1 is the reference period. The pre-election coefficients provide a visual diagnostic for stability/parallel trends.

The event-study coefficients capture changes in expressed sentiment among users who are active in the comment space. As with any social-media-based design, the estimates can be influenced by changing participation (who comments when) and by strategic self-presentation. For this reason, I interpret the event study as evidence about public opinion as expressed in platform discourse during a high-salience international political event, rather than as a direct measure of latent attitudes.

Conditional and unconditional logistic regressions (dampening vs. mobilization). The within-user fixed effects models estimate whether sentiment becomes more or less positive among users who express any relevant sentiment. However, a ratio shift toward less negative sentiment could arise through two distinct mechanisms: *mobilization* of positive voices or *dampening* of negative voices. To distinguish these possibilities, I estimate logistic regressions of the form:

$$\Pr(Y_i = 1) = \frac{1}{1 + \exp(-(\beta_0 + \beta_1 \cdot \text{Post}_i + \beta_2 \cdot \text{Weibo}_i + \beta_3 \cdot \text{Post}_i \times \text{Weibo}_i))},$$

applied to two complementary samples with different outcome definitions.

Conditional models restrict the sample to comments expressing relevant sentiment (democracy view $\in \{-1, 1\}$ or US sentiment $\in \{-1, 1\}$) and define $Y_i = 1$ if the comment is anti-democracy or anti-US. These models estimate whether the *composition* of relevant discourse shifts—that is, whether the share of negative sentiment among those discussing the topic declines after the election. The interaction term β_3 tests whether this compositional shift differs

across platforms.

Unconditional models use the full comment corpus and define Y_i in three ways: (a) anti-democracy/anti-US expression (1 if sentiment = -1 , 0 otherwise), (b) pro-democracy/pro-US expression (1 if sentiment = 1 , 0 otherwise), and (c) any relevant expression (1 if sentiment $\in \{-1, 1\}$, 0 otherwise). The baseline category is all comments, including those with neutral or no relevant sentiment. These models estimate whether the *absolute volume* of positive and negative expression changes after the election.

4 Results

4.1 Note-level Analysis

4.1.1 Platform Differences in Sentiment

Table 4 summarizes the raw platform gradient in expressed sentiment using bivariate ordered logistic regressions. Because outcomes are coded -1 (negative), 0 (neutral), and 1 (positive), the odds ratios describe the relative odds that a Weibo note falls in a more positive sentiment category than a Rednote note. These bivariate differences are descriptive: they conflate any platform-specific dynamics with compositional differences in who posts and what content the platforms surface.

Relative to Rednote, Weibo notes have modestly higher odds of expressing pro-Trump sentiment (OR = 1.27, $p < 0.001$) and lower odds of expressing pro-Harris sentiment (OR = 0.80, $p < 0.001$). The larger cleavage appears in evaluations of the United States and its political system. Weibo notes have roughly half the odds of expressing more pro-U.S. sentiment (OR = 0.52, $p < 0.001$) and more pro-democracy sentiment (OR = 0.55, $p < 0.001$). Substantively, the platforms diverge less in candidate cheerleading than in whether the election is discussed through an anti-/pro-U.S. and anti-/pro-democracy lens. Pro-China sentiment differs only slightly (OR = 1.23, $p < 0.05$). These patterns are consistent with **H1**: Weibo's information environment, shaped by state content moderation and a predominantly domestic user base, produces substantially more negative baseline sentiment toward the United States and democracy than Rednote.

Table 5 reports analogous bivariate models for outgroup hate and misinformation. Weibo notes are more likely to contain outgroup hate (OR = 1.28, $p < 0.05$) and exhibit slightly higher misinformation scores ($\Delta = 0.02$ on a 0–1 index). Despite statistical significance,

	Trump	Harris	China	US	Democracy
<i>Coefficients</i>					
Platform (Weibo = 1)	0.237*** (0.045)	-0.224*** (0.052)	0.203* (0.097)	-0.649*** (0.053)	-0.598*** (0.061)
<i>Cut points</i>					
-1 0	-2.228 (0.045)	-2.052 (0.047)	-5.408 (0.151)	-2.369 (0.051)	-2.845 (0.060)
0 1	1.822 (0.042)	2.523 (0.053)	3.277 (0.087)	2.503 (0.052)	2.718 (0.058)
<i>N</i>	14,606	14,606	14,606	14,606	14,606
AIC	20,844.8	14,652.3	5,856.3	16,993.8	13,396.3
Residual Deviance	20,838.8	14,646.3	5,850.3	16,987.8	13,390.3

Note: Ordered logistic regression with outcomes coded $\{-1, 0, 1\}$. Positive coefficients indicate Weibo users more likely to express positive sentiment. Standard errors in parentheses. * $p < .05$; ** $p < .01$; *** $p < .001$.

Table 4: Bivariate Models: Platform Effects on Sentiment (Ordered Logit)

platform membership alone explains very little variation in misinformation ($R^2 = 0.003$), consistent with the idea that misinformation is better understood as part of a broader content bundle—a co-occurring set of frames and rhetorical styles—than as an attribute mechanically produced by platform membership.

4.1.2 Democracy as U.S. Legitimacy Talk

A key interpretive question is whether “democracy” in this discourse functions as an abstract institutional evaluation or as a proxy for country-specific affect. Table 6 models democracy views using an ordered logistic specification with post-level covariates. The dominant finding is the tight bundling of democracy evaluations with U.S. sentiment. Holding other covariates constant, a one-unit shift toward more positive U.S. sentiment is associated with roughly a 76-fold increase in the odds of expressing a more pro-democracy view ($\beta = 4.33$, $p < 0.001$). This coefficient dwarfs all other predictors, suggesting that references to “democracy” in this corpus often function as shorthand for broader affect toward the United States rather than a separable evaluation of democratic institutions.

Conditional on U.S. sentiment, several additional associations are informative. Pro-China sentiment is positively associated with pro-democracy views ($\beta = 0.55$, $p < 0.001$), indicating that pro-China notes are not uniformly anti-democratic once their stance toward the United States is held constant. Misinformation is strongly associated with more negative democracy

	Outgroup Hate (Logit)	Misinfo Score (OLS)
Intercept	−3.270*** (0.091)	0.036*** (0.002)
Platform (Weibo = 1)	0.249* (0.101)	0.020*** (0.003)
<i>N</i>	14,606	14,606
AIC / R^2	5,283.5	0.003
Null Deviance / F	5,285.8	49.75***
Residual Deviance	5,279.5	—
<i>Odds Ratio for Hate Model</i>		
Platform (Weibo = 1)	1.283 [1.05, 1.57]	—

Note: Standard errors in parentheses; 95% CI in brackets. * $p < .05$; *** $p < .001$.

Table 5: Bivariate Models: Platform Effects on Hate and Misinformation

views ($\beta = -2.06$, $p < 0.001$), consistent with institutional cynicism being concentrated in high-distortion content. Policy interest is positively associated with democracy views ($\beta = 0.65$, $p < 0.001$), suggesting that substantive political discussion frames democratic processes more favorably than slogan-driven commentary.

This bundling has measurement implications. When we observe post-election shifts in “democracy views,” we should interpret them primarily as shifts in expressed evaluations of U.S. democratic legitimacy—not as evidence that Chinese social media users are becoming more or less committed to democracy in the abstract. The country-specific nature of this sentiment is consistent with cross-border political communication operating through nation-state affect rather than ideational diffusion.

4.1.3 Post-Election Shifts

Table 7 reports platform-by-period means. Temporal change is concentrated on Weibo: average Trump sentiment shifts from approximately neutral pre-election ($M = -0.012$) to positive post-election ($M = 0.115$). Over the same window, Weibo notes become less negative toward the United States and toward democracy (U.S. sentiment from -0.175 to -0.099 ; democracy views from -0.129 to -0.046), although mean evaluations remain negative overall. Rednote is comparatively stable across periods. Misinformation declines on both platforms after the election, consistent with reduced uncertainty once the outcome is known.

	Coefficient	SE	t-value
Platform (Weibo = 1)	−0.112	0.080	−1.40
Trump sentiment	−0.099	0.062	−1.61
Harris sentiment	0.783***	0.098	8.01
China sentiment	0.554***	0.119	4.66
US sentiment	4.332***	0.072	59.86
Policy interest	0.653***	0.081	8.06
Informed	−0.050	0.058	−0.86
Misinfo score	−2.059***	0.195	−10.54
<i>Cut points</i>			
−1 0	−4.337	0.112	−38.68
0 1	4.143	0.110	37.69
<i>N</i>	14,606		
AIC	7,217.2		
Residual Deviance	7,195.2		

Note: Positive coefficients indicate more pro-democracy views. *** $p < .001$.

Table 6: Ordered Logistic Regression: Predictors of Democracy Views

		Pre-Election	Election Day	Post-Election
<i>N</i> posts	Rednote	685	333	2,424
	Weibo	1,908	209	9,047
Mean Trump	Rednote	0.035	0.045	0.045
	Weibo	−0.012	0.019	0.115
Mean Harris	Rednote	−0.038	0.009	−0.020
	Weibo	−0.058	−0.034	−0.056
Mean China	Rednote	0.023	0.030	0.035
	Weibo	0.037	0.019	0.042
Mean US	Rednote	−0.010	0.003	−0.012
	Weibo	−0.175	−0.062	−0.099
Mean Democracy	Rednote	−0.012	0.015	0.011
	Weibo	−0.129	−0.086	−0.046
Outgroup hate (%)	Rednote	4.67	3.60	3.38
	Weibo	5.24	2.87	4.57
Mean Misinfo	Rednote	0.052	0.036	0.031
	Weibo	0.075	0.051	0.051

Note: Bold indicates substantial shift. Pre-election = before Nov 5; Post-election = after Nov 5, 2024.

Table 7: Sentiment by Election Period and Platform

Table 8 formalizes these patterns with a difference-in-differences design that excludes Election Day and includes a linear time trend. The interaction term (Weibo $\times$ Post-election) captures whether sentiment shifted more on Weibo than on Rednote at the election cutoff. Weibo exhibits significantly larger post-election increases in Trump sentiment ($\beta = 0.146$, $p < 0.001$), U.S. sentiment ($\beta = 0.102$, $p < 0.001$), and democracy views ($\beta = 0.067$, $p < 0.001$), with no comparable platform-specific shift for China sentiment.

These results support **H2** (legitimacy updating), **H3** (bandwagon), and **H5** (convergence). Interpreted descriptively, the election appears to have reorganized Weibo discourse toward the winning candidate and, in parallel, toward somewhat less negative assessments of the United States and its democratic system. Critically, the convergence pattern—Weibo showing larger shifts than Rednote—suggests that baseline platform differences reflect information environments and user selection rather than immutable dispositional differences. When exposed to the same electoral shock, users on both platforms update in similar directions, with Weibo partially converging toward Rednote’s more favorable baseline. Because these estimates rely on platform-level posting streams (without user fixed effects), they may reflect both attitude updating and changes in who posts or what topics trend; the within-user comment-level analyses below provide a complementary check on timing and robustness.

	Trump	China	US	Democracy
Intercept	−0.008 (0.019)	0.026** (0.008)	−0.046** (0.017)	−0.021 (0.014)
Platform (Weibo = 1)	−0.072** (0.022)	0.016 (0.010)	−0.185*** (0.019)	−0.123*** (0.016)
Post-election	0.080*** (0.023)	0.006 (0.010)	0.058** (0.020)	0.039* (0.017)
Weibo $\times$ Post-election	0.146*** (0.025)	−0.009 (0.011)	0.102*** (0.021)	0.067*** (0.018)
Days from election	−0.009*** (0.001)	0.001 (0.001)	−0.007*** (0.001)	−0.002* (0.001)
R^2	0.013	0.001	0.017	0.013
N	14,064	14,064	14,064	14,064

Note: OLS regression. Election day observations excluded. Days from election is centered at November 5, 2024, so the intercept represents predicted sentiment for Rednote users in the pre-election period on election day. The days from election coefficient captures the linear time trend (daily rate of change in sentiment). Post-election is an indicator for posts after November 5, 2024; its coefficient captures the discontinuous shift at the election. Standard errors in parentheses. * $p < .05$; ** $p < .01$; *** $p < .001$.

Table 8: Difference-in-Differences Estimates: Post-Election Shifts by Platform

4.2 Comment-level Analysis

4.2.1 Within-Thread Sentiment Diffusion

I next examine whether sentiments expressed in posts are echoed in subsequent comments. Figure 1 reports OLS estimates from models that regress comment-level sentiment on the corresponding post-level sentiment, with hour-of-day and day-of-week fixed effects and standard errors clustered at the post level. A positive coefficient indicates greater within-thread alignment between posts and their comment sections. Because users may self-select into commenting on congenial posts and platform ranking may sort users into threads, these estimates are descriptive and should not be interpreted as a pure causal effect of exposure.

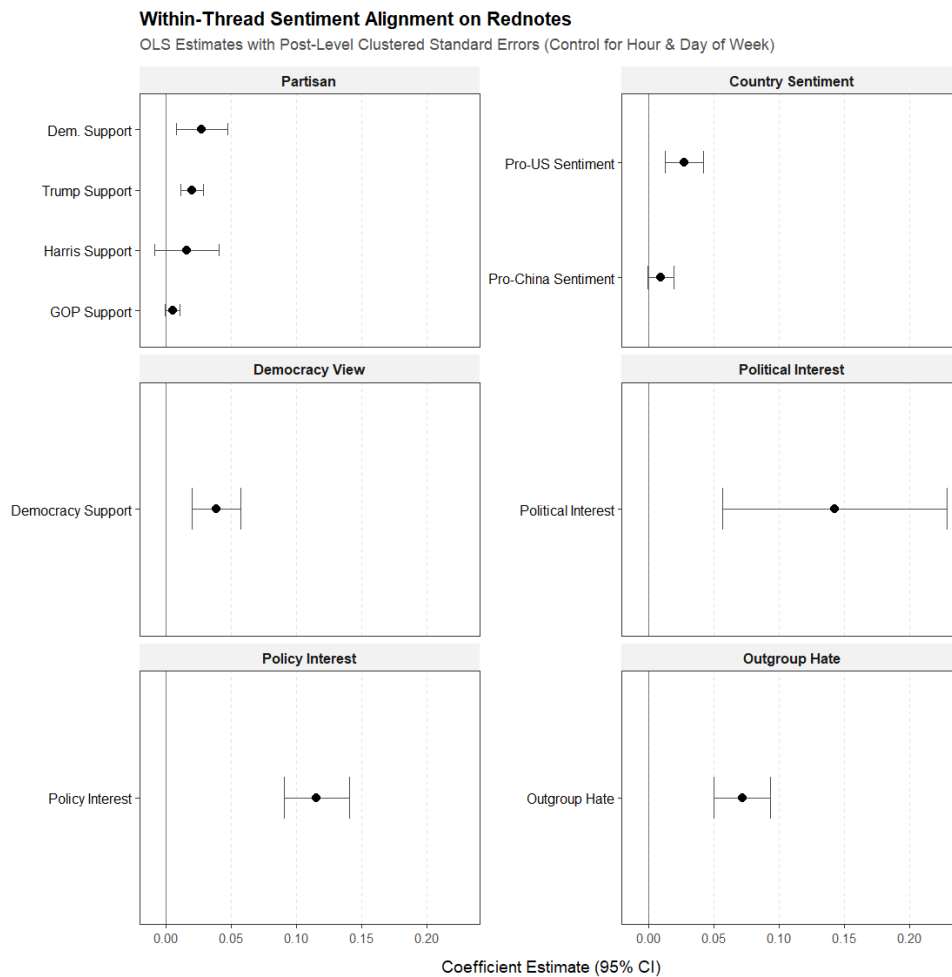


Figure 1: Within Thread Sentiment Alignment on Rednote: Full Sample

Notes: Each point represents an OLS coefficient estimate with 95% confidence intervals. The outcome variable is comment-level sentiment; the predictor is the corresponding post-level sentiment. All models include hour-of-day and day-of-week fixed effects, with standard errors clustered at the post level.

Figure 1 shows that, across outcomes, comment sentiment tends to track the tenor of the originating post on Rednote. The strength of within-thread alignment varies across the election period and across sentiment dimensions. Alignment for pro-China sentiment is stronger before the election ($\beta = 0.046, p < 0.05$) than after ($\beta = 0.007, \text{n.s.}$), consistent with nationalist frames resonating more during the period of electoral uncertainty. Pro-democracy content similarly aligns more strongly before the election ($\beta = 0.068$) than after ($\beta = 0.024$), though both estimates remain statistically significant.

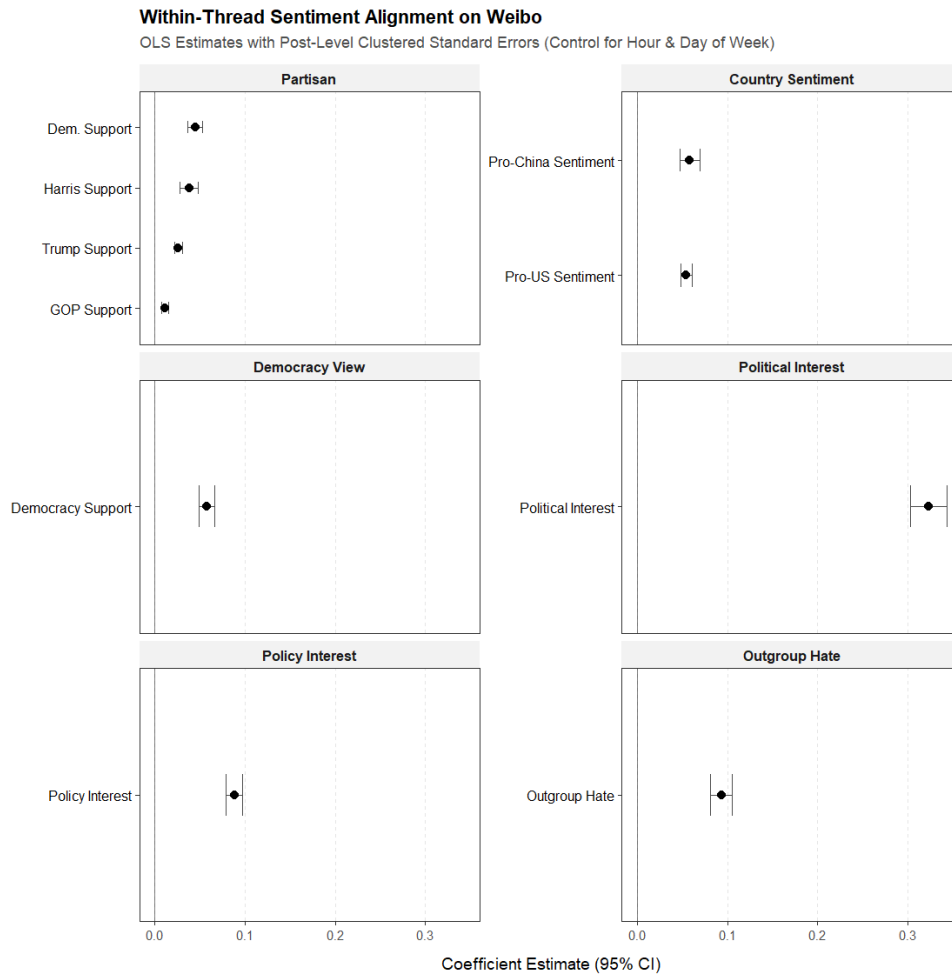


Figure 2: Within Thread Sentiment Alignment on Weibo: Full Sample

Notes: Each point represents an OLS coefficient estimate with 95% confidence intervals. The outcome variable is comment-level sentiment; the predictor is the corresponding post-level sentiment. All models include hour-of-day and day-of-week fixed effects, with standard errors clustered at the post level. “Pre” refers to comments before November 5, 2024; “Post” refers to comments on or after November 5, 2024.

Figure 2 presents within-thread alignment estimates for Weibo. Alignment magnitudes

are substantially larger on Weibo than on Rednote across nearly all outcomes. Political interest shows the strongest alignment ($\beta = 0.323$), roughly double the magnitude observed on Rednote ($\beta = 0.143$). This platform difference is consistent with Weibo's more public, broadcast-oriented architecture producing tighter coupling between posts and comment sections. On Rednote's more intimate, interest-based communities, the same post-level cues may propagate less uniformly.

Taken together, the within-thread alignment results suggest that posts function as cues shaping subsequent discussion. The platform difference in alignment magnitude implies that Weibo's architecture amplifies content diffusion relative to Rednote—a finding with implications for how quickly sentiment shifts can propagate through discourse. The temporal variation in alignment suggests that different frames resonate at different moments: nationalist and pro-democracy frames aligned more strongly during electoral uncertainty, while other forms of alignment persisted or strengthened after resolution.

4.2.2 Within-User Pre/Post Estimates: Rednote

Table 9 presents pre/post difference-in-means estimates from user fixed effects regressions, capturing within-user sentiment changes from the two weeks before the election to the two weeks after. The central finding is that users expressed more favorable evaluations of the United States and its democracy after the election, consistent with **H2**. In the full sample, pro-U.S. sentiment rose by 0.057 ($p < 0.01$) and democracy views by 0.044 ($p < 0.01$). These effects are robust to restricting the sample to pre-existing users, reducing concerns that post-election entrants drive the results. The magnitude is similar across locations. Given the note-level evidence that democracy evaluations closely track U.S. evaluations, these shifts are best read as an improvement in overall assessments of the United States and its electoral process, rather than as an independent endorsement of democratic ideals.

Partisan sentiment shifts are heterogeneous, consistent with **H3**. Trump support increased significantly among Mainland China users ($\beta = 0.035$, $p < 0.01$) but not among U.S.-based users ($\beta = -0.002$, n.s.), consistent with a bandwagon-style update being more pronounced among geographically distant observers who lack strong prior attachments. Harris support showed no significant change in any sample.

Political engagement declined sharply ($\beta = -0.086$, $p < 0.01$), consistent with **H4**: reduced “spectacle” value once electoral uncertainty resolved. China sentiment showed only

Category	Outcome	Sample: Rednote Users			
		Full Sample	Pre-Existing	Mainland China	US
Partisan Sentiment	Trump Support	0.020***	0.020***	0.035***	−0.002
		(0.007)	(0.007)	(0.010)	(0.010)
	Harris Support	[0.001]	[0.001]	[0.002]	[−0.009]
		0.001	0.001	−0.009	0.008
		(0.009)	(0.009)	(0.011)	(0.012)
		[−0.046]	[−0.046]	[−0.040]	[−0.058]
Political Interest	Pol. Engagement	−0.086***	−0.086***	−0.097***	−0.062***
		(0.009)	(0.009)	(0.012)	(0.015)
	Policy Interest	[0.508]	[0.508]	[0.475]	[0.576]
		−0.022**	−0.022**	−0.019**	−0.025
		(0.009)	(0.009)	(0.009)	(0.016)
		[0.120]	[0.120]	[0.099]	[0.180]
Regime Sentiment	China Sentiment	0.006**	0.006**	0.009*	0.003
		(0.003)	(0.003)	(0.005)	(0.004)
	US Sentiment	[0.012]	[0.012]	[0.018]	[−0.002]
		0.057***	0.057***	0.050***	0.071***
	Democracy View	(0.008)	(0.008)	(0.011)	(0.015)
		[−0.097]	[−0.097]	[−0.096]	[−0.100]
		0.044***	0.044***	0.052***	0.050***
		(0.006)	(0.006)	(0.010)	(0.013)
		[−0.078]	[−0.078]	[−0.078]	[−0.072]
<i>N</i> (user-days)		88,756	14,467	63,998	13,326
<i>N</i> (users)		76,998	9,227	57,754	9,439

Notes: OLS estimates with user fixed effects. Standard errors in parentheses, two-way clustered by user and date. Intercept (pre-period mean) in brackets. “Post” is defined as November 5, 2024 onward. “Pre-Existing” restricts to users active before election day. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table 9: Pre/Post Election Difference-in-Means: Comment-Level Sentiment on Rednote

a small increase ($\beta = 0.006$, $p < 0.05$), indicating that the election did not trigger nationalist counter-mobilization. Users updated their evaluations of the United States without revising their orientation toward China.

Event Study of User Sentiment Around 2024 Election

User fixed effects. Two-way clustered SEs. 95% CI shown. Reference = day -1.

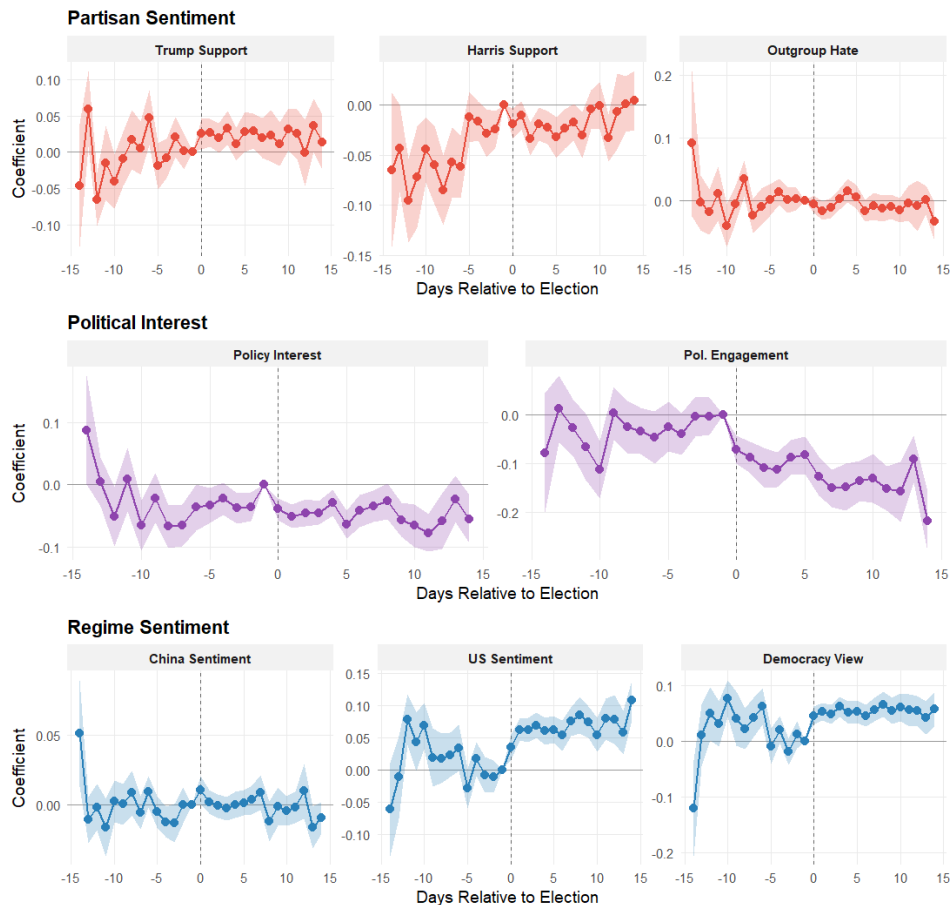


Figure 3: Event Study of User Sentiment Around the 2024 U.S. Presidential Election: Full Sample

Notes: Coefficients represent daily sentiment relative to day -1 (November 4, 2024). All specifications include user fixed effects; two-way clustered standard errors (by user and date). Shaded regions indicate 95% confidence intervals. The dashed vertical line marks Election Day (November 5, 2024). Outcomes are averaged at the user-day level before estimation.

Figure 3 traces the timing of these shifts. Pre-election coefficients fluctuate around zero with no discernible trend, suggesting that post-election shifts reflect the influence of electoral resolution rather than a continuation of pre-existing dynamics. Both pro-U.S. sentiment and democracy views jump on Election Day and rise through day +3 before stabilizing, consistent with updating tied to the resolution of the contest rather than pre-existing trends. The timing

of partisan shifts differs: Trump support increases beginning on Election Day, while Harris support reacts later (largest decline around day +2). Political engagement shows a delayed decline, falling steadily after day +1. Figure 4 suggests that U.S.-based users update more immediately than Mainland China users: pro-U.S. sentiment among U.S.-based users jumps on Election Day and plateaus, whereas Mainland China users show a more gradual rise through day +3. A similar timing gap appears for democracy views. This pattern is consistent with proximity to American media and social networks enabling faster information processing. Because these specifications include user

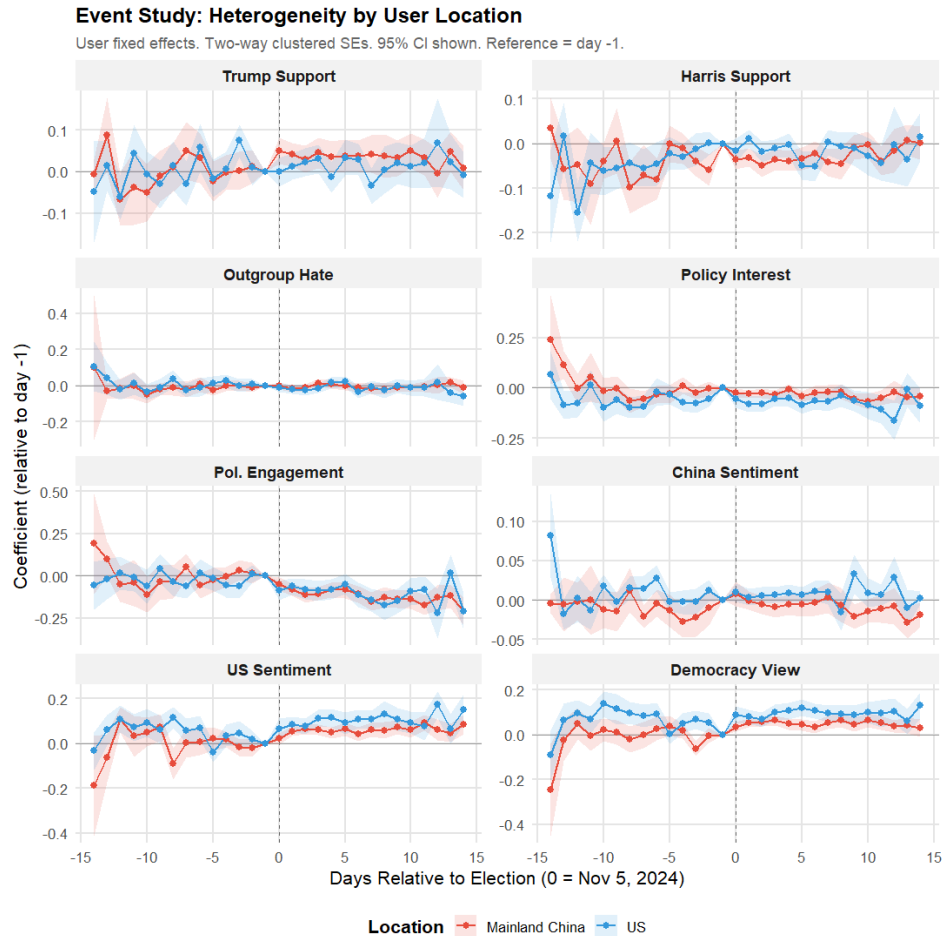


Figure 4: Event Study of User Sentiment Around the 2024 U.S. Presidential Election: Heterogeneity by User Location

Notes: Coefficients represent daily sentiment relative to day -1 (November 4, 2024). All specifications include user fixed effects; two-way clustered standard errors (by user and date). Shaded regions indicate 95% confidence intervals. The dashed vertical line marks Election Day (November 5, 2024).

Figure 4 suggests that U.S.-based users update more immediately than Mainland China users: pro-U.S. sentiment among U.S.-based users jumps on Election Day and plateaus, whereas Mainland China users show a more gradual rise through day +3. A similar timing gap appears for democracy views. This pattern is consistent with proximity to American media and social networks enabling faster information processing. Because these specifications include user

fixed effects, the coefficients capture within-user changes over time and should not be interpreted as differences in baseline sentiment levels across locations.

4.2.3 Within-User Pre/Post Estimates: Weibo

I now turn to Weibo, China’s dominant microblogging platform. Because Weibo’s user base is overwhelmingly domestic—95% of comments in the sample originate from Mainland China IP locations—I focus on overall patterns rather than location heterogeneity.

Category	Outcome	Sample: Weibo Users			
		Full Sample	Pre-Existing	Mainland China	US
Partisan Sentiment	Trump Support	0.022***	0.022***	0.020***	0.117***
		(0.006)	(0.006)	(0.006)	(0.031)
		[−0.015]	[−0.015]	[−0.014]	[−0.115]
	Harris Support	0.001	0.001	0.002	−0.031
		(0.004)	(0.004)	(0.004)	(0.019)
		[−0.031]	[−0.031]	[−0.031]	[−0.034]
Political Interest	Pol. Engagement	−0.040***	−0.040***	−0.042***	−0.007
		(0.008)	(0.008)	(0.009)	(0.032)
		[0.434]	[0.434]	[0.430]	[0.590]
	Policy Interest	0.024***	0.024***	0.023***	0.054*
		(0.004)	(0.004)	(0.003)	(0.027)
		[0.058]	[0.058]	[0.058]	[0.096]
Regime Sentiment	China Sentiment	0.002	0.002	0.002	0.005
		(0.002)	(0.002)	(0.002)	(0.008)
		[0.011]	[0.011]	[0.011]	[0.005]
	US Sentiment	0.049***	0.049***	0.051***	−0.007
		(0.005)	(0.005)	(0.005)	(0.025)
		[−0.104]	[−0.104]	[−0.104]	[−0.067]
	Democracy View	0.053***	0.053***	0.053***	0.013
		(0.005)	(0.005)	(0.005)	(0.020)
		[−0.084]	[−0.084]	[−0.085]	[−0.051]
<i>N</i> (user-days)		167,809	63,932	160,274	2,605
<i>N</i> (users)		114,988	28,885	110,258	1,572

Notes: OLS estimates with user fixed effects. Standard errors in parentheses, two-way clustered by user and date. Intercept (pre-period mean) in brackets. “Post” is defined as November 5, 2024 onward. “Pre-Existing” restricts to users active before election day. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table 10: Pre/Post Election Difference-in-Means: Comment-Level Sentiment on Weibo

Table 10 presents within-user pre/post estimates from user fixed effects regressions. As on Rednote, the central pattern is a post-election increase in pro-U.S. sentiment and democracy views: pro-U.S. sentiment rose by 0.049 ($p < 0.01$) and democracy views by 0.053 ($p < 0.01$), and the results are robust to restricting to pre-existing users. Trump support increased modestly ($\beta = 0.022$, $p < 0.01$), while Harris support showed no significant change. Political engagement declined ($\beta = -0.040$, $p < 0.01$).

A key divergence from Rednote is that policy interest increased on Weibo ($\beta = 0.024$, $p < 0.01$), suggesting that Weibo users maintained substantive attention to the policy implications of the result even after the electoral drama concluded. China sentiment showed no significant change ($\beta = 0.002$, n.s.), again indicating little evidence of nationalist counter-mobilization in response to the election. Consistent with H5, Weibo users updated in the same direction as Rednote users—toward less negative U.S. and democracy sentiment—despite starting from a more hostile baseline.

Figure 5 presents the event study coefficients. Pro-U.S. sentiment and democracy views rise sharply beginning on Election Day and remain elevated through day +14, with no clear pre-election trends. Trump support shows a brief positive spike around day +1 before returning closer to baseline by day +3, consistent with a transient winner-oriented response rather than sustained candidate enthusiasm. One possible interpretation is that Weibo discourse quickly pivots from candidate evaluations to broader commentary on U.S. politics and governance, though identifying the mechanisms would require additional evidence on information exposure.

4.2.4 Dampening of Negative Discourse

The within-user fixed effects models show that evaluations of the United States and democracy became less negative after the election. But do these shifts reflect *mobilization* of positive sentiment or *dampening* of negative sentiment? To distinguish these mechanisms, I examine both conditional models—which subset to democracy-relevant or U.S.-relevant comments only—and unconditional models that predict whether any comment expresses relevant sentiment relative to all comments.

Table 11 presents the results for democracy sentiment. The conditional analysis (Models 1–2) shows that among comments expressing any democracy view, the odds of that view being anti-democracy decreased by approximately 31% after the election (Model 1: OR = 0.69). This

Event Study of User Sentiment on Weibo Around 2024 Election

User fixed effects. Two-way clustered SEs. 95% CI shown. Reference = day -1.

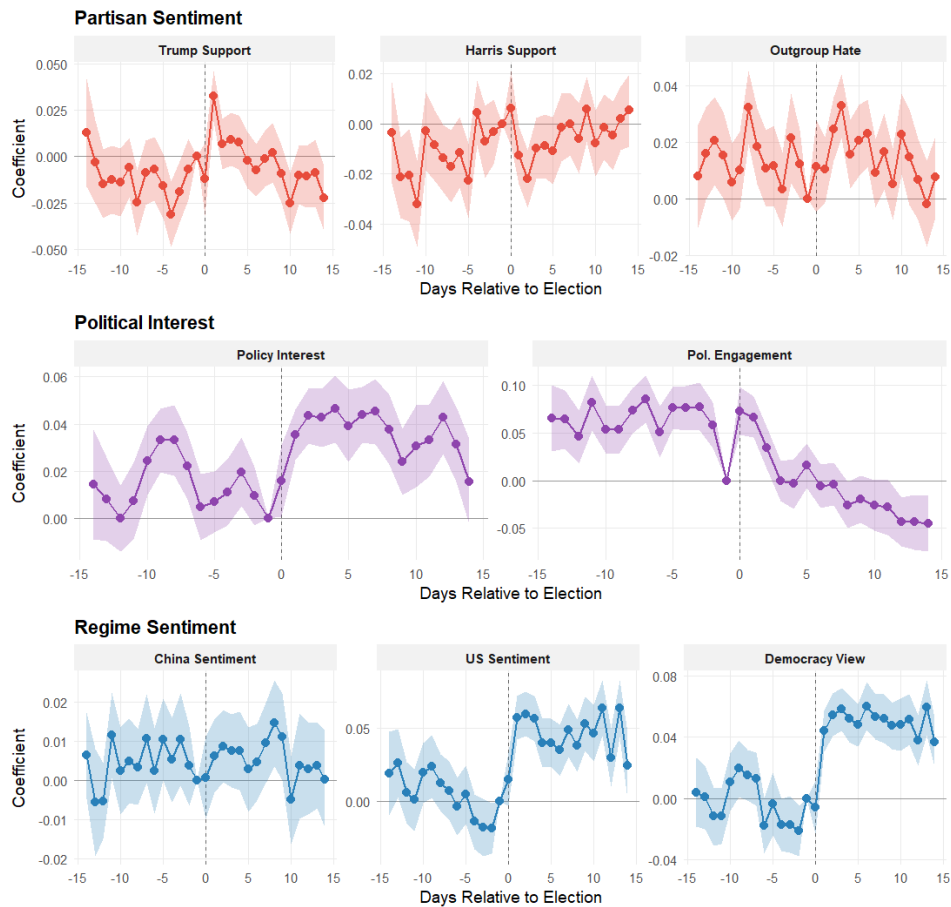


Figure 5: Event Study of User Sentiment Around the 2024 U.S. Presidential Election: Full Sample

Notes: Coefficients represent daily sentiment relative to day -1 (November 4, 2024). All specifications include user fixed effects; two-way clustered standard errors (by user and date). Shaded regions indicate 95% confidence intervals. The dashed vertical line marks Election Day (November 5, 2024). Outcomes are averaged at the user-day level before estimation.

shift is robust to controlling for platform (Model 2), and Weibo comments exhibit substantially higher baseline anti-democracy sentiment than Rednote comments ($\beta = 0.52$, $p < 0.001$). However, the conditional analysis alone cannot distinguish whether the ratio shift reflects (a) genuine attitude change among those discussing democracy, (b) differential exit from the conversation (anti-democracy voices withdrawing more than pro-democracy voices), or (c) differential entry (pro-democracy voices joining more than anti-democracy voices).

The unconditional analysis (Models 3–5) clarifies the mechanism. Both anti-democracy expression (Model 3: OR = 0.42) and pro-democracy expression (Model 4: OR = 0.59) declined in absolute terms after the election, but anti-democracy rhetoric fell more sharply. Overall democracy-relevant discussion contracted by over 50% (Model 5: OR = 0.43). This pattern is consistent with dampening rather than mobilization: the election did not bring new pro-democracy voices into the conversation; rather, it reduced the volume of democracy-related discourse overall, with anti-democracy content declining disproportionately.

	Conditional		Unconditional (All Comments)		
	Anti vs. Pro (1)	Anti vs. Pro (2)	Anti (3)	Pro (4)	Any (5)
Post-Election	−0.374*** (0.047) [0.688]	−0.342*** (0.048) [0.710]	−0.875*** (0.017) [0.417]	−0.520*** (0.045) [0.594]	−0.841*** (0.016) [0.431]
Weibo		0.523*** (0.043) [1.687]	0.319*** (0.017) [1.375]	−0.203*** (0.039) [0.817]	0.241*** (0.016) [1.272]
Intercept	2.106*** (0.041)	1.737*** (0.050)	−2.585*** (0.020)	−4.377*** (0.048)	−2.397*** (0.018)
N	20,055	20,055	379,725	379,725	379,725

Table 11: Logistic Regression: Democracy Sentiment

Note: Standard errors in parentheses; odds ratios in brackets. Models 1–2: DV = anti-democracy (1) vs. pro-democracy (0) among democracy-relevant comments. Models 3–5: DV = respective sentiment (1) vs. all other comments (0). *** $p < 0.001$.

The pattern for U.S. sentiment is parallel (Table 12). The conditional analysis shows that among U.S.-relevant comments, the odds of anti-U.S. sentiment decreased by approximately 34% after the election (Model 1: OR = 0.66). Notably, the platform interaction is statistically significant (Model 2: $\beta = 0.261$, $p < 0.01$): the post-election decline in anti-U.S. sentiment is stronger on Rednote than on Weibo. Among Rednote users, anti-U.S. sentiment dropped

from 88.9% to 82.0% of U.S.-relevant comments (6.9 percentage points); among Weibo users, the decline was smaller, from 91.3% to 88.6% (2.7 percentage points). This asymmetry may reflect Rednote’s diaspora user base, who have direct experience with the United States and may be more responsive to signals that American institutions function effectively.

The unconditional analysis again reveals dampening: anti-U.S. expression declined by 50% (Model 4: OR = 0.50) while pro-U.S. expression declined by only 24% (Model 5: OR = 0.76). Overall U.S.-relevant discussion dropped by 48% (Model 6: OR = 0.52).

	Conditional			Unconditional (All Comments)		
	Anti vs. Pro (1)	Anti vs. Pro (2)	Anti vs. Pro (3)	Anti (4)	Pro (5)	Any (6)
Post-Election	−0.421*** (0.045) [0.656]	−0.563*** (0.080) [0.569]	−0.388*** (0.045) [0.678]	−0.690*** (0.015) [0.502]	−0.278*** (0.043) [0.757]	−0.653*** (0.014) [0.520]
Weibo		0.272** (0.087) [1.312]	0.483*** (0.038) [1.621]	0.309*** (0.015) [1.362]	−0.177*** (0.035) [0.838]	0.245*** (0.014) [1.277]
Post × Weibo		0.261** (0.097) [1.298]				
Intercept	2.276*** (0.040)	2.078*** (0.073)	1.933*** (0.047)	−2.346*** (0.017)	−4.356*** (0.046)	−2.184*** (0.016)
N	27,799	27,799	27,799	379,725	379,725	379,725

Table 12: Logistic Regression: U.S. Sentiment

Note: Standard errors in parentheses; odds ratios in brackets. Models 1–3: DV = anti-U.S. (1) vs. pro-U.S. (0) among U.S.-relevant comments. Models 4–6: DV = respective sentiment (1) vs. all other comments (0). ** $p < 0.01$; *** $p < 0.001$.

This asymmetric dampening suggests that pre-election criticism of American democracy served partly as spectacle—narratives of dysfunction and chaos that entertained audiences during uncertainty but lost salience once the contest concluded decisively. Peaceful electoral resolution does not convert skeptics into democracy supporters; rather, it reduces the supply of negative content that thrives on uncertainty. The mechanism is dampening of criticism rather than generation of enthusiasm.

All five hypotheses receive support from the data. Weibo discourse exhibits substantially more negative baseline sentiment toward the U.S. and democracy than Rednote (**H1**), but eval-

uations on both platforms shift in a less negative direction after the election (**H2**), with Weibo showing larger shifts that partially converge toward Rednote’s baseline (**H5**). Trump support increases after victory, especially among mainland users (**H3**), while political engagement declines once uncertainty resolves (**H4**). China sentiment remains flat, indicating that updating foreign-country evaluations does not require revising domestic attachments.

5 Conclusion

This paper examined how Chinese-speaking social media users responded to the 2024 U.S. presidential election, using over 390,000 posts and comments from Rednote and Weibo. Four main implications emerge.

First, information environments shape baseline priors, but peaceful election resolution produces convergent updating. At baseline, Weibo discourse is substantially more anti-American and anti-democracy than Rednote, consistent with prolonged exposure to state-curated narratives. Yet when the contest resolves peacefully, evaluations shift: difference-in-differences estimates show Weibo converging toward Rednote’s more favorable baseline, with significantly larger post-election increases in pro-U.S. sentiment and democracy views. This convergence pattern suggests that authoritarian narrative environments establish starting points but do not fully determine endpoints—the platform with the most negative baseline also exhibits the largest post-election correction.

Second, “democracy” talk is largely U.S. legitimacy talk in this discourse. Post-level models reveal that U.S. sentiment overwhelmingly predicts democracy views, indicating that institutional language often functions as a symbolic label attached to a specific country rather than as a separable abstract commitment. This has important measurement implications for studying authoritarian publics: apparent changes in “support for democracy” may often reflect changes in perceived legitimacy of a foreign power rather than generalized democratic diffusion. Researchers should be cautious about interpreting expressive data on “democracy” as evidence of ideational conversion when it may instead capture country-specific affect.

Third, sentiment expressed in posts diffuses to comment sections. Within-thread alignment is substantial across outcomes, with comments echoing the tenor of originating posts. Alignment magnitudes are larger on Weibo than on Rednote, consistent with Weibo’s more public, broadcast-oriented architecture producing tighter coupling between posts and replies. This diffusion pattern suggests that posts function as cues shaping subsequent discussion, and

that platform structure moderates how effectively such cues propagate.

Fourth, electoral spectatorship involves three distinct post-election dynamics. The first is a bandwagon effect: Trump support rises after victory, especially among mainland Chinese users, consistent with distant observers updating toward winners. The second is dampening of negative sentiment: both anti-U.S. and anti-democracy expression decline after the election, producing less hostile evaluations even as overall discussion contracts. This suggests that pre-election criticism served partly as spectacle—entertaining audiences with narratives of dysfunction—that lost salience once the contest concluded decisively. Peaceful resolution reduces the supply of negative content rather than converting skeptics into supporters. The third dynamic is attention withdrawal: political engagement declines substantially once uncertainty resolves, consistent with elections functioning partly as entertainment whose value dissipates after resolution. These dynamics coexist: evaluative updating persists even as discussion becomes less frequent. Yet the election does not produce nationalist counter-mobilization: China sentiment remains stable across platforms and designs, suggesting that softening criticism of a foreign democracy does not require revising domestic attachments.

The findings contribute to several literatures. For research on authoritarian public opinion, they demonstrate that state-curated information environments shape baseline priors but do not fully insulate audiences from updating when foreign democratic institutions visibly function. The dampening mechanism—whereby peaceful resolution reduces negative spectacle rather than generating positive conversion—offers a more precise account of how democratic performance might influence foreign publics than simple “demonstration effects” frameworks suggest. For research on cross-border political communication, the within-thread alignment findings highlight how post-level framing can shape comment-level discourse, with platform architecture moderating diffusion. For platform studies, the convergence pattern underscores that baseline differences across information environments need not persist: the same observable event can produce larger updating precisely where priors are most negative.

Several limitations warrant emphasis. The analysis captures expressed discourse among a politically active subset of Chinese-speaking users rather than representative mass opinion or privately held beliefs. Expression may be strategic or performative, shaped by platform norms, perceived audiences, and censorship risk. Key outcomes mix shifts in valence with shifts in salience: the finding that democracy discussion declined overall while becoming proportionally less negative illustrates this complexity. The observational design limits causal inference;

although event-study timing and stable pre-trends provide reassurance, I cannot fully rule out contemporaneous confounds. Separating mechanisms—genuine belief updating, expressive conformity, reduced entertainment value of criticism—would benefit from experimental designs manipulating exposure to election-resolution information.

Future research can extend these findings in several directions. Comparative work across elections would clarify whether dampening operates similarly when contests resolve contentiously rather than decisively, and whether the democracy-country bundling is U.S.-specific or general. Longer time horizons could test whether reduced negative discourse represents durable change or temporary exhaustion that rebounds as new controversies emerge. Individual-level measures of policy exposure, visa status, or travel history could help distinguish stakeholder from spectator responses more precisely. Finally, combining social media data with surveys would help distinguish changes in public expression from changes in underlying belief.

More broadly, foreign elections create brief but consequential windows for cross-border opinion dynamics. When democratic institutions resolve uncertainty peacefully, they can dampen criticism among skeptical audiences—but this operates by reducing negative spectacle rather than generating positive conversion, and it coexists with rapid attention decline once the narrative closes. For democracies concerned with international legitimacy, the implication is cautiously optimistic: functioning institutions can reduce hostile discourse even in unfavorable information environments, though such effects may prove shallow if they depend more on spectacle fatigue than on genuine persuasion.

References

- Bail, Christopher A., Lisa P. Argyle, Taylor W. Brown, John P. Bumpus, Haohan Chen, M. B. Fallin Hunzaker, Jaemin Lee, Marcus Mann, Friedolin Merhout, and Alexander Volfovsky.** 2018. “Exposure to Opposing Views on Social Media Can Increase Political Polarization.” *Proceedings of the National Academy of Sciences*, 115(37): 9216–9221.
- Bartels, Larry M.** 1988. *Presidential Primaries and the Dynamics of Public Choice*. Princeton, NJ: Princeton University Press.
- Bischoff, Ivo, and Henrik Egbert.** 2013. “Social Information and Bandwagon Behavior in Voting: An Economic Experiment.” *Journal of Economic Psychology*, 34: 270–284.
- Brady, William J., Julian A. Wills, John T. Jost, Joshua A. Tucker, and Jay J. Van Bavel.** 2017. “Emotion Shapes the Diffusion of Moralized Content in Social Networks.” *Proceedings of the National Academy of Sciences*, 114(28): 7313–7318.

- Brinks, Daniel, and Michael Coppedge.** 2006. "Diffusion Is No Illusion: Neighbor Emulation in the Third Wave of Democracy." *Comparative Political Studies*, 39(4): 463–489.
- Crabtree, Charles, David Darmofal, and Holger L. Kern.** 2015. "A Spatial Analysis of the Impact of West German Television on Public Opinion in East Germany." *Political Science Research and Methods*, 3(3): 477–498.
- Fang, Songying, Alastair Iain Johnston, and Xiaojun Shen.** 2022. "Chinese Public Opinion about US–China Relations from Trump to Biden." *The Chinese Journal of International Politics*, 15(1): 27–50.
- Gleditsch, Kristian Skrede, and Michael D. Ward.** 2006. "Diffusion and the International Context of Democratization." *International Organization*, 60(4): 911–933.
- Grimmer, Justin, and Brandon M. Stewart.** 2013. "Text as Data: The Promise and Pitfalls of Automatic Content Analysis Methods for Political Texts." *Political Analysis*, 21(3): 267–297.
- Hopkins, Daniel J., and Gary King.** 2010. "A Method of Automated Nonparametric Content Analysis for Social Science." *American Journal of Political Science*, 54(1): 229–247.
- Huang, Haifeng, Chanita Intawan, and Stephen P. Nicholson.** 2023. "In Government We Trust: Implicit Political Trust and Regime Support in China." *Perspectives on Politics*, 21(4): 1357–1375.
- Huntington, Samuel P.** 1991. *The Third Wave: Democratization in the Late Twentieth Century*. University of Oklahoma Press.
- Kern, Holger L., and Jens Hainmueller.** 2009. "Opium for the Masses: How Foreign Media Can Stabilize Authoritarian Regimes." *Political Analysis*, 17(4): 377–399.
- King, Gary, Jennifer Pan, and Margaret E. Roberts.** 2013. "How Censorship in China Allows Government Criticism but Silences Collective Expression." *American Political Science Review*, 107(2): 326–343.
- King, Gary, Jennifer Pan, and Margaret E. Roberts.** 2017. "How the Chinese Government Fabricates Social Media Posts for Strategic Distraction, Not Engaged Argument." *American Political Science Review*, 111(3): 484–501.
- Kramer, Adam D. I., Jamie E. Guillory, and Jeffrey T. Hancock.** 2014. "Experimental Evidence of Massive-Scale Emotional Contagion through Social Networks." *Proceedings of the National Academy of Sciences*, 111(24): 8788–8790.
- Landis, J Richard, and Gary G Koch.** 1977. "The measurement of observer agreement for categorical data." *Biometrics*, 33(1): 159–174.
- Levi, Margaret.** 1997. *Consent, Dissent, and Patriotism*. New York:Cambridge University Press.

- Lu, Yingdan, Jennifer Pan, Xu Xu, and Yiqing Xu.** 2025. “Decentralized Propaganda in the Era of Digital Media: The Massive Presence of the Chinese State on Douyin.” *American Journal of Political Science*, 1–17.
- Mehrabian, Albert.** 1998. “Effects of Poll Reports on Voter Preferences.” *Journal of Applied Social Psychology*, 28(23): 2119–2130.
- Mukerjee, Subhayan, Kokil Jaidka, and Yphtach Lelkes.** 2022. “The Political Landscape of the US Twittersverse.” *Political Communication*, 1–31.
- Mutz, Diana C.** 1997. “Mechanisms of Momentum: Does Thinking Make It So?” *The Journal of Politics*, 59(1): 104–125.
- Prior, Markus.** 2007. *Post-Broadcast Democracy: How Media Choice Increases Inequality in Political Involvement and Polarizes Elections*. New York:Cambridge University Press.
- Qin, Bei, David Strömberg, and Yanhui Wu.** 2024. “Social Media and Collective Action in China.” *Econometrica*, 92(6): 1993–2026.
- Repnikova, Maria.** 2017. *Media Politics in China: Improvising Power under Authoritarianism*. Cambridge:Cambridge University Press.
- Repnikova, Maria.** 2022. *Chinese Soft Power. Elements in Global China*, Cambridge University Press.
- Roberts, Margaret E.** 2018. *Censored: Distraction and Diversion Inside China’s Great Firewall*. Princeton University Press.
- Ruths, Derek, and Jürgen Pfeffer.** 2014. “Social Media for Large Studies of Behavior.” *Science*, 346(6213): 1063–1064.
- Sunstein, Cass R.** 2017. *#Republic: Divided Democracy in the Age of Social Media*. Princeton:Princeton University Press.
- Tufekci, Zeynep.** 2014. “Big Questions for Social Media Big Data: Representativeness, Validity and Other Methodological Pitfalls.” 505–514. AAAI Press.
- Tyler, Tom R.** 2006. *Why People Obey the Law*. . 2nd ed., Princeton, NJ:Princeton University Press. Originally published 1990.
- Weyland, Kurt.** 2014. *Making Waves: Democratic Contention in Europe and Latin America since the Revolutions of 1848*. Cambridge University Press.

Appendices

Table of Contents

A	Ethical Statement	1
B	Pre-registration	1
C	Labeling Accuracy	4
D	Post-level Analysis	6
D.1	UMAP Embedding Projections	6
D.2	Keyword Frequency Shifts	10
D.3	Regression Analysis	13
	Platform Baseline Difference	13
	Multivariate content models.	14
	Engagement Models	16
E	Comment-level Analysis	20
E.1	Comment-Level UMAP Projections	20
E.2	Red Note	21
E.3	Weibo	25

A Ethical Statement

Data collection of this study on a larger scale within and beyond the Chinese-speaking population and a companion survey experimental design (which is not yet a part of this paper) have been approved by the Institutional Review Board at Columbia University (Protocol Number: IRB-AAAV9204).

This study analyzes only secondary data collected from public social media comments and posts and therefore does not involve any interactions with human subjects. Comments and posts from individuals whose privacy settings have been set to “private” cannot be retrieved by the data collection techniques employed in this study. I did not collect data that could possibly contain private identifying information (e.g., profile photos) throughout the course of this study. User handles have been completely removed from the datasets, and each user has been assigned a randomly generated unique identifier so that their presence in the dataset can never be linked back even to their online presence. In this analysis, user locations are aggregated and analyzed at the country level. Only the content of comments and posts, after having been fully de-associated from user handles, are input into LLMs.

B Pre-registration

LLM labeling code and design of this study have been pre-registered at AsPredicted.org (AsPredicted #262,912) before LLM labeling is conducted.

Electoral spectators (#262912)

Author(s)

This pre-registration is currently anonymous to enable blind peer-review.
It has one author.

Pre-registered on:

2025/12/08 08:45 (PT)

1) Have any data been collected for this study already?

It's complicated. We have already collected some data but explain in Question 8 why readers may consider this a valid pre-registration nevertheless.

2) What's the main question being asked or hypothesis being tested in this study?

The hypotheses concern how Chinese social media users' online sentiments change before and after US elections. I pre-register the following bi-directional hypotheses:

H1: If Chinese citizens interpret Trump's victory and the peaceful election resolution as evidence of democratic functionality, pro-US sentiment and democracy views should increase after the election. If instead they interpret the outcome through the lens of state narratives emphasizing democratic dysfunction, these attitudes should decrease or remain unchanged.

H2: If bandwagon effects operate for foreign elections, Trump support should increase after his victory. This effect should be stronger among Mainland China users---who observe from a distance with weaker priors---than among US-based users with more crystallized partisan attachments.

H3: If the election functions as political ``spectacle," engagement should decline once uncertainty resolves. The entertainment value of following the election diminishes after the outcome is known, even if substantive interest in policy implications persists.

H4: If exposure to American democratic events activates nationalist counter-mobilization, China sentiment should increase after the election as users express defensive national pride. If cross-border political information does not trigger such responses, China sentiment should remain stable.

H5: If platform differences reflect user selection rather than platform-specific socialization, Weibo and Rednote users should update in similar directions---even if from different baselines---in response to the common shock of the election outcome.

3) Describe the key dependent variable(s) specifying how they will be measured.

Dependent variables contain the following LLM-labeled outcomes: Trump attitude, Harris attitude, attitude toward the Republican Party, attitude toward the Democratic Party, attitude toward China, attitude toward the United States, views of Democracy, hate toward political outgroups, policy interests, whether the content is political, whether the content shows informedness, and to what extent the content is considered misinformation.

I will be using LLM to label the sentiment of each post and comment. The labeling model will be GPT-4.1. Due to the word limit, I am providing most essential parts of the labeling code.

You are classifying Chinese-language social media posts about the 2024 US Presidential Election (Trump vs Harris). Output ONLY valid JSON with the following fields:

```
{
  "trump": X,
  "harris": X,
  "republican": X,
  "democrat": X,
  "outgroup_hate": X,
  "policy_interest": X,
  "china": X,
  "us": X,
  "democracy_view": X,
  "informed": X,
  "political": X,
  "misinfo": X
}
```

ATTITUDE VARIABLES (1 = positive/supportive, 0 = neutral/not mentioned, -1 = negative/critical):

1. trump: Attitude toward Donald Trump
2. harris: Attitude toward Kamala Harris (or Biden administration before she became nominee)
3. republican: Attitude toward Republican Party (beyond just Trump)
4. democrat: Attitude toward Democratic Party (beyond just Harris)
5. china: Attitude toward China expressed in the post
6. us: Attitude toward the United States as a country/system
7. democracy_view: Attitude toward democracy/electoral system

2

BINARY VARIABLES (1 = present, 0 = absent):

Version of AsPredicted Questions: 2.00
PDF generated: 2025/12/08 08:51 (PT)

Available at <https://aspredicted.org/6fn2u5.pdf>

1/2

- 8. outgroup_hate: Contains hostile, contemptuous, or mocking language toward a political outgroup
- 9. policy_interest: Mentions specific policy issues the author cares about
- 10. political: Is the content politically relevant?

ORDINAL VARIABLE (0 = low, 1 = moderate, 2 = high):

- 11. informed: How politically informed does the content suggest the author is?
- 12. misinfo: Factual accuracy (0 = accurate, 0.25 = minor issues, 0.5 = mixed/misleading, 0.75 = mostly false, 1 = fabricated)

4) How many and which conditions will participants be assigned to?

This study is an observational study that does not involve random assignment of participants and treatment/control conditions. Social media platforms I analyze contain Weibo and Rednote (Xiaohongshu).

5) Specify exactly which analyses you will conduct to examine the main question/hypothesis.

I would conduct analyses of posts and comments I collected from both platforms.

Post analyses involve: 1) how each sentiments of a post are correlated with each other, and what sentiments are more likely to predict other sentiments; 2) what sentiments are more likely to predict more engagement; 3) whether there is a platform baseline difference of sentiment, and whether the election resolution leads to a temporal shift of sentiment, and whether this shift is moderated by platforms.

Comment analyses involve: 1) post-comment diffusion analysis, which is essentially regressing comment sentiment on post sentiment. The standard error is clustered at post level. 2) User-level average sentiment change. This involves a diff-in-means, comparing average sentiment of users before and after election resolution, and a event-study style daily sentiment change. The robust standard error is clustered at user-day level. The analyses will be conducted for both platforms.

6) Describe exactly how outliers will be defined and handled, and your precise rule(s) for excluding observations.

There will be no outliers in this study.

7) How many observations will be collected or what will determine sample size? No need to justify decision, but be precise about exactly how the number will be determined.

On Rednote, there are 3,442 posts and 138,246 comments. On Weibo, there are 11,164 posts and 241,479 comments.

8) Anything else you would like to pre-register? (e.g., secondary analyses, variables collected for exploratory purposes, unusual analyses planned?)

The data was collected during the 2024 US presidential election window, because otherwise these posts and comments might be censored or deleted voluntarily by users. LLM labeling will be performed after this pre-registration.

C Labeling Accuracy

I evaluate the accuracy of LLM-based content classification by comparing labels generated by GPT-4.1 against expert-coded labels for 200 randomly sampled posts and comments from each platform. Tables A1–A3 report accuracy, Matthews Correlation Coefficient (MCC), and Cohen’s Kappa for each outcome variable. MCC provides a balanced measure of classification quality that is robust to class imbalance, while Cohen’s Kappa accounts for chance agreement between raters.

Dataset	Variable	Accuracy	MCC	κ
Rednote Comments	Outgroup hate	0.985	0.863	0.862
	Policy interest	1.000	1.000	1.000
	Political	0.995	0.990	0.990
Rednote Posts	Outgroup hate	0.995	0.923	0.921
	Policy interest	0.995	0.989	0.989
	Political	1.000	1.000	1.000
Weibo Comments	Outgroup hate	0.965	0.780	0.756
	Policy interest	1.000	1.000	1.000
	Political	1.000	1.000	1.000
Weibo Posts	Outgroup hate	0.995	0.864	0.855
	Policy interest	1.000	1.000	1.000
	Political	1.000	1.000	1.000

Note: $N = 200$ for each dataset. Variables coded as $\{0, 1\}$. MCC = Matthews Correlation Coefficient; κ = Cohen’s Kappa (unweighted).

Table A1: LLM Labeling Accuracy: Binary Outcomes

I assess LLM–expert agreement using Matthews Correlation Coefficient (MCC) and Cohen’s Kappa (κ), both of which account for chance agreement and provide robust measures even under class imbalance. For ordinal variables, I report quadratic weighted Kappa (κ_w), which penalizes larger disagreements more heavily than adjacent-category errors. Following Landis and Koch (1977), Kappa values above 0.80 indicate “almost perfect” agreement, 0.61–0.80 “substantial” agreement, and 0.41–0.60 “moderate” agreement.

Classification of binary variables achieves strong to near-perfect agreement. *Political* and *policy interest* show perfect or near-perfect reliability across all datasets (MCC and $\kappa \geq 0.989$). *Outgroup hate* demonstrates substantial to almost perfect agreement ($\kappa = 0.756$ – 0.921), with slightly lower performance on Weibo comments, likely reflecting platform-specific linguistic conventions in expressing hostility.

Sentiment toward political figures shows consistently high reliability. Trump sentiment achieves almost perfect agreement across platforms ($\kappa = 0.874$ – 0.956), as does Harris sentiment ($\kappa = 0.890$ – 1.000). Party-level sentiment (Republican and Democrat) similarly demonstrates substantial to almost perfect agreement ($\kappa = 0.665$ – 1.000). China sentiment shows

Dataset	Variable	Accuracy	MCC	κ
Rednote Comments	Trump sentiment	0.980	0.893	0.887
	Harris sentiment	1.000	1.000	1.000
	Republican sentiment	1.000	1.000	1.000
	Democrat sentiment	0.990	0.926	0.923
	China sentiment	1.000	1.000	1.000
	US sentiment	0.980	0.808	0.790
	Democracy view	0.995	0.864	0.855
Rednote Posts	Trump sentiment	0.970	0.909	0.907
	Harris sentiment	0.990	0.946	0.946
	Republican sentiment	0.990	0.934	0.933
	Democrat sentiment	0.980	0.861	0.859
	China sentiment	0.995	0.911	0.907
	US sentiment	0.960	0.785	0.762
	Democracy view	0.985	0.872	0.864
Weibo Comments	Trump sentiment	0.990	0.957	0.956
	Harris sentiment	0.985	0.891	0.890
	Republican sentiment	0.995	0.706	0.665
	Democrat sentiment	0.990	0.853	0.852
	China sentiment	0.995	0.924	0.921
	US sentiment	1.000	1.000	1.000
	Democracy view	0.980	0.836	0.836
Weibo Posts	Trump sentiment	0.960	0.878	0.874
	Harris sentiment	0.985	0.896	0.890
	Republican sentiment	0.990	0.842	0.829
	Democrat sentiment	0.990	0.927	0.924
	China sentiment	0.995	0.924	0.921
	US sentiment	0.945	0.812	0.795
	Democracy view	0.980	0.909	0.908

Note: $N = 200$ for each dataset. Variables coded as $\{-1, 0, 1\}$ representing negative, neutral, and positive sentiment. MCC = Matthews Correlation Coefficient; κ = Cohen's Kappa (unweighted).

Table A2: LLM Labeling Accuracy: Ternary Sentiment Outcomes

Dataset	Variable	Accuracy	MCC	κ_w
Rednote Comments	Informed	0.995	0.988	0.992
	Misinformation	0.990	0.909	0.878
Rednote Posts	Informed	0.985	0.978	0.989
	Misinformation	0.990	0.939	0.965
Weibo Comments	Informed	0.985	0.956	0.939
	Misinformation	0.995	0.941	0.901
Weibo Posts	Informed	0.950	0.918	0.943
	Misinformation	0.990	0.963	0.981

Note: $N = 200$ for each dataset. *Informed* coded as $\{0, 1, 2\}$; *Misinformation* coded as $\{0, 0.25, 0.5, 0.75, 1\}$. MCC = Matthews Correlation Coefficient; κ_w = Cohen’s Kappa with quadratic weights.

Table A3: LLM Labeling Accuracy: Ordinal Outcomes

almost perfect agreement ($\kappa = 0.907$ – 1.000), indicating clear linguistic markers when users express views toward China. US sentiment achieves substantial to almost perfect agreement ($\kappa = 0.762$ – 1.000), with Weibo comments showing perfect agreement and other datasets ranging from 0.762 to 0.812. Democracy view similarly demonstrates almost perfect agreement across all datasets ($\kappa = 0.836$ – 0.908).

Both ordinal variables achieve strong reliability. *Informed*—measuring the degree to which content reflects substantive political knowledge—shows almost perfect weighted agreement ($\kappa_w = 0.939$ – 0.992). *Misinformation* also achieves almost perfect agreement ($\kappa_w = 0.878$ – 0.981), a notable result given the five-level ordinal scale and the inherent difficulty of distinguishing degrees of misinformation even among human coders. The high MCC values (0.909–0.963) for misinformation further confirm that disagreements, when they occur, do not represent systematic bias.

These validation results strongly support the use of LLM-based classification for my main analyses. All outcome variables achieve almost perfect agreement ($\kappa > 0.80$). For variables where measurement error exists, such error would attenuate rather than inflate estimated effects, providing a conservative test of my hypotheses.

D Post-level Analysis

D.1 UMAP Embedding Projections

To examine whether political discourse about the U.S. election differs systematically across platforms and time periods, I project text embeddings into a two-dimensional space using Uniform Manifold Approximation and Projection (UMAP). I embed all political posts using OpenAI’s `text-embedding-3-small` model, then reduce dimensionality with UMAP (`n_neighbors=20`, `min_dist=0.05`, cosine metric). To quantify separability beyond visual inspection, I train logis-

tic regression classifiers on the full-dimensional embeddings using 5-fold cross-validation.

Table A4 reports cross-validated classification performance for distinguishing (1) platform (Rednote vs. Weibo) and (2) period (pre- vs. post-election).

Classification Task	AUC	Accuracy
Platform (Rednote vs. Weibo)	0.898 (0.009)	0.826 (0.010)
Period (Pre vs. Post)	0.815 (0.013)	0.766 (0.005)

Table A4: Embedding-Based Separability of Platform and Period

Note: Standard deviations from 5-fold cross-validation in parentheses. Classifiers are logistic regression models trained on 1,536-dimensional embeddings. Sample includes 7,359 political posts (Rednote: 2,530; Weibo: 4,829).

The high platform separability (AUC = 0.898) indicates that Chinese-speaking users on Rednote and Weibo discuss the U.S. election in systematically different ways. This divergence may reflect differences in platform demographics, content moderation regimes, or community norms. Rednote’s younger, more internationally-oriented user base may engage with American politics through different frames than Weibo’s broader domestic audience. The result suggests that platform context shapes how foreign political events are interpreted and discussed, even among users sharing a common language and cultural background.

The substantial period separability (AUC = 0.815) demonstrates that the election outcome produced a detectable shift in discourse. Post-election content occupies partially distinct regions of the embedding space, suggesting the emergence of new topics (e.g., reactions to Trump’s victory, discussions of policy implications), shifts in emotional valence, or changes in dominant frames. This temporal separation is consistent with the difference-in-differences results reported in the main text, providing converging evidence that the election served as an exogenous shock to political expression.

Notably, platform identity explains more variance in embedding space than temporal period, suggesting that cross-platform differences in how Chinese-speaking audiences engage with U.S. partisan content are robust features of the information environment rather than artifacts of particular political moments.

Figure A1 displays the two-dimensional UMAP projection colored by platform. While substantial overlap exists in the central region—indicating shared baseline discourse about American politics—platform-specific clusters emerge at the periphery. Weibo posts (orange) concentrate in the upper-left and lower-left regions, while Rednote posts (blue) dominate portions of the central and lower areas.

Figure A2 shows the same projection colored by period. Pre-election posts (orange) cluster more heavily in the lower-left quadrant, while post-election posts (blue) expand into the upper and central regions. This visual pattern corroborates the classifier-based finding that the election induced measurable shifts in discourse content.

Several limitations warrant consideration. First, UMAP projections are sensitive to hyperparameter choices; alternative configurations may yield different visual patterns, though the

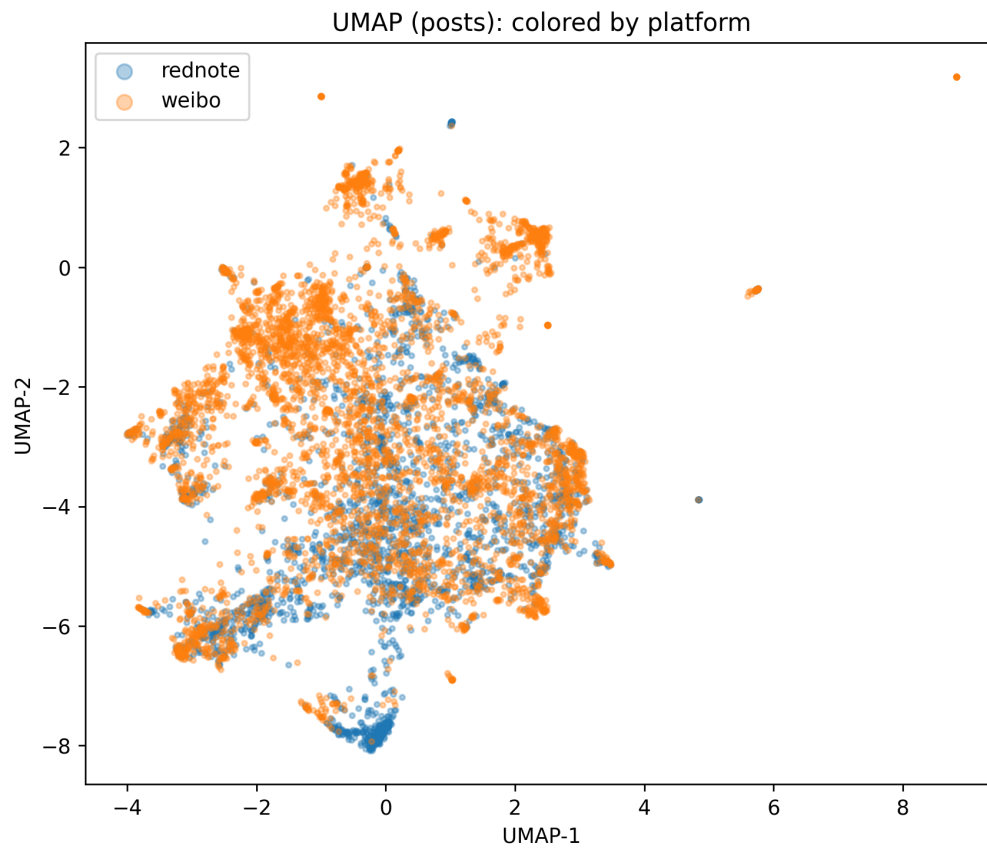


Figure A1: UMAP Projection of Political Posts by Platform

Note: Each point represents a political post embedded using `text-embedding-3-small` and projected via UMAP. Blue = Rednote; Orange = Weibo.

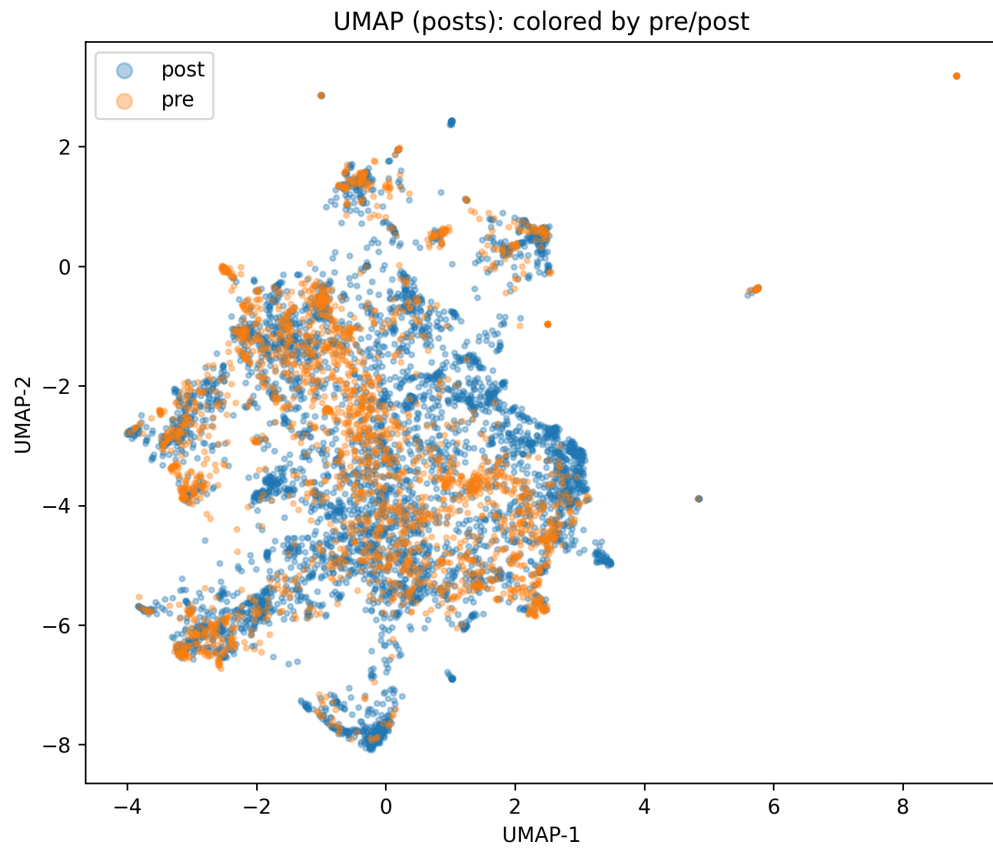


Figure A2: UMAP Projection of Political Posts by Period

Note: Same projection as Figure A1, colored by period. Blue = post-election; Orange = pre-election.

classifier-based metrics are more robust. Second, high separability does not identify *which* linguistic features drive differentiation—topic modeling or qualitative analysis would be needed to characterize the substantive differences. Third, the embeddings capture semantic similarity but may conflate distinct phenomena (e.g., topic, sentiment, style). Future work could decompose these dimensions more systematically.

D.2 Keyword Frequency Shifts

To complement the embedding-based analysis, I conduct a keyword frequency analysis that provides interpretable insights into the specific terms dominating discourse on each platform. We tokenize all political posts using jieba for Chinese text segmentation, apply comprehensive stopword filtering, and compute term frequencies by platform and period. Figure ?? and Figure ?? present cross-platform comparisons of the top keywords before and after the election.

Both platforms exhibit similar baseline shifts reflecting the transition from campaign to outcome. On both Rednote and Weibo, terms associated with the losing candidate and campaign dynamics decline sharply: “Harris”, “Democratic Party”, “Vote”, “Voters”, “Poll”, “Swing state”, and “General election” all show substantial negative shifts. Meanwhile, “Trump” increases dramatically on both platforms as discourse reorients toward the winner.

The Rednote shift pattern reveals a distinctive pivot toward immigration and personal-stakes discourse. The top-increasing keywords—“China”, “Visa”, “International students”, “Study abroad”, “Application”, “Work”, “Tariff”, “Trade”, and “Restrict”—collectively paint a picture of diaspora users rapidly reassessing how Trump’s victory will affect their visas, employment prospects, and ability to remain in the United States. Terms “Professional”, “Employment”, and “Difficulty” further suggest anxiety about career pathways under the new administration. This pattern is consistent with Rednote’s user base of Chinese international students and immigrants for whom U.S. policy carries direct personal consequences.

Weibo’s shift pattern emphasizes geopolitical implications and electoral outcomes rather than personal stakes. The top-increasing terms include “Ukraine”, “Europe”, “Russia”, and “Government”, reflecting domestic Chinese audiences’ interest in how Trump’s return might reshape international relations—particularly regarding the Russia-Ukraine conflict and U.S.-China tensions. Outcome-oriented vocabulary also surges: “Win”, “Victory”, “Elected”, “White House”, “Appointment”, and “Congratulations” capture the shift from speculative campaign coverage to post-hoc analysis of results. Notably, “Electoral votes” increases substantially as Weibo users engaged with the mechanics of Trump’s decisive victory.

These shift patterns underscore an asymmetry in how the two platforms’ user bases process the same political event. Rednote users, facing direct policy exposure, immediately pivot to practical concerns—visas, work authorization, and tariffs. Weibo users, observing from geographic and political distance, frame Trump’s victory through a geopolitical lens, emphasizing implications for international relations and great-power competition. This divergence validates the theoretical distinction between “electoral spectators” engaging as distant observers versus diaspora communities engaging as affected stakeholders.

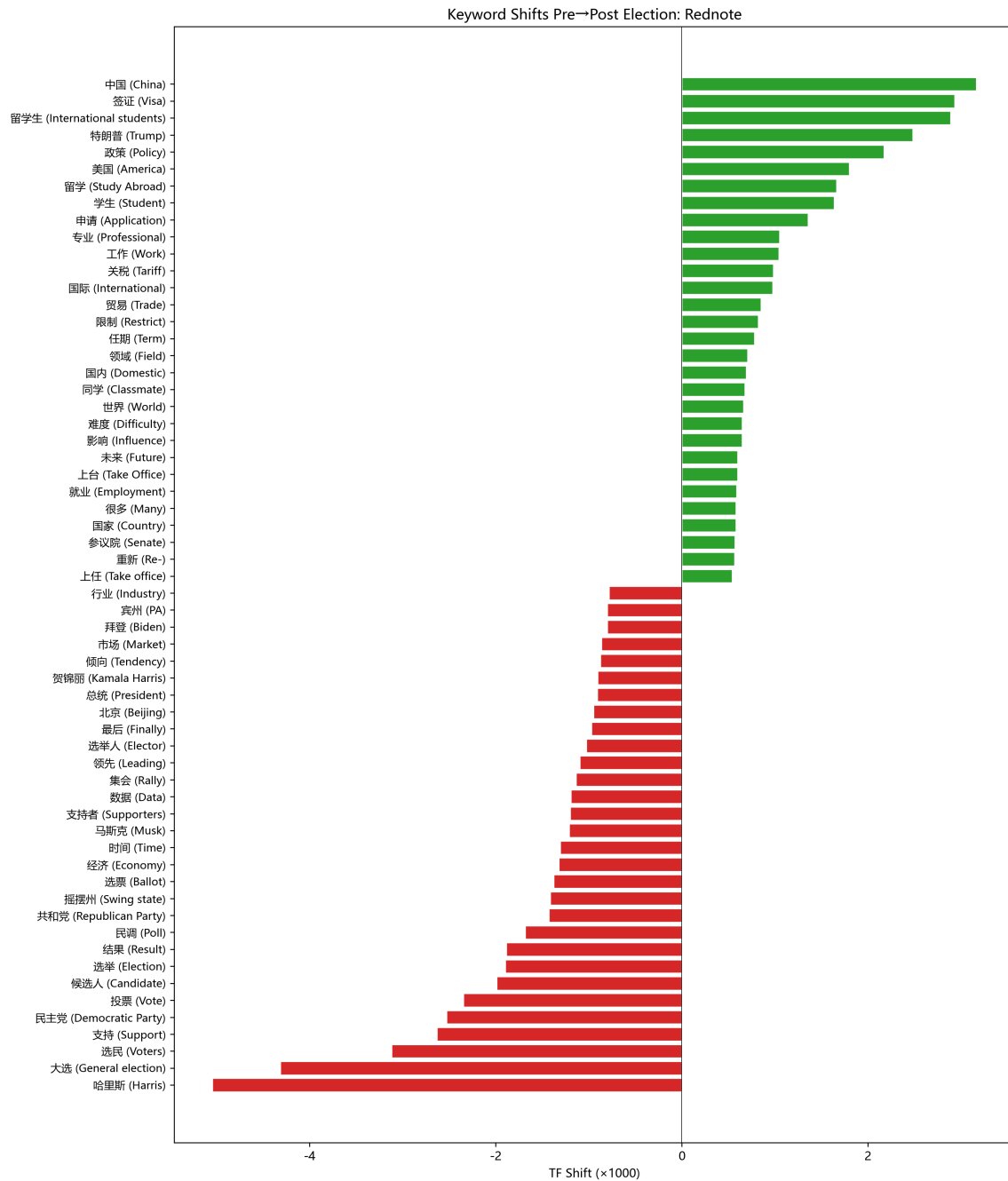


Figure A3: Keyword Frequency Shifts Pre→Post Election: Rednote

Note: Horizontal bars show the change in term frequency ($\times 1000$) from pre-election to post-election periods. Green bars indicate keywords that increased in frequency after the election; red bars indicate keywords that decreased. Only keywords appearing at least 20 times are included.

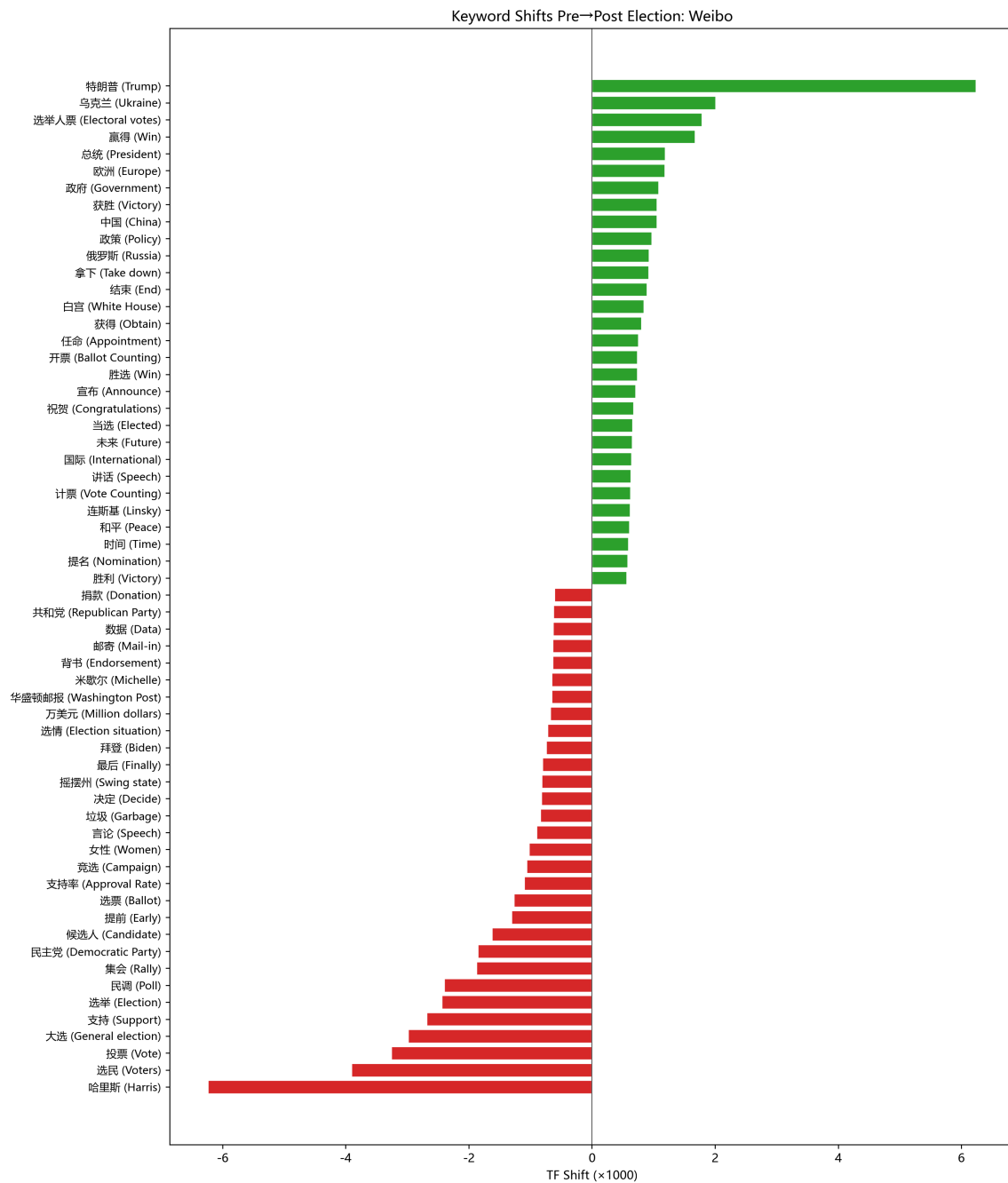


Figure A4: Keyword Frequency Shifts Pre→Post Election: Weibo

Note: Horizontal bars show the change in term frequency ($\times 1000$) from pre-election to post-election periods. Green bars indicate keywords that increased in frequency after the election; red bars indicate keywords that decreased. Only keywords appearing at least 20 times are included.

D.3 Regression Analysis

Platform Baseline Difference Table A5 confirms that most platform differences are statistically significant but vary in substantive size. The largest differences are concentrated in **information and salience dimensions**: the informed rating ($V = 0.199$) and whether content is classified as political ($V = 0.139$). Attitudinal differences are smaller but nontrivial for U.S. sentiment ($V = 0.104$) and democracy views ($V = 0.084$). By contrast, candidate sentiment differences are modest in effect size ($V \approx 0.05$ – 0.06), and China sentiment does not differ reliably across platforms ($p = .064$). This pattern foreshadows a key theme: the platforms diverge most sharply in *how* election content is framed and with what informational tone, not simply in candidate cheerleading.

Variable	χ^2	df	p-value	Cramér's V
Trump sentiment	52.0	2	<.001	0.060
Harris sentiment	38.9	2	<.001	0.052
China sentiment	5.5	2	.064	0.019
US sentiment	158.0	2	<.001	0.104
Democracy view	103.0	2	<.001	0.084
Informed	576.0	2	<.001	0.199
Outgroup hate	5.9	1	.016	0.020
Policy interest	14.6	1	<.001	0.032
Political content	282.0	1	<.001	0.139

Note: Cramér's V measures effect size (0.1 = small, 0.3 = medium, 0.5 = large).

Table A5: Chi-Square Tests for Platform Differences

Variable	Mean (Rednote)	Mean (Weibo)	Difference	t-statistic	p-value
Trump sentiment	0.043	0.091	−0.049	−5.38	<.001
Harris sentiment	−0.021	−0.056	0.035	5.78	<.001
China sentiment	0.032	0.041	−0.009	−2.21	.027
US sentiment	−0.010	−0.111	0.101	12.60	<.001
Democracy view	0.007	−0.061	0.068	9.85	<.001
Informed	1.09	1.17	−0.080	−5.16	<.001
Misinfo score	0.036	0.055	−0.020	−7.74	<.001
Outgroup hate	0.037	0.047	−0.010	−2.62	.009
Policy interest	0.347	0.312	0.035	3.79	<.001

Note: Positive difference indicates Rednote mean is higher. All variables treated as continuous for t-test purposes.

Table A6: T-Tests for Platform Mean Differences

The mean comparisons in Table A6 sharpen the descriptive story. The largest gap is **U.S. sentiment**: the average Weibo post is notably more negative toward the United States than the average Rednote post (difference = 0.101, $t = 12.60$), and democracy views are similarly more negative on Weibo (difference = 0.068, $t = 9.85$). Candidate sentiment differences

are directionally consistent with Table 3 (Weibo more pro-Trump and more anti-Harris), but smaller in magnitude. In parallel, Weibo has modestly higher misinformation and hate, while Rednote shows slightly higher policy interest. Taken together, the descriptive results suggest that the platforms differ less in *whether* they discuss the election and more in *the evaluative lens* brought to the discussion: Weibo discourse tilts more strongly toward negatively valenced interpretations of the U.S. and democracy.

Multivariate content models. To identify which discourse features travel together—and to assess whether bivariate platform gaps are primarily compositional—I estimate multivariate models that include platform alongside candidate sentiment, country sentiment, and information quality. For outgroup hate, I estimate:

$$\log \left(\frac{\Pr(\text{Hate}_i = 1)}{1 - \Pr(\text{Hate}_i = 1)} \right) = \beta_0 + \beta_1 \cdot \text{Weibo}_i + \beta_2 \cdot \text{Trump}_i + \beta_3 \cdot \text{Harris}_i \\ + \beta_4 \cdot \text{China}_i + \beta_5 \cdot \text{US}_i + \beta_6 \cdot \text{Democracy}_i \\ + \beta_7 \cdot \text{PolicyInterest}_i + \beta_8 \cdot \text{Informed}_i + \beta_9 \cdot \text{Misinfo}_i. \quad (4)$$

For democracy views, I estimate an ordered logistic model with the same covariates:

$$\Pr(\text{Democracy}_i \leq j) = \frac{1}{1 + \exp(-(\alpha_j - \mathbf{X}_i' \boldsymbol{\beta}))}. \quad (5)$$

These models are descriptive: they identify bundles of discourse features (e.g., misinformation, hostility, and institutional cynicism) rather than causal pathways.

Table A7 reports a multivariate logistic regression that treats platform, political sentiment, and discourse-quality indicators as joint correlates of outgroup hate. Three patterns stand out.

First, misinformation is the strongest correlate in the model. A one-unit increase in the misinformation index (from 0 to 1) is associated with a 6.6-fold increase in the odds of outgroup hate (OR = 6.63, $p < 0.001$); equivalently, a 0.10-point increase corresponds to about 21% higher odds (OR $\approx e^{0.189} = 1.21$). Second, informational quality is protective: each one-point increase in the informed score is associated with 33% lower odds of hate (OR = 0.67, $p < 0.001$). Together, these associations locate hostile rhetoric in a low-information, high-distortion subset of notes rather than uniformly across election discussions.

Third, the bivariate Weibo–Rednote gap in hate appears largely compositional. Once misinformation, informational quality, and political framing are included, the platform coefficient attenuates to OR = 1.10 and is statistically indistinguishable from zero, suggesting that higher hate prevalence on Weibo reflects the kinds of content that co-occur with hate rather than an irreducible platform effect.

Finally, country- and institution-oriented sentiment is more tightly linked to hate than candidate sentiment. Pro-U.S. sentiment is strongly associated with lower odds of hate (OR = 0.39), whereas pro-China sentiment is associated with higher odds (OR = 1.65). Candidate

sentiment is asymmetric: Trump sentiment is not meaningfully related to hate, while pro-Harris sentiment is strongly negatively associated (OR = 0.23). One interpretation is that pro-Harris notes tend to appear in conversational contexts with less outgroup targeting, rather than that candidate evaluations mechanically generate hostility.

	Coefficient	SE	Odds Ratio	95% CI
Intercept	−3.529***	0.119	0.029	[0.02, 0.04]
Platform (Weibo = 1)	0.093	0.108	1.097	[0.89, 1.36]
Trump sentiment	−0.034	0.080	0.967	[0.83, 1.13]
Harris sentiment	−1.481***	0.103	0.227	[0.19, 0.28]
China sentiment	0.500**	0.157	1.649	[1.21, 2.23]
US sentiment	−0.950***	0.121	0.387	[0.31, 0.49]
Democracy view	−0.152	0.133	0.859	[0.66, 1.11]
Policy interest	0.467***	0.111	1.595	[1.28, 1.98]
Informed	−0.398***	0.075	0.672	[0.58, 0.78]
Misinfo score	1.891***	0.207	6.629	[4.40, 9.92]
<i>N</i>			14,606	
AIC			4,538.7	
Null Deviance			5,285.8	
Residual Deviance			4,518.7	

Note: ** $p < .01$; *** $p < .001$.

Table A7: Logistic Regression: Predictors of Outgroup Hate

	Coefficient	SE
Intercept	−3.016***	0.181
Platform (Weibo)	0.121	0.208
Post-election	−0.336	0.213
Weibo × Post-election	0.191	0.242
<i>N</i>		14,064
AIC		5,124.0
Null Deviance		5,126.1
Residual Deviance		5,116.0

Note: *** $p < .001$. No coefficients except intercept are significant.

Table A8: Difference-in-Differences: Outgroup Hate (Logistic Regression)

Table A8 shows no evidence that outgroup hate changes systematically after Election Day, either overall or differentially by platform. This temporal null result complements the multivariate findings above: hostility appears less responsive to electoral “events” and more tied to the underlying informational and framing characteristics of posts.

I next examine whether each platform differentially amplifies problematic content. Here, the descriptive theory is explicitly about **attention**: even modest differences in the prevalence

of hate or misinformation can matter if platform engagement disproportionately rewards those posts.

Engagement Models Table A9 shows that Weibo posts receive higher likes on average (and a much higher median), consistent with platform scale and exposure differences. On both platforms, engagement is highly right-skewed (large mean–median gaps), so a small number of viral posts can disproportionately shape users’ experienced information environment.

	Rednote	Weibo
Mean likes	240	354
Median likes	14	67
Mean comments	51	37
Median comments	2	13

Table A9: Engagement Summary Statistics by Platform

Table A10 reveals the most consequential platform divergence in the results: **hate is rewarded on Weibo but not on Rednote**. On Weibo, hateful posts receive substantially more likes than non-hateful posts (491 vs. 353). On Rednote, hateful posts receive fewer likes (146 vs. 243). This contrast indicates that even if the *generation* of hate is largely explained by content composition (Table A7), the *visibility* of hate may still differ sharply because engagement concentrates attention differently across platforms.

A similar directional contrast appears for anti-U.S. and anti-democracy content, which receives higher likes on Weibo than pro-positions, but not on Rednote. By contrast, misinformation does not display a consistent engagement advantage on either platform, suggesting that misinformation’s influence may operate indirectly (e.g., via its co-occurrence with hostility) rather than through direct engagement rewards.

Table A11 examines which content features predict engagement (log-transformed likes). Hate is positively associated with likes in the pooled model, but this relationship is platform-contingent: the Weibo $\times$ Hate interaction is positive and significant ($\beta = 0.53$, $p < 0.01$), while the hate main effect becomes insignificant. Hostility is rewarded on Weibo but not on Rednote.

Policy-focused notes receive fewer likes ($\beta = -0.32$, $p < 0.001$), indicating that substantive discussion is less likely to go viral. Pro-Trump sentiment is associated with lower engagement ($\beta = -0.09$, $p < 0.01$), suggesting that the platform-level pro-Trump skew in posting does not translate into engagement advantages for pro-Trump content.

The regression results in Table A11 reinforce the descriptive engagement pattern. In Model 1, hate is positively associated with likes. Model 2 shows that this relationship is **platform-contingent**: the Weibo $\times$ Hate interaction is positive and significant, while the hate main effect becomes small and statistically insignificant (with Rednote as the reference). Substantively, hate is not generally “popular” across contexts; rather, it is *rewarded* in the Weibo attention environment.

	Rednote	Weibo
<i>By Trump Sentiment</i>		
Pro-Trump	155	277
Anti-Trump	166	402
Neutral	261	366
<i>By China Sentiment</i>		
Pro-China	170	278
Anti-China	590	184
Neutral	241	358
<i>By US Sentiment</i>		
Pro-US	184	203
Anti-US	141	366
Neutral	256	359
<i>By Democracy View</i>		
Pro-democracy	210	358
Anti-democracy	147	407
Neutral	249	350
<i>By Discourse Quality</i>		
High misinfo	168	333
Low misinfo	242	355
Outgroup hate	146	491
No hate	243	353

Note: High misinfo defined as score ≥ 0.5 ; low misinfo as < 0.5 .

Table A10: Mean Likes by Content Characteristics and Platform

	Model 1 (Main Effects)	Model 2 (Interactions)
Intercept	3.497*** (0.069)	3.199*** (0.031)
Platform (Weibo)	1.280*** (0.032)	1.244*** (0.034)
Trump sentiment	−0.087** (0.029)	−0.093** (0.029)
Harris sentiment	−0.036 (0.046)	−0.038 (0.046)
China sentiment	0.087 (0.065)	0.085 (0.065)
US sentiment	−0.034 (0.044)	−0.032 (0.044)
Democracy view	−0.032 (0.051)	−0.031 (0.051)
Outgroup hate	0.457*** (0.068)	0.021 (0.153)
Misinfo score	0.195 (0.102)	0.322 (0.236)
Policy interest	−0.321*** (0.030)	−0.342*** (0.030)
Political	−0.347*** (0.071)	—
Weibo × Hate	—	0.529** (0.170)
Weibo × Misinfo	—	−0.161 (0.256)
R^2	0.113	0.112
N	14,582	14,582

Note: Dependent variable is $\log(1 + \text{likes})$. Standard errors in parentheses. ** $p < .01$; *** $p < .001$.

Table A11: OLS Regression: Predictors of Log-Transformed Engagement

Two other engagement patterns are notable. Policy-focused posts receive fewer likes, indicating that substantive discussion is less likely to go viral. Meanwhile, pro-Trump sentiment is associated with *lower* engagement, suggesting that platform-level pro-Trump skew in posting (Table 3) does not automatically translate into engagement advantages for pro-Trump posts.

	Rednote	Weibo
<i>Amplification Ratios</i>		
Hate / No-hate likes	0.61	1.39
High-misinfo / Low-misinfo likes	0.70	0.94
<i>Within-Platform t-test: Hate vs. No-Hate</i>		
<i>p</i> -value	0.014	0.033

Note: Amplification ratio = mean likes for posts with attribute / mean likes for posts without. Values > 1 indicate amplification; < 1 indicate suppression. The *t*-tests assess whether the difference in likes between hate and no-hate posts is statistically significant within each platform.

Table A12: Content Amplification Ratios by Platform

Table A12 summarizes this engagement contrast: On both platforms, the difference in engagement between hateful and non-hateful posts is statistically significant ($p = .014$ for Rednote; $p = .033$ for Weibo). However, the direction differs: on Rednote, hateful posts receive significantly **fewer** likes (ratio = 0.61), while on Weibo, hateful posts receive significantly **more** likes (ratio = 1.39). This divergent pattern suggests fundamentally different community norms or algorithmic amplification across platforms.

	Coefficient	SE	IRR	95% CI
Intercept	5.589***	0.030	267.4	[252.6, 283.3]
Platform (Weibo)	0.402***	0.032	1.494	[1.40, 1.59]
Trump sentiment	-0.254***	0.029	0.776	[0.73, 0.82]
Harris sentiment	0.044	0.046	1.045	[0.95, 1.14]
China sentiment	-0.099	0.066	0.906	[0.79, 1.04]
US sentiment	-0.171***	0.044	0.843	[0.76, 0.93]
Democracy view	0.047	0.052	1.048	[0.94, 1.16]
Outgroup hate	0.334***	0.069	1.397	[1.22, 1.61]
Misinfo score	-0.043	0.103	0.957	[0.79, 1.17]
Policy interest	-0.448***	0.030	0.639	[0.60, 0.68]
<i>N</i>	14,582			
θ (dispersion)	0.373 (SE = 0.004)			
AIC	182,827			

Note: IRR = Incidence Rate Ratio. Values > 1 indicate more likes. *** $p < .001$.

Table A13: Negative Binomial Regression: Predictors of Engagement (Likes)

The negative binomial model in Table A13 confirms the main engagement conclusions while accounting for overdispersed counts. Hateful posts receive 40% more likes (IRR = 1.40),

and policy-focused posts receive substantially fewer likes ($IRR = 0.64$). Misinformation does not independently predict likes, again implying that misinformation’s main consequence is its tight linkage to hostility (Table A7) rather than direct engagement rewards.

E Comment-level Analysis

E.1 Comment-Level UMAP Projections

To complement the post-level analysis in Appendix D.1, I extend the UMAP embedding analysis to comment-level data. Comments represent a distinct mode of engagement—reactive, conversational, and typically shorter than original posts. Analyzing comments separately allows us to assess whether platform and temporal patterns observed in posts generalize to user responses. I embed 138,066 political comments (after filtering) using OpenAI’s `text-embedding-3-small` model. For visualization and separability analysis, I sample 5,000 comments per platform-period group (20,000 total), stratified to ensure balanced representation. We apply the same UMAP parameters as the post-level analysis (`n_neighbors=20`, `min_dist=0.05`, cosine metric) and compute cross-validated logistic regression classifiers on the full-dimensional embeddings. Table A14 reports classification performance for platform and period distinctions at the comment level.

Classification Task	AUC	Accuracy
Platform (Rednote vs. Weibo)	0.835 (0.003)	0.746 (0.004)
Period (Pre vs. Post)	0.664 (0.006)	0.618 (0.006)

Table A14: Embedding-Based Separability: Comments

Note: Standard deviations from 5-fold cross-validation in parentheses. Sample includes 20,000 comments stratified by platform and period. Total embedded comments: 138,066.

Table A15 compares separability metrics between posts and comments. Several pattern emerges from the comparison:

Content Type	Classification Task	AUC	Accuracy
Posts	Platform	0.898 (0.009)	0.826 (0.010)
Posts	Period	0.815 (0.013)	0.766 (0.005)
Comments	Platform	0.835 (0.003)	0.746 (0.004)
Comments	Period	0.664 (0.006)	0.618 (0.006)

Table A15: Embedding Separability: Posts vs. Comments

Note: Standard deviations in parentheses. Posts sample: 7,359; Comments sample: 20,000.

First, platform separability remains substantial but attenuated. Comments exhibit lower platform separability ($AUC = 0.835$) compared to posts ($AUC = 0.898$), though the difference remains meaningful. This attenuation likely reflects the nature of comment discourse: comments are shorter, more reactive, and less likely to contain the extended argumentation or framing

that distinguishes platform cultures. Nevertheless, the persistent platform effect suggests that even brief user responses carry detectable signatures of their originating community.

Second, temporal separability is markedly weaker for comments. The most striking difference appears in period classification. While posts show strong pre/post separability ($AUC = 0.815$), comments approach chance levels ($AUC = 0.664$). This divergence admits several interpretations. First, the election's discursive impact may manifest primarily in original content creation rather than reactive commentary—users may craft new posts to process the outcome while their commenting behavior remains relatively stable. Second, comments' brevity may limit the semantic signal available for temporal classification. Third, comments often respond to the parent post's framing, creating a dependency structure that partially decouples comment content from the broader temporal context.

These findings suggest that platform identity shapes user expression more consistently than temporal context, particularly for reactive content. The robust platform separability across both posts and comments indicates that community-specific norms, audiences, and affordances leave detectable traces even in brief interactions. For researchers studying imported partisanship, this implies that platform selection may be a more reliable indicator of discourse patterns than timing relative to political events, at least for comment-level engagement.

Figure A5 displays the two-dimensional UMAP projection of comments colored by platform. Compared to the post-level visualization (Figure A1), comments exhibit greater overall mixing, consistent with the lower separability metrics. However, platform-specific clustering remains visible, with Rednote comments (blue) concentrating in certain regions distinct from Weibo comments (orange).

Figure A6 shows the same projection colored by period. The visual overlap between pre-election (orange) and post-election (blue) comments is substantially greater than observed for posts, corroborating the weak temporal separability. This suggests that the semantic shift induced by the election outcome is less pronounced—or less detectable—in reactive comment discourse.

The comment-level analysis inherits the limitations noted for the post-level analysis. Additionally, comments present unique challenges: their brevity reduces semantic signal, their dependence on parent posts introduces confounding, and their higher volume may include more noise (e.g., spam, off-topic responses). The lower period separability should be interpreted cautiously—it may reflect genuine homogeneity in comment discourse or methodological limitations in detecting subtle temporal shifts in short texts.

E.2 Red Note

Table A16 presents the full set of content diffusion estimates, expanding on the results discussed in the main text. The table reports how post-level sentiment predicts comment-level sentiment across all classified outcomes, with heterogeneity by election period and user location.

Several patterns beyond the main findings are notable. First, diffusion effects are strongest

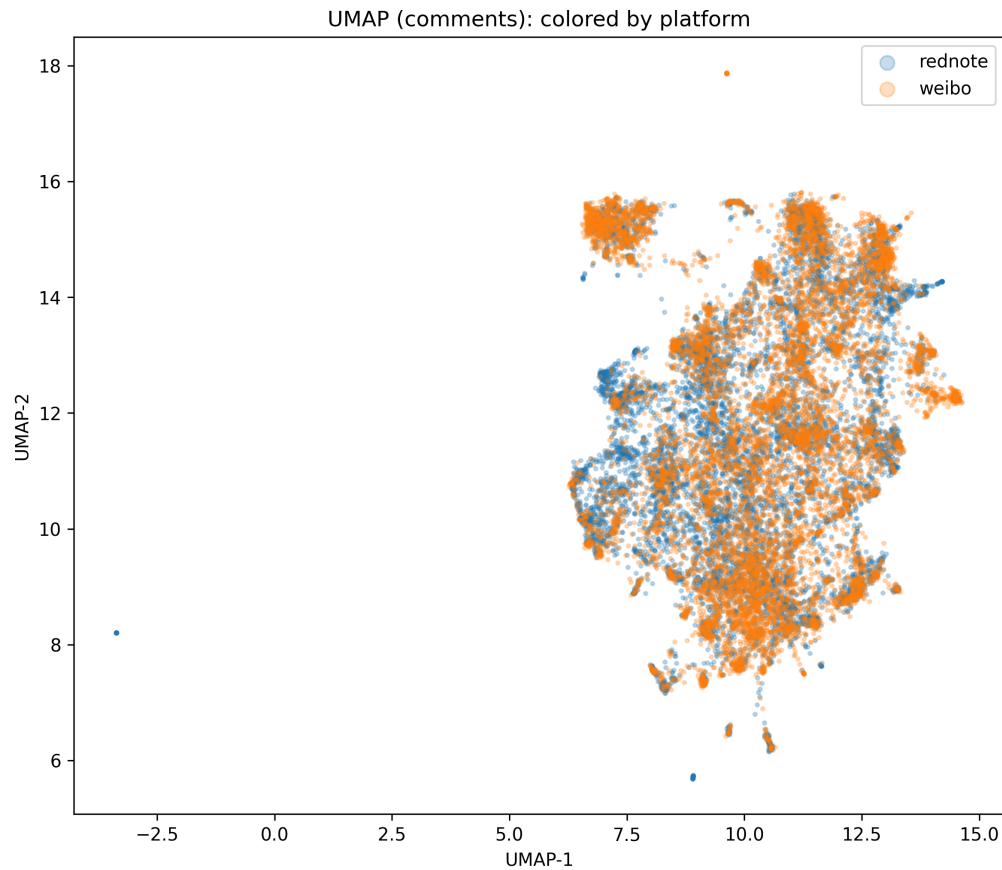


Figure A5: UMAP Projection of Political Comments by Platform

Note: Each point represents a political comment embedded using `text-embedding-3-small` and projected via UMAP. Blue = Rednote; Orange = Weibo. Sample: 20,000 comments.

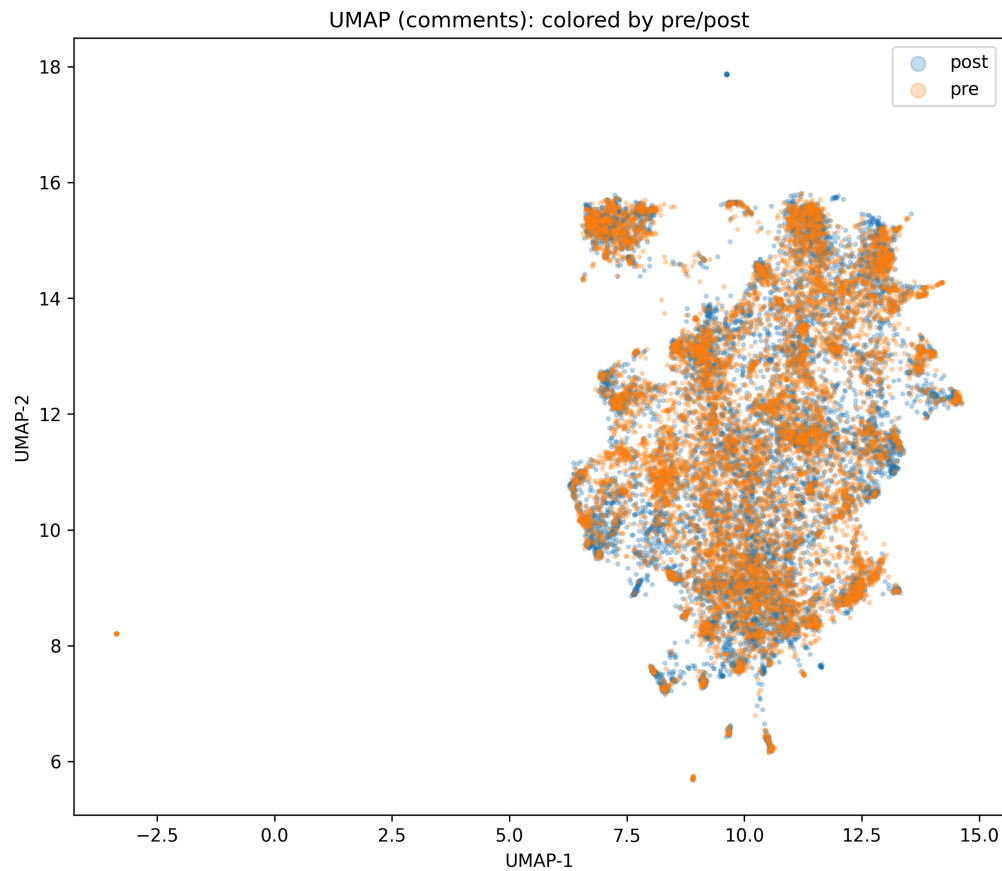


Figure A6: UMAP Projection of Political Comments by Period

Note: Same projection as Figure A5, colored by period. Blue = post-election; Orange = pre-election.

for political and policy interest ($\beta = 0.143$ and $\beta = 0.116$), indicating that engagement-focused content spreads more readily than attitudinal content. Second, partisan diffusion is asymmetric: Democratic Party support diffuses significantly ($\beta = 0.028$) while Republican Party support shows weaker effects ($\beta = 0.005$). Third, outgroup hate diffuses strongly in the post-election period ($\beta = 0.082$) but not before ($\beta = 0.021$), suggesting that hostile rhetoric became more contagious after the election outcome was known. Fourth, pro-China sentiment diffuses only among Mainland China users ($\beta = 0.012$) and not among US users ($\beta = -0.000$), while pro-democracy content shows the opposite pattern—stronger diffusion among US users ($\beta = 0.056$) than Mainland China users ($\beta = 0.029$).

Category	Outcome	Overall	Election Period		User Location	
			Pre	Post	China	US
Partisan	Trump Support	0.020*** (0.004)	0.008 (0.013)	0.022*** (0.004)	0.015*** (0.005)	0.031*** (0.008)
	Harris Support	0.016 (0.012)	0.033** (0.016)	0.009 (0.015)	0.019 (0.013)	0.016 (0.012)
	GOP Support	0.005* (0.003)	0.008 (0.005)	0.005 (0.003)	-0.001 (0.003)	0.009* (0.005)
	Dem. Support	0.028*** (0.010)	0.020** (0.010)	0.027** (0.012)	0.027** (0.010)	0.014 (0.010)
Country Sentiment	Pro-China Sentiment	0.010* (0.005)	0.046** (0.018)	0.007* (0.004)	0.012* (0.006)	-0.000 (0.008)
	Pro-US Sentiment	0.028*** (0.008)	0.042** (0.018)	0.022*** (0.007)	0.023*** (0.009)	0.032*** (0.009)
Democracy View	Democracy Support	0.039*** (0.009)	0.068*** (0.019)	0.024*** (0.009)	0.029*** (0.009)	0.056*** (0.013)
Political Interest	Political Interest	0.143*** (0.044)	0.032 (0.061)	0.139*** (0.043)	0.122*** (0.045)	0.108** (0.048)
Policy Interest	Policy Interest	0.116*** (0.013)	0.112*** (0.031)	0.117*** (0.013)	0.100*** (0.014)	0.134*** (0.015)
Outgroup Hate	Outgroup Hate	0.072*** (0.011)	0.021 (0.028)	0.082*** (0.009)	0.076*** (0.013)	0.056*** (0.015)

Notes: OLS estimates with post-level clustered standard errors in parentheses. All models include hour-of-day and day-of-week fixed effects. “Overall” reports the full-sample effect. “Pre” and “Post” columns show marginal effects from an interaction model, where “Pre” refers to comments before November 5, 2024. “China” and “US” columns report effects estimated separately by user IP address location. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table A16: Content Diffusion from Posts to Comments on Rednote

Table A17 reports the full set of daily event study coefficients underlying Figure ?? in the main text. Panel A presents pre-election and Election Day estimates (days -7 through 0); Panel B presents post-election estimates (days 1 through 7). Day -1 serves as the reference

period.

The coefficients support the parallel trends assumption for the key outcomes. For US sentiment and democracy views, pre-election estimates (days -7 through -2) fluctuate around zero without a clear trend, and most are statistically insignificant. The sharp increase on day 0 (Election Day) and sustained elevation thereafter is consistent with the election outcome—rather than pre-existing trends—driving the observed shifts.

For partisan sentiment, Harris support shows some negative pre-election coefficients (e.g., $\beta = -0.058$ on day -7), suggesting possible anticipation effects or differential sample composition in the early period. However, the sharpest decline occurs on day $+2$ ($\beta = -0.034$), when the election result was unambiguous. Trump support increases more gradually, with significant positive coefficients appearing on Election Day and persisting through day $+7$.

Political engagement (Pol. Engagement) shows the clearest temporal break: estimates are close to zero before the election but turn sharply negative beginning on day 0 ($\beta = -0.071$) and decline steadily to -0.151 by day $+7$. This pattern is consistent with the resolution of electoral uncertainty reducing users' motivation to engage with political content.

China sentiment remains flat throughout the observation window, with no coefficient exceeding 0.01 in absolute value. The election did not trigger a nationalist response in either direction.

E.3 Weibo

Table A18 presents content diffusion estimates for Weibo, providing a comparison to the Rednote results. Several differences emerge. First, diffusion effects on Weibo are substantially larger across nearly all outcomes compared to Rednote. Political interest shows the strongest diffusion ($\beta = 0.323$), roughly double the magnitude observed on Rednote ($\beta = 0.143$). This suggests that Weibo's platform architecture—with its more public, broadcast-oriented structure—amplifies content contagion relative to Rednote's more intimate community-based design.

Second, the temporal dynamics differ markedly. On Weibo, pro-US sentiment and democracy views diffuse *more strongly before* the election (pro-US: $\beta = 0.084$ pre vs. $\beta = 0.043$ post; democracy: $\beta = 0.093$ pre vs. $\beta = 0.037$ post), whereas on Rednote the pattern was reversed or stable. This may reflect Weibo's more politically engaged user base, who were already primed to discuss American democracy before the election, with diffusion effects attenuating once the outcome was known. In contrast, outgroup hate shows the same pattern as Rednote: stronger diffusion post-election ($\beta = 0.100$) than pre-election ($\beta = 0.069$).

Third, location heterogeneity is less pronounced on Weibo than on Rednote. Pro-China sentiment diffuses similarly among Mainland China users ($\beta = 0.058$) and US users ($\beta = 0.070$), unlike on Rednote where it diffused only among Mainland users. This likely reflects Weibo's overwhelmingly domestic user base (95% Mainland China), making US-based users on Weibo a more selected sample. Harris support fails to diffuse among US users on Weibo

Panel A: Pre-Election Period

Category	Outcome	−7	−6	−5	−4	−3	−2	0
Partisan Sentiment	Trump Support	0.006 (0.018)	0.048** (0.018)	−0.019 (0.016)	−0.008 (0.013)	0.021 (0.014)	0.001 (0.012)	0.025** (0.010)
	Harris Support	−0.058*** (0.017)	−0.061*** (0.015)	−0.012 (0.012)	−0.016 (0.010)	−0.028** (0.012)	−0.024** (0.009)	−0.019** (0.007)
	Outgroup Hate	−0.024* (0.013)	−0.010 (0.014)	0.001 (0.012)	0.013 (0.010)	0.001 (0.010)	0.002 (0.009)	0.006 (0.007)
Political Interest	Policy Interest	−0.065*** (0.016)	−0.036* (0.018)	−0.032** (0.014)	−0.022* (0.013)	−0.037*** (0.012)	−0.035*** (0.011)	−0.039*** (0.009)
	Pol. Engagement	−0.033 (0.023)	−0.046* (0.026)	−0.024 (0.025)	−0.039* (0.021)	−0.003 (0.020)	−0.003 (0.019)	−0.071*** (0.014)
Regime Sentiment	China Sentiment	−0.005 (0.006)	0.009* (0.006)	−0.005 (0.005)	−0.012** (0.005)	−0.013* (0.007)	−0.000 (0.006)	0.011** (0.005)
	US Sentiment	0.024 (0.017)	0.035* (0.018)	−0.028* (0.015)	0.018 (0.013)	−0.007 (0.013)	−0.019 (0.011)	0.046*** (0.010)
	Democracy View	0.042** (0.017)	0.062*** (0.016)	−0.009 (0.015)	0.020* (0.012)	−0.019 (0.011)	0.011 (0.012)	0.046*** (0.010)

Panel B: Post-Election Period

Category	Outcome	1	2	3	4	5	6	7
Partisan Sentiment	Trump Support	0.027*** (0.009)	0.020* (0.010)	0.034*** (0.011)	0.011 (0.011)	0.028** (0.013)	0.029** (0.011)	0.020 (0.014)
	Harris Support	−0.011 (0.007)	−0.034*** (0.008)	−0.019** (0.008)	−0.023*** (0.008)	−0.032*** (0.010)	−0.024** (0.010)	−0.018* (0.008)
	Outgroup Hate	−0.016** (0.007)	−0.012 (0.008)	0.003 (0.008)	0.015 (0.009)	0.006 (0.009)	−0.017* (0.009)	−0.009 (0.011)
Political Interest	Policy Interest	−0.050*** (0.009)	−0.045*** (0.009)	−0.045*** (0.009)	−0.028** (0.009)	−0.063*** (0.011)	−0.042*** (0.012)	−0.034*** (0.011)
	Pol. Engagement	−0.087*** (0.016)	−0.109*** (0.020)	−0.112*** (0.017)	−0.087*** (0.018)	−0.082*** (0.019)	−0.126*** (0.019)	−0.151*** (0.011)
Regime Sentiment	China Sentiment	0.002 (0.004)	−0.001 (0.004)	−0.001 (0.004)	0.001 (0.004)	0.004 (0.005)	0.004 (0.005)	0.009 (0.006)
	US Sentiment	0.063*** (0.008)	0.062*** (0.009)	0.070*** (0.009)	0.063*** (0.010)	0.054*** (0.011)	0.045*** (0.010)	0.076*** (0.011)
	Democracy View	0.053*** (0.008)	0.048*** (0.008)	0.061*** (0.008)	0.050*** (0.008)	0.052*** (0.010)	0.045*** (0.010)	0.056*** (0.011)

Notes: Two-way clustered standard errors (by user and date) in parentheses. Day −1 is the omitted reference period. All specifications include user fixed effects. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table A17: Event Study Estimates: Effects Relative to Day −1

($\beta = 0.006$, n.s.), echoing the Rednote pattern where pro-Harris content resonated less with overseas audiences.

Category	Outcome	Overall	Election Period		User Location	
			Pre	Post	China	US
Partisan	Trump Support	0.026*** (0.002)	0.024*** (0.006)	0.026*** (0.003)	0.026*** (0.002)	0.042*** (0.014)
	Harris Support	0.038*** (0.005)	0.020** (0.008)	0.049*** (0.006)	0.038*** (0.005)	0.006 (0.013)
	GOP Support	0.011*** (0.002)	0.022*** (0.005)	0.008*** (0.002)	0.011*** (0.002)	0.034*** (0.012)
	Dem. Support	0.045*** (0.004)	0.025*** (0.006)	0.054*** (0.005)	0.045*** (0.004)	0.043** (0.017)
Country Sentiment	Pro-China Sentiment	0.058*** (0.006)	0.052*** (0.015)	0.059*** (0.006)	0.058*** (0.006)	0.070** (0.031)
	Pro-US Sentiment	0.054*** (0.003)	0.084*** (0.007)	0.043*** (0.003)	0.054*** (0.003)	0.035*** (0.013)
Democracy View	Democracy Support	0.058*** (0.005)	0.093*** (0.009)	0.037*** (0.004)	0.057*** (0.005)	0.060*** (0.012)
Political Interest	Political Interest	0.323*** (0.010)	0.308*** (0.033)	0.319*** (0.010)	0.321*** (0.010)	0.387*** (0.064)
Policy Interest	Policy Interest	0.088*** (0.005)	0.066*** (0.007)	0.093*** (0.005)	0.087*** (0.005)	0.113*** (0.017)
Outgroup Hate	Outgroup Hate	0.093*** (0.006)	0.069*** (0.010)	0.100*** (0.007)	0.092*** (0.006)	0.095*** (0.024)
<i>N</i> (comments)		241,479	46,917	194,561	230,229	3,901
<i>N</i> (posts)		8,818	1,809	7,416	8,766	1,465

Notes: OLS estimates with post-level clustered standard errors in parentheses. All models include hour-of-day and day-of-week fixed effects. “Overall” reports the full-sample effect. “Pre” and “Post” columns show marginal effects from an interaction model, where “Pre” refers to comments before November 5, 2024. “China” and “US” columns report effects estimated separately by user IP address location. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table A18: Content Diffusion from Posts to Comments on Weibo

Table A19 reports the full set of daily event study coefficients for Weibo. The patterns differ notably from Rednote. For US sentiment and democracy views, pre-election coefficients show more variation, but the post-election increase is similarly pronounced and sustained. US sentiment rises sharply on Election Day ($\beta = 0.057$) and remains elevated through day +7 ($\beta = 0.024$). Democracy views show a parallel pattern, increasing from $\beta = 0.058$ on day 0 to sustained levels around 0.05–0.06 throughout the post-election period.

Trump support on Weibo shows a different pattern than on Rednote. Pre-election coef-

ficients are mostly negative or near zero, with a brief positive spike on day +1 ($\beta = 0.032$, $p < 0.01$) that quickly dissipates. This may reflect Weibo's predominantly domestic Chinese user base, where bandwagon effects toward the winner are weaker or more transient than on the more internationally oriented Rednote platform.

Political engagement declines post-election on Weibo, though less dramatically than on Rednote. Outgroup hate increases modestly after the election, with significant positive coefficients emerging by day +5 ($\beta = 0.021$) and persisting through day +7. Policy interest shows more variability on Weibo, with significant positive coefficients appearing in the post-election period.

Panel A: Pre-Election Period

Category	Outcome	−7	−6	−5	−4	−3	−2	0
Partisan Sentiment	Trump Support	−0.009 (0.009)	−0.007 (0.008)	−0.016 (0.007)	−0.031*** (0.009)	−0.019* (0.008)	−0.007 (0.008)	−0.012 (0.012)
	Harris Support	−0.017** (0.007)	−0.011 (0.007)	−0.023*** (0.007)	0.004 (0.006)	0.007 (0.007)	−0.003 (0.006)	−0.013** (0.006)
	Outgroup Hate	−0.014** (0.007)	0.011 (0.007)	0.011 (0.007)	0.021*** (0.006)	0.008 (0.006)	0.012 (0.006)	0.012 (0.010)
Political Interest	Policy Interest	0.022** (0.009)	0.033*** (0.007)	0.051*** (0.007)	0.077*** (0.011)	0.011 (0.011)	0.007 (0.007)	0.016* (0.009)
	Pol. Engagement	0.074** (0.012)	0.053** (0.012)	0.074** (0.013)	0.033*** (0.010)	0.022** (0.007)	0.043*** (0.012)	−0.004 (0.011)
Regime Sentiment	China Sentiment	0.011* (0.006)	0.002 (0.004)	0.010 (0.006)	0.005 (0.005)	0.011* (0.005)	0.018** (0.006)	0.046 (0.041)
	US Sentiment	−0.004 (0.006)	−0.014 (0.010)	−0.018** (0.009)	−0.014 (0.018)	−0.017* (0.008)	−0.021** (0.006)	0.057*** (0.007)
	Democracy View	0.013 (0.008)	0.013 (0.009)	−0.018* (0.010)	−0.009 (0.008)	−0.003 (0.011)	−0.017** (0.007)	0.058*** (0.007)

Panel B: Post-Election Period

Category	Outcome	1	2	3	4	5	6	7
Partisan Sentiment	Trump Support	0.032*** (0.007)	0.007 (0.007)	0.008 (0.007)	−0.009 (0.006)	−0.002 (0.007)	−0.007 (0.007)	−0.001 (0.007)
	Harris Support	0.022*** (0.006)	−0.010 (0.006)	−0.001 (0.006)	−0.009 (0.006)	−0.011 (0.007)	−0.002 (0.007)	0.000 (0.006)
	Outgroup Hate	0.025** (0.010)	0.033*** (0.006)	0.016** (0.006)	0.021*** (0.006)	0.039*** (0.007)	0.023*** (0.006)	0.017** (0.007)
Political Interest	Policy Interest	0.036** (0.011)	0.024*** (0.007)	0.045*** (0.007)	0.038*** (0.007)	0.039*** (0.007)	0.044*** (0.006)	0.045*** (0.006)
	Pol. Engagement	0.034** (0.011)	−0.004 (0.011)	−0.006 (0.011)	0.016 (0.012)	−0.026** (0.011)	−0.020 (0.012)	−0.004 (0.011)
Regime Sentiment	China Sentiment	0.005 (0.005)	0.005 (0.005)	0.010* (0.005)	0.015** (0.005)	0.003 (0.005)	0.005 (0.005)	0.000 (0.005)
	US Sentiment	0.025*** (0.007)	0.057*** (0.007)	0.040*** (0.008)	0.035*** (0.009)	0.049*** (0.009)	0.038*** (0.009)	0.053*** (0.009)
	Democracy View	0.058*** (0.007)	0.048*** (0.007)	0.060*** (0.007)	0.053*** (0.007)	0.052*** (0.007)	0.047*** (0.008)	0.051*** (0.007)

Notes: Two-way clustered standard errors (by user and date) in parentheses. Day −1 is the omitted reference period. All specifications include user fixed effects. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table A19: Event Study Estimates on Weibo: Effects Relative to Day −1