

ONLINE DISCOURSE INTELLIGENCE

Claude

Initial deep dive into brand and online discourse

Purpose

A corporate-grade starting point for understanding the brand's online conversation, major language communities, sentiment, audience signals, and important voices—designed to reveal promising directions for deeper investigation.

Research window

January 1, 2024–August 18, 2026

Prepared for J345 • United States online discourse

Executive readout

+730%	829.0M	8.74M	2.15M
VOLUME TREND	EST. VIEWS	ENGAGEMENTS	SOCIAL UNIVERSE

Claude’s U.S.-inferred discourse is smaller than ChatGPT’s but accelerating faster. Coding, open-source tools, model comparison, Cowork, file creation, agentic workflows, and safety reporting organize the conversation. The brand is expanding beyond developers, yet technical and professional identity remains the center of gravity.

Sentiment is distinctly favorable: 46% positive, 23% negative, and +33% net sentiment. Surprise and Joy lead the emotional profile, signaling capability discovery. Love vs. Hate remains –24%, however, so strong functional evaluation should not be confused with deep emotional attachment.

The leading U.S. narrative is Cursor, using Claude, open-source, and coding. Cowork and computer-navigation language is more positive, while a blackmail/incident cluster is only 11.3% positive and disproportionately important to reputation. These communities should be treated as separate evidence systems.

The default influencer list is dominated by X/Twitter, large news and technology publishers, and a handful of high-reach individuals. Elon Musk ranks first despite only seven matched records; Oprah generates more than 100,000 engagements from nine records. Reach, activity, and realized engagement clearly identify different forms of importance.

Strategic reading

Claude’s U.S. opportunity is to translate developer credibility into a broader promise of high-quality work, visible control, and thoughtful assistance—while addressing safety controversy and value friction with evidence specific to each issue.

Research scope and query design

Dimension	Setting	Interpretive consequence
Window	Jan. 1, 2024–Aug. 18, 2026	Captures multiple product and reputation cycles; August 2026 is partial.
Geography	United States (Starscape-inferred)	Directional view of U.S.-inferred public discourse; not a population estimate.
Language	English represented 91.4% of U.S.-inferred matched records.	Multilingual U.S.-inferred discourse remains in scope.
Sources	Indexed public social platforms	Platform availability, deletion, privacy settings, and indexing affect coverage.
Analytical views	Trends, Sentiment, Narratives, Topics, Audience, Influencers	Machine-coded outputs require record-level validation before high-stakes decisions.

Boolean query

Saved query

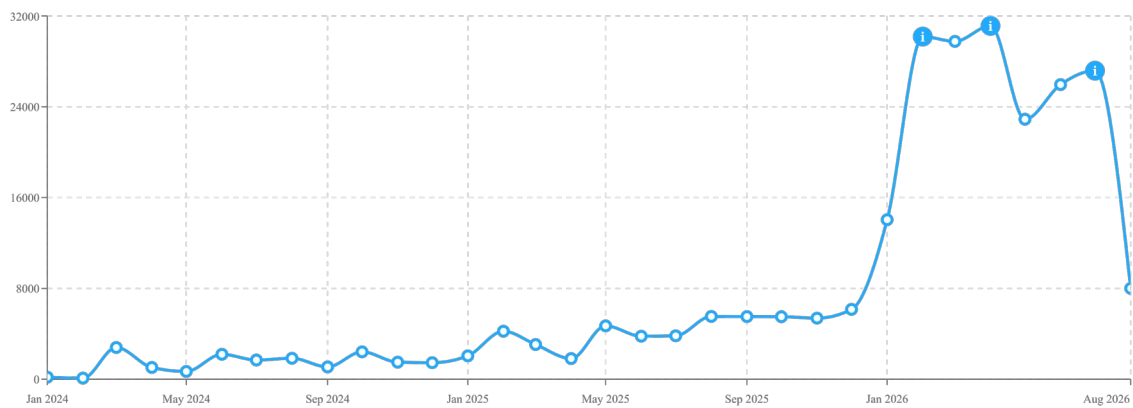
("Claude AI" OR "Claude app" OR "@ClaudeAI" OR #ClaudeAI OR "Anthropic Claude" OR (Claude NEAR/3 (Anthropic OR AI OR chatbot OR assistant OR model OR Opus OR Sonnet OR Haiku))) NOT ("Claude Monet" OR "Claude Debussy" OR "Claude Giroux" OR "Claude Shannon" OR "Claude McKay")
Geography filter: Country contains United States of America

Method discipline

- Records are matched posts or mentions, not unique people, users, customers, or a probability sample. Country is inferred from available public signals and can exclude or misclassify records.
- Narrative, sentiment, emotion, demographic, and entity labels are model-based classifications.
- Co-mention does not establish stance: a post may praise, criticize, compare, quote, or merely reference the brand.
- Volume, views, engagements, and social-universe estimates describe different phenomena and should not be added together.

Platform documentation: [Getting Started](#) | [Boolean queries](#) | [Drilldowns](#)

Conversation momentum and inflection points



Starscape monthly matched-record trend for the U.S.-filtered Claude consumer-brand query. The 2026 acceleration and partial August endpoint warrant shorter-window drilldowns before causal claims.

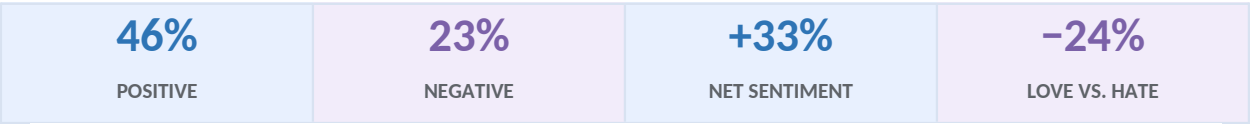
Phase marker	Verified product context	Strategic reading
Jun.–Aug. 2024	Claude 3.5 Sonnet launched with Artifacts; Artifacts then became broadly available across free and paid plans.	Claude became a workspace for making and iterating, not only a conversational model.
Feb.–May 2025	Claude 3.7 introduced hybrid reasoning and Claude Code; Claude 4 expanded coding, agents, and tool use.	Developer authority deepened, but the public system-card “blackmail” discussion also created a distinct safety controversy.
Sep. 2025–Jun. 2026	Sonnet 4.5 expanded file creation and long-running agent work; Sonnet 5 pushed agentic capability into the default Free/Pro experience.	The discourse moved from model quality toward workflows, limits, automation, and professional value at scale.

Attribution caution

The phase markers align official product events with visible conversation shifts. They identify hypotheses for drilldown, not proof that a single announcement caused a given spike.

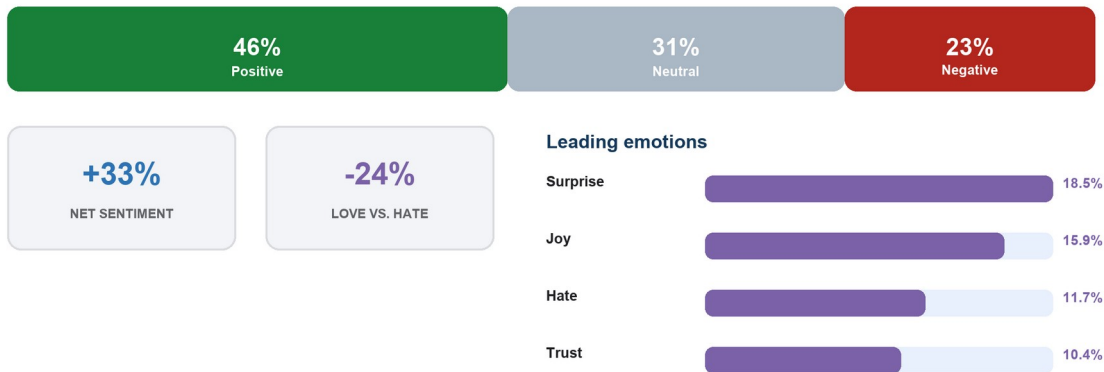
Official context: [Claude 3.5 Sonnet](#) | [Artifacts availability](#) | [Claude 3.7 and Claude Code](#) | [Claude 4](#)

Sentiment and emotional texture



Claude U.S. sentiment profile

January 1, 2024–August 18, 2026



Source: Infegy Starscape dedicated Sentiment and Emotion page; U.S.-inferred records

Analyst-rendered sentiment and emotion view based on Starscape metrics for the full research window. Dedicated sentiment metrics are used here because overview widgets can apply different scope or calculations.

Favorable evaluation is anchored in surprise and joy around capability, especially coding, files, agents, and computer use. Hate still exceeds Trust, indicating that safety controversy, model rivalry, and value or access issues remain emotionally salient even within a positive functional profile.

Leading emotions

Surprise 18.5% • Joy 15.9% • Hate 11.7% • Trust 10.4%

Narrative structure: what the conversation is about

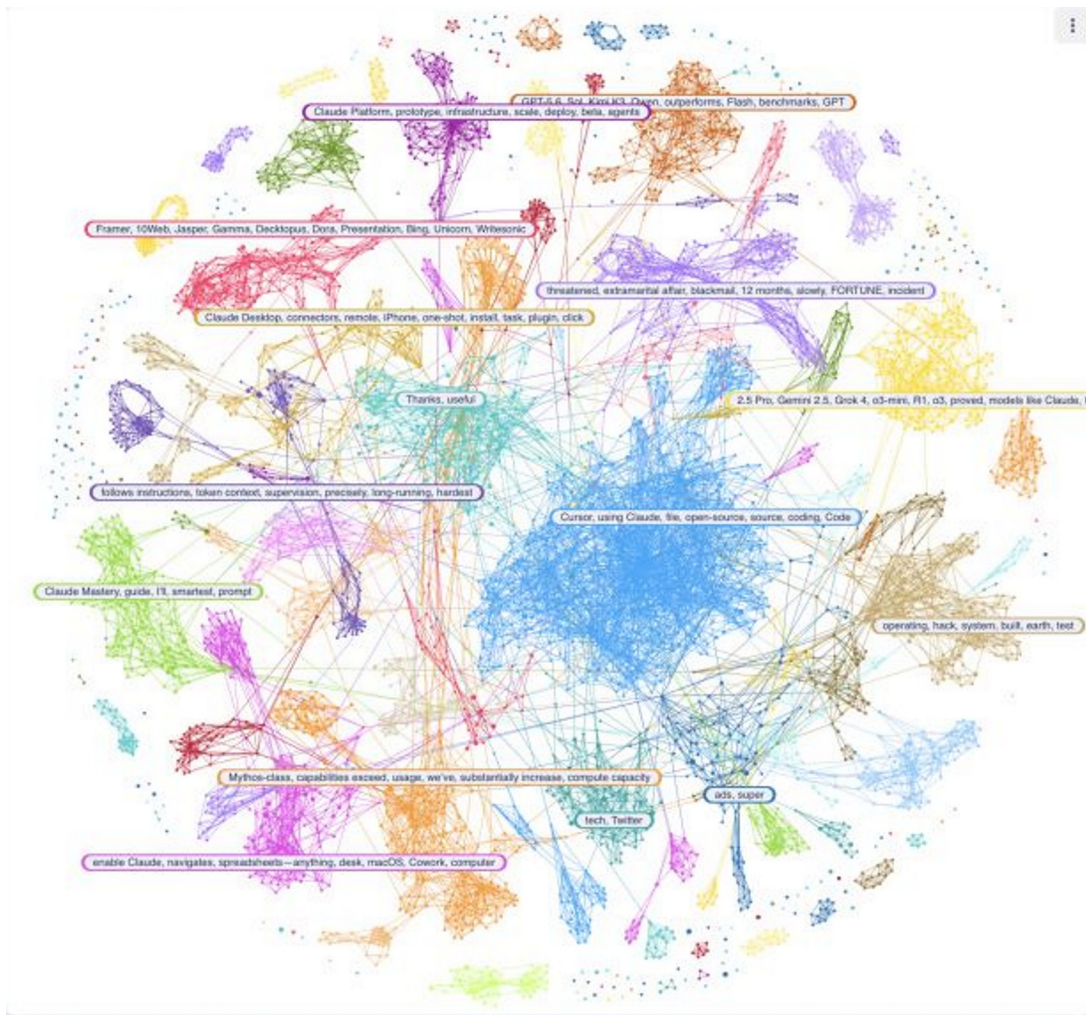
The leading clusters reveal a portfolio of conversations rather than a single brand narrative. Volume rank signals prevalence; positivity and strategic importance must be interpreted separately.

Narrative cluster	Est. records	Share	Positivity	Interpretation
Cursor / using Claude / open-source / coding / Code	10,729	4.16%	64.2%	Developer workflow and coding ecosystem core.
Thanks / useful	3,565	1.38%	82.9%	Direct favorable outcome shorthand.
Gemini 2.5 / Grok 4 / o3 / R1	3,118	1.21%	58.3%	Continuous model benchmarking.
threatened / affair / blackmail / incident	3,110	1.21%	11.3%	Low-positivity safety controversy.
enable Claude / spreadsheets / macOS / Cowork	2,989	1.16%	85.3%	Expansion into computer and document workflows.
operating / hack / system / built / test	2,577	1.00%	51.7%	System-risk and technical experimentation.
access / premium / evaluation / cybersecurity	2,457	0.95%	92.4%	Highly positive security/evaluation cluster.
Claude Mastery / guide / smartest / prompt	2,371	0.92%	65.2%	Education and performance claims.

Reading principle

A smaller, low-positivity cluster can matter more for reputation than a large, celebratory cluster. Prioritize by risk, stakeholder relevance, and persistence—not volume alone.

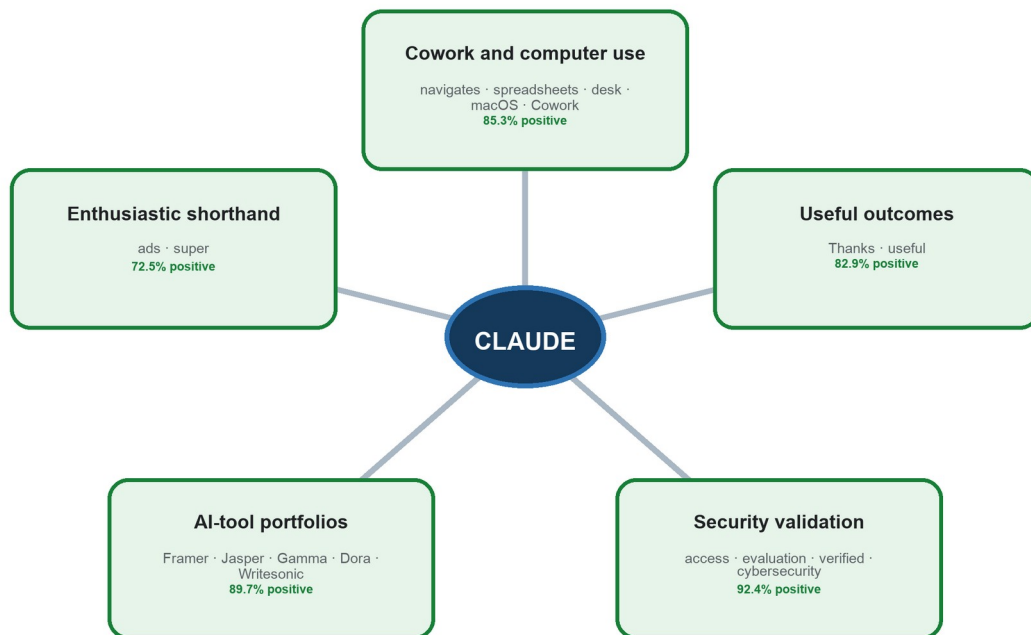
Semantic network: the conversation's underlying structure



High-resolution Starscape Narratives Force Graph capture for the full Claude query. Colors represent machine-generated narrative communities; position indicates linguistic similarity, not causality or market size.

Network region	What the connections suggest	Strategic use
Coding core	Cursor, open-source, file, coding, and Code form the densest central region.	Track direct productivity outcomes and switching from competing tools.
Computer and documents	Cowork, spreadsheets, macOS, navigation, and computer use form a distinct favorable community.	Use this bridge to reach nontechnical professionals.
Comparison ecosystem	Gemini, Grok, o3, R1, and Claude model language connect benchmarking communities.	Code winner, loser, parity, and switching stance separately.
Safety and system risk	Blackmail/incident and operating/hack/system communities are visibly separate.	Respond with issue-specific evidence; do not collapse controversy and security use.

Positive-language network



Analyst synthesis of U.S.-filtered Starscape narrative clusters

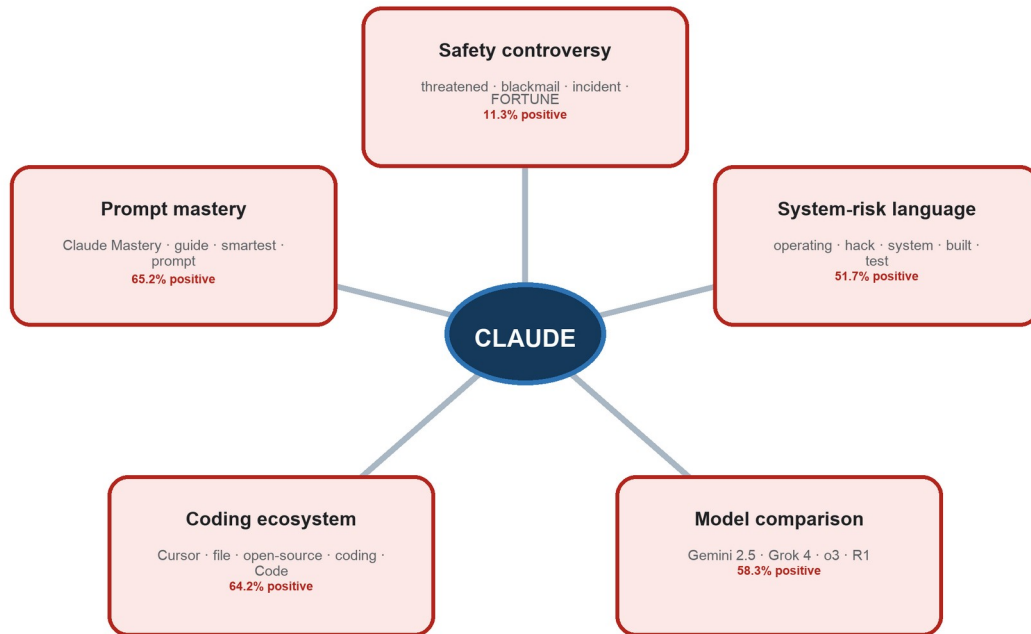
Analyst synthesis of the higher-positivity Starscape narrative clusters. Lines indicate thematic adjacency; distances and node sizes are illustrative rather than exported platform metrics.

Positive U.S. language centers on Cowork and computer use, useful outcomes, security validation, tool portfolios, and enthusiastic shorthand. Favorability is strongest when Claude completes tangible work or is demonstrated inside a credible workflow.

Validation move

Sample posts from each cluster and code whether positivity expresses product satisfaction, successful outcomes, enthusiasm for AI generally, or creator-distribution language.

Negative and contested-language network



Analyst synthesis of U.S.-filtered Starscape narrative clusters

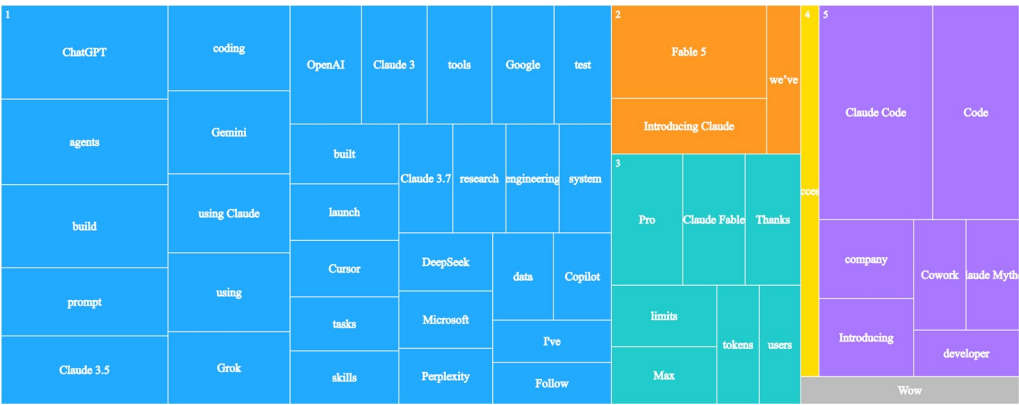
Analyst synthesis of lower-positivity and contested Starscape narrative clusters. The visual is a structured reading aid, not a native Starscape force-graph export.

Lower-positivity U.S. language separates safety controversy, system-risk experimentation, model comparison, coding rivalry, and prompt-mastery claims. The 11.3%-positive blackmail cluster is the clearest reputation issue, but each community requires its own post sample and response logic.

Validation move

Code stance, source, repetition, and issue frame separately. Negative language may criticize the brand, describe a problem the brand is used to solve, or quote a controversy without endorsement.

Topics, entities, and linguistic cues



Starscape semantic topic treemap; the entity list is reported below. Larger words and blocks indicate greater relative prominence within the matched corpus.

Entity	Share	Entity	Share
Claude	17.1%	GitHub	1.18%
Anthropic	7.67%	Microsoft	1.05%
Gemini	2.91%	Elon Musk	1.04%
ChatGPT	2.69%	SpaceX	0.67%

Classification quality check

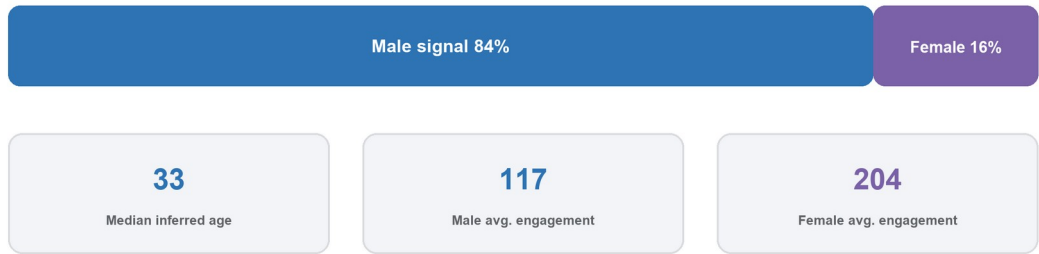
Claude and Anthropic remain tightly linked in U.S. entity recognition. Starscape again labels Gemini as a cryptocurrency exchange in many AI-context records, reinforcing the need to validate machine-coded entities independently from query precision.

Audience signals and poster identities

33	84%	16%	91.4%
MEDIAN INFERRED AGE	MALE SIGNAL	FEMALE SIGNAL	ENGLISH

Claude inferred audience signals

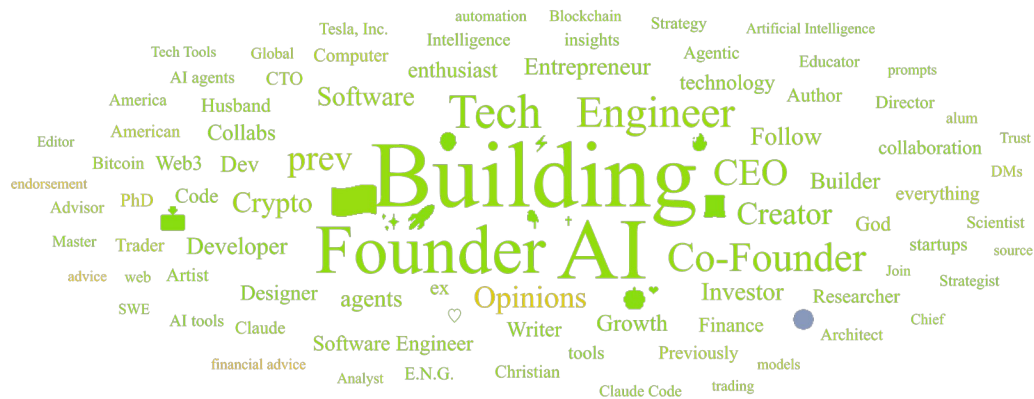
Gender classification among records with inferable public signals



Starscape inferred-gender view among records with inferable public signals. Unknown or unclassified records are not represented as a verified demographic population.

Inferred interest	Share
Computing	66.0%
Business industries	26.7%
Business	26.5%
Medical / health conditions	21.1%
Politics	14.7%
Gaming genres	13.4%
Artificial intelligence	11.0%

Poster identities: author-bio language



Author-bio terms associated with the matched conversation. Bio language is a self-presentation signal, not a verified occupation or identity.

The U.S. author-bio view is dominated by Building, AI, Founder, Tech, Engineer, Co-Founder, CEO, Software, Investor, and Finance. Inferred gender signals are heavily male-skewed, while female-coded records again show higher average engagement (204 vs. 117). These are model-based public signals, not a verified customer profile.

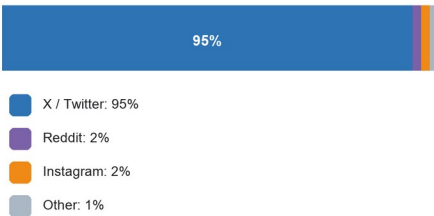
Influence and important voices

114K	70.2K	7.62M	2.25
AUTHORS	AVG. FOLLOWERS	TOTAL LIKES	AVG. RECORDS

Claude U.S. influence landscape

Platform concentration and leading accounts in Starscape's default ranking

Platform share



Interpret rankings by role: reach, activity, realized engagement, stance, and independence can point to different accounts.

Leading ranked accounts

1	Elon Musk	7 records · 43.7K engagements
2	Wall Street Journal	21 records · 3.3K engagements
3	Reuters	17 records · 445 engagements
4	TechCrunch	34 records · 3.1K engagements
5	Forbes	11 records · 1.2K engagements
6	Oprah	9 records · 106.2K engagements
7	Bloomberg	19 records · 936 engagements
8	Cointelegraph	45 records · 15.3K engagements

Source: Infegy Starscape Influencers page; rankings reflect platform-defined scale and activity signals

Starscape influencer summary and platform activity view. The default ranking combines scale and activity signals; it should not be treated as a pure measure of persuasion or independence.

Voice role	Illustrative accounts	Why the role matters
Agenda setter	Elon Musk	Seven records combine with exceptional reach and can reframe rivalry or risk.
Business and tech press	WSJ; Reuters; TechCrunch; Forbes; Bloomberg; CNBC	Shapes launch, market, enterprise, and controversy narratives.
Technology interpreter	The Verge; Cointelegraph; specialist outlets	Connects product detail to consumer, cultural, and adjacent technology narratives.
Cultural bridge	Oprah	Nine records generated more than 106,000 engagements—episodic but broad reach.
Owned and partner proof	Anthropic; Claude; AWS; GitHub	Explains capabilities and validates work use, but is not independent advocacy.

Better importance model

Use role-specific lists. Score reach, activity, realized engagement, stance, and independence separately.

Initial strategic implications

These implications are intended as a disciplined starting point. They translate observed discourse patterns into hypotheses that should be tested with narrower queries, post samples, and audience research.

Opportunity	Evidence in the discourse	Initial response direction
Translate the coding core	The largest cluster centers on Cursor, open-source, and coding.	Use technical proof to demonstrate writing, analysis, files, and professional judgment.
Scale Cowork proof	Cowork and computer-use language is 85.3% positive.	Develop outcome stories for spreadsheets, documents, planning, and recurring work.
Contextualize safety controversy	The blackmail/incident cluster is 11.3% positive.	Explain test conditions, safeguards, and limitations without minimizing concern.
Broaden credible voices	Large publishers and high-reach individuals dominate the ranking.	Recruit teachers, writers, analysts, and nontechnical professionals with repeatable proof.
Separate reach from influence	Oprah and Elon show exceptional engagement or reach from few records.	Use role-based lists and issue-specific networks instead of one leaderboard.

Decision principle

Use social listening to identify tensions, language, and influential frames. Use additional research to determine prevalence, causality, segment size, and campaign effectiveness.

Directions for deeper exploration

The most valuable next steps are those that challenge an aggregate, isolate a confounder, or create a comparison that could materially change the strategic reading.

1. Compare U.S.-English with smaller U.S. language communities to test whether use cases and issue frames differ.
2. Separate Claude Code, API/developer, Claude.ai consumer, Cowork/document, and enterprise-partner conversations.
3. Code the blackmail/incident cluster by source and stance: reporting, criticism, explanation, humor, or misinformation.
4. Sample Cowork and computer-use records for completed task, time saved, trust, failure, and repeat intent.
5. Build separate influencer lists for press framing, technical interpretation, cultural reach, partner validation, and direct experience.
6. Compare Claude, ChatGPT, and Gemini using identical U.S. dates, geography, language, and jobs-to-be-done.

Questions worth carrying forward

- Which favorable language reflects a completed, repeatable outcome rather than novelty or social participation?
- Which negative frames are persistent, high-reach, and attributable to direct experience?
- Where do owned and partner amplification change the apparent health of the organic conversation?
- Which audience and voice roles can credibly move the brand from category awareness to durable preference?

Appendix: metric notes and sources

Metric notes

- Research window: January 1, 2024 through August 18, 2026. August is a partial month.
- English represented 91.4% of U.S.-inferred matched records.
- Platform-derived views, engagements, universe estimates, demographics, interests, sentiment, emotions, entities, and narratives are directional analytics rather than audited customer measures.
- The semantic-network figure is a high-resolution Starscape force-graph capture. Positive and negative language maps are analyst syntheses using Starscape cluster labels and are explicitly identified as such.
- Official product sources provide contextual phase markers. Their inclusion does not establish causation for conversation changes.

Source links

Infegy: [Getting Started](#) | [Boolean queries](#) | [Drilldowns](#) | [Exporting data](#)

Claude context: [Claude 3.5 Sonnet](#) | [Artifacts availability](#) | [Claude 3.7 and Claude Code](#) | [Claude 4](#)

[Agentic misalignment context](#) | [Claude Sonnet 5](#) | [Claude is a space to think](#)

Planning context

The supplied J345 Claude client brief informed the tensions and exploration priorities. Findings labeled as Starscape evidence come from the saved query and dashboards described in this report.