

Borrowed Battlegrounds: Global Audiences and Online Discussion of American Political Violence

Muzhi Liu *Columbia University*

Why do foreign audiences mobilize around American political violence, and does their engagement mirror domestic partisan dynamics? I analyze 1.2 million posts on X across five languages (Chinese, Japanese, Korean, French, Russian) surrounding two exogenous shocks: the Trump assassination attempt (July 2024) and the Kirk assassination (September 2025). Both events generate 50–100× increases in event-specific discussion, with non-U.S. users showing larger proportional volume gains than U.S. users—consistent with dramatic shocks activating otherwise-dormant distant audiences. Sentiment analysis reveals that the clearest effects appear in directional stance: violence produces pro-target sympathy shifts, with asymmetric spillover to adjacent partisan objects (e.g, Trump to GOP but not Kirk to Trump). Affective intensity and hostility decline on average, reflecting fatigue among returning users partially offset by positively-selected new entrants. U.S.-based users exhibit generalized activation across all partisan topics; non-U.S. users respond specifically to focal events. American partisan symbols are globally legible, but cross-border engagement remains episodic rather than systemic—activated by dramatic focal points without importing the connective tissue of the domestic partisan landscape.

Word Count: 9,649

Muzhi Liu is a PhD student at Columbia University, 420 W. 118th Street, New York, NY, 10027 (ml4967@columbia.edu).

I thank Shigeo Hirano, Donald Green, Eunji Kim, Yamil Velez, Carlo Prato, Naoki Egami, Junyan Jiang, Fred Gui, Patrick Liu, Kirill Chmel, Ronald Wu, Hanying Wei, Yinghui Zhou, Lesley Zhang, Alec Ewig, Daniel Markovits, and Emma Swanson for valuable comments and suggestions, as well as participants in the Kimuli Kasara Comparative Politics Seminar, Graduate Student Seminar, American Politics Graduate Seminar, and Mini-APSA at Columbia University. Any errors remain my own.

1 INTRODUCTION

On July 13, 2024, gunshots rang out at a Donald Trump rally in Butler, Pennsylvania. Within hours, social media filled with commentary from around the world: users writing in Chinese, Korean, Japanese, French, and Russian debated whether the attack was staged, analyzed video frames, and circulated memes portraying Trump as heroic or divinely protected. Many of these commentators were not U.S. citizens and had no vote in the upcoming election. Yet their reactions resembled domestic partisan conflict: rapid stance-taking, coalition signaling, and moralized interpretation of a violent event. This pattern raises a puzzle: why do geographically distant audiences mobilize around American political violence, and does their engagement replicate the systemic dynamics of U.S. partisanship or remain episodic and event-bound?

This paper develops and tests a two-stage framework of cross-border partisan expression. In Stage 1, salient U.S. events reach foreign-language communities through global media coverage amplified by diaspora-mediated transmission within language networks. In Stage 2, exposed users decide whether to participate and what to express based on how American partisan symbols resonate with their own identities. The framework generates a key distinction: users with material stakes in American politics—those who vote, pay taxes, or are directly affected by U.S. policy—are embedded in a connected partisan environment where shocks to one target can activate engagement across the entire landscape. Users without such stakes can engage intensely but episodically, responding to focal events without generalizing to broader partisan involvement. I call this contrast *scope of activation*: whether political shocks produce narrow, target-specific responses or broader engagement across related partisan objects.

I test this framework using 1.2 million posts on X surrounding two high-salience shocks: the assassination attempt on Donald Trump (July 2024) and the assassination of conservative activist Charlie Kirk (September 2025). For each event, I collect posts in a ± 7 day window across five language communities, constructing parallel corpora for event-specific discussion, entity-focused discussion of the focal target (Trump or GOP), cross-target discussion of adjacent partisan objects (GOP or Trump), and a placebo corpus centered on Biden and Democrats. I estimate causal effects using difference-in-differences designs that compare shifts in treatment corpora against placebo

trends, separately analyzing volume (corpus–language–day cells) and sentiment (user–day panels with LLM-coded measures of stance and affective tone).

Four findings emerge. First, both events generate massive cross-border mobilization—50- to 100-fold increases in event-specific discussion—with non-U.S. users showing larger proportional volume gains than U.S. users. Dramatic shocks serve as focal points that activate distant audiences who otherwise remain below the participation threshold. Second, violence produces detectable shifts in directional stance: sympathy effects toward targets, and in some cases, spillover to politically adjacent figures. The Trump assassination attempt activates pro-Trump sentiment that generalizes to broader GOP discussion; the Kirk assassination produces more localized effects that do not spill over to Trump. Third, affective tone does not mechanically intensify with attention. Average intensity and hostility decline in entity-focused discourse relative to placebo, reflecting a combination of returning-user fatigue and compositional change as new entrants—positively selected on identity resonance—partially offset but do not fully reverse the cooling effect among incumbents. Fourth, U.S. and non-U.S. users differ in scope of activation: U.S. users show generalized engagement that elevates even placebo discussion of Democrats after violence against Republicans, while non-U.S. users respond specifically to focal targets without such spillover.

The core finding is that American partisan symbols are globally legible—foreign audiences engage as active evaluators who take sides, express emotion, and signal coalition membership. Yet cross-border engagement is structurally shallower than domestic participation. Non-U.S. users exhibit high volume elasticity to dramatic events (they need focal points to enter the conversation) but narrow scope of activation (they respond to specific events without generalizing across the partisan landscape). U.S. users show the reverse pattern: lower volume elasticity (they engage with American politics routinely) but broader scope (shocks reverberate across the connected partisan system). American political conflict has been exported, but the connective tissue that links partisan symbols into a unified system remains specific to those with material stakes.

This paper contributes to research on political communication, comparative politics, and transnational information flows. First, it provides a theoretical framework that links cross-border diffusion to distinct behavioral margins—participation versus expression, focal versus spillover effects—showing that imported partisanship is not merely visible but evaluative. Second, it

introduces scope of activation as a diagnostic concept, demonstrating how differential placebo dynamics can reveal whether engagement is systemic or episodic. Third, it offers a scalable methodological approach for studying cross-lingual political behavior, combining large-scale multilingual data collection with validated LLM-based coding of stance and tone across five languages and over one million posts.

2 LITERATURE AND THEORY

Political Shocks and Cross-Border Attention

Sudden, high-salience events provide a window into how information environments respond to exogenous shocks. The focusing-events literature emphasizes that unexpected, highly visible shocks can reorganize attention and elevate issues on public agendas (Birkland 1998; Kingdon 1984). In domestic settings, political shocks can produce detectable changes even when baseline attitudes are entrenched (Krosnick and Petty 1995; Sears 1993). Recent work confirms this pattern for contemporary events: the 2020 George Floyd protests shifted police favorability among low-prejudice Americans (Reny and Newman 2021), while the July 2024 Trump assassination attempt reduced Republicans' support for partisan violence (Holliday et al. 2024).

Political information increasingly traverses national and linguistic borders via digital platforms (Bail 2021; Tucker et al. 2017). This paper extends the shock-response logic to an increasingly important but understudied setting: *foreign-language discourse about U.S. politics*. When U.S. political violence occurs, it can trigger immediate attention well beyond U.S. borders, but we know less about whether that attention translates into durable partisan-style engagement or remains concentrated on the event itself.

Existing research on transnational political influence often emphasizes elites, institutions, or coordinated campaigns (Elkins and Simmons 2005; King et al. 2017; Bradshaw and Howard 2019). State-sponsored disinformation campaigns have received substantial scholarly attention (Rid 2020; Starbird et al. 2019), but organic engagement by ordinary users remains understudied. This paper centers such engagement, examining how users become *distant witnesses* to foreign political conflict (Chouliaraki 2016; Robertson et al. 2023). The key outcome is not only whether

distant audiences talk about U.S. events, but whether they take *directional stances* that resemble partisan evaluation and whether those stances spill over across targets.

Stage 1: Exposure through Diaspora and Network Bridges

Cross-border political information does not diffuse uniformly. Network structure shapes what information reaches which communities and how quickly (Centola 2010; González-Bailón et al. 2011). Diaspora populations are plausibly important bridges because they are simultaneously embedded in U.S. political information environments and connected to heritage-language networks (Adamson 2016; Brinkerhoff 2009; Levitt 2001).

Research on diaspora behavior documents that immigrant communities maintain engagement with home-country politics even after extended residence abroad (Guarnizo et al. 2003; Waldinger 2015). Digital platforms reduce the costs of maintaining these transnational ties (Diminescu 2008; Komito 2011), increasing opportunities for U.S. political frames to appear in foreign-language discussions. Information can flow bidirectionally: diaspora users consume host-country political content and may introduce it into heritage-language networks (Eleta and Golbeck 2012).

The United States hosts substantial diaspora populations across the language communities in this study. When high-salience U.S. events occur, diaspora users encounter this information early—often through English-language sources—and may frame events in ways that shape how co-linguistic audiences understand American politics. The volume and speed of cross-border information flow thus depend on diaspora size, concentration, and platform activity, implying variation across language communities in baseline exposure to U.S. partisan content.

At the same time, very high-salience events can saturate global media and reduce the marginal role of diaspora channels (Wu 2000; Galtung and Ruge 1965). This suggests an important empirical tension: diaspora-linked exposure may matter most for lower-salience shocks or less globally recognized targets, while blockbuster events may diffuse broadly regardless of diaspora structure. That tension motivates cross-event and cross-language comparisons in both volume and content.

Stage 2: Identity-Based Expression, Stance, and Spillover Across Targets

Exposure does not determine expression. Conditional on encountering a shock, users still decide whether to participate and what to say. In U.S. contexts, partisanship operates as a social identity that shapes attention, evaluation, and expression (Green et al. 2002; Campbell et al. 1960; Huddy et al. 2015). Partisan identity fosters in-group cohesion while generating out-group animosity (Mason 2015; Iyengar and Westwood 2015; Iyengar et al. 2019), and creates a “perceptual screen” that makes individuals differentially attentive to identity-consistent information (Taber and Lodge 2006; Kunda 1990).

Identity-based politics is not only about greater participation; it also changes the *direction* of expressed evaluations. Strong partisan affect predicts more extreme expression: individuals holding intense negative feelings toward political opponents are more likely to share misinformation to attack out-groups (Osmundsen et al. 2021), use hostile language (Brady et al. 2017; Rathje et al. 2021), and express contempt rather than measured disagreement. Affective polarization—the tendency to view political opponents with fear and loathing rather than mere disagreement—has intensified these patterns (Iyengar et al. 2012; Finkel et al. 2020).

A central question for cross-border politics is whether U.S. partisan symbols function as portable identity markers. Scholars have noted that figures like Trump carry ideological meanings that transcend borders—attitudes toward immigration, nationalism, cultural traditionalism, and elite institutions (Norris and Inglehart 2019; Mudde 2019). This allows foreign users to express and perform identity by taking stances toward American figures, linking cross-border engagement to expressive theories of participation (Hillman 2010). Political expression can yield intrinsic benefits even without electoral stakes.

Whether these U.S.-derived patterns extend to foreign contexts is an open question. Kobayashi et al. (2024) find that users in Japan and Hong Kong do not exhibit the selective exposure patterns observed in America, attributing this to weaker affective polarization. However, their experiments use domestic party cues. When foreign audiences encounter American partisan cues—a globally legible symbolic system—different dynamics may emerge.

Instrumental Stakes and the Scope of Activation

The distinction between instrumental and expressive motivation (Downs 1957; Riker and Ordeshook 1968) helps clarify why cross-border engagement may be intense yet structurally different from domestic participation. Instrumental motivation arises from material stakes: voters seek policy outcomes, activists seek career advancement, citizens seek to influence collective decisions. Expressive motivation arises from identity: participation is intrinsically valuable as a signal of who one is, regardless of consequences (Hillman 2010).

U.S.-based users (including diaspora) have instrumental stakes: they vote in U.S. elections, pay U.S. taxes, and are directly affected by U.S. policy. Their engagement with American partisan conflict is embedded in a broader context of material interest and social coordination. Non-U.S. users generally lack these stakes; their participation is purely expressive. They cannot vote, are minimally affected by U.S. policy, and have no instrumental reason to engage with American political events.

The absence of instrumental motivation among distant witnesses does not imply weaker engagement—indeed, users who participate despite lacking material stakes may be especially strongly identified with the symbolic content of American partisanship (Passini 2017). However, it may imply *structurally different* engagement. Instrumental participants must navigate the entire partisan landscape continuously; expressive participants can engage episodically, responding to specific events without maintaining ongoing engagement with the broader political system (Schudson 1998; Zaller 1992).

This paper therefore emphasizes *scope of activation*: whether a shock targeting one partisan side spills into broader engagement with other partisan targets (including the opposing side). Empirically, the placebo corpus provides a meaningful benchmark for detecting scope differences: if a shock against Republicans increases discussion of Democrats among some users, that indicates generalized activation beyond the focal target. Conversely, stable placebo dynamics alongside sharp focal responses indicates targeted, event-bound engagement. This scope distinction organizes the heterogeneity analyses by user location and motivates the spillover tests across corpora.

A Two-Stage Model of Cross-Border Partisan Expression

The literatures reviewed above motivate a two-stage account of cross-border partisan expression. I model how a salient U.S. political shock changes (i) *who posts* (participation) and (ii) *what they say* (affective tone and directional stance).

Timing and shock. Time is indexed by $t \in \{-7, \dots, 7\}$, with an unexpected political shock at $t = 0$ of magnitude $S > 0$. Let $\mathbb{1}_{post,t} \equiv \mathbb{1}[t \geq 0]$ denote the post-shock period.

Users. Users are indexed by i . Each user belongs to a language community $l \in \{zh, ja, ko, fr, ru\}$ and a location group $g \in \{US, Non-US\}$. User i has:

- *Partisan resonance* $\theta_i \in [0, 1]$: how identity-relevant U.S. partisan conflict is to user i .
- *Baseline evaluations* m_{ik}^{pre} for each political target k .

Corpus topics vs. political targets. I distinguish between *corpus topics* (what a post is about) and *political targets* (whom the post evaluates):

- **Corpus topics** $j \in \{E, F, R, O\}$: event-specific discussion (E), focal entity (F), related entity (R), and opposing/placebo topic (O).
- **Political targets** $k \in \{F, R, O\}$: focal, related, and opposing targets toward which users may express directional stances.

A post in corpus j may express stances toward multiple targets k . Let $w_{jk} \geq 0$ denote the salience weight of target k within corpus topic j .

Exposure. The shock reaches user i with probability

$$\Pr(Z_i = 1) = \pi_{il} = \pi(D_l, S), \quad \frac{\partial \pi}{\partial D_l} > 0, \quad \frac{\partial \pi}{\partial S} > 0, \quad (1)$$

where D_l measures diaspora presence in language community l . If $Z_i = 0$, the user does not respond to the shock (i.e., their perceived environment remains at the pre-shock level). For

compactness, define

$$\mathbb{1}_{it} \equiv \mathbb{1}_{post,t} \cdot Z_i.$$

Topic salience. Let S_{igjt} denote the salience of corpus topic j for user i (in group g) at time t :

$$S_{igjt} = S_j^{pre} + \eta_{jg} S \cdot \mathbb{1}_{it}, \quad (2)$$

where $\eta_{Eg} = 1$, $\eta_{Fg}, \eta_{Rg} \in (0, 1)$, and $\eta_{Og} = \sigma_g \geq 0$ captures the *scope of activation* into placebo/opposing topics, potentially varying by location group g .

Stage 2: Expression Conditional on Posting Conditional on posting in corpus topic j at time t , user i chooses (a) an affective tone intensity $e \geq 0$ and (b) directional stances $s_k \in \{-1, 0, +1\}$ toward each target $k \in \{F, R, O\}$.

Affective tone. Tone utility is

$$U_{igjt}^{tone}(e) = (\alpha_g + \beta \theta_i S_{igjt}) e - \frac{\gamma}{2} e^2 - c, \quad (3)$$

where α_g is a group-specific baseline incentive, $\beta > 0$ scales how resonance amplifies topic salience, $\gamma > 0$ governs the marginal cost of intensity, and $c > 0$ is a fixed cost of producing a post.

The optimal tone is

$$e_{igjt}^* = \max \left\{ 0, \frac{\alpha_g + \beta \theta_i S_{igjt}}{\gamma} \right\}. \quad (4)$$

Substituting (4) into (3) yields the indirect tone payoff

$$V_{igjt}^{tone} = \frac{(\max\{0, \alpha_g + \beta \theta_i S_{igjt}\})^2}{2\gamma} - c. \quad (5)$$

Directional stance. The shock shifts evaluations of political targets:

$$m_{ikt} = m_{ik}^{pre} + (\chi_k + \lambda_{F \rightarrow k} \chi_F) S \cdot \mathbb{1}_{it}, \quad \lambda_{F \rightarrow F} = 0. \quad (6)$$

Here χ_k is a direct effect on target k , and $\lambda_{F \rightarrow k}$ captures affective spillover from the focal target.

Conditional on posting in corpus j , stance utility is separable across targets:

$$U_{ijt}^{stance}(\mathbf{s}) = \sum_{k \in \{F, R, O\}} [w_{jk} s_k m_{ikt} - \kappa \mathbb{1}[s_k \neq 0]], \quad (7)$$

where $\kappa > 0$ is the fixed cost of taking a non-neutral stance toward a target.

Maximization is target-by-target. The optimal stance toward k (when posting in corpus j) is

$$s_{ijkt}^* = \begin{cases} +1 & \text{if } w_{jk} m_{ikt} > \kappa, \\ 0 & \text{if } |w_{jk} m_{ikt}| \leq \kappa, \\ -1 & \text{if } w_{jk} m_{ikt} < -\kappa. \end{cases} \quad (8)$$

Stage 1: Participation and Topic Choice

Posting rule. Conditional on exposure ($Z_i = 1$), user i posts at time t if and only if the best available corpus topic delivers nonnegative indirect payoff:

$$\text{Post}_{it} = 1 \iff \max_{j \in \{E, F, R, O\}} V_{igjt}^{tone} \geq 0, \quad (9)$$

and when posting the user chooses

$$j_{it}^* \in \arg \max_{j \in \{E, F, R, O\}} V_{igjt}^{tone}.$$

Users with $Z_i = 0$ behave as in the pre-shock environment (i.e., $\mathbb{K}_{it} = 0$).

Resonance cutoff. For a given topic j , a sufficient and necessary condition for $V_{igjt}^{tone} \geq 0$ is

$$\alpha_g + \beta \theta_i S_{igjt} \geq \sqrt{2\gamma c}.$$

Equivalently, posting in j is feasible if

$$\theta_i \geq \bar{\theta}_{igjt} \equiv \max \left\{ 0, \frac{\sqrt{2\gamma c} - \alpha_g}{\beta S_{igjt}} \right\}, \quad (10)$$

(where the outer $\max\{\cdot, 0\}$ handles the case $\alpha_g \geq \sqrt{2\gamma c}$, in which the cutoff is nonbinding).

Comparative Statics and Empirical Predictions The model yields five predictions, each tied to interpretable parameters.

Prediction 1: Volume increases with shock magnitude and topic salience. From (10), the resonance cutoff $\bar{\theta}_{igjt}$ falls as $\mathcal{S}_{igjt}$ rises. For exposed users in the post period ($\kappa_{it} = 1$), (2) implies $\mathcal{S}_{igjt}$ increases in S , hence participation expands:

$$\frac{\partial \bar{\theta}_{igjt}}{\partial S} = -\frac{\eta_{jg}(\sqrt{2\gamma c} - \alpha_g)}{\beta \mathcal{S}_{igjt}^2} < 0 \quad \text{whenever } \sqrt{2\gamma c} > \alpha_g \text{ and } \eta_{jg} > 0.$$

Because only a fraction π_{il} are exposed, reduced-form effects are attenuated by π_{il} .

Prediction 2: Event-specific discussion dominates. Since $\eta_{Eg} = 1 > \eta_{Fg}, \eta_{Rg}$, the shock-induced increase in $\mathcal{S}_{igjt}$ (and thus in V_{igjt}^{tone}) is largest for $j = E$. Post-shock posting therefore concentrates most strongly in event-specific discussion.

Prediction 3: Extensive-margin selection on resonance (and ambiguous average intensity). Posting is increasing in θ_i , but the shock *lowers* the cutoff (10), bringing in users with θ_i between the old and new thresholds. Thus, compared to pre-shock posters, *marginal entrants have lower θ_i* (though still higher than persistent non-posters).

For intensity, (4) shows e_{igjt}^* rises mechanically with $\mathcal{S}_{igjt}$ for inframarginal posters, while entrants tend to have lower θ_i . The net effect on average tone is therefore composition-dependent.

Prediction 4: Location heterogeneity in scope of activation. The scope parameter σ_g governs how much the shock activates the placebo/opposing corpus O via $\eta_{Og} = \sigma_g$. If $\sigma_{\text{US}} > \sigma_{\text{Non-US}}$, then U.S. users exhibit *generalized activation* (volume increases even in O), while non-U.S. users exhibit more *targeted activation* concentrated in E, F, R .

Prediction 5: Stance shifts and cross-target spillover. From (6) and (8), the shock moves users across stance thresholds $\pm\kappa/w_{jk}$. Three parameters govern the pattern:

- $\chi_F > 0$: increased sympathy toward the focal target.
- $\lambda_{F \rightarrow R} > 0$: spillover to related targets.
- $\lambda_{F \rightarrow O} \approx 0$: little/no spillover to opposing targets.

3 DATA AND METHODS

Data Collection

This study analyzes how non-U.S. audiences engage with U.S. partisan politics online following major political shocks. I focus on cross-lingual discourse on X, a platform where users broadcast short messages and interact via replies, quotes, and reposts without requiring mutual consent for network formation (Bossetta 2018; Theocharis et al. 2022). This setting is well-suited for observing “borrowed” political conflict: users who cannot vote in U.S. elections can nonetheless perform partisan alignment, target political outgroups, and repurpose U.S. events to comment on politics elsewhere.

I examine two high-salience events that are closely related but analytically distinct: the assassination attempt on Donald Trump on July 13, 2024, and the assassination of Charlie Kirk on September 10, 2025. The first event centers on a single political figure (Trump), whereas the second plausibly activates broader partisan identity associated with the Republican coalition. Studying both allows me to distinguish candidate-centered reactions (stance toward Trump) from party-centered reactions (stance toward Republicans as a group), and—critically—to test whether shocks to one target spill over to affect sentiment toward the other.

For each event, I collect posts from a symmetric window of ± 7 days around the event date (15 days total, including the event day), using X’s full-archive search API.¹ Data collection occurred on December 20, 2025, yielding approximately 1.2 million posts across both events. The unit of collection is a language–day–query–family cell. I study five language communities: Chinese (zh), Korean (ko), Japanese (ja), French (fr), and Russian (ru). These languages were selected to vary on two theoretically relevant dimensions. First, they represent diverse geopolitical relationships to the United States—from close allies (Japan, South Korea, France) to strategic competitors

¹I use X’s Pro tier, which provides full-archive search and a monthly cap of 1 million posts. When this cap was exceeded, I purchased additional retrieval quota through X’s top-up service.

(China, Russia)—which may shape how U.S. partisan symbols are interpreted. Second, they differ substantially in U.S. diaspora size: Chinese speakers constitute the largest group (5.4 million), followed by French (2.1 million), Japanese (1.5 million), Korean (1.1 million), and Russian (0.9 million). This variation enables direct tests of the diaspora exposure mechanism outlined in the theory.

Within each language and day, I execute four query families designed to capture event-specific attention, baseline discussion of the focal entity, cross-target spillover, and a placebo comparison:

1. **Event corpus:** Posts mentioning both event-related terms (e.g., “assassination,” “shooting”) and the focal entity (Trump for the 2024 attempt; Kirk or the Republican Party for the 2025 assassination). This corpus captures direct engagement with the political shock.
2. **Entity-focal corpus:** Posts mentioning the focal entity but explicitly excluding event-related terms, capturing baseline partisan discussion of the primary target within the same calendar window.
3. **Entity-cross corpus:** Posts mentioning the non-focal entity (GOP for the Trump event; Trump for the Kirk event), excluding event-related terms. This corpus enables testing of cross-target spillover: whether a candidate-centered shock activates party-level sentiment, and vice versa.
4. **Placebo corpus:** Posts mentioning a comparison topic (Biden or the Democratic Party) to assess whether observed changes reflect a general attention shock rather than an event-specific reaction.

Table 1 summarizes the query structure. The symmetric design—collecting both Trump and GOP discussion for each event—allows me to estimate whether political shocks produce narrow, target-specific effects or broader activation of partisan sentiment across related political objects.

To ensure reproducibility under API rate limits, I retrieve posts at the day level and de-duplicate by post ID within each UTC day. I retain original posts, replies, and quote posts, which are central to measuring targeted stances and outgroup hostility, and I exclude simple retweets that contain no original text. When daily volume exceeds a pre-specified operational threshold of 15,000 posts, I employ stratified time-slice sampling: I randomly draw eight one-hour bins per day, allocating two bins to each of four six-hour blocks (00:00–05:59, 06:00–11:59, 12:00–17:59, 18:00–23:59

TABLE 1. Query Families by Event

Query Family	Trump Attempt (Jul. 2024)	Kirk Assassination (Sep. 2025)
Event	$\text{EVENT} \wedge \text{TRUMP}$	$\text{EVENT} \wedge (\text{KIRK} \vee \text{GOP})$
Entity-focal	$\text{TRUMP} \wedge \neg \text{EVENT}$	$\text{GOP} \wedge \neg \text{EVENT} \wedge \neg \text{KIRK}$
Entity-cross	$\text{GOP} \wedge \neg \text{EVENT}$	$\text{TRUMP} \wedge \neg \text{EVENT} \wedge \neg \text{KIRK}$
Placebo	BIDEN/DEMS	BIDEN/DEMS

Notes: Each query is executed separately by language (`lang:zh`, `lang:ko`, etc.) and UTC day. `EVENT` denotes assassination- or attack-related terms. `TRUMP`, `KIRK`, and `GOP` denote name variants for each entity across languages. `BIDEN/DEMS` includes Biden name variants and Democratic Party terms. Negation ($\neg$) is implemented term-by-term rather than on grouped expressions to ensure correct API behavior. Full multilingual term lists appear in Appendix D.

UTC). This procedure reduces the risk that sampling systematically over-represents particular times of day.

Because the research question concerns engagement outside the United States, language serves as a useful but imperfect proxy for user location. Diaspora users—Chinese speakers in California, Russian speakers in Brooklyn—post in their heritage language but reside in the U.S. information environment. I therefore treat location as a measured attribute rather than an assumption. X users optionally self-report location in free text, yielding genuine geographic strings (e.g., “Paris”), missing entries, and fabricated values. I standardize user-reported locations using an LLM-assisted procedure and classify them into US and non-US IP categories. This classification serves two purposes: it enables descriptive comparisons of diaspora versus non-diaspora users, and it supports robustness checks excluding users with US locations.

Sample Descriptive Statistics

The final sample comprises 1,204,745 posts from 366,045 unique users across both events. Table 2 presents summary statistics by event, language, and corpus type.

The Trump event generates approximately 50% more posts than the Kirk event, consistent with Trump’s greater global salience. Japanese-language discussion dominates the sample, followed by French and Chinese. Language communities differ markedly in diaspora composition: Chinese-language posts show the highest US share (16.7%), reflecting the large Chinese-speaking population in the United States, while Japanese-language posts show the lowest (1.1%). Approx-

TABLE 2. Sample Descriptive Statistics

	Posts	Users	% US	% Non-US
<i>Panel A: By Event</i>				
Trump Assassination Attempt (Jul 2024)	724,754	252,640	5.0	36.7
Kirk Assassination (Sep 2025)	479,991	163,239	4.3	32.3
<i>Panel B: By Language</i>				
Japanese	521,198	178,223	1.1	36.7
French	337,772	111,056	3.3	40.2
Chinese	205,007	45,867	16.7	21.9
Korean	99,057	19,521	2.7	32.1
Russian	41,711	12,378	6.6	41.0
<i>Panel C: By Corpus</i>				
Event Corpus	145,483	66,100	4.4	39.8
Entity-Focal	351,780	163,117	6.0	34.2
Entity-Cross	327,125	143,381	4.1	35.3
Placebo (Biden/Dems)	380,357	140,970	3.9	34.3

Notes: % US and % Non-US reflect location classification from user profiles; remaining users have unknown or ambiguous locations. Entity-Cross combines GOP corpus (Trump event) and Trump corpus (Kirk event).

imately 35% of posts come from users with identifiable non-US locations, 5% from US-based users, and 60% from users with unknown or ambiguous location information. The high share of unknown locations is typical of X data and motivates the location-based heterogeneity analysis as a robustness check rather than a primary specification.

The event corpus exhibits the sharpest pre-post contrast. Across both events and all languages, event-specific discussion is nearly absent in the pre-period (averaging fewer than 100 posts per language-day) and surges immediately following each shock. The placebo corpus, by contrast, shows relatively stable volume across the event window, supporting its use as a counterfactual for difference-in-differences estimation.

Constructing Sentiment Variables

The research design requires both volume and sentiment outcomes. Volume outcomes (post counts, unique users) are computed directly from the full sample of approximately 1.2 million posts. Sentiment outcomes require text-based content coding, which I implement using large language models (LLMs) on a stratified subsample.

Stratified Sampling for Sentiment Coding. To balance statistical power against coding costs, I draw a stratified random sample for LLM labeling. Strata are defined by the cells of the research design: event (2) $\times$ language (5) $\times$ corpus (4) $\times$ day (15: days -7 to $+7$) $\times$ user location (3: US/non-US/Unknown), yielding 1,350 strata. I sample up to 300 posts per stratum, or all available posts if the stratum contains fewer than 300. This procedure yields approximately 300,000 posts for sentiment coding, ensuring balanced representation across events, languages, corpora, days, and user types. The daily sampling structure supports event-study analyses that trace the evolution of sentiment before and after each shock.

Coding Framework. The goal of automated coding is not to recover a single “overall sentiment” score, but to measure targeted partisan expression on dimensions that map directly to the theoretical model. As a result, the coding scheme captures both the direction of partisan stance and the intensity of expression. A single post may praise Trump while criticizing other Republicans, or express anti-Democratic hostility without explicitly referencing Trump. The coding scheme is therefore multi-target: separate variables capture stance toward each political object.

Table 3 summarizes the coding framework. Partisan stance is coded on a directional scale ($-1/0/+1$) to capture both valence and neutrality. Expression intensity and outgroup hostility are coded on ordinal scales (0–3) to capture variation in the strength of partisan expression—a key prediction of the model is that users with higher partisan resonance θ_i express more intensely.

Implementation and Validation. Automated coding is inherently imperfect for nuanced political language, especially in a cross-lingual setting where sarcasm, slang, and context-specific references are common. I mitigate these risks through conservative label definitions, language-specific prompts that include culturally appropriate examples, and validation against expert annotation.

I employ GPT-4.1 for automated content coding. Each post is processed with a prompt that provides the post text, the event context (Trump attempt or Kirk assassination), and coding instructions with few-shot examples in the post’s language. The model returns structured output for all variables. To validate LLM labels, I draw a separate random sample of 200 posts per language per event (2,000 posts total) for expert coding. Validation results show strong agreement between LLM and expert labels, with exact match accuracy exceeding 94% and Cohen’s κ ranging

TABLE 3. LLM Content Coding Framework

Variable	Coding Rules
<i>Partisan Stance (Direction)</i>	
stance_trump	-1 = anti-Trump; 0 = neutral/not mentioned; +1 = pro-Trump
stance_gop	-1 = anti-Republican; 0 = neutral/not mentioned; +1 = pro-Republican
stance_dem	-1 = anti-Democrat; 0 = neutral/not mentioned; +1 = pro-Democrat
<i>Expression Intensity</i>	
intensity	Strength of expressed sentiment regardless of direction: 0 = neutral/factual reporting; 1 = mild opinion; 2 = strong opinion; 3 = extreme/inflammatory
<i>Outgroup Hostility</i>	
hostility	Hostility toward political opponents: 0 = none; 1 = mild criticism; 2 = strong criticism/mockery; 3 = dehumanizing or hateful language
<i>Content Type</i>	
US_relevant	1 if post actually discusses U.S. politics; 0 otherwise

Notes: Each post receives codes for all variables. Stance variables are directional (-1/0/+1); intensity and hostility are ordinal (0-3). The scheme is multi-target: a single post can receive non-zero stance codes for multiple political objects (e.g., pro-Trump and anti-Democrat).

from 0.88 to 0.95 across all variables and languages.²

Identification Strategy

I test the model's predictions through difference-in-differences designs that exploit exogenous shocks to estimate treatment effects on participation (volume) and expression (sentiment). This section describes the core estimand, the four main analyses, and the validity of the placebo corpus.

Core Estimand The key estimand across all specifications is the *relative* change in a treatment corpus compared to a politically adjacent placebo corpus (Biden/Democrats) around an exogenous shock. For any outcome Y (log volume or user-day sentiment), define the pre/post change within corpus $c \in \{T, P\}$ as

$$\Delta^c \equiv \bar{Y}_{\text{post}}^c - \bar{Y}_{\text{pre}}^c,$$

²Complete implementation details, prompts, few-shot examples, and validation results are provided in the Appendix. The LLM labeling pipeline has been pre-registered at AsPredicted.org (AsPredicted #266,281).

FIGURE 1. Scope-of-activation schematic. When the placebo corpus moves (generalized activation), the DiD coefficient captures a *relative* shift between treated and placebo topics; when the placebo corpus is stable (targeted activation), DiD more closely approximates the absolute change in the treated topic.

	Placebo corpus (Biden/Dems)	Treatment corpus (Event/Entity)
Targeted activation (Non-US)	$\Delta_{\text{Non-US}}^P \approx 0$	$\Delta_{\text{Non-US}}^T \neq 0$
Generalized activation (US)	$\Delta_{\text{US}}^P \neq 0$	$\Delta_{\text{US}}^T \neq 0$
	$\text{DiD}_g = \Delta_g^T - \Delta_g^P$	$\beta_4 = \text{DiD}_{\text{Non-US}} - \text{DiD}_{\text{US}}$

where T denotes a treatment corpus (event or entity) and P denotes the placebo corpus. The canonical difference-in-differences estimand is

$$\text{DiD} \equiv \Delta^T - \Delta^P.$$

Thus, the DiD coefficient identifies *how much more* the treated topic shifts relative to the placebo topic over the same window, under the assumption of parallel trends.

Importantly, placebo corpus movement is not nuisance variation—it is a diagnostic of *scope of activation*. Whether an event activates only the focal target (targeted activation) or also adjacent partisan topics (generalized activation) is central to this paper’s theory. This interpretation is most transparent in the location heterogeneity design. Let $g \in \{\text{US}, \text{Non-US}\}$ index user location groups. The triple-difference coefficient can be written as

$$\beta_4 \equiv \text{DiD}_{\text{Non-US}} - \text{DiD}_{\text{US}} = \left(\Delta_{\text{Non-US}}^T - \Delta_{\text{Non-US}}^P \right) - \left(\Delta_{\text{US}}^T - \Delta_{\text{US}}^P \right).$$

If treatment-corpus changes Δ_g^T are similar across groups, then β_4 is driven primarily by differences in placebo movement Δ_g^P —reflecting different scope of activation across user types. Figure 1 illustrates this logic.

Specifications I estimate four main analyses, progressing from baseline event effects to heterogeneity by location, language, and cross-target spillover. Throughout, I use two treatment corpora: the *event corpus* (posts mentioning assassination-related terms alongside the focal entity) and the *entity corpus* (baseline discussion of Trump or GOP without event-related terms). The event corpus exhibits near-zero pre-period volume and massive post-event spikes; the entity corpus

provides more balanced samples for parallel trends assessment.³

Analysis 1: Baseline Event Effects. For volume, the unit of analysis is the corpus–language–day cell:

$$\log(N_{jlt} + 1) = \alpha + \beta_1 \text{Treat}_j + \beta_3 (\text{Treat}_j \times \text{Post}_t) + \lambda_l + \tau_t + \epsilon_{jlt}, \quad (11)$$

where N_{jlt} is the count of posts in corpus j , language l , day t ; Treat_j indicates the treatment corpus; Post_t indicates days after the shock; and λ_l and τ_t are language and day fixed effects. Standard errors are clustered by language and day. The model predicts $\beta_3 > 0$: higher salience lowers the participation threshold, drawing more users into discourse.

For sentiment, the unit of analysis is the user-day within the entity corpus:

$$Y_{ilt} = \alpha + \beta_1 \text{Treat}_j + \beta_3 (\text{Treat}_j \times \text{Post}_t) + \lambda_l + \tau_t + \mathbf{X}'_{it} \boldsymbol{\gamma} + \epsilon_{ilt}, \quad (12)$$

where Y_{ilt} is intensity (0–3), hostility (0–3), or stance (–1/0/+1) for user i in language l on day t ; and $\mathbf{X}_{it}$ includes logged follower count and account age. Standard errors are clustered at the user level. The model predicts $\beta_3 < 0$ if fatigue among returning users dominates, but β_3 could be attenuated or positive if new entrants with high intensity raise the average.

Analysis 2: Location Heterogeneity. The model distinguishes U.S.-based users (instrumental motivation) from non-U.S. users (expressive motivation only). I test for differential responses using triple-difference specifications that map directly onto the scope-of-activation logic in Figure 1.

For volume, the unit of analysis is the corpus–language–day–location cell:

$$\log(N_{jltg} + 1) = \alpha + \beta_3 (\text{Treat}_j \times \text{Post}_t) + \beta_4 (\text{Treat}_j \times \text{Post}_t \times \text{NonUS}_g) + \boldsymbol{\delta}' \mathbf{Z} + \lambda_l + \tau_t + \epsilon_{jltg}, \quad (13)$$

where $g \in \{\text{US}, \text{Non-US}\}$ indexes user location and $\mathbf{Z}$ includes all lower-order interactions. The coefficient β_3 captures the DiD effect for US-based users; β_4 captures the *additional* effect for non-US users. The model predicts $\beta_4 > 0$: dramatic events disproportionately activate distant

³Volume results use the event corpus in the main text (where the dramatic spike is substantively important) and both corpora for robustness. Sentiment results use the entity corpus throughout due to pre-period sample limitations in the event corpus.

audiences who otherwise remain below the participation threshold, while US-based users engage more routinely.

For sentiment, the specification parallels the volume model at the user-day level:

$$Y_{ilt} = \alpha + \beta_3(\text{Treat}_j \times \text{Post}_t) + \beta_4(\text{Treat}_j \times \text{Post}_t \times \text{NonUS}_i) + \delta'Z + \lambda_l + \tau_t + X'_{it}\gamma + \epsilon_{ilt}. \quad (14)$$

The interpretation of β_4 here concerns *scope*: if US users exhibit generalized activation (moving the placebo corpus), their *relative* sentiment shift will be smaller than that of non-US users who exhibit targeted activation (stable placebo corpus).

Analysis 3: Language Heterogeneity. Diaspora size varies across language communities (Chinese: 5.4M; French: 2.1M; Japanese: 1.5M; Korean: 1.1M; Russian: 0.9M in the U.S.). I estimate language-specific effects using French as the baseline:

$$\log(N_{jlt} + 1) = \alpha + \beta_3^{\text{fr}}(\text{Treat}_j \times \text{Post}_t) + \sum_{l \neq \text{fr}} \beta_{3,l}(\text{Treat}_j \times \text{Post}_t \times \mathbf{1}[l]) + \delta'Z + \tau_t + \epsilon_{jlt}. \quad (15)$$

The total effect for language l is $\beta_3^{\text{fr}} + \beta_{3,l}$. I correlate these language-specific effects with pre-period diaspora activity on X to test the exposure mechanism. The model predicts that high-diaspora communities show larger volume increases, though this relationship may be attenuated for highly salient events that achieve saturation coverage through mainstream media.

Analysis 4: Cross-Target Spillover. The symmetric design—collecting both Trump and GOP discussion for each event—enables tests of whether shocks to one political object spill over to related objects. For each event, I estimate separate DiD models for the focal corpus (direct target) and the cross corpus (related partisan object):

$$\log(N_{jlt}^{\text{focal}} + 1) = \alpha + \beta_3^{\text{focal}}(\text{Focal}_j \times \text{Post}_t) + \lambda_l + \tau_t + \epsilon_{jlt}, \quad (16)$$

$$\log(N_{jlt}^{\text{cross}} + 1) = \alpha + \beta_3^{\text{cross}}(\text{Cross}_j \times \text{Post}_t) + \lambda_l + \tau_t + \epsilon_{jlt}. \quad (17)$$

For the Trump event, focal is Trump and cross is GOP; for Kirk, the reverse. Asymmetric spillover—Trump events affecting GOP discourse ($\beta_3^{\text{cross}} > 0$) but Kirk events not affecting Trump discourse ($\beta_3^{\text{cross}} \approx 0$)—would indicate that cross-border engagement is candidate-centered rather than coalition-general, anchored by globally salient individuals.

Validity of the Placebo Corpus The difference-in-differences design assumes that the placebo corpus (Biden/Democrats) provides a valid counterfactual for how treatment corpus discussion would have evolved absent the event. Several design features support this assumption: First, the placebo corpus is thematically related but causally distant: Biden and Democrats are not directly implicated in assassination attempts against Republicans, yet they occupy the same broad domain (American politics) and thus absorb general fluctuations in cross-border political attention. A completely unrelated placebo would fail to control for politically relevant time trends. Second, the query structure explicitly excludes event-related terms from the placebo corpus, ensuring that posts directly referencing the assassination do not contaminate the control group. Third, the symmetric ± 7 day window and day fixed effects absorb day-specific shocks affecting all corpora equally, such as platform-wide activity changes or unrelated news events.

The key identifying assumption is parallel trends: absent the event, treatment and placebo corpora would have evolved similarly. For sentiment analysis, I assess this through event-study specifications estimating day-specific treatment effects relative to the pre-period average (days -7 to -1). This reference choice is more robust to single-day noise than the conventional day -1 reference. For volume analysis, day -1 is set as the reference day. Pre-event coefficients clustering near zero support parallel trends; systematic pre-trends would raise concerns.

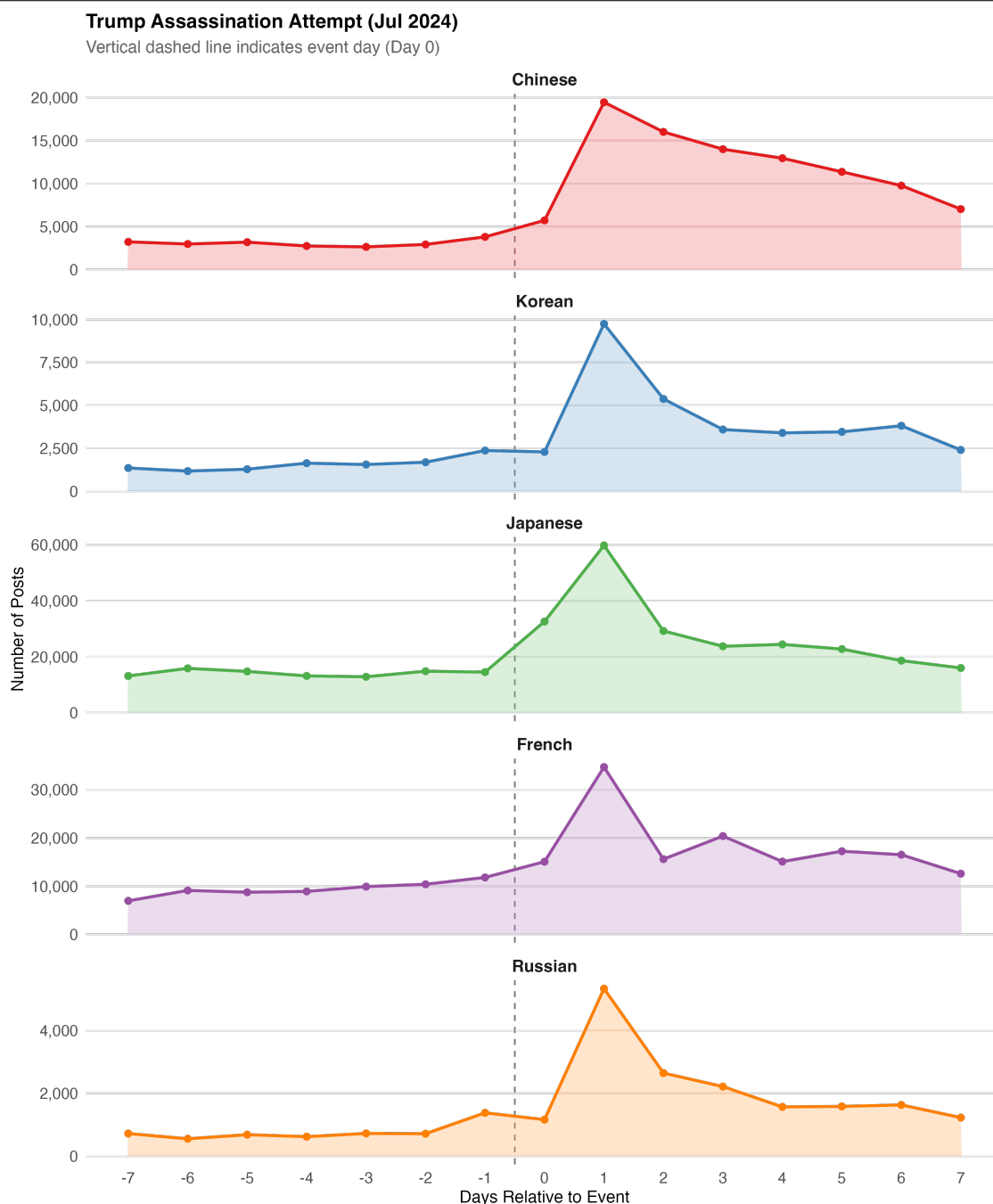
4 RESULTS

Time-series Patterns

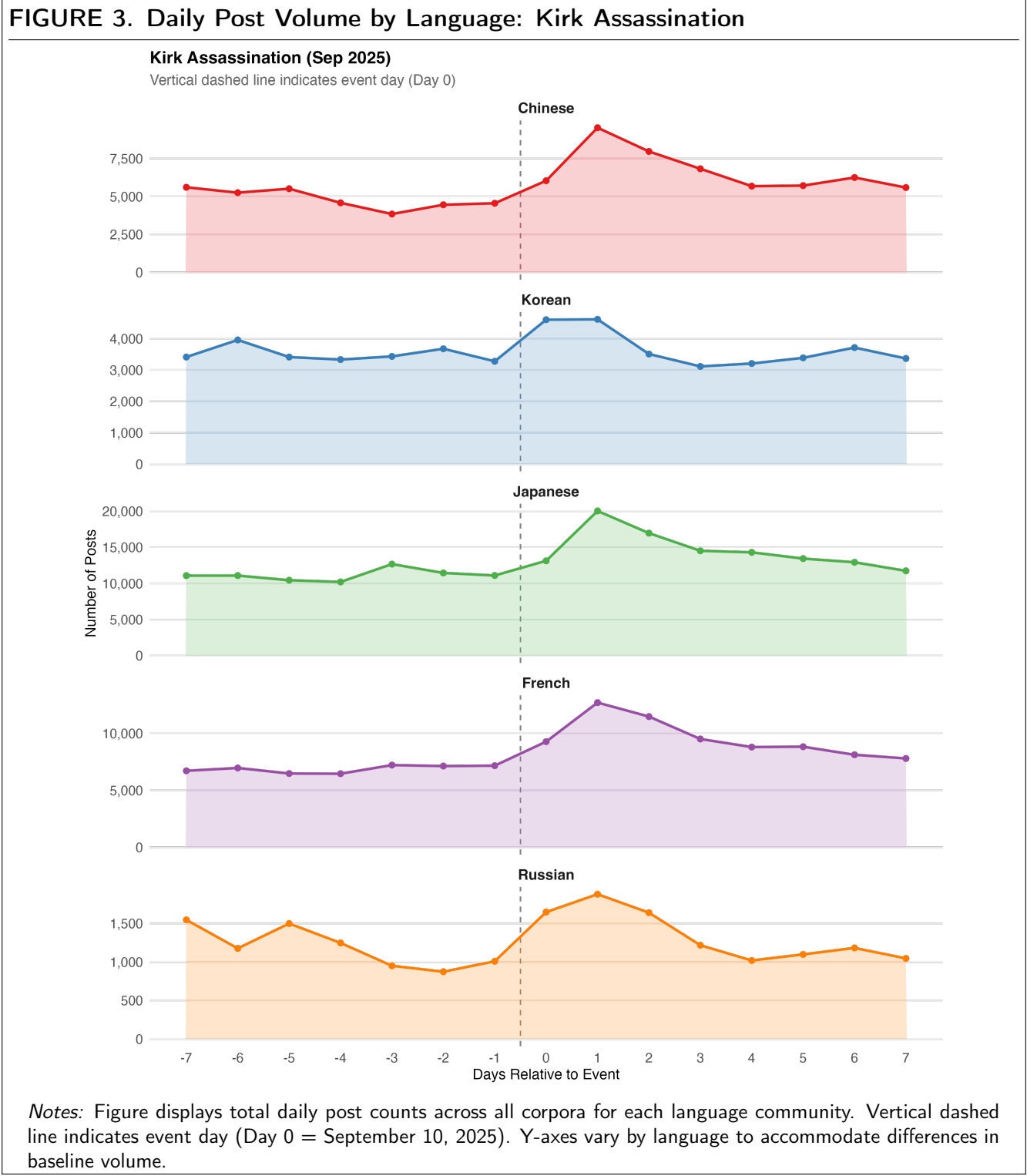
Before turning to causal estimates, I present time-series patterns in post volume across language communities. Figures 2 and 3 display daily post counts for each language group in the ± 7 day window around each event.

Several patterns emerge. First, the Trump event generates sharp volume spikes across all five

FIGURE 2. Daily Post Volume by Language: Trump Assassination Attempt



Notes: Figure displays total daily post counts across all corpora (event, entity-focal, entity-cross, and placebo) for each language community. Vertical dashed line indicates event day (Day 0 = July 13, 2024). Y-axes vary by language to accommodate differences in baseline volume.



language communities. Japanese-language discussion exhibits the largest absolute increase, rising from approximately 15,000 posts per day to over 60,000 at peak—a four-fold increase. Chinese, French, and Korean communities show similar proportional responses, while Russian-language volume increases modestly from a lower baseline. Second, the Kirk event produces noticeably weaker responses. Japanese and French communities show discernible spikes, but Korean-language volume remains essentially flat, suggesting that Kirk—a figure prominent in U.S. conservative media but largely unknown internationally—failed to penetrate Korean information networks. Third, decay patterns differ across events: the Trump spike dissipates gradually over the post-event week, while the Kirk spike attenuates more rapidly, consistent with lower sustained interest in a less globally salient figure.

These descriptive patterns preview the formal results. The cross-language variation in responsiveness, particularly Korean’s non-response to Kirk, suggests that diaspora-mediated exposure plays a key role in determining which events reach which communities. The volume analysis that follows isolates these effects using the placebo corpus as a counterfactual.

Volume Analysis

I first examine volume outcomes using the full sample of approximately 1.2 million posts. Table 4 presents the main difference-in-differences estimates.

Event Effects. Both events generate massive increases in event-specific discussion. The DiD estimates for the event corpus are 4.05 log points for Trump and 4.62 log points for Kirk (both $p < 0.001$), corresponding to roughly 50- to 100-fold increases in daily post volume relative to placebo. Effects on the entity-focal corpus are smaller but significant: 1.20 log points for Trump ($p < 0.01$) and 0.28 log points for Kirk ($p < 0.01$). The larger entity spillover for Trump suggests that his global brand recognition translates event-specific attention into broader political engagement, whereas Kirk—less known internationally—activates event-specific discussion without substantial spillover.

Figure 4 presents event-study estimates. Pre-event coefficients are stable near zero, supporting the parallel trends assumption required for causal identification. Both events show sharp disconti-

TABLE 4. Event Effects on Volume (Difference-in-Differences)

	Trump Event		Kirk Event	
	Event Corpus	Entity-Focal	Event Corpus	Entity-Focal
Treat	-5.142*** (0.297)	-0.555 (0.278)	-5.602*** (0.570)	-1.838* (0.644)
Treat × Post	4.048*** (0.399)	1.200** (0.183)	4.622*** (0.447)	0.280** (0.037)
Observations	145	150	140	150
R^2 (within)	0.912	0.392	0.902	0.579
Language FE	Yes	Yes	Yes	Yes
Day FE	Yes	Yes	Yes	Yes

Notes: Unit of analysis is corpus–language–day cell. Outcome is $\log(\text{posts} + 1)$. Placebo corpus (Biden/Democrats) serves as control group. All specifications include language and day fixed effects. Standard errors clustered by language and day in parentheses. $^+p < 0.1$; $^*p < 0.05$; $^{**}p < 0.01$; $^{***}p < 0.001$.

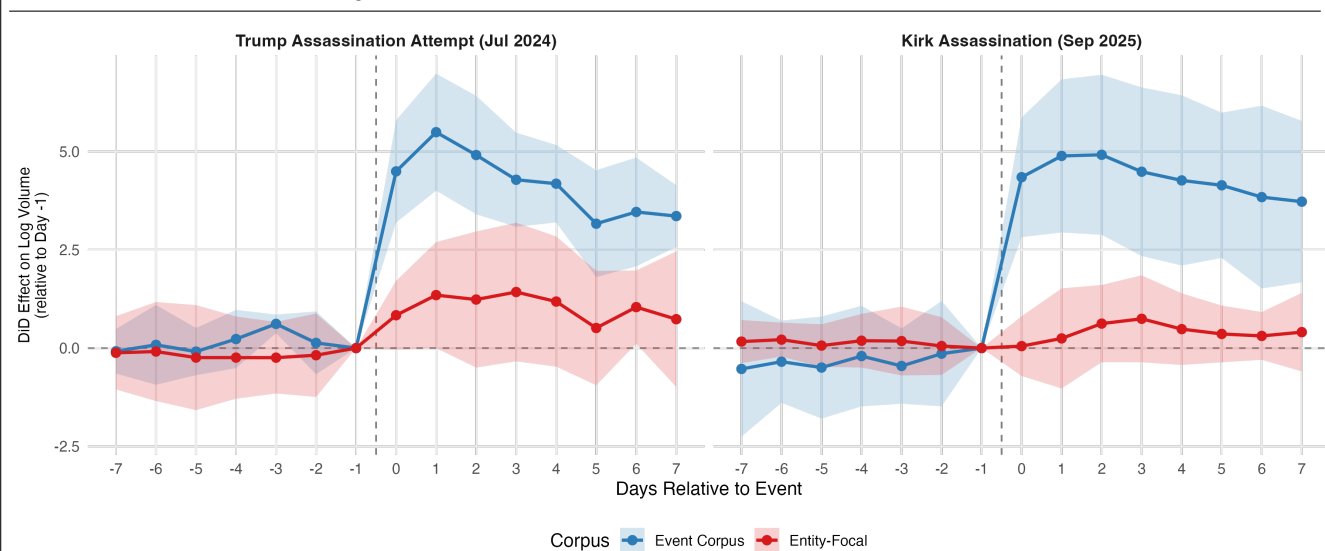
nities at Day 0, peak effects at Days 1–2, and gradual decay thereafter. The event corpus (blue) shows dramatically larger effects than the entity-focal corpus (red), confirming that the shocks activate event-specific discussion far more than general political discourse.

User Location. Table 5 presents estimates interacting treatment with user location, comparing both event-specific and entity corpora to placebo. The results reveal a striking pattern consistent with differential scope of activation.

For the *event corpus*, non-U.S. users exhibit significantly larger volume responses than U.S.-based users. The triple interaction is positive and significant for the Trump event ($\beta_4 = 1.31$, $p < 0.05$) and positive though imprecisely estimated for Kirk ($\beta_4 = 1.39$, $p > 0.10$). However, for the *entity corpus*—which captures broader discussion of political figures rather than the focal event itself—the location differential attenuates substantially. The Trump entity corpus shows only a marginally significant differential ($\beta_4 = 0.34$, $p < 0.10$), while the Kirk entity corpus shows essentially no location heterogeneity ($\beta_4 = 0.12$, $p > 0.10$).

This pattern illuminates the mechanism underlying imported partisanship. Raw counts clarify the extensive margin driving these results: in the Trump event corpus, non-U.S. users increase from approximately 182 unique posters pre-event to over 19,000 post-event (a 106-fold increase), while U.S.-based users increase from 26 to 1,279 (a 49-fold increase). Crucially, this differential

FIGURE 4. Event Study: Difference-in-Differences Effects on Post Volume



Notes: Figure plots daily DiD coefficients from the specification: $\log(N_{jlt} + 1) = \alpha + \beta_1 \text{Treat}_j + \sum_{\tau \neq -1} \gamma_\tau \mathbf{1}[t = \tau] + \sum_{\tau \neq -1} \delta_\tau (\text{Treat}_j \times \mathbf{1}[t = \tau]) + \lambda_l + \epsilon_{jlt}$, where N_{jlt} is post count in corpus j , language l , day t . Coefficients δ_τ represent the differential change in treatment corpus volume relative to placebo, with Day -1 as the reference period. Shaded regions indicate 95% confidence intervals based on standard errors clustered by language. Horizontal dashed line marks zero; vertical dashed line marks event day.

TABLE 5. Volume Effects by User Location (Triple Difference)

	Trump		Kirk	
	Event	Entity	Event	Entity
Treat	-3.775*** (0.235)	-0.282 (0.178)	-2.936** (0.353)	-1.134** (0.156)
Non-US	2.045* (0.613)	2.045* (0.612)	2.217* (0.699)	2.217* (0.698)
Treat × Post (US baseline)	2.544*** (0.283)	0.960* (0.239)	2.470** (0.422)	0.134 (0.187)
Treat × Non-US	-1.060* (0.282)	-0.349+ (0.145)	-1.711 (0.924)	-0.627 (0.564)
Post × Non-US	-0.127+ (0.055)	-0.127+ (0.056)	-0.163 (0.083)	-0.163+ (0.063)
Treat × Post × Non-US (β_4)	1.306* (0.330)	0.340+ (0.124)	1.390 (0.735)	0.121 (0.071)
Observations	271	300	249	296
Language FE	Yes	Yes	Yes	Yes
Day FE	Yes	Yes	Yes	Yes

Notes: Unit of analysis is corpus–language–day–location cell. Outcome is $\log(\text{posts} + 1)$. Each column compares the listed corpus to placebo. Non-US indicates users with non-U.S. profile locations. All specifications include language and day fixed effects. Standard errors clustered by language and day in parentheses. + $p < 0.1$; * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

concentrates in event-specific discussion. When we examine the entity corpus—ongoing discussion of Trump or GOP more broadly—U.S. users show robust spillover effects ($\text{Treat} \times \text{Post} = 0.96$, $p < 0.05$ for Trump), while the Non-US differential shrinks.

These results support the scope-of-activation interpretation: dramatic events serve as focal points that disproportionately draw distant witnesses into *event-specific* discourse, but this activation does not generalize as broadly to the wider partisan landscape. U.S.-based users, already embedded in American political discussion, show smaller event-specific surges but broader spillover to entity-level engagement. Non-U.S. users require salient shocks to overcome their higher participation threshold, and their engagement remains more narrowly tethered to the precipitating event.

Cross-Target Spillover. Table 6 examines whether shocks to one political target spill over to related targets. The Trump event generates significant spillover to GOP discussion ($\beta = 0.98$, $p < 0.01$): an assassination attempt on Trump activates broader Republican-related discourse. The Kirk event, however, shows *negative* spillover to Trump discussion ($\beta = -0.41$, $p < 0.001$). This asymmetry suggests attention crowding—the Kirk event absorbs attention that might otherwise go to Trump, rather than activating a broader partisan network.

Summary. The volume analysis confirms that U.S. political shocks generate substantial cross-border engagement, with effects concentrated on event-specific discussion but spilling over into broader political discourse—particularly for globally salient figures like Trump. The finding that non-U.S. users show larger proportional increases (significantly so for the Trump event) sets up a key test for the sentiment analysis: if these newly activated distant witnesses also express more intensely, this would confirm that events disproportionately mobilize high- θ individuals who were previously below the participation threshold.

Sentiment Analysis

I next examine how major U.S. political violence shapes the *direction* and *tone* of foreign-language political discourse. All sentiment outcomes are measured in an LLM-coded sample and analyzed in a user–day panel. The main identification strategy compares the *entity-focal* corpus (discussion of the focal political target) to a placebo corpus (Biden/Democrats) before vs. after each event,

TABLE 6. Cross-Target Spillover Effects

	Trump Event		Kirk Event	
	Focal (Trump)	Cross (GOP)	Focal (GOP)	Cross (Trump)
Treat (Focal)	−0.555 (0.278)		−1.838* (0.644)	
Treat (Cross)		−2.925* (0.860)		1.226+ (0.500)
Treat (Focal) × Post	1.200** (0.183)		0.280** (0.037)	
Treat (Cross) × Post		0.979** (0.157)		−0.410*** (0.045)
Observations	150	150	150	150
R^2 (within)	0.392	0.684	0.579	0.504
Language FE	Yes	Yes	Yes	Yes
Day FE	Yes	Yes	Yes	Yes

Notes: Unit of analysis is corpus–language–day cell. Outcome is log(posts + 1). Placebo corpus serves as control. Focal corpus contains the event's direct target; Cross corpus contains the related target (GOP for Trump event; Trump for Kirk event). All specifications include language and day fixed effects. Standard errors clustered by language and day in parentheses.

+ $p < 0.1$; * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

including language and relative-day fixed effects, and clustering standard errors by user.⁴

Directional and affective effects. Table 7 summarizes the treatment-on-treated contrast of interest (Treat × Post) across all primary sentiment outcomes. To keep the main text readable, I report only the focal coefficient for each outcome; full regression tables with all coefficients, fit statistics, and raw-means decompositions appear in the Appendix.

Directional consequences: stance. I lead with stance because it maps most directly onto politically consequential changes in evaluation. The Trump attempt generates a modest but precisely estimated pro-Trump shift in entity-focused discourse ($\beta = 0.058$, $p < 0.001$). Interestingly, stance toward Democrats also shifts positive in the same comparison ($\beta = 0.066$, $p < 0.001$), consistent with a broad anti-violence or sympathy reaction that is not confined to the focal target. In contrast, the Kirk assassination produces a smaller pro-Trump shift ($\beta = 0.025$, $p < 0.05$) alongside a significant anti-Democratic shift ($\beta = -0.031$, $p < 0.01$), suggesting a more partisan pattern of

⁴I focus on the entity-focal corpus in the main text because the event-corpus has sparse pre-event observations. Event-corpus estimates, event-study diagnostics (including Wald pre-trend tests), and placebo-only diagnostics are reported in the Appendix.

TABLE 7. Headline Sentiment Effects (Entity vs. Placebo): Treat $\times$ Post

	Trump Attempt (Jul 2024)	Kirk Assassination (Sep 2025)
<i>Panel A: Directional stance (−1 anti, 0 neutral, +1 pro)</i>		
Stance toward Trump	0.058*** (0.008)	0.025* (0.011)
Stance toward GOP	−0.005 (0.006)	−0.005 (0.016)
Stance toward Democrats	0.066*** (0.008)	−0.031** (0.011)
<i>Panel B: Affective tone (0–3)</i>		
Intensity	−0.082*** (0.012)	−0.054** (0.019)
Hostility	−0.172*** (0.014)	−0.054** (0.021)
User-day observations	69,979	31,616

Notes: Each entry reports the Treat $\times$ Post coefficient from a separate user–day difference-in-differences regression comparing the entity-focal corpus to the placebo corpus. All models include language and relative-day fixed effects and control for log followers and log account age. Standard errors clustered by user in parentheses. ⁺ $p < 0.1$; * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

attribution or backlash in the later event. Across both events, average stance toward the GOP in the entity-focal comparison does not significantly change.

Cross-target spillover in stance. A central implication of the theoretical framework is that shocks can generalize from the focal target to politically adjacent targets. Table 8 estimates parallel entity-vs-placebo comparisons for stance outcomes in focal and cross-target corpora. Two patterns stand out. First, the Trump event exhibits broader coalition linkage: effects are detectable not only in Trump-focused discussion but also in GOP-related discourse (e.g., GOP-focused discussion becomes more pro-GOP). Second, the Kirk event is more localized: in cross-target Trump discussion, stance toward Trump is near zero, consistent with weaker generalization to Trump when the focal victim is a less globally salient GOP figure.

Affective tone as mechanism: fatigue and selective entry. While stance shifts indicate directional political consequences, intensity and hostility help diagnose *how* discourse changes. Despite enormous volume increases, average affect in the entity-focal corpus declines relative to placebo after the Trump attempt (Table 7, Panel B). The Kirk event shows smaller negative differences.

To probe mechanisms, I decompose composition changes in the entity-focal corpus. Table 9

TABLE 8. Cross-Target Spillover Effects (Stance Outcomes)

	Trump Attempt		Kirk Assassination	
	Focal (Trump)	Cross (GOP)	Focal (GOP)	Cross (Trump)
<i>Panel A: Stance toward Trump</i>				
Treat × Post	0.058*** (0.008)	−0.016+ (0.009)	0.025* (0.011)	0.004 (0.009)
<i>Panel B: Stance toward GOP</i>				
Treat × Post	−0.005 (0.006)	0.072*** (0.018)	−0.005 (0.016)	−0.063*** (0.010)
User-day obs.	69,979	42,442	31,616	56,973

Notes: Each column is a separate regression comparing the indicated corpus to placebo. All models include language and relative-day fixed effects and control for log followers and log account age. Standard errors clustered by user in parentheses. ⁺ $p < 0.1$; * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

classifies users into returning participants (present pre and post), new entrants (appear only post), and churned users (appear only pre). Two patterns emerge for tone. First, returning users exhibit *fatigue*: their average intensity declines from pre to post. Second, new entrants are not affectively neutral; they are often *more* intense than returning users in the post-period, indicating positive selection into participation rather than simple dilution.⁵

TABLE 9. Decomposition by User Type: Entity Corpus

Event	User Type	Intensity		Hostility	
		Pre	Post	Pre	Post
Trump	Returning (N=2,671)	1.24	1.12	0.96	0.77
	New Entrant (N=16,337)	—	1.19	—	0.72
	Churned (N=8,758)	1.29	—	0.94	—
Kirk	Returning (N=708)	1.28	1.27	1.12	1.10
	New Entrant (N=6,209)	—	1.39	—	1.20
	Churned (N=2,828)	1.25	—	1.05	—

Notes: Returning users posted at least once in both pre and post periods. New entrants appear only post-event; churned users appear only pre-event. N indicates unique users.

Table 10 formalizes this logic by comparing the observed post-minus-pre change to a counterfactual in which only returning users contribute in the post period. For the Trump event, the counterfactual decline is larger than the observed decline, implying that new entrants partially

⁵Appendix tables replicate this decomposition for *stance* outcomes. Directional selection is also evident: new entrants can materially contribute to observed stance shifts even when returning users move little or in the opposite direction.

offset returning-user fatigue. For the Kirk event, new entrants account for a large share of the observed post-period increase in intensity within the entity-focal corpus.

TABLE 10. Counterfactual Decomposition of Intensity Changes

Event	Actual Δ	Counterfactual Δ	New Entrant Contribution
Trump Attempt	-0.097	-0.149	+0.052
Kirk Assassination	+0.114	+0.013	+0.101

Notes: Actual Δ is the observed pre-to-post change in mean intensity (entity corpus). Counterfactual Δ is the change that would have occurred if only returning users posted in the post-period. New entrant contribution = Actual – Counterfactual; positive values indicate new entrants raised average intensity.

Scope of activation: U.S. vs. Non-U.S. users (stance and tone). The identification strategy treats the placebo corpus as a proxy for time-varying baseline engagement. This matters especially when assessing location heterogeneity: if U.S. users become more activated across *all* partisan topics, placebo outcomes can move mechanically, changing the baseline against which entity-focal effects are measured.

Table 11 reports triple-difference estimates interacting $\text{Treat} \times \text{Post}$ with a Non-U.S. indicator, now including both stance and affective outcomes. For the Trump event, Non-U.S. users exhibit a larger pro-Trump stance shift relative to U.S. users (positive and significant interaction for stance toward Trump) and substantially smaller declines in intensity and hostility (positive and significant interactions for both tone measures). For the Kirk event, location heterogeneity is weaker overall; the clearest difference is suggestive evidence that Non-U.S. users become less pro-GOP relative to U.S. users (negative interaction for stance toward GOP).

Table 12 decomposes the location differential for affective tone into raw mean changes by corpus and location. The key empirical pattern is that U.S. users' placebo discourse shifts more strongly than Non-U.S. users' placebo discourse, mechanically changing the baseline against which entity-focal effects are measured.⁶

Cross-language heterogeneity. Table 13 documents heterogeneity across language communities in event-corpus responses, including a directional outcome (stance toward Trump) alongside tone.

⁶Appendix tables report the analogous raw-mean decompositions for stance outcomes, which clarify whether location differentials are driven primarily by entity-side shifts, placebo-side shifts, or both.

TABLE 11. Differential Response by User Location (US vs. Non-US)

	Trump Attempt (Jul 2024)	Kirk Assassination (Sep 2025)
<i>Panel A: Stance toward Trump</i>		
Treat × Post (US baseline)	0.019 (0.022)	0.035 (0.037)
Treat × Post × Non-US (β_4)	0.052* (0.026)	−0.038 (0.041)
<i>Panel B: Stance toward GOP</i>		
Treat × Post (US baseline)	−0.017 (0.017)	0.058 (0.055)
Treat × Post × Non-US (β_4)	0.016 (0.020)	−0.115+ (0.062)
<i>Panel C: Stance toward Democrats</i>		
Treat × Post (US baseline)	0.075*** (0.021)	−0.006 (0.036)
Treat × Post × Non-US (β_4)	−0.004 (0.024)	−0.010 (0.041)
<i>Panel D: Affective tone</i>		
Intensity: Treat × Post (US baseline)	−0.193*** (0.030)	−0.148* (0.066)
Intensity: Treat × Post × Non-US (β_4)	0.124*** (0.036)	0.053 (0.074)
Hostility: Treat × Post (US baseline)	−0.328*** (0.033)	−0.169* (0.072)
Hostility: Treat × Post × Non-US (β_4)	0.180*** (0.040)	0.105 (0.080)
User-day obs.	39,595	14,969
R^2	0.278	0.266
<i>Notes:</i> Triple-difference estimates from entity corpus vs. placebo, interacted with user location. Sample restricted to users with identifiable US or Non-US location. All models include language and relative-day fixed effects and control for log followers and log account age. Standard errors clustered by user in parentheses. +$p < 0.1$; *$p < 0.05$; **$p < 0.01$; ***$p < 0.001$.		

TABLE 12. Decomposition of Location Differential in Tone: Raw Mean Changes

Event	Component	Intensity		Hostility	
		US	Non-US	US	Non-US
Trump	Entity Δ	-0.17	-0.09	-0.33	-0.19
	Placebo Δ	+0.04	+0.01	+0.02	-0.03
	DiD	-0.21	-0.10	-0.35	-0.16
Kirk	Entity Δ	+0.06	+0.08	+0.05	+0.11
	Placebo Δ	+0.21	+0.15	+0.22	+0.15
	DiD	-0.15	-0.07	-0.17	-0.04
β_4 (DiD _{Non-US} - DiD _{US})		+0.12 / +0.08		+0.19 / +0.14	

Notes: Δ denotes post minus pre change in raw means. DiD = Entity Δ - Placebo Δ . Positive β_4 indicates Non-US users show less negative DiD than US users. Table reports tone outcomes (intensity/hostility) to clarify baseline dynamics.

Given the small number of language groups, these estimates are best interpreted as descriptive evidence that motivates mechanism discussion (e.g., media environment and diaspora composition) rather than as a high-powered test of subgroup theory. For readers interested in substantively interpreting language differences, Appendix tables report total language-specific effects (baseline + interaction) with correctly propagated standard errors.

TABLE 13. Cross-Language Heterogeneity: Event Corpus vs. Placebo

	Trump Attempt			Kirk Assassination		
	Stance (Trump)	Intensity	Hostility	Stance (Trump)	Intensity	Hostility
Treat $\times$ Post (French)	0.042* (0.019)	-0.147*** (0.027)	-0.252*** (0.026)	0.001 (0.032)	-0.191*** (0.054)	-0.321*** (0.056)
$\times$ Japanese	-0.020 (0.025)	0.144*** (0.040)	0.294*** (0.041)	0.037 (0.045)	0.122 (0.087)	0.096 (0.086)
$\times$ Korean	0.033 (0.076)	0.082 (0.078)	0.391*** (0.092)	-0.068+ (0.041)	0.528*** (0.064)	0.898*** (0.055)
$\times$ Russian	-0.003 (0.091)	0.345*** (0.100)	0.252** (0.080)	0.121* (0.060)	-0.122 (0.108)	-0.352*** (0.088)
$\times$ Chinese	0.076** (0.029)	0.024 (0.050)	0.276*** (0.064)	0.035 (0.053)	0.514*** (0.145)	0.522*** (0.141)

Notes: Estimates from event corpus vs. placebo, interacted with language (French is the omitted baseline). All models include relative-day fixed effects and control for log followers and log account age. Standard errors clustered by user in parentheses. + $p < 0.1$; * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

Summary. Taken together, the sentiment evidence complements the volume results in three ways. First, the clearest political consequences appear in stance: violence shocks move directional

evaluations in ways consistent with sympathy and coalition linkage, particularly for globally salient targets. Second, affective tone does not simply intensify with attention; instead, average intensity/hostility reflect an interaction of returning-user fatigue and positively selected entry by newly activated participants (with analogous selection patterns in stance documented in the Appendix). Third, heterogeneity analyses clarify that “activation” can be generalized (moving placebo baselines) for U.S. users but more targeted for distant observers, and that both direction and tone vary meaningfully across language communities.

5 DISCUSSION

This paper examined how American partisan conflict travels globally: why foreign-language audiences mobilize around U.S. political violence, and whether their engagement reproduces the structure of domestic partisan dynamics. Across five language communities and two exogenous violence shocks, I find a consistent pattern. Participation is highly elastic abroad—dramatic events trigger extraordinary volume spikes among users who otherwise remain disengaged from American politics. But expression does not simply scale domestic polarization. Instead, users with and without material stakes in U.S. politics differ systematically in the *scope* of their activation.

The volume results establish that American political violence generates massive cross-border engagement. Both events produce 50- to 100-fold increases in event-specific discussion, with non-U.S. users showing larger proportional gains than U.S. users. This pattern is consistent with a focal-point logic: U.S.-based users engage with American politics routinely, while non-U.S. users require high-salience shocks to cross the participation threshold. Spillover effects are asymmetric. The Trump event activates broader Republican-related discussion, consistent with coalition-level mobilization around a globally salient symbol. The Kirk event shows negative spillover to Trump-focused volume, suggesting attention crowding rather than coalition-wide expansion. Diaspora presence does not robustly predict volume responses for the globally saturated Trump event but appears more relevant for the less universally salient Kirk shock—consistent with diaspora transmission mattering most when mainstream exposure is incomplete.

On the sentiment side, the clearest political consequences appear in directional stance rather than affective tone. Stance shifts are modest in magnitude but precisely estimated, revealing

how foreign audiences evaluate American partisan actors rather than merely how emotionally charged their language becomes. Two patterns stand out. First, the Trump attempt produces both a pro-Trump shift and a pro-Democratic shift in entity-focused discourse, consistent with a broad anti-violence or sympathy response not confined to the targeted actor. Second, the Kirk assassination exhibits a more partisan pattern: a smaller pro-Trump shift alongside an anti-Democratic movement, suggesting attribution or backlash dynamics that differ qualitatively from the earlier event. Cross-target spillovers reinforce this interpretation. Trump's event generates coalition-linked effects that extend into GOP-related discourse; Kirk's event remains localized, with limited evidence of shifting Trump evaluations. Global partisan engagement appears candidate-centered: it does not automatically generalize across adjacent targets unless the focal figure is sufficiently salient to anchor coalition-wide meaning.

Affective tone provides a complementary diagnostic of how discourse changes. Despite massive increases in posting, average intensity and hostility decline in the entity-focal corpus relative to placebo after the Trump attempt. Decompositions by user type clarify the mechanism: returning users exhibit fatigue, moderating their tone post-event, while new entrants are positively selected and often more intense than incumbents. New entrants therefore partially offset—but do not reverse—incumbent fatigue. This composition logic matters substantively: high-salience shocks expand the participant pool, but intensive-margin moderation among returning users can still move average tone downward. The result is not "more attention, more anger" but rather a reshuffling of who speaks and how.

The clearest structural divergence between U.S. and non-U.S. users emerges in scope of activation. U.S. users respond to violence in ways that generalize across partisan topics, including the opposing-party placebo. Non-U.S. users respond more narrowly to event-relevant discussion without comparable spillover. This difference appears in the triple-difference designs: relative to placebo, non-U.S. users show smaller declines in affective tone and, for the Trump event, larger pro-Trump stance shifts than U.S. users in the same contrast. For the Kirk event, location differentials are weaker and noisier, suggesting that generalized activation is most pronounced when the focal shock involves a globally dominant symbol receiving saturation coverage.

These results reconcile two facts that might otherwise seem in tension. Non-U.S. users are

more responsive on the extensive margin—they enter in larger proportional numbers when shocks occur—but less generalized on the scope margin—their engagement concentrates on the focal event rather than diffusing across the partisan landscape. U.S. users show the opposite profile: smaller proportional entry but broader activation across partisan topics. The implication is that American partisan symbols are globally legible, but the connective tissue that links these symbols into a unified political system remains rooted in material stakes. Foreign audiences engage as active participants rather than detached observers; they mobilize at scale and take directional stances toward U.S. partisan actors. Yet their engagement is episodic rather than systemic. Dramatic events create focal points that temporarily lower participation thresholds, but absent material stakes, engagement does not consistently generalize across the interconnected terrain that structures domestic conflict.

Several limitations qualify the interpretation. User-profile location is measured with error, and the analysis cannot perfectly separate diaspora users within the U.S. from non-diaspora users abroad. The study covers five languages and two violence shocks, limiting generalizability to other platforms, issues, and political contexts. While LLM-coded sentiment measures are validated against expert annotation, performance may vary across languages and rhetorical styles. Observational data cannot definitively identify the mechanism of exposure—diaspora transmission versus direct global media—or the microfoundations of fatigue among returning users. Future work could extend the design to a wider set of political shocks, incorporate richer measures of cross-border network diffusion, and test mechanisms directly via experiments on exposure and identity resonance.

The broader contribution is methodological as much as substantive. Cross-lingual social media data combined with validated content coding can reveal not only whether global audiences attend to American politics, but how that attention translates into stance-taking and coalition linkage—or fails to do so. The vocabulary of American partisanship travels globally; its systemness does not.

REFERENCES

Adamson, Fiona B. (2016). The growing importance of diaspora politics. *Current History* 115(784), 291–297.

- Bail, Christopher A. (2021). *Breaking the Social Media Prism: How to Make Our Platforms Less Polarizing*. Princeton: Princeton University Press.
- Birkland, Thomas A. (1998). Focusing events, mobilization, and agenda setting. *Journal of Public Policy* 18(1), 53–74.
- Bossetta, Michael (2018). The digital architectures of social media: Comparing political campaigning on facebook, twitter, instagram, and snapchat in the 2016 u.s. election. *Journalism and Mass Communication Quarterly* 95(2), 471–496.
- Bradshaw, Samantha and Philip N. Howard (2019). The global disinformation order: 2019 global inventory of organised social media manipulation. *Oxford Internet Institute Working Paper*.
- 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.
- Brinkerhoff, Jennifer M. (2009). *Digital Diasporas: Identity and Transnational Engagement*. Cambridge University Press.
- Campbell, Angus , Philip E. Converse, Warren E. Miller, and Donald E. Stokes (1960). *The American Voter*. New York: Wiley.
- Centola, Damon (2010). The spread of behavior in an online social network experiment. *Science* 329(5996), 1194–1197.
- Chouliaraki, Lilie (2016). Human rights in an age of distant witnesses: remixed lives, reincarnated images and live-streamed co-presence. *The International Journal of Human Rights* 20(7), 929–944.
- Diminescu, Dana (2008). The connected migrant: An epistemological manifesto. *Social Science Information* 47(4), 565–579.
- Downs, Anthony (1957). *An Economic Theory of Democracy*. New York: Harper and Row.
- Eleta, Irene and Jennifer Golbeck (2012). Bridging languages in social networks: How multilingual users of twitter connect language communities. In *Proceedings of the American Society for Information Science and Technology*, Volume 49, pp. 1–4.

-
- Elkins, Zachary and Beth Simmons (2005). On waves, clusters, and diffusion: A conceptual framework. *The Annals of the American Academy of Political and Social Science* 598(1), 33–51.
- Finkel, Eli J. , Christopher A. Bail, Mina Cikara, Peter H. Ditto, Shanto Iyengar, Samara Klar, Lilliana Mason, Mary C. McGrath, Brendan Nyhan, David G. Rand, et al. (2020). Political sectarianism in America. *Science* 370(6516), 533–536.
- Galtung, Johan and Mari Holmboe Ruge (1965). The structure of foreign news: The presentation of the Congo, Cuba and Cyprus crises in four Norwegian newspapers. *Journal of Peace Research* 2(1), 64–90.
- González-Bailón, Sandra , Javier Borge-Holthoefer, Alejandro Rivero, and Yamir Moreno (2011). The dynamics of protest recruitment through an online network. *Scientific Reports* 1, 197.
- Green, Donald , Bradley Palmquist, and Eric Schickler (2002). *Partisan Hearts and Minds: Political Parties and the Social Identities of Voters*. Yale University Press.
- Guarnizo, Luis Eduardo , Alejandro Portes, and William Haller (2003). Assimilation and transnationalism: Determinants of transnational political action among contemporary migrants. *American Journal of Sociology* 108(6), 1211–1248.
- Hillman, Arye L. (2010). Expressive behavior in economics and politics. *European Journal of Political Economy* 26(4), 403–418.
- Holliday, D. E. , Y. Lelkes, and S. J. Westwood (2024). The July 2024 trump assassination attempt was followed by lower in-group support for partisan violence and increased group unity. *Proceedings of the National Academy of Sciences* 121(49), e2414689121.
- Huddy, Leonie , Lilliana Mason, and Lene Aarøe (2015). Expressive partisanship: Campaign involvement, political emotion, and partisan identity. *American Political Science Review* 109(1), 1–17.
- Iyengar, Shanto , Yphtach Lelkes, Matthew Levendusky, Neil Malhotra, and Sean J. Westwood (2019). The origins and consequences of affective polarization in the United States. *Annual Review of Political Science* 22, 129–146.
- Iyengar, Shanto , Gaurav Sood, and Yphtach Lelkes (2012). Affect, not ideology: A social identity perspective on polarization. *Public Opinion Quarterly* 76(3), 405–431.

- Iyengar, Shanto and Sean J. Westwood (2015). Fear and loathing across party lines: new evidence on group polarization. *American Journal of Political Science* 59(3), 690–707.
- 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. Publisher's version available at <https://tinyurl.com/ycvo9zog>.
- Kingdon, John W. (1984). *Agendas, Alternatives, and Public Policies*. Boston: Little, Brown.
- Kobayashi, Tetsuro , Zhifan Zhang, and Ling Liu (2024). Is partisan selective exposure an american peculiarity? a comparative study of news browsing behaviors in the united states, japan, and hong kong. *Communication Research* 0(0).
- Komito, Lee (2011). Social media and migration: Virtual community 2.0. *Journal of the American Society for Information Science and Technology* 62(6), 1075–1086.
- Krosnick, Jon A. and Richard E. Petty (1995). Attitude strength: Antecedents and consequences.
- Kunda, Ziva (1990). The case for motivated reasoning. *Psychological Bulletin* 108(3), 480–498.
- Landis, J Richard and Gary G Koch (1977). The measurement of observer agreement for categorical data. *Biometrics* 33(1), 159–174.
- Levitt, Peggy (2001). *The Transnational Villagers*. Berkeley: University of California Press.
- Mason, Lilliana (2015). ‘i disrespectfully agree’: The differential effects of partisan sorting on social and issue polarization. *American Journal of Political Science* 59(1), 128–145.
- Mudde, Cas (2019). *The Far Right Today*. Cambridge: Polity Press.
- Norris, Pippa and Ronald Inglehart (2019). *Cultural Backlash: Trump, Brexit, and Authoritarian Populism*. Cambridge: Cambridge University Press.
- Osmundsen, Mathias , Alexander Bor, Puk Bødker Vahlstrup, Anja Bechmann, and Michael Bang Petersen (2021). Partisan polarization is the primary psychological motivation behind political fake news sharing on Twitter. *American Political Science Review* 115(3), 999–1015.
- Passini, Stefano (2017). From global to glocal: The development of a sense of global citizenship. *Peace and Conflict: Journal of Peace Psychology* 23(3), 265–273.

-
- Rathje, Steve , Jay J. Van Bavel, and Sander van der Linden (2021). Out-group animosity drives engagement on social media. *Proceedings of the National Academy of Sciences* 118(26), e2024292118.
- Reny, Tyler T. and Benjamin J. Newman (2021). The opinion-mobilizing effect of social protest against police violence: Evidence from the 2020 george floyd protests. *American Political Science Review* 115(4), 1499–1507.
- Rid, Thomas (2020). *Active Measures: The Secret History of Disinformation and Political Warfare*. New York: Farrar, Straus and Giroux.
- Riker, William H. and Peter C. Ordeshook (1968). A theory of the calculus of voting. *American Political Science Review* 62(1), 25–42.
- Robertson, Ronald E. , Jon Green, Damian J. Ruck, Katherine Ognyanova, Christo Wilson, and David Lazer (2023). Users choose to engage with more partisan news than they are exposed to on Google Search. *Nature* 618, 342–348.
- Schudson, Michael (1998). *The Good Citizen: A History of American Civic Life*. New York: Free Press.
- Sears, David O. (1993). Symbolic politics: A socio-psychological theory.
- Starbird, Kate , Ahmer Arif, and Tom Wilson (2019). Disinformation as collaborative work: Surfacing the participatory nature of strategic information operations. *Proceedings of the ACM on Human-Computer Interaction* 3(CSCW), 1–26.
- Taber, Charles S. and Milton Lodge (2006). Motivated skepticism in the evaluation of political beliefs. *American Journal of Political Science* 50(3), 755–769.
- Theocharis, Yannis , Shelley Boulianne, Karolina Koc-Michalska, and Bruce Bimber (2022). Platform affordances and political participation: How social media reshape political engagement. *West European Politics* 46(4), 788–811.
- Tucker, Joshua A. , Yannis Theocharis, Margaret E. Roberts, and Pablo Barberá (2017). From liberation to turmoil: Social media and democracy. *Journal of Democracy* 28(4), 46–59.
- Waldinger, Roger (2015). *The Cross-Border Connection: Immigrants, Emigrants, and Their Homelands*. Cambridge, MA: Harvard University Press.

- Wu, H. Denis (2000). Systemic determinants of international news coverage: A comparison of 38 countries. *Journal of Communication* 50(2), 110–130.
- Zaller, John R. (1992). *The Nature and Origins of Mass Opinion*. Cambridge: Cambridge University Press.

Appendix

A	Ethical Statement	1
B	Pre-registration	1
C	Data Collection Pipeline Using X's API	4
C.1	API Access and Endpoints	4
C.2	Query Construction	4
C.3	Tweet and User Fields Retrieved	5
C.4	Sampling Strategy for High-Volume Days	5
C.5	Rate Limiting and Retry Logic	5
C.6	Deduplication	5
C.7	Checkpointing and Resumption	6
C.8	Output Schema	6
C.9	Reproducibility	6
D	Multilingual Term Lists (ASCII Transliteration)	6
D.1	Chinese (lang:zh) — Pinyin without tone marks	8
D.2	Korean (lang:ko) — Revised Romanization (ASCII)	8
D.3	Japanese (lang:ja) — Hepburn romanization (ASCII)	8
D.4	French (lang:fr) — ASCII (no accents)	8
D.5	Russian (lang:ru) — Latin transliteration (ASCII)	8
D.6	How these lists map to query families	8
E	LLM Labeling Accuracy	8
F	Supplementary Results for Volume Analysis	11
F.1	Event Study Results: US v. Non-US Triple Diff-in-Diff	11
F.2	Full Results: Diff-in-Diff by Language Groups	12
G	Supplementary Results for Sentiment Analysis	14
G.1	Full DiD Results: Entity vs. Placebo	14
G.2	Raw Means: Entity vs. Placebo	16
G.3	Event Corpus Results	16
G.4	Triple-Difference by User Location	16
G.5	Raw Means by User Location	20
G.6	Cross-Language Heterogeneity	20
G.7	Cross-Target Spillover	20
G.8	Cross-Event Comparison	23
G.9	Robustness: Clean Placebo	23
G.10	Event Corpus vs. Entity Corpus Comparison	24

A ETHICAL STATEMENT

Data collection of this study on X 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, Y01M01).

B PRE-REGISTRATION

The LLM Labeling pipeline employed by this study has been pre-registered at AsPredicted.org before the labeling process starts. A PDF version of the pre-registration detail is inserted below.

Borrowed Battlegrounds - Jan 4 (#266281)

Author(s)

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

Pre-registered on:

2026/01/04 13:14 (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 study examines how US political polarization operates as a globalized phenomenon.

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

The key dependent variables involved in this pre-registration are LLM labeled sentiments and attributes for each tweet collected from X, using their official API. These sentiments and attributes include: whether a post is US-relevant, whether a post talks about foreign politics, political stances of each post, including stances on Trump, the Republican Party in the United States, and the Democratic party, as well as intensity for these partisan expressions and expression of hate toward political outgroups.

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

This study is an observational study, so no conditions under which the participants will be assigned to.

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

I will conduct LLM labeling of social media posts collected from X (Twitter) that contain keywords related to two U.S. political events: the Trump assassination attempt (July 2024) and the Charlie Kirk assassination (September 2025). Posts are collected in five non-English languages (Chinese, Japanese, Korean, French, Russian) across a 15-day window (−7 to +7 days) around each event.

LLM Coding Scheme. Each post is coded by GPT-4.1 on seven dimensions:

us_relevant (0/1): Whether the post actually discusses U.S. politics (vs. keyword coincidence)

stance_trump (−1/0/+1): Stance toward Donald Trump

stance_gop (−1/0/+1): Stance toward Republican Party/conservatives

stance_dem (−1/0/+1): Stance toward Democratic Party/liberals

intensity (0–3): Strength of political expression

hostility (0–3): Hostility toward political opponents

foreign_proxy (0/1): Whether the post uses U.S. politics to comment on the author's own country's politics

User locations are classified as US/Non-US/Unknown using LLM-based geocoding of profile location strings.

Analyses. Using the coded data, I will conduct five pre-specified analyses:

Event Effects: Difference-in-differences comparing treatment corpora (event-focal, entity-focal) to placebo corpus, estimating effects on both volume (post counts) and expression (intensity, hostility, stance).

Instrumental vs. Expressive Users: Test for differential responses by user location (US vs. Non-US).

Cross-Language Heterogeneity: Estimate language-specific treatment effects and correlate with diaspora size.

Cross-Target Spillover: Compare effects on direct targets (Trump for Trump event; GOP for Kirk event) versus related targets (GOP for Trump event; Trump for Kirk event) to assess whether shocks spill over to associated political objects.

Cross-Event Comparison: Compare effect sizes across events.

All sentiment analyses condition on us_relevant=1. Standard errors are clustered at the user level for post-level outcomes.

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

Non-US-relevant posts (us_relevant=0): Posts identified by the LLM as not actually about U.S. politics (e.g., cryptocurrency tokens named "\$TRUMP", other countries' "Republican" parties, keyword coincidences) are excluded from all sentiment analyses. These posts are retained in the dataset but coded outcomes are treated as not applicable.

Coding errors (coding_error=TRUE): Posts where the LLM API failed after maximum retries are excluded from sentiment analyses. I will report the coding error rate; if it exceeds 5%, I will investigate systematic patterns.

Retweets (is_retweet=TRUE): Pure retweets without additional commentary are excluded from sentiment analyses to avoid double-counting content originating from other users. Quote tweets with original commentary are retained.

Duplicate posts: Posts appearing multiple times across collection windows (identified by tweet_id) are deduplicated, retaining the first instance.

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.

Population: All posts are collected via X API using keyword searches in five languages (Chinese, Japanese, Korean, French, Russian) across a 15-day window (days -7 to +7) around each event. The total collected population is approximately 1.16 million posts across both events.

Stratified Sampling: I draw a stratified random sample with a target of 500 posts per stratum. Strata are defined by the full cross of:

Event (2): Trump assassination attempt, Kirk assassination

Language (5): Chinese, Japanese, Korean, French, Russian

Corpus (3): Placebo (Biden/Democrat keywords), Entity-focal (direct target), Entity-cross (related target)

Day (15): Days -7 through +7 relative to event

User location (3): US, Non-US, Unknown

This yields a maximum of $2 \times 5 \times 3 \times 15 \times 3 = 1,350$ strata.

Sample Size Determination:

For strata with ≥ 300 posts: Randomly sample exactly 300 posts (seed: 20251220)

For strata with < 300 posts: Include all available posts

The final labeled sample size is determined by: (a) the number of non-empty strata in the collected data, and (b) the within-stratum post counts.

Maximum possible sample size is $1,350 \times 300 = 405,000$ posts; actual size will be smaller due to sparse strata.

Full Dataset Retention: All unlabeled posts are retained in the final dataset with missing values in coding columns. This preserves population structure for weighting or alternative analyses.

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

Although data associated with this study had been collected, the LLM labeling process hasn't started yet. The only purpose of this pre-registration is to pre-register my LLM labeling pipeline.

This pre-registration is bundled with an earlier pre-registration with the same content but a different sample size of randomized posts to be labeled by LLM in each event-language-corpus-day-location strata.

BUNDLE

This pre-registration is part of a bundle which includes:

#265,768 - <https://aspredicted.org/ai3zj8.pdf> - Title: 'Borrowed Battlegrounds'

C DATA COLLECTION PIPELINE USING X'S API

This appendix documents the engineering of the X (Twitter) API v2 data collection pipeline used in this study. The pipeline is implemented in Python and designed to handle full-archive search, rate limiting, high-volume sampling, and resumable execution.

C.1 API Access and Endpoints

The pipeline uses X's API v2 with Pro-tier access, which provides full-archive search capabilities through two primary endpoints:

- `/tweets/search/all`: Full-archive search returning tweets matching a query within a specified time range.
- `/tweets/counts/all`: Returns tweet counts (without content) for volume estimation prior to retrieval.

At initialization, the pipeline performs a capability check by issuing test requests to both endpoints. If full-archive access is unavailable (HTTP 403), the pipeline falls back to recent-search endpoints (`/tweets/search/recent`, `/tweets/counts/recent`), which cover only the past seven days. For historical event windows (2024–2025), the pipeline requires full-archive access and will terminate with an informative error if only recent endpoints are available.

C.2 Query Construction

Each API request is defined by a query string that specifies search terms, language filters, and content-type exclusions. Query construction must respect two constraints: (1) X's query syntax rules and (2) a maximum query length of 1,024 characters.

Syntax for Negation. X's search API does not reliably support negating parenthesized groups (e.g., `-(term1 OR term2)`). The pipeline therefore implements term-by-term negation: to exclude a set of terms, each term is negated individually (e.g., `-term1 -term2 -term3`). This ensures correct Boolean evaluation across all query families.

Query Length Management. When the combined length of OR-grouped terms exceeds the 1,024-character limit, the pipeline automatically splits the term list into chunks and generates multiple query variants. Each variant is executed separately, and results are deduplicated by tweet ID. The chunking algorithm allocates character budgets proportionally across term groups and forms the Cartesian product of chunks when multiple groups require splitting.

Query Families. For each event–language–day cell, the pipeline executes three query families:

1. **Event corpus:** `(event_terms) AND (entity_terms) lang:{lang} -is:retweet`
2. **Entity-only corpus:** `(entity_terms) -event_term1 -event_term2 ... lang:{lang} -is:retweet`
3. **Placebo corpus:** `(placebo_terms) lang:{lang} -is:retweet`

where `event_terms`, `entity_terms`, and `placebo_terms` are language-specific keyword lists (see Appendix D). The `-is:retweet` operator excludes simple retweets, which contain no original text. Replies and quote tweets are retained, as they represent original expression and are central to measuring targeted stances.

C.3 Tweet and User Fields Retrieved

Each API request retrieves the following tweet-level fields:

- `id`, `text`, `created_at`, `author_id`, `lang`, `source`
- `public_metrics`: like count, retweet count, reply count, quote count
- `conversation_id`, `in_reply_to_user_id`
- `referenced_tweets`: array indicating whether the tweet is a reply, quote, or retweet, including the ID of the referenced tweet
- `entities`: hashtags, mentions, URLs

User-level fields are retrieved via the `author_id` expansion:

- `id`, `username`, `name`, `created_at`, `description`
- `location`: free-text self-reported location
- `verified`: verification status
- `public_metrics`: follower count, following count, tweet count

C.4 Sampling Strategy for High-Volume Days

For some event–language–day–family cells, tweet volume may exceed operational thresholds or API rate limits. The pipeline implements a stratified time-slice sampling procedure for days with estimated volume exceeding 15,000 tweets.

Prior to retrieval, the pipeline queries the counts endpoint to estimate daily volume. If the estimate exceeds the threshold, the pipeline samples eight one-hour bins from the 24-hour UTC day using stratified random sampling: two bins are drawn from each of four six-hour blocks (00:00–05:59, 06:00–11:59, 12:00–17:59, 18:00–23:59). This stratification ensures temporal coverage across the full day and mitigates bias from time-of-day variation in posting activity.

Sampled observations are flagged in the output data (`sampled = TRUE`) to enable sensitivity analyses that compare sampled and fully enumerated cells.

C.5 Rate Limiting and Retry Logic

X's API enforces rate limits of approximately 300 requests per 15-minute window for full-archive search (Pro tier). The pipeline manages rate limits through two mechanisms:

1. **Proactive delays:** A configurable delay (default: 2–3 seconds) is inserted between consecutive API requests to avoid triggering rate limits.
2. **Reactive backoff:** If the API returns HTTP 429 (rate limit exceeded), the pipeline reads the `x-rate-limit-reset` header, sleeps until the reset time plus a small buffer, and retries the request.

Transient server errors (HTTP 500, 502, 503, 504) trigger an exponential backoff retry (up to 3 attempts with delays of 1, 2, and 4 seconds). Persistent errors terminate the current task and log the failure.

C.6 Deduplication

Deduplication occurs at two levels:

1. **Within-task:** When a query is split into multiple variants or time bins, the pipeline deduplicates by tweet ID before saving results.

2. **Across-task:** A daily ledger (stored as a text file per UTC day) tracks all tweet IDs retrieved for that day. Before adding a tweet to the output, the pipeline checks whether its ID already exists in the ledger. This prevents double-counting when tweets match multiple query families (e.g., a tweet mentioning both Trump and Biden).

C.7 Checkpointing and Resumption

To support long-running collection jobs and recovery from interruptions, the pipeline implements task-level checkpointing. Each task is defined by a unique key: `event_id|language|window_day|family`. Upon successful completion of a task, its key is appended to a JSON checkpoint file (`completed_tasks.json`).

At startup, the pipeline loads the checkpoint file and skips any tasks already marked complete. This enables resumption from the point of failure without re-fetching previously collected data. To force a full re-run, the user deletes the checkpoint file.

C.8 Output Schema

Each task produces a CSV file with UTF-8-SIG encoding (to ensure correct rendering of CJK characters in Excel). The output directory structure is:

```
{output_root}/
  {event_id}/
    {lang}/
      event/
        {event_id}_{lang}_event_day{+/-N}_{count}.csv
      entity/
        {event_id}_{lang}_entity_day{+/-N}_{count}.csv
      placebo/
        {event_id}_{lang}_placebo_day{+/-N}_{count}.csv
  ledger_daily_ids/
    {YYYY-MM-DD}.txt
  run_state.json
  completed_tasks.json
```

Table A1 describes the columns in each output CSV file.

C.9 Reproducibility

To ensure reproducibility, the pipeline uses a fixed random seed (`RANDOM_SEED = 20251219`) for all sampling operations. The `run_state.json` file logs daily unique tweet counts and cumulative totals. All query strings are recorded in the output data (`query_used` column), enabling exact replication of API calls.

The pipeline requires Python 3.8+ and the following dependencies: `requests`, `pandas`, `tqdm`. The optional `retry` package provides automatic retry with exponential backoff; if unavailable, a minimal fallback implementation is used.

D MULTILINGUAL TERM LISTS (ASCII TRANSLITERATION)

This appendix reports the multilingual term lists used to construct the query families. For maximum portability across $\text{\LaTeX}$ compilers, all non-Latin scripts are presented as ASCII-only transliterations. The replication code contains the native-script terms used in API queries.

TABLE A1 . Output Data Schema

Column	Type	Description
<i>Provenance</i>		
event_id	string	Event identifier (trump_attempt_2024 or kirk_assassination_2025)
family	string	Query family (event, entity, placebo)
lang	string	Language tag (zh, ko, ja, fr, ru)
window_day	integer	Days relative to event (−7 to +7)
sampled	boolean	Whether stratified sampling was applied
<i>Tweet-Level Fields</i>		
tweet_id	string	Unique tweet identifier
created_at	datetime	Tweet creation timestamp (UTC)
tweet_text	string	Full text content
source	string	Posting client (e.g., “Twitter for iPhone”)
conversation_id	string	ID of the root tweet in the thread
in_reply_to_user_id	string	User ID being replied to (if reply)
in_reply_to_tweet_id	string	Tweet ID being replied to (if reply)
quoted_tweet_id	string	Tweet ID being quoted (if quote tweet)
is_retweet	boolean	Whether tweet is a retweet
is_quote	boolean	Whether tweet is a quote tweet
is_reply	boolean	Whether tweet is a reply
hashtags	string	Comma-separated hashtags
like_count	integer	Number of likes
retweet_count	integer	Number of retweets
reply_count	integer	Number of replies
quote_count	integer	Number of quote tweets
<i>User-Level Fields</i>		
user_id	string	Unique user identifier
username	string	Handle (e.g., @example)
name	string	Display name
user_created_at	datetime	Account creation date
user_verified	boolean	Verification status at collection time
user_location	string	Self-reported location (free text)
user_description	string	Profile bio
user_followers_count	integer	Follower count
user_following_count	integer	Following count
user_tweet_count	integer	Total tweets posted

TABLE A2 . Chinese Term Lists (ASCII-only transliteration)

Category	Terms (comma-separated)
Trump (name variants)	telangpu, chuanpu, dongna-de telangpu, telangpu zongtong, chuanpu zongtong, trump, donald trump
Republicans / GOP	gonghedang, gonghedangren, gonghepai, gop, republican party, republican
Charlie Kirk	charlie kirk, cha-li ke-ke, ke-ke
Biden / Democrats (placebo)	baideng, qiaozhi baideng, minzhudang, minzhudangren, democrats, democratic party, biden, joe biden
Event terms (Trump attempt)	ansha, cisha, yuci, qiangji, shaji, juji, chongji, gongji, qiangshou, qiangsheng, jiuhui, jizhi, pennsylvania, butler
Event terms (Kirk assassination)	ansha, cisha, yuci, qiangji, shaji, gongji, qiangsheng, yanjiang, xiaoyuan, utah

D.1 Chinese (lang:zh) — Pinyin without tone marks**D.2 Korean (lang:ko) — Revised Romanization (ASCII)****TABLE A3 . Korean Term Lists (ASCII-only transliteration)**

Category	Terms (comma-separated)
Trump (name variants)	teureompeu, donaldu teureompeu, trump, donald trump
Republicans / GOP	gonghwadang, gonghwadangwon, gop, republican, republican party
Charlie Kirk	challyeokeokeu, chareli keokeu, charlie kirk, keokeu
Biden / Democrats (placebo)	baideun, jo baideun, minjudang, democrats, democratic party, biden, joe biden
Event terms (Trump attempt)	amsal, pigeok, chonggyeok, chonggyeok sageon, jeogyek, gonggyeok, jiphoe, rally, pensilbeinia, beoteulreo, butler
Event terms (Kirk assassination)	amsal, pigeok, chonggyeok, gonggyeok, yeonseol, kaempeoseu, utah

D.3 Japanese (lang:ja) — Hepburn romanization (ASCII)**D.4 French (lang:fr) — ASCII (no accents)****D.5 Russian (lang:ru) — Latin transliteration (ASCII)****D.6 How these lists map to query families**

For each language and UTC day, I construct three query families: (i) **Event corpus**: EVENT TERMS AND FOCAL ENTITY TERMS; (ii) **Entity-only corpus**: FOCAL ENTITY TERMS AND NOT(EVENT TERMS) using term-by-term negation; (iii) **Placebo corpus**: BIDEN/DEMOCRATS TERMS. All queries include the language restriction (e.g., lang:zh) and exclude simple retweets (e.g., -is:retweet).

E LLM LABELING ACCURACY

To validate the accuracy of our LLM-based content classification, we conducted expert hand-coding on a stratified random sample of 2,000 social media posts across five languages (French, Japanese,

TABLE A4 . Japanese Term Lists (ASCII-only transliteration)

Category	Terms (comma-separated)
Trump (name variants)	torampu, donarudo torampu, trump, donald trump
Republicans / GOP	kyowato, kyowato-in, gop, republican, republican party
Charlie Kirk	charlie kirk, chaarii kaaku, kaaku
Biden / Democrats (placebo)	baiden, jo baiden, minshuto, democrats, democratic party, biden, joe biden
Event terms (Trump attempt)	ansatsu, juugeki, shageki, sotsugeki, kougeki, jiken, shuusai, rally, pen-shirubenia, batoraa, butler
Event terms (Kirk assassination)	ansatsu, juugeki, shageki, kougeki, enzetsu, kyanpasu, utah

TABLE A5 . French Term Lists (ASCII-only)

Category	Terms (comma-separated)
Trump (name variants)	trump, donald trump
Republicans / GOP	republicain, republicains, parti republicain, GOP, republican party
Charlie Kirk	charlie kirk, kirk
Biden / Democrats (placebo)	biden, joe biden, democrates, parti democrate, democratic party, democrats
Event terms (Trump attempt)	tentative d'assassinat, assassinat, attaque, tir, coups de feu, fusillade, rally, meeting, pennsylvanie, butler
Event terms (Kirk assassination)	assassinat, attaque, tir, coups de feu, fusillade, discours, campus, utah

TABLE A6 . Russian Term Lists (ASCII-only transliteration)

Category	Terms (comma-separated)
Trump (name variants)	tramp, donald tramp, trump, donald trump
Republicans / GOP	respublikantsy, respublikanskaya partiya, GOP, republican party, republican
Charlie Kirk	charli kirk, charlie kirk, kirk
Biden / Democrats (placebo)	baiden, dzho baiden, demokraty, demokraticheskaya partiya, democrats, democratic party, biden, joe biden
Event terms (Trump attempt)	pokushenie, ubiistvo, popytka ubiistva, vystrel, strelba, napadenie, atentat, miting, rally, pensilvaniya, batler, butler
Event terms (Kirk assassination)	ubiistvo, pokushenie, vystrel, strelba, napadenie, vystuplenie, kampus, utah

Korean, Russian, and Chinese) and two events (the 2024 Trump assassination attempt and the 2025 Kirk assassination). For each post, expert coders independently classified stance toward Trump, the GOP, and the Democratic Party (each on a scale from -1 to 1), as well as emotional intensity and hostility (each on a scale from 0 to 3). We then compared these expert labels against the LLM-generated classifications.

Table A7 presents the validation results. Panel A reports overall accuracy metrics across all 2,000 posts. The LLM achieves exact match accuracy exceeding 94% for all five variables, with stance toward Trump reaching 96.6% agreement. When allowing for minor discrepancies of ± 1 on the ordinal scales, accuracy rises to between 98.2% and 99.9%. Mean absolute error remains below 0.06 for all variables. Cohen's κ coefficients range from 0.882 to 0.950, all exceeding the conventional threshold of 0.80 for “almost perfect” agreement (Landis and Koch 1977).

Panel B disaggregates exact match accuracy by language. Performance is consistently high across all five languages, ranging from 92.0% to 98.2%. Chinese-language posts exhibit the highest accuracy (97.5%–98.2%), likely reflecting greater availability of Chinese-language training data. French shows slightly lower accuracy for stance toward the GOP (92.0%), though this still represents strong agreement. Panel C reports Cohen's κ by language, with all values exceeding 0.83 and the majority surpassing 0.90. The consistency of high κ values across typologically diverse languages suggests that the LLM captures meaningful variation in political sentiment rather than exploiting superficial linguistic patterns.

TABLE A7 . LLM Classification Validation Results

<i>Panel A: Overall Accuracy</i>					
Variable	<i>N</i>	Exact Match	Within ± 1	MAE	Cohen's κ
Stance: Trump	2,000	96.6%	99.5%	0.039	0.917
Stance: GOP	2,000	94.7%	99.9%	0.055	0.882
Stance: Dem	2,000	96.2%	99.9%	0.039	0.908
Intensity	2,000	96.2%	98.8%	0.054	0.950
Hostility	2,000	96.4%	98.2%	0.059	0.945
<i>Panel B: Exact Match Accuracy by Language (%)</i>					
	French	Japanese	Korean	Russian	Chinese
Stance: Trump	96.2	96.5	96.0	96.8	97.5
Stance: GOP	92.0	96.2	93.8	93.8	97.5
Stance: Dem	93.8	97.5	95.2	96.5	98.0
Intensity	95.5	95.0	95.2	97.2	98.2
Hostility	94.0	95.5	97.2	96.8	98.2
<i>Panel C: Cohen's κ by Language</i>					
	French	Japanese	Korean	Russian	Chinese
Stance: Trump	0.915	0.904	0.883	0.913	0.953
Stance: GOP	0.863	0.897	0.839	0.834	0.946
Stance: Dem	0.849	0.938	0.883	0.927	0.940
Intensity	0.919	0.934	0.941	0.966	0.980
Hostility	0.926	0.924	0.951	0.932	0.983

Note: Validation based on 2,000 randomly sampled posts across five languages and two events. Stance variables coded on a -1 to 1 scale; intensity and hostility coded on a 0 to 3 scale. “Within ± 1 ” indicates the proportion of cases where LLM and expert codes differ by at most one unit. MAE = Mean Absolute Error.

In sum, the validation exercise demonstrates that LLM-based classification achieves accuracy levels comparable to or exceeding those typically reported for human coders in political communication research (??). The strong performance across five typologically diverse languages—spanning Romance, Slavic, Sinitic, Japonic, and Koreanic language families—provides confidence that our automated coding captures meaningful cross-national variation in political sentiment.

F SUPPLEMENTARY RESULTS FOR VOLUME ANALYSIS

F.1 Event Study Results: US v. Non-US Triple Diff-in-Diff

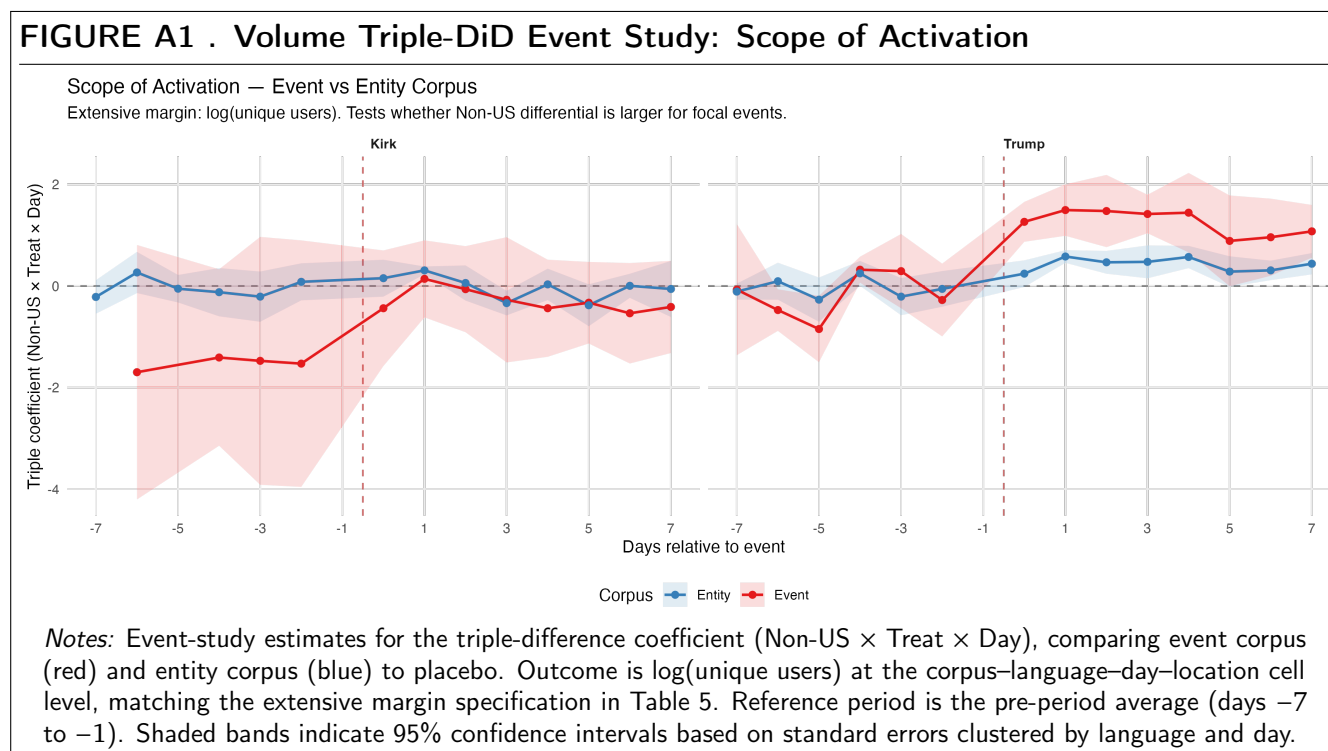


Figure A1 presents event-study estimates for the volume triple-difference, separately for event-specific and entity corpora. The triple coefficient captures whether Non-US users show a larger treatment effect (corpus vs. placebo) than US users on each day relative to the pre-period average.

For the Trump event, the event corpus (red) shows a sharp positive shift at day 0, sustained through the post-period at approximately 1.0–1.5 log points. This indicates that Non-US users disproportionately entered event-specific discussions relative to US users. In contrast, the entity corpus (blue) remains flat near zero throughout, indicating no differential location response for broader political discussion. This pattern supports the scope-of-activation hypothesis: Non-US users exhibit high volume elasticity for focal events but do not generalize to the wider partisan landscape. The Kirk event shows a noisier pattern with the event corpus below zero in the pre-period, converging toward the entity corpus at the event date. The pre-period divergence complicates causal interpretation for this event, though both corpora show similar post-period levels—consistent with narrower scope of activation among Non-US users. Due to the small number of clusters inherent in cell-level volume data (approximately 75 language–day combinations), I present visual evidence rather than formal pre-trend tests.

F.2 Full Results: Diff-in-Diff by Language Groups

Table A8 presents the difference-in-differences estimates allowing for heterogeneous effects across language communities. French serves as the baseline category. The key parameters are the three-way interactions (Treat $\times$ Post $\times$ Language), which capture how each language community's response differs from the French baseline.

TABLE A8 . Cross-Language Heterogeneity in Volume Effects

	Trump Event	Kirk Event
Treat $\times$ Post (French baseline)	4.269*** (0.278)	4.219*** (0.136)
Treat $\times$ Post $\times$ Japanese	0.409* (0.108)	1.681*** (0.069)
Treat $\times$ Post $\times$ Korean	0.527** (0.095)	0.345+ (0.143)
Treat $\times$ Post $\times$ Russian	-1.435*** (0.071)	-2.277*** (0.083)
Treat $\times$ Post $\times$ Chinese	-0.664*** (0.039)	0.800** (0.146)
Observations	145	140
R^2	0.973	0.987
Day FE	Yes	Yes
Language main effects	Yes	Yes
Two-way interactions	Yes	Yes

Notes: Unit of analysis is corpus–language–day cell. Outcome is $\log(\text{posts} + 1)$. French is the baseline language. Full model includes language main effects and all two-way interactions (Treat $\times$ Language, Post $\times$ Language). Standard errors clustered by language and day in parentheses.

+ $p < 0.1$; * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

Table A9 summarizes the total DiD effects by language, calculated as the sum of the baseline effect and the language-specific interaction term.

Japanese and Korean speakers show larger effects than French for both events, with Japanese showing particularly strong responses to the Kirk event (+1.68 log points). Russian speakers consistently show smaller effects (−1.44 for Trump, −2.28 for Kirk), likely reflecting platform restrictions. Chinese speakers show mixed patterns: smaller effects for Trump (−0.66) but larger for Kirk (+0.80). Notably, cross-language variation does not correlate with U.S. diaspora presence on X ($r = -0.08$ for Trump, $r = 0.27$ for Kirk), suggesting diaspora-mediated transmission is not the primary driver of cross-border engagement.

Language Heterogeneity in Volume. Figures A2 and A3 decompose the volume difference-in-differences by language community. Each panel plots the DiD in $\log(\text{unique users})$ between the treatment corpus and placebo, separately for event-specific content (red) and entity content (blue). The figures serve as a visual diagnostic to assess whether the scope-of-activation pattern documented in Table 5 holds consistently across language communities or is driven by a single outlier.

For both events, the pattern is strikingly consistent across all five languages: the event corpus shows a sharp upward spike at day 0, converging toward or exceeding the entity corpus, while

TABLE A9 . Total DiD Effects by Language Community

Language	Log Diaspora	Trump Event		Kirk Event	
		Effect	SE	Effect	SE
French (baseline)	7.54	4.269***	(0.278)	4.219***	(0.136)
Japanese	5.94	4.678***	(0.298)	5.900***	(0.153)
Korean	5.05	4.796***	(0.294)	4.564***	(0.197)
Russian	5.42	2.834***	(0.287)	1.942***	(0.159)
Chinese	7.35	3.605***	(0.281)	5.019***	(0.199)
Correlation with Log Diaspora		$r = -0.08$		$r = 0.27$	

Notes: Total effects calculated as baseline (French) coefficient plus language-specific interaction. Log Diaspora measures the log number of U.S.-based users posting in each language community during the pre-event period. Standard errors approximated as $\sqrt{SE_{\text{base}}^2 + SE_{\text{interaction}}^2}$. All effects significant at $p < 0.001$.

FIGURE A2 . Volume Scope of Activation by Language — Trump Event

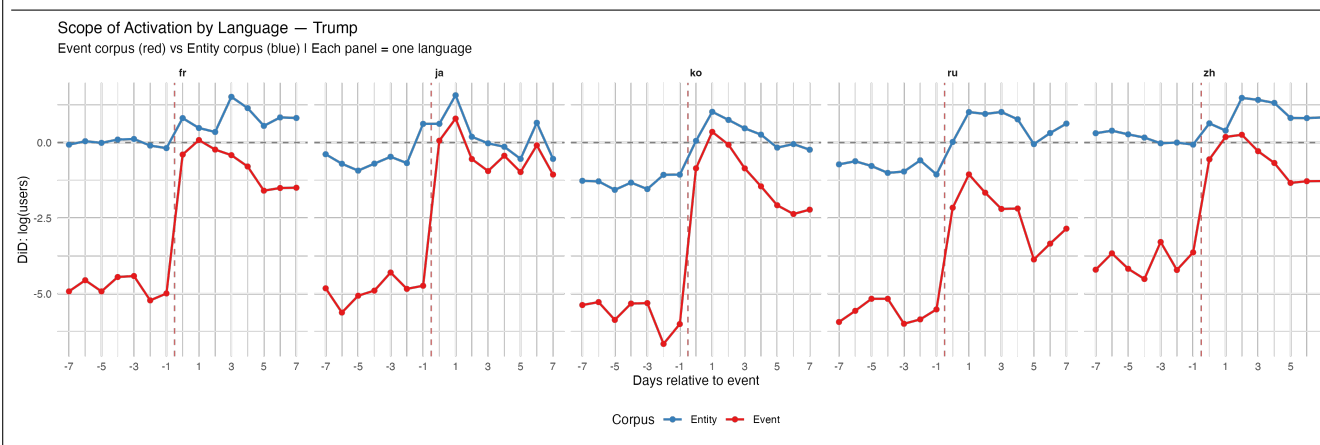
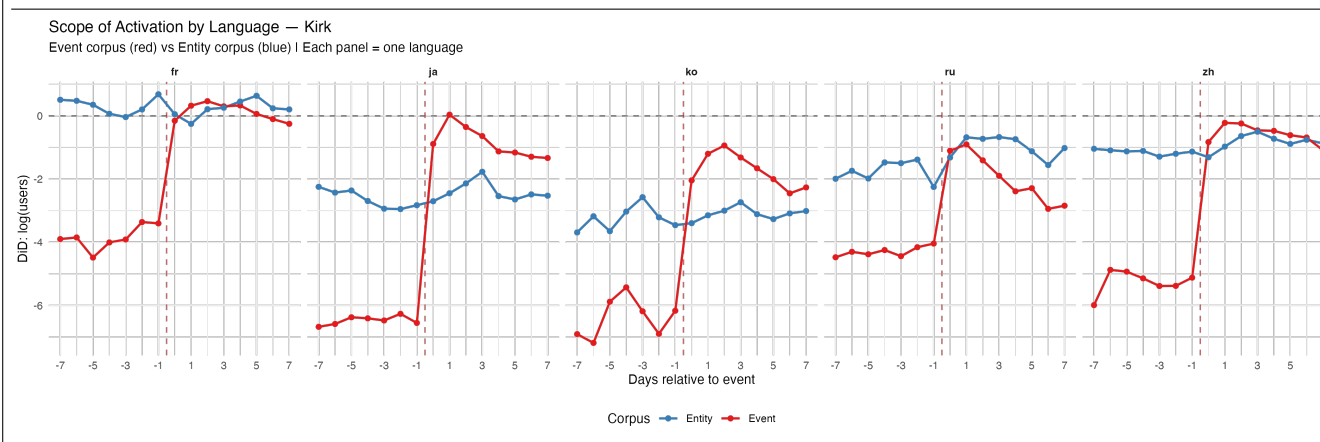


FIGURE A3 . Volume Scope of Activation by Language — Kirk Event



the entity corpus remains relatively stable throughout the window. This universality suggests that the scope-of-activation mechanism—whereby political shocks draw users into event-specific discourse rather than broader partisan discussion—operates similarly across French, Japanese, Korean, Russian, and Chinese language communities. No single language appears to drive the aggregate results, supporting the external validity of the main findings.

G SUPPLEMENTARY RESULTS FOR SENTIMENT ANALYSIS

This appendix provides full regression output, raw means, and robustness checks for the sentiment analysis. Tables are numbered A1–A14 to correspond with main text references.

G.1 Full DiD Results: Entity vs. Placebo

Tables A10 and A11 report full regression output for the main entity-vs-placebo DiD specifications summarized in the main text.

TABLE A10 . Full DiD Results (Entity vs. Placebo): Stance Outcomes

	Trump Attempt			Kirk Assassination		
	Trump	GOP	Dem	Trump	GOP	Dem
Treat	−0.027*** (0.007)	0.036*** (0.005)	0.391*** (0.007)	−0.100*** (0.009)	−0.366*** (0.014)	0.520*** (0.009)
Treat × Post	0.058*** (0.008)	−0.005 (0.006)	0.066*** (0.008)	0.025* (0.011)	−0.005 (0.016)	−0.031** (0.011)
Log Followers	0.009*** (0.001)	0.004*** (0.001)	0.009*** (0.001)	0.004* (0.001)	0.021*** (0.002)	0.002 (0.003)
Log Account Age	−0.034*** (0.002)	−0.016*** (0.002)	0.009*** (0.002)	−0.018*** (0.002)	−0.039*** (0.004)	0.019*** (0.003)
User-day obs.	69,979	69,979	69,979	31,616	31,616	31,616
R^2	0.015	0.006	0.170	0.014	0.078	0.200

Notes: Stance coded as −1 (anti), 0 (neutral), +1 (pro). All models include language and day fixed effects. Standard errors clustered by user in parentheses. ⁺ $p < 0.1$; ^{*} $p < 0.05$; ^{**} $p < 0.01$; ^{***} $p < 0.001$.

Sentiment Event Studies. Figures A4 and A5 present event-study estimates for the difference-in-differences comparing entity corpus to placebo, with coefficients referenced to the pre-period average (days −7 to −1) rather than a single reference day. This approach is more robust to day-to-day fluctuations in sentiment driven by unrelated news events.

For stance outcomes (Figure A4), the pre-period coefficients are centered around zero with no systematic directional trend, supporting the parallel trends assumption. The Trump event shows a sharp positive discontinuity at day 0 for stance toward Trump, consistent with a sympathy effect following the assassination attempt. The Kirk event shows noisier dynamics but a similar centering of pre-period coefficients.

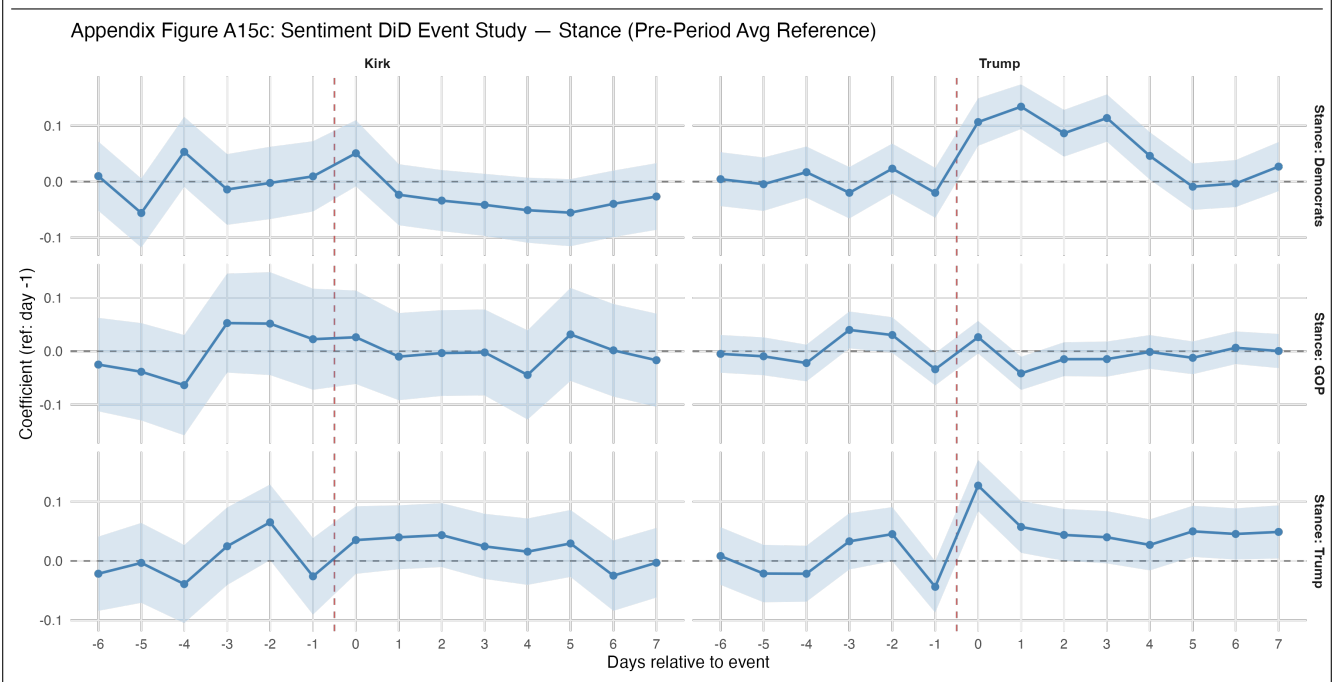
For tone outcomes (Figure A5), intensity and hostility show a notable *decline* following day 0, particularly for the Trump event. This pattern—whereby initial emotional responses moderate over the post-event window—is consistent with the fatigue and selection dynamics discussed in the main text. The pre-period variation, while present, does not exhibit systematic trending toward the event date.

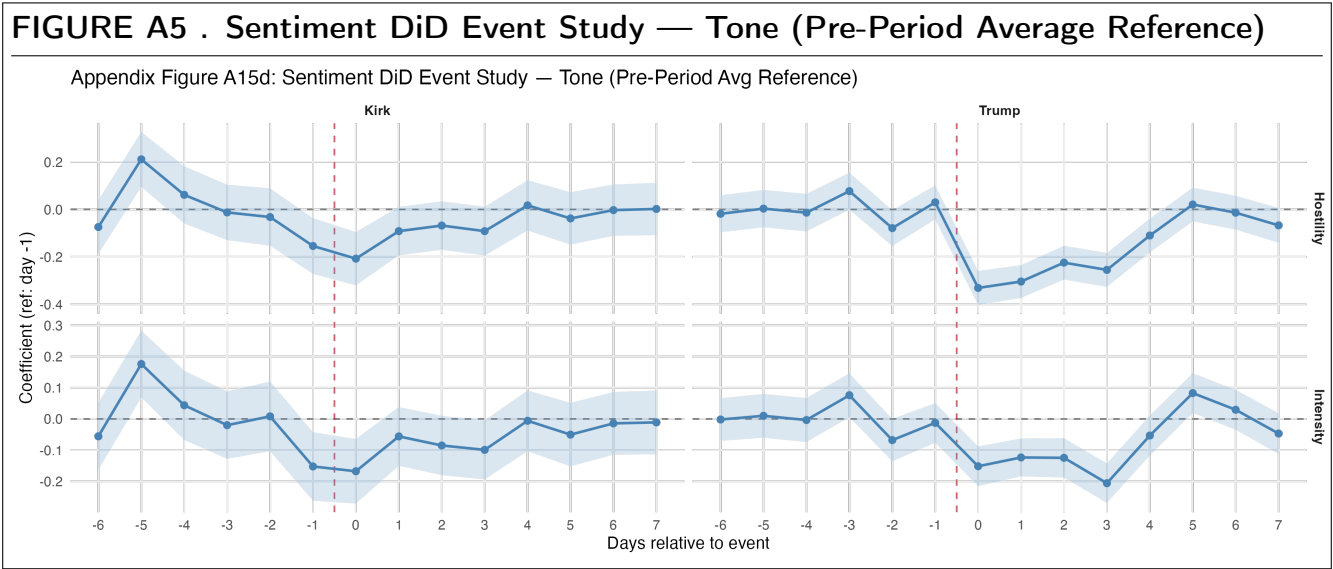
TABLE A11 . Full DiD Results (Entity vs. Placebo): Affective Tone Outcomes

	Trump Attempt		Kirk Assassination	
	Intensity	Hostility	Intensity	Hostility
Treat	−0.087*** (0.010)	−0.270*** (0.012)	−0.259*** (0.016)	−0.327*** (0.018)
Treat × Post	−0.082*** (0.012)	−0.172*** (0.014)	−0.054** (0.019)	−0.054** (0.021)
Log Followers	−0.058*** (0.002)	−0.050*** (0.002)	−0.032*** (0.006)	−0.031*** (0.006)
Log Account Age	−0.011** (0.004)	−0.003 (0.004)	−0.014* (0.007)	−0.015* (0.007)
User-day obs.	69,979	69,979	31,616	31,616
R^2	0.064	0.097	0.080	0.091

Notes: Intensity and hostility coded 0–3. All models include language and day fixed effects. Standard errors clustered by user in parentheses. ⁺ $p < 0.1$; * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

FIGURE A4 . Sentiment DiD Event Study — Stance (Pre-Period Average Reference)





Taken together, these event studies support the identifying assumption underlying the difference-in-differences design: pre-period coefficients fluctuate around zero without directional bias, while sharp discontinuities emerge precisely at the event date.

G.2 Raw Means: Entity vs. Placebo

Table A12 reports raw pre/post means for each corpus and event, providing context for interpreting DiD coefficients.

TABLE A12 . Raw Means by Corpus, Event, and Period							
Event	Corpus	Intensity		Hostility		Stance: Trump	
		Pre	Post	Pre	Post	Pre	Post
Trump	Entity	1.27	1.18	0.94	0.73	−0.03	0.05
	Placebo	1.36	1.36	1.23	1.19	−0.01	0.02
Kirk	Entity	1.26	1.37	1.07	1.19	−0.07	−0.08
	Placebo	1.60	1.73	1.49	1.63	0.02	−0.00

Notes: Raw means computed at user-day level. Δ Entity – Δ Placebo approximates the DiD coefficient.

G.3 Event Corpus Results

The event corpus—posts explicitly discussing assassination attempts—suffers from severe pre-period sample limitations. Table A13 documents the sample sizes: the event corpus contains only 72–546 pre-event observations, compared to 4,066–14,162 for the entity corpus.

Tables A14 and A15 report regression results using the event corpus. Coefficients are larger in magnitude but less precisely estimated due to the limited pre-period samples.

G.4 Triple-Difference by User Location

Table A16 reports the full triple-difference specification interacting treatment effects with user location for all sentiment outcomes.

TABLE A13 . Sample Sizes by Corpus, Event, and Period

Event	Corpus	Pre	Post
Trump Attempt	Event	546	15,438
	Entity	14,162	24,151
	Placebo	12,823	18,848
Kirk Assassination	Event	72	13,361
	Entity	4,066	8,367
	Placebo	7,512	11,675

Notes: Cell entries are user-day observations. Event corpus has severely limited pre-period observations (72–546), making DiD estimates less reliable than those using the entity corpus.

TABLE A14 . DiD Results (Event Corpus vs. Placebo): Stance Outcomes

	Trump Attempt			Kirk Assassination		
	Trump	GOP	Dem	Trump	GOP	Dem
Treat	0.110*** (0.027)	0.080*** (0.021)	0.329*** (0.021)	0.105+ (0.057)	−0.149 (0.099)	0.486*** (0.052)
Treat × Post	−0.076** (0.027)	−0.064** (0.021)	0.133*** (0.021)	−0.112* (0.057)	0.126 (0.100)	0.023 (0.052)
Log Followers	0.004*** (0.001)	0.003*** (0.001)	0.014*** (0.001)	0.003* (0.001)	0.014*** (0.002)	0.005* (0.002)
Log Account Age	−0.017*** (0.002)	−0.010*** (0.001)	0.011*** (0.003)	−0.013*** (0.002)	−0.034*** (0.004)	0.017*** (0.003)
User-day obs.	47,654	47,654	47,654	32,618	32,618	32,618
R ²	0.014	0.007	0.166	0.006	0.013	0.215

Notes: Stance coded as −1 (anti), 0 (neutral), +1 (pro). All models include language and day fixed effects. Standard errors clustered by user in parentheses. Caution: Pre-period event corpus contains only 546 (Trump) and 72 (Kirk) observations. + $p < 0.1$; * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

TABLE A15 . DiD Results (Event Corpus vs. Placebo): Affective Tone Outcomes

	Trump Attempt		Kirk Assassination	
	Intensity	Hostility	Intensity	Hostility
Treat	−0.136*** (0.036)	−0.360*** (0.039)	−0.285** (0.096)	−0.356*** (0.107)
Treat × Post	−0.302*** (0.037)	−0.250*** (0.040)	−0.302** (0.097)	−0.357*** (0.107)
Log Followers	−0.063*** (0.003)	−0.053*** (0.002)	−0.041*** (0.004)	−0.037*** (0.004)
Log Account Age	−0.016*** (0.004)	−0.020*** (0.004)	−0.018** (0.006)	−0.014* (0.006)
User-day obs.	47,654	47,654	32,618	32,618
R^2	0.116	0.149	0.138	0.164

Notes: All models include language and day fixed effects. Standard errors clustered by user in parentheses. Caution: Pre-period event corpus contains only 546 (Trump) and 72 (Kirk) observations. $^+p < 0.1$; $^*p < 0.05$; $^{**}p < 0.01$; $^{***}p < 0.001$.

TABLE A16 . Differential Response by User Location (US vs. Non-US): All Outcomes

	Trump Attempt					Kirk Assassination				
	Trump	GOP	Dem	Int.	Host.	Trump	GOP	Dem	Int.	Host.
Treat × Post (US baseline)	0.019 (0.022)	−0.017 (0.017)	0.075*** (0.021)	−0.193*** (0.030)	−0.328*** (0.033)	0.035 (0.037)	0.058 (0.055)	−0.006 (0.036)	−0.148* (0.066)	−0.169* (0.072)
× Non-US (β_4)	0.052* (0.026)	0.016 (0.020)	−0.004 (0.024)	0.124*** (0.036)	0.180*** (0.040)	−0.038 (0.041)	−0.115+ (0.062)	−0.010 (0.041)	0.053 (0.074)	0.105 (0.080)
User-day obs.	39,595	39,595	39,595	39,595	39,595	14,969	14,969	14,969	14,969	14,969
R^2	0.021	0.012	0.158	0.076	0.098	0.017	0.082	0.196	0.080	0.089

Notes: Entity vs. placebo, interacted with user location. Sample restricted to users with identifiable US or Non-US location. All models include language and day fixed effects. Standard errors clustered by user in parentheses. $^+p < 0.1$; $^*p < 0.05$; $^{**}p < 0.01$; $^{***}p < 0.001$.

FIGURE A6 . Triple-DiD Event Study — Stance (Non-US × Treat)

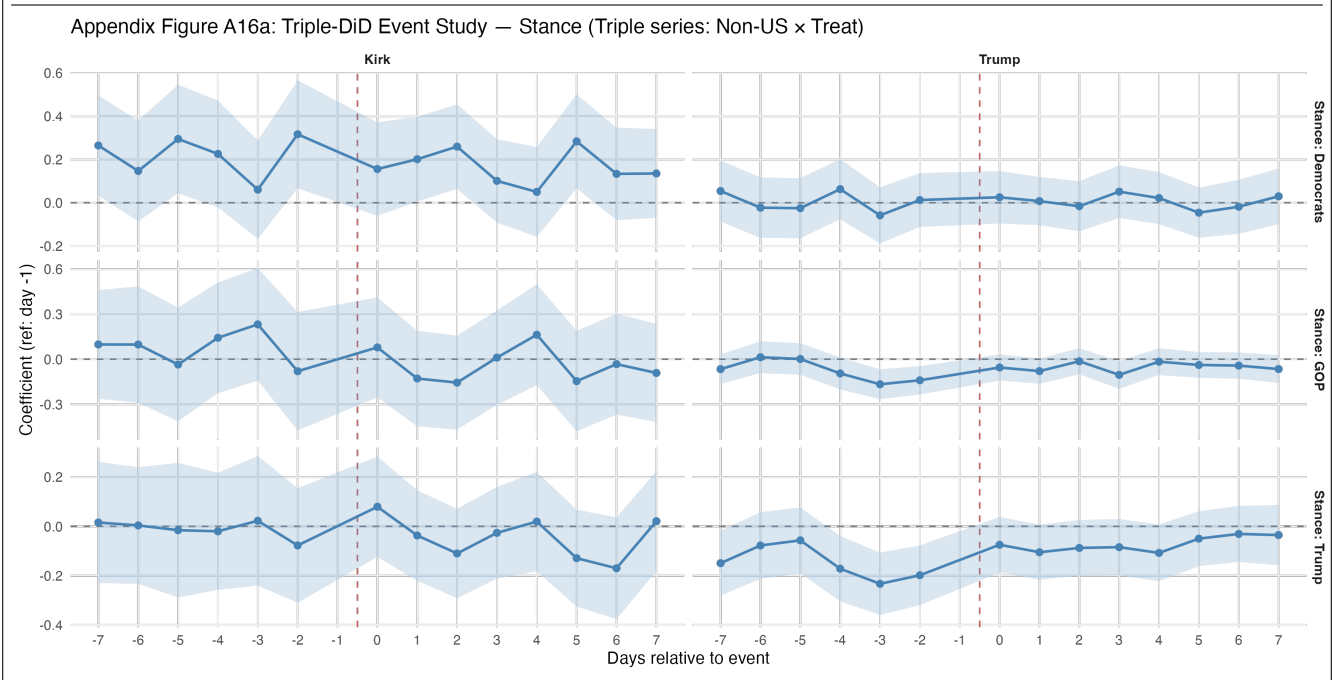
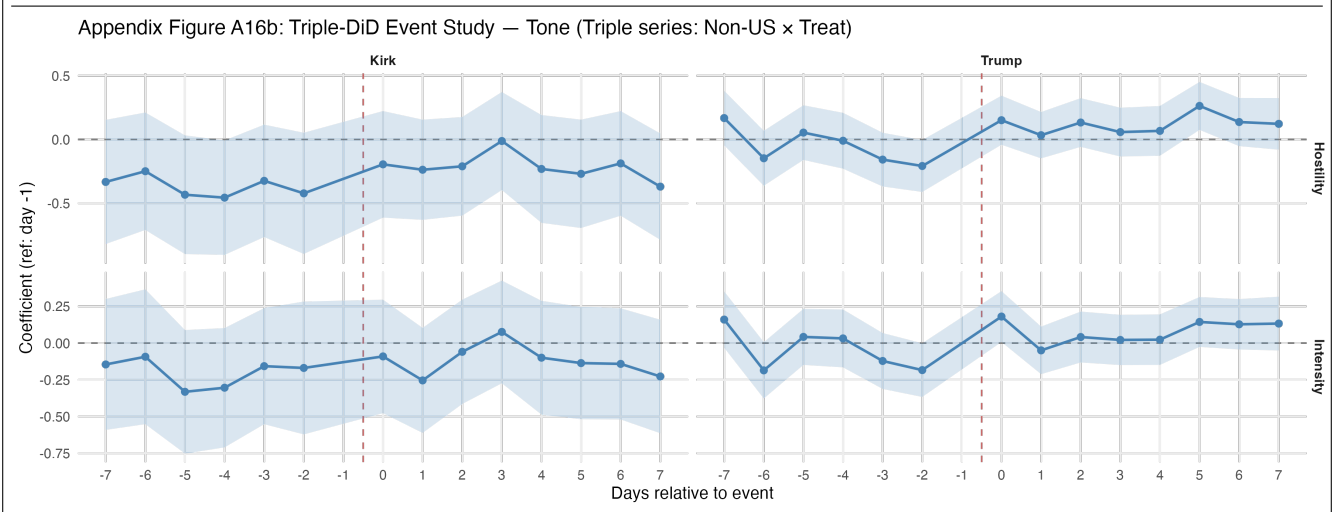


FIGURE A7 . Triple-DiD Event Study — Tone (Non-US × Treat)



Location Heterogeneity Event Studies. Figures A6 and A7 present event-study estimates for the triple-difference coefficient ($\text{Non-US} \times \text{Treat} \times \text{Day}$), which captures whether Non-US users respond differentially to the treatment relative to US users on each day. Coefficients are referenced to day -1 .

For the Kirk event, pre-period coefficients are relatively stable across all outcomes, with confidence intervals generally encompassing zero. This supports the parallel trends assumption for the triple-difference specification and validates the location heterogeneity findings in Table ???. The Trump event shows somewhat more pre-period variation, particularly for stance toward Trump, though the coefficients remain centered near zero without systematic directional trends.

Notably, post-period coefficients for the triple interaction are generally flat and close to zero across both events and all outcomes. This indicates that *conditional on posting*, Non-US and US users exhibit similar sentiment responses to the events. Combined with the volume results showing that Non-US users are disproportionately drawn into event-specific discourse (Table 5), this suggests that the key location difference operates through the *extensive margin* of participation rather than the *intensive margin* of sentiment expression. Once users engage, their expressed attitudes do not systematically differ by location.

G.5 Raw Means by User Location

Table A17 reports raw means by user location, corpus, and period. The key pattern is that US users' placebo corpus shifts more strongly than Non-US users' placebo corpus, mechanically changing the baseline against which entity-focal effects are measured.

TABLE A17 . Raw Means by User Location: Entity Corpus vs. Placebo

Event	Location	Corpus	Intensity		Hostility		Stance: Trump	
			Pre	Post	Pre	Post	Pre	Post
Trump	US	Entity	1.40	1.23	1.12	0.79	0.06	0.10
		Placebo	1.37	1.41	1.20	1.21	-0.03	-0.01
	Non-US	Entity	1.19	1.10	0.86	0.67	-0.06	0.04
		Placebo	1.25	1.26	1.13	1.11	-0.01	0.03
Kirk	US	Entity	1.38	1.44	1.19	1.24	-0.04	-0.06
		Placebo	1.60	1.81	1.49	1.71	0.07	0.02
	Non-US	Entity	1.23	1.31	1.03	1.14	-0.07	-0.08
		Placebo	1.52	1.67	1.41	1.56	-0.00	-0.01

Table A18 manually decomposes the triple-difference to show that positive β_4 arises because US users decline more than Non-US users.

G.6 Cross-Language Heterogeneity

Tables A20 and A21 report language-specific heterogeneity in event effects. French is the omitted baseline. Table A21 computes total effects for each language with correctly propagated standard errors using the full variance-covariance matrix.

G.7 Cross-Target Spillover

Tables A22 and A23 report cross-target spillover effects, comparing focal targets (Trump for Trump event; GOP for Kirk event) to cross-targets (GOP for Trump event; Trump for Kirk event).

TABLE A18 . Manual Triple-Difference Decomposition

Event	Outcome	ΔEntity_{US}	$\Delta\text{Placebo}_{US}$	DiD_{US}	DiD_{Non-US}	β_4
Trump	Intensity	-0.17	+0.04	-0.21	-0.09	+0.12
	Hostility	-0.33	+0.02	-0.35	-0.16	+0.19
Kirk	Intensity	+0.06	+0.21	-0.15	-0.07	+0.08
	Hostility	+0.05	+0.22	-0.17	-0.04	+0.13

Notes: Δ denotes post – pre change. $\text{DiD} = \Delta\text{Entity} - \Delta\text{Placebo}$. $\beta_4 = \text{DiD}_{Non-US} - \text{DiD}_{US}$. Positive β_4 indicates Non-US users show smaller declines than US users.

TABLE A19 . Sample Sizes by User Location, Corpus, and Period

Event	Location	Corpus	Pre	Post
Trump Attempt	US	Entity	2,292	4,226
		Placebo	2,164	3,382
	Non-US	Entity	5,462	9,662
		Placebo	4,962	7,446
Kirk Assassination	US	Entity	284	808
		Placebo	838	1,829
	Non-US	Entity	1,336	2,780
		Placebo	2,691	4,403

Notes: Cell entries are user-day observations.

TABLE A20 . Cross-Language Heterogeneity: Interaction Coefficients (Event Corpus vs. Placebo)

	Trump Attempt			Kirk Assassination		
	Intensity	Hostility	Stance: Trump	Intensity	Hostility	Stance: Trump
Treat $\times$ Post (French)	-0.388*** (0.081)	-0.483*** (0.086)	-0.000 (0.068)	-0.437* (0.175)	-0.618*** (0.176)	0.067 (0.045)
$\times$ Japanese	0.237* (0.097)	0.455*** (0.102)	-0.150* (0.075)	0.132 (0.223)	0.216 (0.234)	-0.306** (0.112)
$\times$ Korean	0.224 (0.300)	0.279 (0.297)	0.153 (0.163)	0.873** (0.331)	1.271*** (0.185)	-0.192 (0.136)
$\times$ Russian	0.529 (0.436)	0.873+ (0.512)	-0.370+ (0.207)	-0.684*** (0.205)	-0.690*** (0.205)	0.902*** (0.055)
$\times$ Chinese	-0.122 (0.111)	0.061 (0.128)	-0.010 (0.096)	0.201 (0.256)	0.644* (0.303)	-0.271+ (0.150)
User-day obs.	47,654	47,654	47,654	32,618	32,618	32,618

Notes: French is the omitted baseline. All models include day fixed effects. Standard errors clustered by user in parentheses. + $p < 0.1$; * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

TABLE A21 . Language-Specific Total Effects with Correct Standard Errors

Event	Language	Intensity		Hostility	
		Effect	SE	Effect	SE
Trump Attempt	French (baseline)	-0.388	(0.081)	-0.483	(0.086)
	Japanese	-0.151	(0.052)	-0.028	(0.054)
	Korean	-0.165	(0.289)	-0.204	(0.285)
	Russian	+0.141	(0.428)	+0.390	(0.504)
	Chinese	-0.510	(0.076)	-0.422	(0.095)
Kirk Assassination	French (baseline)	-0.437	(0.175)	-0.618	(0.176)
	Japanese	-0.305	(0.139)	-0.402	(0.154)
	Korean	+0.435	(0.280)	+0.653	(0.055)
	Russian	-1.122	(0.091)	-1.308	(0.090)
	Chinese	-0.236	(0.189)	+0.026	(0.248)

Notes: Total effects for non-French languages computed as $\hat{\beta}_{\text{French}} + \hat{\beta}_{\text{French} \times \text{Language}}$. Standard errors computed using $\text{Var}(a + b) = \text{Var}(a) + \text{Var}(b) + 2\text{Cov}(a, b)$.

TABLE A22 . Cross-Target Spillover Effects: Stance Outcomes

	Stance toward Trump				Stance toward GOP			
	Trump Event		Kirk Event		Trump Event		Kirk Event	
	Focal	Cross	Focal	Cross	Focal	Cross	Focal	Cross
Treat $\times$ Post	0.058*** (0.008)	-0.016+ (0.009)	0.025* (0.011)	0.004 (0.009)	-0.005 (0.006)	0.072*** (0.018)	-0.005 (0.016)	-0.063*** (0.010)
User-day obs.	69,979	42,442	31,616	56,973	69,979	42,442	31,616	56,973
R ²	0.015	0.011	0.014	0.031	0.006	0.039	0.078	0.023

Notes: Each column is a separate regression comparing the indicated corpus to placebo. Focal = Trump corpus (Trump event) or GOP corpus (Kirk event); Cross = GOP corpus (Trump event) or Trump corpus (Kirk event). All models include language and day fixed effects. Standard errors clustered by user in parentheses. + $p < 0.1$; * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

TABLE A23 . Cross-Target Spillover Effects: Affective Tone Outcomes

	Intensity				Hostility			
	Trump Event		Kirk Event		Trump Event		Kirk Event	
	Focal	Cross	Focal	Cross	Focal	Cross	Focal	Cross
Treat $\times$ Post	-0.082*** (0.012)	-0.113*** (0.022)	-0.054** (0.019)	0.008 (0.014)	-0.172*** (0.014)	-0.128*** (0.024)	-0.054** (0.021)	0.017 (0.015)
User-day obs.	69,979	42,442	31,616	56,973	69,979	42,442	31,616	56,973
R ²	0.064	0.096	0.080	0.110	0.097	0.104	0.091	0.135

Notes: Each column is a separate regression comparing the indicated corpus to placebo. All models include language and day fixed effects. Standard errors clustered by user in parentheses. + $p < 0.1$; * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

G.8 Cross-Event Comparison

Table A24 tests whether effects differ significantly between the Trump and Kirk events by pooling data and interacting treatment effects with an event indicator.

TABLE A24 . Difference in Effects: Kirk vs. Trump Event

	Stance: Trump	Stance: GOP	Stance: Dem	Intensity	Hostility
Treat × Post × Kirk	−0.029* (0.013)	−0.005 (0.017)	−0.097*** (0.014)	0.036 (0.023)	0.122*** (0.025)

Notes: Coefficient on Treat × Post × Kirk tests whether the DiD effect differs between events. Negative values indicate smaller effects for Kirk; positive values indicate larger effects. * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

G.9 Robustness: Clean Placebo

Tables A25 and A26 report robustness checks that drop placebo posts containing violence-related terms (e.g., “assassination,” “shooting,” “attack”) to address concerns about spillover contamination. Results are qualitatively unchanged.

TABLE A25 . Robustness: Drop Placebo Posts With Violence Terms (Stance)

	Trump Attempt			Kirk Assassination		
	Trump	GOP	Dem	Trump	GOP	Dem
Treat × Post	0.062*** (0.008)	−0.001 (0.006)	0.064*** (0.008)	0.021* (0.011)	−0.007 (0.016)	−0.019+ (0.011)
User-day obs.	68,812	68,812	68,812	30,530	30,530	30,530
R^2	0.015	0.006	0.169	0.015	0.080	0.207

Notes: Same specification as main DiD, but placebo posts containing violence/assassination terms are dropped. All models include language and day fixed effects. Standard errors clustered by user in parentheses. + $p < 0.1$; * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

TABLE A26 . Robustness: Drop Placebo Posts With Violence Terms (Tone)

	Trump Attempt		Kirk Assassination	
	Intensity	Hostility	Intensity	Hostility
Treat × Post	−0.078*** (0.012)	−0.165*** (0.014)	−0.044* (0.019)	−0.046* (0.021)
User-day obs.	68,812	68,812	30,530	30,530
R^2	0.063	0.096	0.077	0.088

Notes: Same specification as main DiD, but placebo posts containing violence/assassination terms are dropped. All models include language and day fixed effects. Standard errors clustered by user in parentheses. + $p < 0.1$; * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

G.10 Event Corpus vs. Entity Corpus Comparison

Table A27 compares DiD estimates across corpus definitions. Event corpus effects are larger in magnitude but noisier due to limited pre-period samples.

TABLE A27 . Comparison of DiD Estimates Across Corpus Definitions

Event	Outcome	Event Corpus		Entity Corpus	
		β	SE	β	SE
Trump Attempt	Intensity	−0.302***	(0.037)	−0.082***	(0.012)
	Hostility	−0.250***	(0.040)	−0.172***	(0.014)
	Stance: Trump	−0.076**	(0.027)	0.058***	(0.008)
	Stance: GOP	−0.064**	(0.021)	−0.005	(0.006)
	Stance: Dem	0.133***	(0.021)	0.066***	(0.008)
Kirk Assassination	Intensity	−0.302**	(0.097)	−0.054**	(0.019)
	Hostility	−0.357***	(0.107)	−0.054**	(0.021)
	Stance: Trump	−0.112*	(0.057)	0.025*	(0.011)
	Stance: GOP	0.126	(0.100)	−0.005	(0.016)
	Stance: Dem	0.023	(0.052)	−0.031**	(0.011)

Notes: Event corpus effects are larger in magnitude but estimated with substantial noise due to limited pre-period samples (72–546 obs). Entity corpus provides more reliable estimates. ⁺ $p < 0.1$; * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.