



The Threats Hiding in Plain Sight: Coded Violent Rhetoric Targeting Big Tech

Alethea Risk Intelligence Brief · September 2026

Executive Summary

Reports on a rise in violent threats against AI executives have prompted a broader debate across mainstream and fringe platforms, with a convergence on the shared framing that executives, investors, and physical infrastructure are legitimate targets for holding AI companies accountable for public grievances, an Alethea analysis found.

Using the Artemis for Security platform and threats of violence model, Alethea reviewed 36.8K posts referencing AI executives across X, Reddit, Bluesky, and news sources between January and August 2026, finding that content placing violent language in proximity to a named AI executive ran at roughly one post per day from mid-January through mid-July.

On July 18, it jumped to 57 per day, and on August 11, it reached 255 per day, **marking a 250-fold increase over baseline in less than a month**. Notably, the rise in press coverage about threats against AI executives appears to have played a role in the escalation in volume, as users responded to the coverage with approval and satisfaction that the executives were afraid.

EXHIBIT 1

Timeline — Post Volume



Source: Alethea, Artemis for Security platform.

By volume, this is a predictable rise in hostility as AI reaches further into jobs, utility bills, and local land use ahead of the midterms. By substance, it's a different story; a blind spot in traditional detection, which was built to flag words rather than pattern.

Threatening rhetoric increasingly does not announce itself, naming a tactic instead of an act, with users invoking a "Sam Altman special," saying a target looks like he wants "Molotov time," or writing that they are whistling the Luigi's Mansion theme, all phrasing that contains no violent verb, no stated intent, and nothing explicitly violent.

High profile attacks on executives often spawn new or novel vocabulary. Without updating monitoring after these events, real escalation can read as ordinary criticism.

This environment is difficult to monitor for several reasons:

- **Event-generated vocabulary:** the phrasing changes and evolves after every incident.
- **Dispersion:** the content shows up in communities with no connection to AI companies and executives.
- **Low keyword yield:** the most target-specific material (imagery, eponyms, historical references) contains no keywords that a conventional system would flag.
- **Portable:** the language can transfer easily, cross-industry, from executives to employees to local officials to family members

Because of these factors, a low keyword match rate is not the same as a low threat level. **Increasingly, the content that may look like noise can be where the most specific targeting lives.**

Alethea also observed an underlying shift from how users feel about AI executives to what they picture happening to them: naming them dead, staging it in imagery or coded wording, and treating it as a foreseeable and deserved outcome. The rhetoric has also extended beyond executives, targeting engineers, facility staff, and local officials who approve permits.

KEY TAKEAWAY

None of this is confined to technology leaders. Any executive whose company deploys AI into a contested public domain inherits the same structure: an identifiable individual, a grievance with a body count attached, and a vocabulary of coded violence already built and waiting to be pointed at them.

The Coded Language Environment

Alethea's review found that the majority of hostile content directed at AI executives contained little to no threatening language. It instead names a tactic rather than an act, substitutes a proper noun or historical reference for a verb, or arrives as an image. That's exactly the content keyword-based monitoring is built to miss, and precisely what our pattern analysis is designed to surface.

Names Function as Verbs

The Luigi Mangione reference established the pattern of supplying the shorthand, and remains well in use nearly two years after the killing. Users write that Sam Altman “needs to get Luigi'd,” “needs to meet a Luigi real bad,” or “Luigi the Sam Altman and people will literally celebrate you as the hero for history for centuries ahead.”

Incidents Become Eponyms

Notably, the April attack on Altman's home generated its own eponym within days, applied to tech executives with no connection to OpenAI, demonstrating how any incident against an AI exec that receives coverage is a candidate to become the next referenceable tactic, invokable against any executive by name without describing an act or stating any intent.

One X user warned that a target was going to get “LUIGI'd... Or at least a Sam Altman special” while a Bluesky user said they hoped Zuckerberg would get a Molotov thrown at his house “next.” An X user posted that Microsoft's CEO “looked like he wants Molotov time,” while another asked “why doesn't anyone ever firebomb” the house of Nvidia's CEO.

French Revolution Language Runs Across Every Grievance

The largest coded register Alethea identified casts AI executives as aristocrats and users as peasants. One user complaining about model access limits called Anthropic “an arrogant monarchy — issuing the latest decree from Emperor Dario to us peasants,” and asked for images of Amodei saying “let them eat cake.” Residents opposing data centers have called Altman a “modern day marie antoinette.”

The same French Revolution wording is used for calls to kill them: variations of “[CEO] deserves the guillotine” were observed targeting Sam Altman, Jensen Huang, Dario Amodei, and more. One account turned it into a slogan, stating, “#bringbacktheguillotine #abolishbillionaires.” Another user said “mark zuckerberg... the last thing you ever see will be the blade of a guillotine.”

On Nvidia, one user suggested AI companies “need some guillotine insurance for the societal turbulence ahead.” A Bluesky user asked “What size guillotine do you recommend for Jeff Bezos, Sam Altman, Peter Thiel, et al.,” providing a named target list framed as a sizing question.

Online Rhetoric Turns Into Offline Action

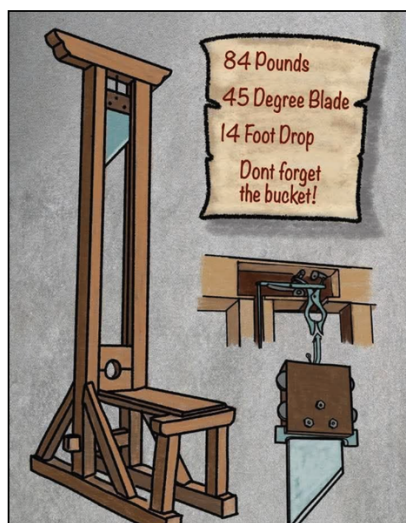
Notably, on August 6, a resident towed a roughly one-meter replica guillotine behind a bicycle outside a Salem, Oregon city council hearing on a proposed \$5.1 billion AI data center, resulting in company reps leaving the meeting due to safety concerns.

Within 48 hours of the incident, online users converted the symbol into a method. One post stated the guillotine “needs to be brought to every meeting where data centers will be discussed nationwide.” Another framed it as instruction:

“Want the Secret to Stop AI Data Center Developments? Show up to the City Council Hearings with a Guillotine.”

— Online user, following the Salem, Oregon incident

One user proposed cities “buy guillotines to make the data centers die” and store them “for the next time,” while a Salt Lake City user suggested replicating it locally.



Source: Sample images depicting French revolution imagery from various social platforms

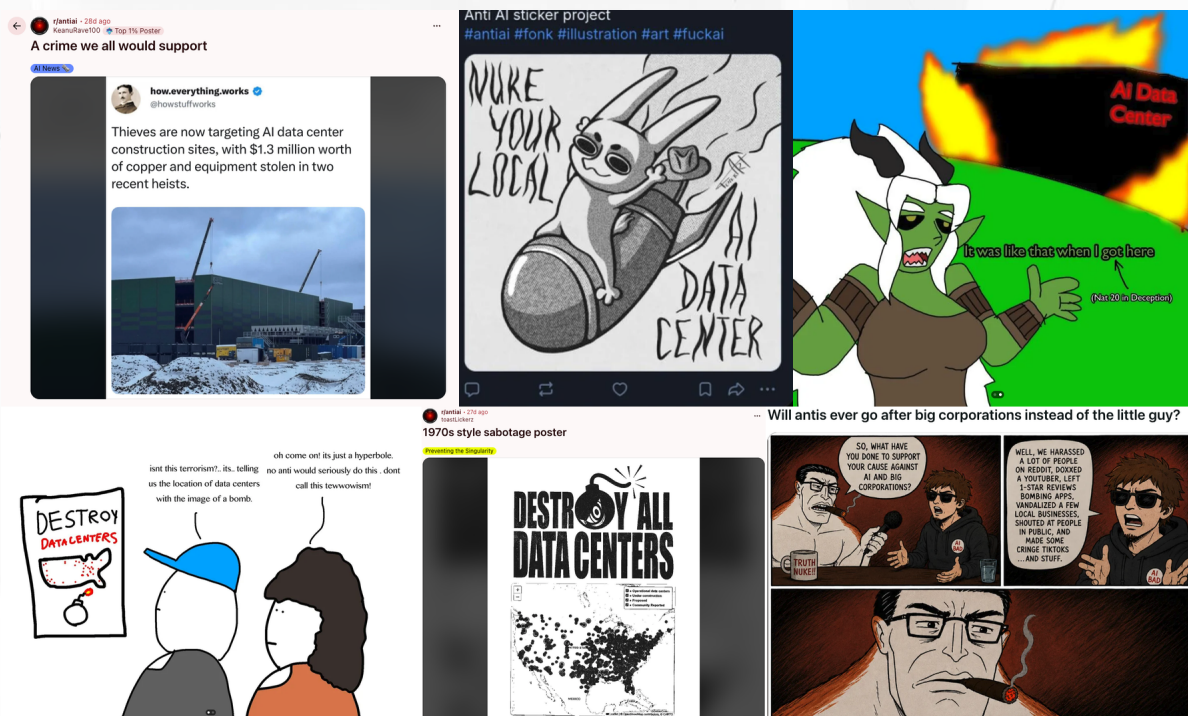


Imagery and Meme Formats

Beyond the coded language described above, a second and less visible register carries the same content through images and text-based memes, which traditional keyword-based monitoring cannot read, as captions of images are frequently innocuous, while the threat sits entirely in the accompanying graphic.

EXHIBIT 2

Representative imagery circulating across Reddit and X



Source: Alethea, Artemis for Security platform

Alethea found that images, particularly on Reddit, are where the most operationally specific content in this environment appears. A 1970s-style sabotage poster circulating under the heading “Preventing the Singularity” pairs the text “DESTROY ALL DATA CENTERS” rendered around a lit bomb, with a map of the United States plotting facilities by category: operational, under construction, proposed, and community reported.

Other related images include an anti-AI sticker depicting a cartoon figure riding a falling bomb captioned “NUKE YOUR LOCAL AI DATA CENTER,” an illustration of a data center in flames, and a text-based image about \$1.3 million in copper and equipment stolen from construction sites under the title “A crime we all would support.”

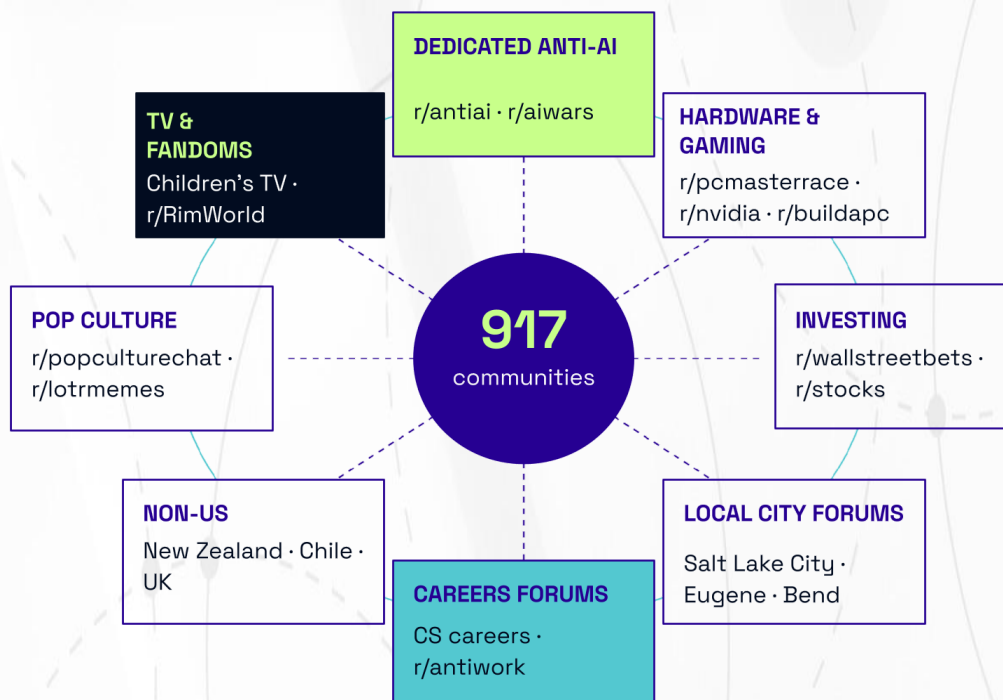
Read individually, the images appear as jokes, merchandise, or commentary. However, read together, they show a community establishing sabotage as a shared reference point, and the imagery pairing a bomb with the locations of named facilities functions as reconnaissance material regardless of the intent behind it.

This content also persists in ways text does not. An image circulates as a file, gets reposted, screenshotted, and printed onto merchandise, surviving long after the originating account is suspended or the post removed while carrying no searchable string to flag.

For security teams, the practical consequence is that the most target-specific material in this environment is the material least likely to surface in monitoring built around language.

Where These Conversations Are Happening

Alethea's review found that hostile content directed at AI executives is not confined to dedicated anti-AI communities. While spaces such as the subreddits [r/antiai](#) and [r/aiwars](#) produce a steady volume of hostile content, the majority of the content surfaced in communities with no connection to AI.



For example, a since-deleted Reddit [post](#) deliberating which tactics would most effectively [intimidate tech CEOs](#), including whether more extreme violence would be more effective than the attempts on Altman, appeared in a computer science careers forum, while a post suggesting that Altman [buy a bulletproof vest](#) because someone is going to shoot him appeared in a subreddit for an animated children's series.

Guillotine rhetoric aimed at named executives surfaced in communities for [PC hardware](#), [self-hosting](#), [celebrity gossip](#), [personal finance](#), and local city subreddits including [Salt Lake City](#), alongside non-US communities in [New Zealand](#) and [Chile](#).

The risk is that monitoring scoped to obvious venues will miss a high volume of rhetoric, as hostile content can arrive as a comment inside an unrelated conversation, produced by a user with no anti-AI posting history, in a forum with no apparent connection to the subject.

Beyond Executives: Staff, Facilities, and Local Officials

Coded violent rhetoric in this environment is not confined to named principals. Alethea observed the same vocabulary applied to engineering staff, facility representatives, and local officials. This extends exposure past the executives that a traditional monitoring posture is typically built around.

One account combined all three tiers in a single post, stating: “We will kill all the AI engineers, kill all the computer scientists & kill all the ai billionaire oligarchs. Reign of Terror French Revolution is human nature.” Another user, responding to a Meta push poll about data centers, wrote that “AI infrastructure should be burnt to the ground and the programmers should meet the guillotine.” Similarly, a Reddit post said, “Anger is typically directed only towards mega corps but the scientists, engineers, and techbros who are actually building it are never called out for their complicity.”

Users have also discussed whether employees could sabotage systems from inside their own companies, and one post anticipated that departing staff would not escape “the guillotines when the blades start eating the wealthy.”

This demonstrates how a grievance directed at a company resolves into named people at every level of it, ranging from the chief executive to the engineer to the official who approved the permit.

Why This Matters

The rhetoric in this environment is changing faster than the methods used to detect it. Volume is rising while content carrying the most specific information about targets, tactics, and locations is the content least likely to contain a keyword that is flagged as a threat.

This type of rhetoric is challenging to detect without the right tools. The vocabulary changes after every incident and cannot be defined in advance. It is dispersed, appearing in communities with no topical connection to AI and among users with no prior history on the subject. And it is portable, transferring intact from executives to engineers to local officials, and from AI to any other industry that puts automated decision-making into a contested public setting.

For CSOs and Protective Intelligence Teams

RECOMMENDATIONS

Read for indicators, not just violent language. A named eponym, a historical reference, a sizing question about equipment, or a map paired with a bomb graphic carries more information than a keyword match.

Retune after every incident, not on a schedule. Each publicized attack or protest generates new shorthand within days. A term list built before April would not have caught "Sam Altman special"; one built before August would not have caught the Salem guillotine as a tactic.

Scope coverage by target and language, not by venue. Watchlists organized around known anti-AI forums will miss the majority of hostile content. Coverage should follow the principal's name and the coded vocabulary wherever they appear, including in hobbyist, local, and non-English communities.

Extend the protected population. Engineers, facility representatives, and permitting officials appear in the same rhetoric as execs and are named as complicit rather than incidental.

Treat imagery as evidence. Text-based collection cannot resolve the highest-specificity material in this environment. Image and video review belongs in standing coverage, and artifacts should be captured on discovery, since they persist and recirculate long after the originating account is removed.

For Communications and Security Teams Working Jointly

RECOMMENDATIONS

Coverage of the threat propagates the threat. The largest escalations Alethea observed followed press cycles about violence against tech leaders rather than any executive action, and the dominant register in response was satisfaction and celebration of the news coverage rather than sympathy. Security and communications need a shared view before a principal discusses personal safety publicly.

Public caution is not protective cover. Executives who have positioned themselves as the industry's most careful actors appear in this rhetoric alongside everyone else, indicating that a safety-forward public posture does not always reduce exposure.

Not all hostility is a threat, and treating it as such carries its own cost. Much of this content is criticism, complaint, or dark humor, and forums include people arguing explicitly against violence. Separating lawful criticism from targeting is what allows a team to act on the material that matters, ignore the noise, and respond proportionately when something genuine appears.

How Alethea Helps

Alethea gives security, protective intelligence, and communications teams the intelligence to distinguish a coded threat from a complaint, and to know which function should own the response.

Unlike keyword-based monitoring, Alethea's approach is AI-backed and based on decades of analyst experience and labeled data sets. That means cross-platform coverage spanning mainstream, niche, fringe, and non-English communities; image and video review as part of standing coverage; monitoring that extends to workforce and facility-level exposure alongside named execs; actor attribution and coordinated inauthentic behavior detection; coordinated takedown support; and strategic communications counsel, including the assessment of when public engagement will increase an executive's exposure rather than reduce it.

Above all, it means a monitoring posture that is re-tuned against the language each new incident produces to ensure the most accurate, holistic picture of risk. If your organization wants a baseline read of its executive threat environment, contact us to learn more.

Methodology note: This analysis draws on Alethea's monitoring of mainstream, niche, broadcast, and foreign-language coverage between January 2026 and August 2026, using its Artemis for Security platform's proprietary threats of violence model.