

Threats to AI Infrastructure and Executives

Source: Liferaft, Data Centers / Domestic Sabotage Query

Period: Feb to Jun 2026 (120 days)

Sample: 7,714 posts across 13 platforms

Threats to AI Infrastructure & Executives

Source Liferaft, Data Centers / Domestic Sabotage query **Period** 17 Feb to 17 Jun 2026 (120 days)

Sample 7,714 posts across 13 platforms

AT A GLANCE.

Online chatter about sabotaging AI data centers and targeting AI executives **peaked at roughly 7x daily volume in May** (3,441 posts) before showing initial signs of moderation in June (1,073 posts in 17 days, equating to ~4x the February baseline). Sentiment remains overwhelmingly hostile, with a **mean score of -66 on a -100 to +100 scale**. Threat-themed language (arson, sabotage, Molotov, “blow up”) appears in **roughly 1 in 3 posts**.

Source: Liferaft, **Period** 17 Feb to 17 Jun 2026 (120 days) **Sample** 7,714 posts across 13 platforms

Volume & Escalation

Monthly volume escalated from sub-200 posts in late February to a peak of **3,441 posts in May**, a ~7x increase in daily rate from the February baseline. June (17 days captured) is on pace for approximately 1,900 posts, suggesting the steep escalation is moderating but remains well above pre-April baseline. The two highest-volume weeks (6 to 12 April, and 11 to 17 May) together account for **25% of the entire 120-day dataset**.

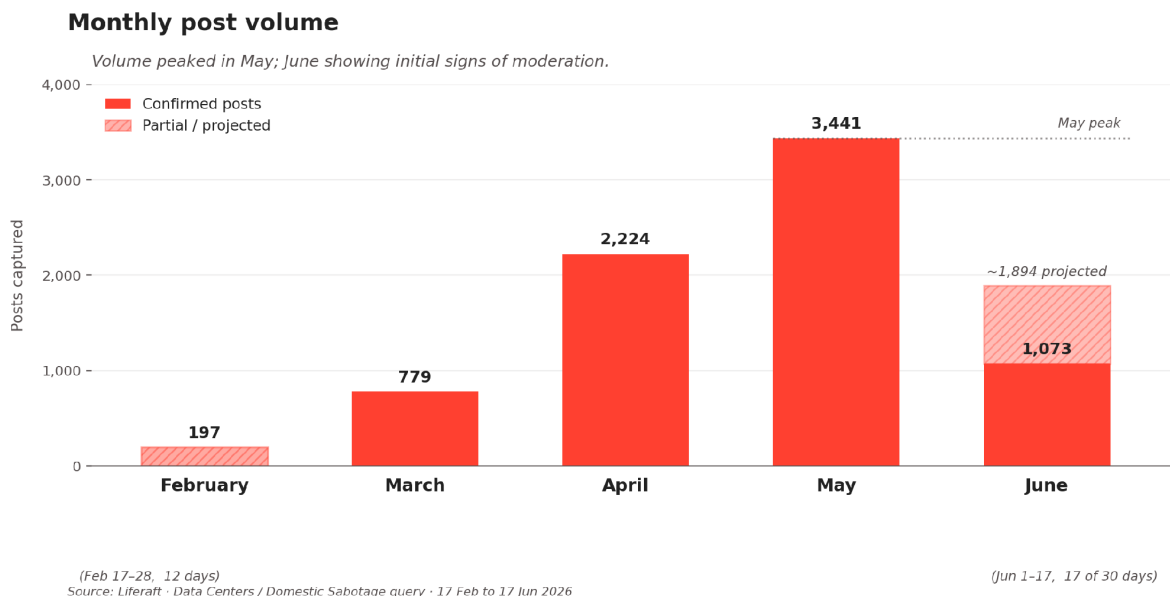


Figure 1. Monthly post volume. Hatched fill indicates partial months and projection.

Executive Perspective

“What we are seeing is not evidence of a coordinated plot, but it is a clear **warning signal**. When hostile rhetoric starts to cluster around specific infrastructure, local grievances, and action-oriented language, security teams need to treat the online environment as an early indicator of real-world risk.”

— Eduardo Capouya, co-founder and Chief Innovation Officer, Liferaft

Executive Targeting

Sam Altman is named in **450 posts**, roughly **8x** the next-most-named figure. **58% of Altman mentions remain concentrated in a single week** (6 to 12 April). This continues to reflect reaction volume around one viral claim rather than 450 discrete threats.

Threat rhetoric, by share of posts referencing each theme.

- **Arson / “burn down”** 23%
- **Sabotage variants** 21%
- **Molotov / firebomb** 11%
- **“Blow up” / explosive / bomb** 6%
- **“Target list” / “hit list” / “kill list”** 3%

Platform distribution.

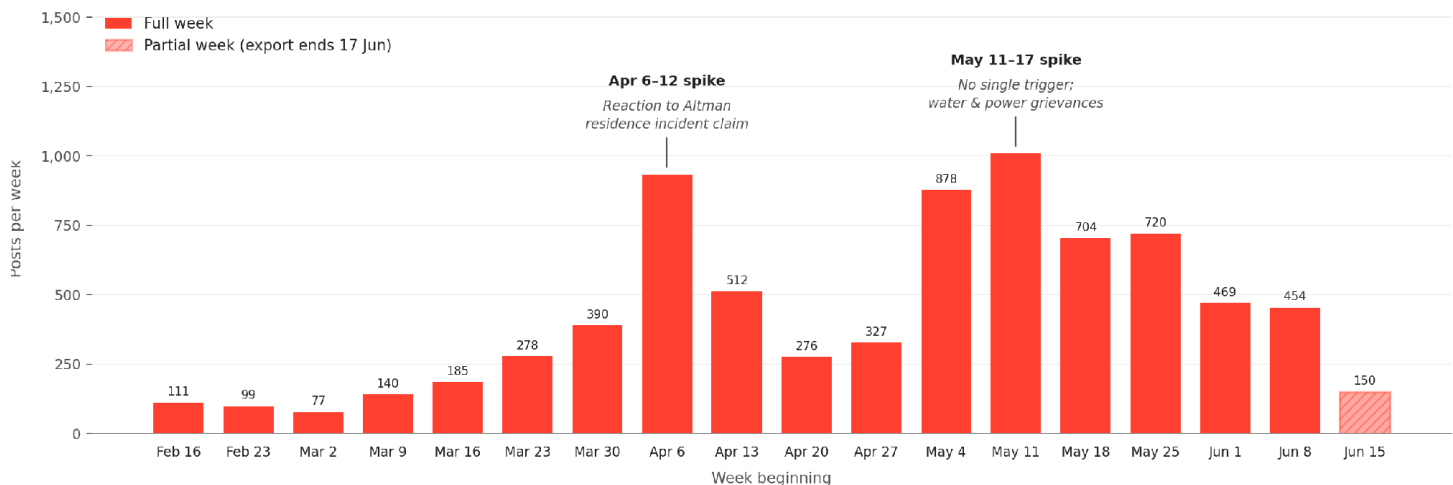
Bluesky represents just **12% of related posts** but accounts for **54% of high-risk-flagged posts**.

Weekly trend, 120-day view.

A weekly breakdown surfaces two distinct spike weeks. The **6 to 12 April spike** was reactive and coherent: 28% of posts that week named Sam Altman and 33% referenced “Molotov,” driven by viral circulation of a claim about an attack on his residence. The **11 to 17 May spike** is structurally different: no single triggering event, and chatter distributed across water and power grievances (Lake Tahoe-area data centers feature heavily), anti-surveillance rhetoric, and broad anti-AI sentiment. Weekly volume has stepped down since late May, settling near 450 posts per week in June versus 700 to 1,000 across May.

Weekly post volume

Two spike weeks driven by different dynamics; weekly volume moderating since late May.



Source: Liferaft • Data Centers / Domestic Sabotage query • 17 Feb to 17 Jun 2026

Figure 2. Weekly post volume across the 120-day window.

Stated Intent

A companion query isolated posts containing first-person and group intent constructions, then filtered for sabotage-relevant content. The resulting set: **199 intent-bearing posts** across 17 April to 17 June 2026. **22% of intent posts are flagged High risk**, versus 1.2% across the broader dataset. Intent language correlates strongly with operational concern. Intent volume grew from 22 posts in April to 135 in May; June (17 days) has logged 42.

Across April to June, the mix of intent constructions **shifted from individual to collective**. First-person language dominated April (50% of intent posts that month). By May, group / collective rhetoric surged to 36%, reflecting a shift in framing from personal frustration to coordinated action language. Executive-specific intent remains rare. Intent is overwhelmingly directed at infrastructure, not named individuals.

First-person individual. 80 posts · 40% of intent-relevant posts

Constructions: "I will," "I want to," "I'm going to," "I'm gonna," "I need to." Representative examples:

"Same like I want to start burning down data centers" - Twitter (X), High risk

"I will personally torch every data center with my own two hands." - Bluesky, High risk

Group / collective. 62 posts · 31% of intent-relevant posts

Constructions: "we will," "we're going to," "let's," "we need to," "we're gonna." Representative examples:

"We are going to blow up the data centres :)" - Twitter (X), High risk

"We are going to see massive crowds of people setting data centers ablaze. It's just a matter of when." - Twitter (X), Medium risk

Anticipatory/conditional. 12 posts · 6% of intent-relevant posts

Constructions: "someone should," "can't wait until," "would be a shame if." Representative examples:

"Someone needs to blow up the data centers" - Twitter (X), High risk

"Man I can't WAIT for the data centers to burn down" - Bluesky, Medium risk

NOTES ON METHODOLOGY.

This brief draws on three Liferaft queries: a main *Data Centers / Domestic Sabotage* query (7,714 posts), a companion *Instructional Material / Guides* query (114 posts), and an *Intent Fusion* query (277 posts, filtered to 199 sabotage-relevant). All posts are captured from public sources. Author handles and direct links have been withheld in this brief.

Three caveats matter for editorial use: (1) the main dataset captures rhetoric and sentiment, not credible plots. Most posts are venting or political commentary rather than operational planning; (2) spike-day totals are inflated by reactions to a small number of viral posts, and the April spike claim about an attack on Sam Altman's residence warrants; (3) the intent query covers narrower windows than the main query.

We're here to help.

AI infrastructure is becoming a flashpoint for online hostility, local grievances, and action-oriented threat rhetoric. While most online chatter does not indicate a credible plot, escalating volume, hostile sentiment, and intent-bearing language can create meaningful early warning signals for security teams. You can't protect what you can't see. Proactively monitoring OSINT for emerging narratives, infrastructure-specific threats, and executive targeting should be part of any organization's physical security and intelligence practices.

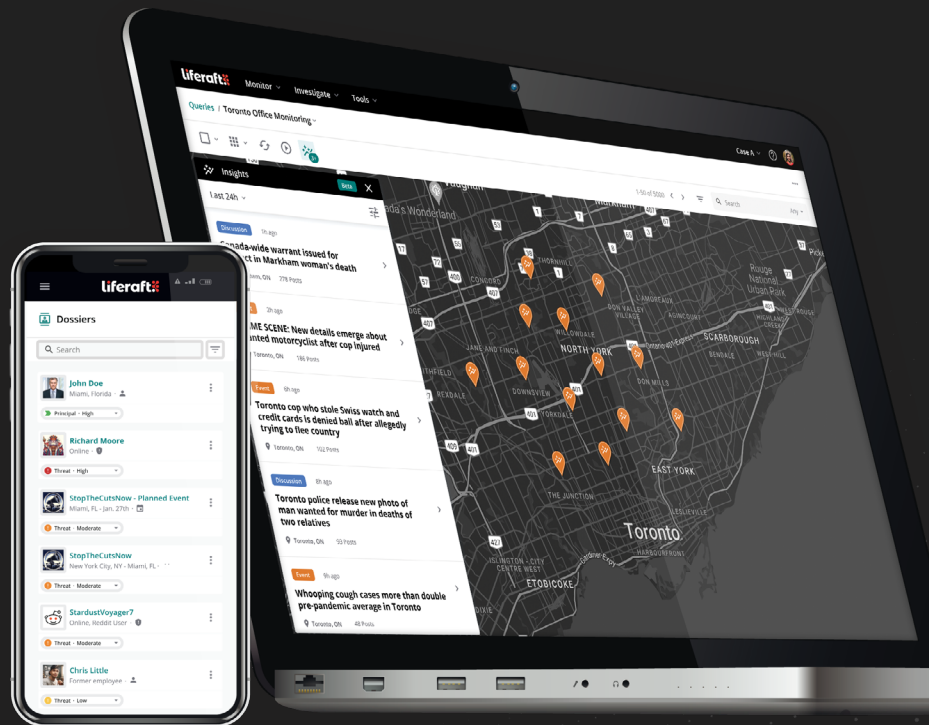
Why Liferaft?

In an environment where online rhetoric can quickly shift from frustration to operational concern, security teams need more than data, they need context.

Liferaft is a decision enablement platform that helps teams turn open-source intelligence into timely, actionable insight.

By surfacing relevant signals across social, alternative, deep, and dark web sources, Liferaft helps analysts detect emerging threats, assess risk around critical infrastructure and executives, and understand where online activity may indicate real-world security concerns.

With intuitive workflows, flexible integrations, and expert support, Liferaft helps security and intelligence teams move from noise to informed action, faster.



Liferaft
A Securitas company

Protect what **matters**.

Visit liferaftlabs.com/demo to schedule a discovery call.