



Siftsy Score & Sentiment Scoring Methodology



A New Paradigm in Comment Analysis for the AI Age

Comment sections represent a unique form of interaction – they're not just standalone text, but reactions and responses to online content. So, to understand comments you have to *understand their context*. That's why Siftsy's approach combines this **crucial context** with **cutting-edge AI** to deliver unprecedented insight into audience engagement.

The Context + AI Advantage

Modern Large Language Models (LLMs) have achieved remarkable capabilities in understanding human communication, even surpassing human analysts if they lack full contextual awareness. Siftsy harnesses this power through a unique approach:

Our Core Focus: Context-Enhanced Analysis

Every comment is analyzed by Siftsy within its context:

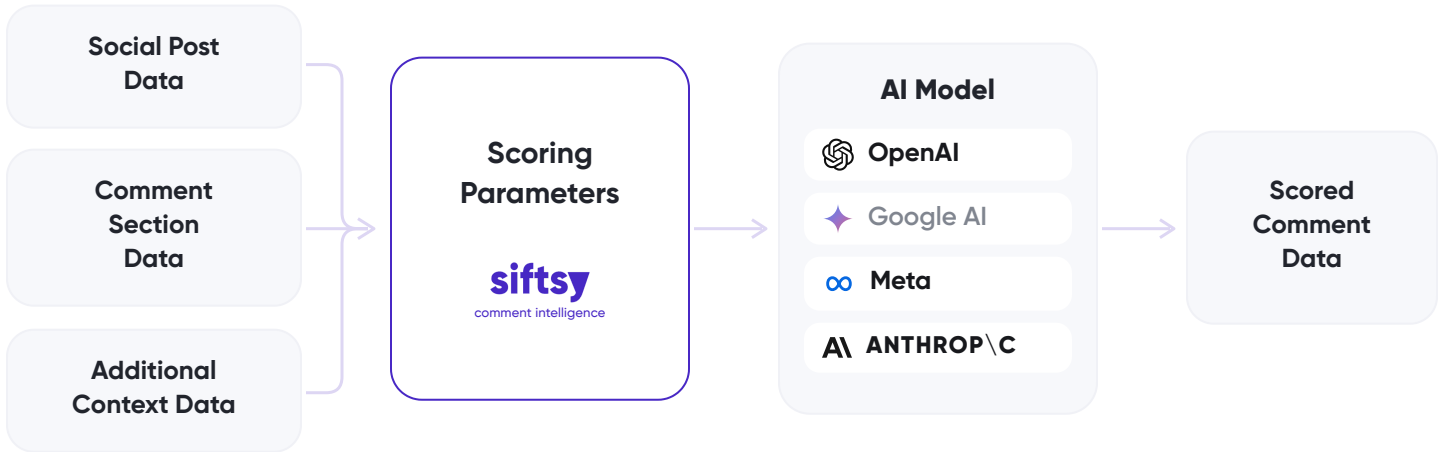
- **Social Post Data:** description, transcript (when available), additional metadata, etc.
- **Comment Section Data:** comments are analyzed with other comments simultaneously.
- **Additional Context Data:** data from search partners (like Perplexity) for up-to-date relevant information.

The AI Model-Agnostic Architecture

Siftsy's system is designed for continuous evolution:

- Compatible with any modern AI model (OpenAI, Anthropic, Google, Meta, etc...).
- Implements Siftsy's scoring parameters (and our updates of these parameters).
- Implements additional context when needed.

Diagram A: Siftsy Scoring Workflow



Scoring Parameters & Definitions

What Siftsy Score measures

Siftsy Score is a **1–10 temperature check** of a comment section in response to content. It is not a like count, not a virality metric, and not “how people feel about the post” alone.

It blends three independent 1–10 signals so a loud, off-topic, or fighting thread cannot look healthy just because one axis is high.

Siftsy Score

Comments are rated on three signals – Sentiment, Relevance, and Consensus.

Siftsy Score = (Sentiment × w_S) + (Relevance × w_R) + (Consensus × w_C)

- **Version 1 · Equal Average** – Prioritizes holistic engagement. All three signals count equally.

$$w_S = w_R = w_C = \frac{1}{3}$$

- **Version 2 · Sentiment-biased** – Prioritizes sentiment. Sentiment counts double (2 : 1 : 1), so the mood of the section moves the score most.

$$w_S = 0.66, \quad w_R = w_C = 0.165$$

- **Version 3 · Split sentiment** (recommended) – Prioritizes targeted sentiment. The Version 2 blend, with Sentiment split evenly into Content and Topic. Also adds reply hierarchy, a negativity cap, and gentler engagement weighting.

$$\begin{aligned} w_S &= w_{\text{Content}} + w_{\text{Topic}} \\ w_{\text{Content}} = w_{\text{Topic}} &= \frac{w_S}{2} \\ w_S &= 0.66, \quad w_R = w_C = 0.165 \end{aligned}$$

The complete formulas behind each version are documented in **Appendices A–D**.

Scoring Breakdown

Sentiment (1-10)

How commenters feel about the post, creator, or topic:

- **Low scores:** expresses intentionally negative, bitter, angry, mean, or insulting tones.
- **Moderate scores:** expresses a neutral, matter-of-fact, observational tone.
- **High scores:** expresses praise, admiration, or good humor in tone.

For Version 3, Sentiment is a weighted blend of Content and Topic. Versions 1 and 2 use the direct comment-sentiment average. How each comment is scored – and what the model considers – is on [How Sentiment Is Scored](#).

Sentiment = (Content × w_{Content}) + (Topic × w_{Topic})

$$w_{\text{Content}} = w_{\text{Topic}} = \frac{1}{2}$$

Relevance (1-10)

Assesses how closely comments relate to the original post content:

- **Low scores:** Completely off-topic to the post and conversation about the post.
- **Moderate scores:** General commentary not explicitly about the post, but applicable.
- **High scores:** Directly responding to or paraphrasing something about the post.

Consensus (1-10)

Evaluates how commenters interact with **each other** – not whether they agree with the post:

- **Low scores:** Explicitly arguing with other commenters.
- **Moderate scores:** Sharing mixed opinions to other commenters.
- **High scores:** In complete agreement and sharing mutual support with other commenters.

A thread can be uniformly negative toward the post and still have high Consensus.

Accuracy

Siftsy aims to rank and score comments within a Siftsy Score accuracy of ±1 point of error compared to a how a human user would rank based on the scale definitions above.

What the Scores Mean

Siftsy Score is best read in two passes.

First, the temperature check. The score is a fast, standardized read on the overall vibe of a comment section reacting to content – immaculate, mixed, or off. Use it to rank posts, compare them against each other, and spot the ones that deserve a closer look, using the bands below.

Then, the substance. The number tells you *that* something is happening, not *what*. The breakouts show where the temperature is coming from – Sentiment for mood, Relevance for how on-topic the conversation is, Consensus for infighting, and in Version 3, Content versus Topic for what people are actually reacting to. From there, use Search and Insights to read the conversations behind the number before drawing conclusions.

Versions 1 and 2

Range	Label	Meaning
7.0+	Immaculate	The comments are positive and supportive.
5.0 – 6.9	Good	Most of the conversation is in good faith.
3.0 – 4.9	Sub-Par	Something here isn't landing.
< 3	Off	Negativity and infighting run the comments.

Version 3

Range	Label	Meaning
7.0+	Immaculate	The comments are positive and supportive.
4.5 – 6.9	Mixed	The comments are split and attitudes are mixed.
< 4.5	Off	The comments are substantially negative and critical.

Reading the score in context

The same number means different things depending on the job:

- **Campaign performance:** Track scores across posts, creators, and flights to see which activations landed. A rising trend across a campaign says more than any single post's number.
- **Creative reviews:** A high Content score with a low Topic score – or the reverse – separates reaction to the execution from reaction to the subject, and tells you whether to revisit the creative or the message.
- **Consumer research:** Mid-range scores are often the most informative. High Relevance with mixed Sentiment usually marks a section full of unfiltered opinions about the product itself.
- **Brand reputation:** Watch the trend, not the level. An account that normally scores steady and drops two points is a louder signal than one low-scoring post.
- **Press & communications:** A low score with high Consensus means the audience has settled on a shared narrative; a low score with low Consensus means the conversation is still contested. The two call for different responses.

Content Score & Topic Score

Traditional sentiment analysis collapses “I hate this video” and “I hate this subject” into one scale. With Content and Topic sentiment differentiation, we can measure the sentiment towards what exactly the audience is reacting to. Content and Topic Scores are produced only when Version 3 is selected.

Content Score (1–10)

How commenters feel about **this post and its creator**.

Range	Label	Meaning
7+	Supportive	Strongly positive toward the post or creator
5.6 – 6.9	Favorable	Reacting in good faith toward the post or creator
4.5 – 5.5	Neutral	Not reacting directly to the post or creator
	Divided	Reactions to the post or creator are polarized. At least 25% of Content readings lean positive and 25% lean negative.
3.0 – 4.4	Critical	Raising critiques of the post or creator
< 3	Hostile	Aggressively against the post or creator

Topic Score (1–10)

How commenters feel about **the topics being discussed**.

Range	Label	Meaning
7+	Positive	Strongly positive toward the topics
5.6 – 6.9	Positive Lean	Mostly positive toward the topics
4.5 – 5.5	Neutral	Not reacting directly to the topics
	Divided	Reactions to the topics are polarized. At least 25% of Topic readings lean positive and 25% lean negative.
3.0 – 4.4	Negative Lean	Mostly negative toward the topics
< 3	Negative	Aggressively against the topics in the post

Overall Sentiment Score

How commenters feel about the post and its topics.

$$\text{Sentiment} = \frac{\text{Content} + \text{Topic}}{2}$$

Range	Label	Meaning
7+	Positive	Strongly positive about the post overall
5.6 – 6.9	Positive Lean	Mostly positive about the post overall
4.5 – 5.5	Neutral	Reactions are generally neutral. Average Content and Topic scores are within 2 points of each other, giving a neutral reaction overall.
	Divided	Reactions are generally divided. Content and Topic scores diverge by 2+ points, or at least 25% of comments lean positive and 25% lean negative. Sentiment is falling in the middle because of a polarized response.
3.0 – 4.4	Negative Lean	Mostly negative about the post overall
< 3	Negative	Strongly negative about the post overall

How Neutral vs Divided is decided. Only when that score sits in the mid band (4.5–5.5). Overall Sentiment is **Divided** if either test fires. Content and Topic use the split test on their own readings. Otherwise the label stays **Neutral**.

- Content/Topic divergence** (overall Sentiment only): **|Content – Topic| ≥ 2**. With a mid-band average, a 2-point gap puts one dimension at ≥ 5.5 and the other at ≤ 4.5 – a positive read on one and a negative read on the other.
- Split distribution: min(posSide, negSide) ≥ 25%**. posSide is positive (7+) plus positive-lean (5.6–6.9); negSide is negative (< 3) plus negative-lean (3.0–4.4). Both camps must each hold at least a quarter of readings, so the neutral bucket is at most 50%. Overall Sentiment pools Content and Topic (each comment counted once per dimension). Content and Topic use the same split test on their own readings.

How Sentiment Is Scored

Sentiment is the mood signal inside the Siftsy Score, and it is also reported on its own: a 1–10 score per comment, rolled up into a section average and a positive / neutral / negative breakdown. Each comment is scored by context-aware AI applying Siftsy’s scoring parameters, not by keyword lists, and never on the comment text alone. This page is a working overview of how that scoring behaves; the complete parameters are proprietary and continuously refined.

How a comment gets scored

- The post is analyzed first.** Before any comment is read, the model builds a picture of the post itself from the caption, transcript, title, account bio, and retrieved outside context: who the creator is, the editorial stance they are taking, the messages they want the audience to walk away with, and the topics commenters are likely to react to.
- Comments are scored against that picture, together.** Comments are evaluated in the context of the post analysis and alongside the rest of the comment section, so replies, running jokes, and pile-ons are read the way a human moderator would read them: with the thread in view, not one comment at a time.
- Each comment gets two sentiment readings.** The model scores the comment’s attitude toward the **post and creator** and, separately, toward the **topics being discussed**. Each reading uses the 1–10 scale and is paired with a confidence level. Version 3 surfaces these as Content Score and Topic Score; Versions 1 and 2 blend them into a single Sentiment score.

What is considered

The post

Caption, transcript, title, who posted it, and the creator’s stance: what they are arguing, celebrating, or announcing. The same comment can score differently under a different post.

The conversation

Comments are scored together, not one-by-one in a vacuum. A reply is resolved through the parent comment it answers, so “exactly” and “nah” carry the meaning of the exchange they sit in.

Language as people use it

Slang, emoji, sarcasm, and idiom are scored for what they mean in use. Non-English comments are read in their own language and cultural register, not through a literal translation.

Outside context when needed

Recent events, controversies, and background retrieved from search partners. This is the context that makes a reaction legible when the post alone doesn’t explain it.

Reading the 1–10 scale

The scale is anchored to **intent**, not word choice. Per-comment scores are continuous (a comment can score 7.2), and the same bands apply to both sentiment readings:

Range	Read	What lands here
7+	Positive	Praise, excitement, buying intent, admiration, delight
5.6 – 6.9	Positive Lean	Interest, genuine questions, friendly humor, good-faith engagement
4.5 – 5.5	Neutral	Matter-of-fact, informational, or not addressing this at all
3.0 – 4.4	Negative Lean	Skepticism, critique, disappointment, unfavorable comparisons
< 3	Negative	Hostility, insults, mockery, protest

Neutral means “no read”, not “mixed feelings.” A comment that never addresses the creator or the topic sits at 5 rather than pulling the average in either direction. Genuinely mixed comments (“love the concept, not the execution”) are weighted toward their dominant take instead of defaulting to the middle. At the section level, a 4.5–5.5 average can still be **Divided** when the readings are polarized – see Content Score & Topic Score.

These are the same bands behind the Content and Topic Score labels and the positive / neutral / negative breakdown percentages.

How intent is read

- Alignment with the post’s stance, not word polarity.** On a post criticizing something, “this is terrible, I hate it” agrees with the creator, so it counts as positive engagement toward the creator even though the words are negative. Disagreeing with praise works the same way in reverse.
- Intended meaning over literal words.** Sarcasm, mock praise, and rhetorical questions are scored for what the commenter meant. Playful teasing and laughing-with are distinguished from mockery and laughing-at.
- Replies inherit their thread.** A short reply agreeing with a critical comment carries that criticism; a reply pushing back against it flips the stance. Commenter-versus-commenter fights lower the section’s Consensus rather than distorting sentiment.
- Participation counts as engagement.** Answering a call-to-action, tagging friends, or dropping a hype reaction is treated as real positive engagement rather than discarded as contentless.
- Redirected praise is read as deflection.** Celebrating a competitor, predecessor, or alternative on a post about something else is not support for that post, however positive the words sound.
- Story and real world are scored separately.** Hating a fictional villain is engagement with the work; criticizing a real actor, brand, or production choice is scored literally as criticism.

Confidence and consistency

Every score carries a 1–10 confidence level. Ambiguous comments (subtle irony, rare dialects, garbled text) lower confidence rather than forcing a certain read, and spam is excluded entirely. Per-comment scores then average into the section-level Sentiment score, with each comment counting once in the positive / neutral / negative breakdown. The roll-up formulas are in **Appendices A–D**.

Why Use Siftsy Score?

Go Beyond Traditional Sentiment Analysis

Traditional sentiment analysis only scratches the surface. While Siftsy offers sentiment-only analysis for compatibility with existing workflows, our Siftsy Score represents a fundamental evolution in comment understanding.

The Holistic Advantage

- Sentiment analysis tells you if comments are positive, neutral, or negative in tone and intent.
- Siftsy Scores tell you **why**, **how it relates**, and **what it means** for your content.
- Version 3 adds **what they are reacting to** – the post and creator (Content) versus the subjects being discussed (Topic).

Internal Intelligence

Even if your external reporting requires traditional metrics, we strongly recommend tracking Siftsy Score internally because it:

- reveals patterns invisible to sentiment-only analysis.
- provides a faster way to track and nail down audience reception trends.
- offers a standardized way to quickly compare audience reception across posts.
- lets teams switch score versions instantly without losing history.

Practical Impact

The difference between sentiment and Siftsy Scores in real scenarios:

- **Marketing:** Understanding not just if people like your ad/content, but exactly how they received it.
- **Product Launches:** Distinguishing between artificial praise, real feedback, and genuine concerns.
- **Content Strategy:** Identifying which posts truly resonated vs. which simply avoided negative reactions.
- **Community Building:** Recognizing authentic engagement versus superficial reactions.
- **Crisis / controversy:** Seeing whether people hate *this video* or hate *the topic* – and whether commenters are aligned with each other or fighting.

What is not in the score

- Like count, view count, or follower count as a direct input. Likes and replies only matter if engagement weighting is on, and even then they only scale comment *weights*.
- "Agrees with the post" as Consensus.
- Hidden-from-analysis comments.
- Custom sentiment rules change per-comment scores *before* this math; they are not a separate headline formula.

Score bands and how to read them are on [What the Scores Mean](#). For final reporting, use Search and Insights to dive into the conversations the score is based on.

Comparing Siftsy Score Versions

Version 3 is an evolution, not a replacement of Siftsy Score. It uses the same scale, signals, and weights as V2, but adds a Content/Topic split and guardrails so reply pile-ons, hostile minorities, and viral outliers can’t write the headline.

Switching versions is a re-calculation, not a re-run!

Select V3 as the team score version and every score updates instantly, including posts already analyzed. You can switch back to V2 or V1 at any time. Nothing is lost, because all versions are built from the same per-comment AI scores.

	Version 1	Version 2	Version 3
Version Label	Equal Average	Sentiment-biased	Split sentiment-biased
Prioritizes	Holistic reactions	Reaction Sentiment	Targeted reaction sentiment
Scale	1–10	1–10	1–10
Signals	Sentiment, Relevance, Consensus	Same as V1	Same as V1
Signal weights	33.3% / 33.3% / 33.3% 33.3 33.3 33.3	66% / 16.5% / 16.5% 66 16.5 16.5	Same as V2 (66 / 16.5 / 16.5) 66 16.5 16.5
Sentiment	One 1–10 reading per comment	Same as V1	Content + Topic, averaged 50/50 into Sentiment
Per-comment scores	Context-aware AI; versions are rollups only	Same as V1	Same as V1
Score bands	Immaculate 7.0+ Good 5.0–6.9 Sub-Par 3.0–4.9 Off <3	Same as V1	Immaculate 7.0+ Mixed 4.5–6.9 Off <4.5
Top Level vs. Reply	Every comment counts as 1	Same as V1	Top Level = 1. Combined Replies = 1
Negativity cap	None	None	Yes
Engagement Weighting	Optional	Optional	Default

What V3 adds to V2

Addition	What it adds
 Content vs Topic	A split read: reaction to the post and creator versus the subjects discussed
 Mixed band	A first-class label for mid-range scores as mixed reactions, not mild failures
 Negativity cap	A ceiling that gives a hostile minority visibility even when the average is high
 Engagement weighting	Now the default, a 1× floor so quieter comments still count, with likes lifting Sentiment only
 Top Level vs. Reply	Full weight for top-level comments; replies share a capped pool so the score tracks breadth

Real-World Validation

These real-world cases show the Siftsy Score in the context of a real comment section and the post the comments were reacting to.

Scores in these case studies were produced with **Version 1 (Equal Average)** for demonstrative purposes.

1 A comedian's Vitamin Water partnership gets rave reviews

Account	Platform	Comments	Siftsy Score	Sentiment
 veronika_iscool @veronika_iscool	TikTok	758	8.6	88% 12%

- Description:** Comedian @veronika_iscool makes a comedy music video featuring Vitamin Water and the NYC summer heat.
- Audience reception:** With a Siftsy Score of **8.6** and a Sentiment Score of **8.3/10**, the video received high praise across almost all commenters for its humor, catchy tune, and it's charming integration of Vitamin Water.
- Example comments:**

 **Simone** 
Okay but why does this EAT tho 🔥🔥🔥🔥🔥  1.8K

 **Sara Tea** 
First ad I ever watched all the way through 🥺👉  933

 **greatauntmuriel**
@vitaminwater I have never been more influenced to buy a product in my life  59

2 Video app promotion gets positive interest with a side of skepticism

Account	Platform	Comments	Siftsy Score	Sentiment
 Ben Kaluza Creator @benkaluza	TikTok	128	6.4	40% 42% 18%

- Description:** Tech Creator Ben Kaluza demonstrates the Shade app, a video file app for creators.
- Audience reception:** Reception of the post and app promoted was generally positive engagement and interest, with **40.2%** of comments being positive, however questions and skepticism about data privacy and local AI functionality brought the overall sentiment down.
- Example Comments:**

 **ioan.stauffer**
really cool tool, but what about the data safety ? does it mean this ai have access to your computer ? or do you < send > specific files to it ?  42

 **Gabriel S**
man this is a godsend. apple's search is so crap I literally search the word design and it gives me all root files of every app 🙏🥲  2

 **yellowtoaster7**
not installing spyware bro  2

3 Famous entrepreneur's product launch gets critique from fans expecting more

Account	Platform	Comments	Siftsy Score	Sentiment
---------	----------	----------	--------------	-----------

 Sneex @sneex	TikTok	613	4.4	15% 17% 67%
--	--------	-----	-----	---

- **Description:** Entrepreneur Sara Blakely introduces Sneex, a luxury high-heel for women aiming to provide comfort with sneaker styling.
- **Audience reception:** A Siftsy Score of **4.4** indicates mixed to negative commentary, as many users expressed excitement about the concept of a luxury high-heel sneaker, but many expressed concerns regarding the design and price.
- **Example comments:**

 **Heather**  2
Ranging in price from \$395-\$595 🙄

 **Chris Belin**  334
I love Sarah. I love Spanx. I love the thought behind this. LOVE. But, wasn't this focus group tested?

 **Kayleigh** ✨  11
Following for the JUST KIDDING video that surely has to be coming next 🤔🤔

4 A large TV brand tries to do sports marketing, fails miserably

Account	Platform	Comments	Siftsy Score	Sentiment
---------	----------	----------	--------------	-----------

	Instagram	4	1.7	100%
---	-----------	---	-----	------------------------------

- **Description:** A TV brand promotes it's brand with a post about an international sporting event on it's own Instagram account.
- **Audience reception:** Commenters express significant dissatisfaction and concerns regarding the quality of the tech brands products and TVs. The international sporting event is not referenced.
- **Example Comments:**

 **neolandes**  4
Crappy brand that won't stand by it's crappy products. Fix your TVs!

 **kzw7033**  36
Can't enjoy my game day cuz y'all can't fix your damn TVs

 **ra'sodope**  2
Bro yall tv is so slow just freezes out of nothing is trash !! Horrible !! Terrible !! Wifi just stops working and it can barely airplay 📶📶📶📶📶

FAQs

Sentiment is important to understand for my organization. Does Siftsy report on sentiment?

Yes. Siftsy calculates sentiment scores along with its Siftsy Score so you can view and export sentiment analysis and breakdown (% Positive, Neutral, Negative). In Version 3 you also get Content Score and Topic Score as separate breakouts.

What’s the difference between Content Score and Topic Score?

Content is how commenters feel about **this post and its creator**. Topic is how they feel about **the subjects being discussed**. Both are produced only in Version 3, where headline Sentiment is the average of the two.

Does “high Consensus” mean people agree with my post?

No. Consensus measures how commenters react to **other commenters**. A thread that is uniformly negative toward the post can still have high Consensus. Low Consensus means arguing, fighting, or dunking on other people in the section.

How do likes and replies affect the score?

Only if the team setting is **Weighted**. Equal mode gives every comment the same weight (plus Version 3 reply hierarchy). Weighted mode uses the formulas in **Appendix C**. Distribution bars never use likes – they always count each comment as 1.

Does Siftsy understand sarcasm?

Yes, in both positive and negative connotations. In this case study example, Siftsy identified the layered sarcasm which ultimately gave it a negative leaning score, which is what we expect.

K

Kayleigh 🙌
Following for the JUST KIDDING video that surely has to be coming next 🙄🙄

♡ 11

SENTIMENT

3.8

Negative Lean

Comment on entrepreneur Sara Blakely’s launch video for Sneex, a luxury high-heel sneaker.

I'm looking to only understand the positive/negative reactions to content. Can I use Siftsy for that?

Yes, you can use Siftsy search & filter to view and interpret comments for the sentiment you are researching.

Does Siftsy understand the tone of emojis?

Yes, and we find that the AI models are typically as aware of what emojis mean in the context they are used as normal language, whether they are standalone emoji comments or an emoji added to text for emphasis.

In this case study example, the emojis are understood in the context of the overall positive message towards the app involved and the comment is ranked as positive.

G

Gabriel S
man this is a godsend. apple's search is so crap I literally search the word design and it gives me all root files of every app 🙌🙄

♡ 2

SENTIMENT

7.2

Very Positive

Comment on a TikTok demo of Shade, a video file app for creators.

Does Siftsy understand modern internet slang?

In most cases yes, and we are working on bringing even more up-to-date context awareness to Siftsy by searching for recent references not yet seen by the AI.

In this example, Siftsy understood that saying someone “ate” or that something “eats” is a compliment in modern slang and ranked this comment highly positively in sentiment as such.

S

Simone 🦋
Okay but why does this EAT tho 🔥🔥🔥🔥🔥

SENTIMENT

9.1

Very Positive

Comment on a comedy music video featuring Vitamin Water.

Supported Languages

Best Support

- English
- Spanish
- French
- German

- High accuracy in sentiment analysis and context understanding.
- Excellent at interpreting nuances, sarcasm, and idiomatic expressions.
- Well-equipped to handle diverse comment styles (formal, informal, slang).
- Strong availability of training data for various contexts.
- Strong for brand, product, or service mentions and reactions.

Good Support

- | | |
|----------------------|----------|
| Chinese (Simplified) | Dutch |
| Italian | Russian |
| Portuguese | Japanese |

- Generally reliable for sentiment analysis, but may struggle with very subtle nuances.
- Good handling of casual language but can miss some cultural context.
- Adequate training data, though some specific expressions may not be well-represented.
- Adequate for brand, product, or service mentions and reactions.

Med Support

- | | |
|------------|------------|
| Korean | Norwegian |
| Arabic | Danish |
| Turkish | Thai |
| Hindi | Bulgarian |
| Vietnamese | Ukrainian |
| Swedish | Serbian |
| Czech | Croatian |
| Finnish | Slovak |
| Hungarian | Lithuanian |
| Romanian | Latvian |

- Moderate performance in sentiment analysis, with potential misinterpretation of tone.
- May struggle with slang, regional dialects, or informal language.
- Limited availability of training data can affect the model's ability to understand context.
- Good for brand, product, or service mentions and reactions.

Low Support

- | | | |
|--------------------|-----------|-------------|
| Swahili | Malayalam | Hmong |
| Hebrew | Tamil | Amharic |
| Persian (Farsi) | Telugu | Mongolian |
| Indonesian | Kannada | Kazakh |
| Malay | Bengali | Uzbek |
| Filipino (Tagalog) | Punjabi | Tatar |
| Slovenian | Urdu | Azerbaijani |
| Icelandic | Sinhala | Somali |
| Basque | Pashto | Quechua |
| Catalan | Kurdish | Guarani |
| Macedonian | Armenian | Maori |
| Estonian | Georgian | Galician |
| Tigrinya | | |

- High likelihood of errors in understanding sentiment and context.
- Frequent misinterpretations of informal language or cultural references.
- Limited training data leads to vague or generic responses, which can hinder analysis quality.

Appendix A · How the Score is Calculated

Every comment already has Sentiment, Relevance, and Consensus. Those averages get blended into the 1–10 Siftsy Score you see. Below are the weights for each version, and a worked example you can follow through by hand.

Every scored comment already has Sentiment, Relevance, and Consensus on a 1–10 scale (and in V3, Content and Topic as well). Hidden-from-analysis comments are skipped. The remaining comments are averaged, then blended:

$$\text{Siftsy Score} = (\text{Sentiment} \times w_S) + (\text{Relevance} \times w_R) + (\text{Consensus} \times w_C)$$

Display values are **floored to one decimal** (5.94 becomes **5.9**). Version 3 may then limit the result to the negativity-cap maximum for that very-negative share (see Appendix D).

Version weights

Version	Sentiment	Relevance	Consensus	Label
1.0 Equal Average	33.3%	33.3%	33.3%	1:1:1
2.0 Sentiment-biased	66%	16.5%	16.5%	2:1:1
3.0 Split sentiment	66% (33% Content + 33% Topic)	16.5%	16.5%	2:1:1

Worked example

Sentiment **6.0**, Relevance **7.0**, Consensus **5.0**. In Version 3, Sentiment splits into Content **8.0** and Topic **4.0**, which average back to **6.0**.

Version	Blend	Cap	Displayed
1.0	$(6 + 7 + 5) / 3 = 6.00$	none	6.0
2.0	$0.66 \times 6 + 0.165 \times 7 + 0.165 \times 5 = 5.94$	none	5.9
3.0	same 5.94	8.5 at 12% very-neg	5.9

The cap only binds when the blend is above the maximum. 5.94 is already under 8.5, so it still displays as **5.9**. A 9.2 blend with the same 12% very-negative share would display as **8.5**.

Appendix B · How Comments Roll Up

The Siftsy Score is a composite score built on blending individual comment scores.

1. **Per-comment scores** – Sentiment, Relevance, Consensus, and (V3) Content + Topic.
2. **Hide rules** – comments hidden from analysis are skipped.
3. **Structural weight (h)** – V1 / V2: every comment = 1. V3: top-level = 1; replies share a pool (see Appendix D).
4. **Engagement (optional)** – if Weighted, likes and replies scale the mean.
5. **Axis means** – weighted average per signal.
6. **Headline blend** – apply the version weights. In V3, Sentiment is 50 / 50 Content + Topic first.
7. **Version 3 cap** – limit the score to the maximum for the very-negative share (Appendix D).
8. **Display** – floor to one decimal.

Breakdown percentages are never engagement-weighted. Each comment counts as 1 in the distribution bars.

Engagement weighting

- **Equal** – every comment counts the same (plus V3 reply hierarchy).
- **Weighted** – comments with more likes and replies pull harder. Strength **w** is 0–1.

Symbol	Definition
e	$\max(\text{likes} + \text{replies}, 1)$
eMax	Highest e in this comment set
w	Team setting, clamped to 0–1
h	Structural weight (V3 replies share 0.5 or 1.0)

$$\text{weight} = h \times \text{multiplier}(e, e_{\max}, w)$$
$$\text{axis mean} = \frac{\sum(\text{score} \times \text{weight})}{\sum \text{weight}}$$

If $w = 0$, the multiplier is skipped and $\text{weight} = h$. A comment with zero likes and zero replies still has $e = 1$.

Appendix C · Engagement Weighting Formulas

If Weighted mode is on, comments with more likes and replies count for more. Version 3 uses a gentle linear-log lift so quieter comments still have weight and the most-engaged comments are lifted, not used as the only weight.

Version 3 – gentle linear-log lift

Keeps a **1× floor** so quiet comments still have weight. Severity constant is 4. Range is **1× to (1 + 4w)×**. At $w = 1$ the top comment is **5×**; at $w = 0.5$ the top comment is **3×**.

$$\text{normalized} = \frac{\ln(1 + e)}{\ln(1 + e_{\max})}$$
$$\text{multiplier} = 1 + 4w \cdot \text{normalized}$$

Which axes get the multiplier

Axis	Version 1 / 2	Version 3
Sentiment / Content / Topic	Yes, when $w > 0$	Yes, when $w > 0$
Relevance	Yes, when $w > 0$	No – structural weight only
Consensus	Yes, when $w > 0$	No – structural weight only
Bucket % bars	Never	Never

Multipliers when eMax = 1,000 and w = 1

e	V1 / V2	V3	V1 / V2 vs top	V3 vs top
1 (floor)	$\sim 10^{-10} \times$	1.40×	$\sim 0\%$	28%
10	0.00002×	2.39×	$\sim 0\%$	48%
50	0.004×	3.28×	0.4%	66%
100	0.018×	3.67×	1.8%	73%
250	0.11×	4.19×	11%	84%
500	0.35×	4.60×	35%	92%
1,000 (max)	1.00×	5.00×	100%	100%

Versions 1 and 2 used a steep power curve. The top comment's multiplier is always **1×**; everyone else is a fraction of that.

$$\text{logRatio} = \frac{\ln(1 + e)}{\ln(1 + e_{\max})}$$
$$\text{multiplier} = (\text{logRatio}^{10})^w = \text{logRatio}^{10w}$$

Appendix D • Version 3 Extras

Version 3 adds an update to comment reply hierarchy and a negativity cap to keep the score reflective of the overall temperature of comment sections across the most number of normal and edge cases. Reply hierarchy keeps a long reply chain from outweighing the parent thread. The negativity cap adjusts the maximum score of the Siftsy Score scale when a sizeable percentage of the comment sentiment is negative.

Reply hierarchy

- **Top-level comment weight = 1.0**
- **Direct replies share a pool:** 0.5 total if there are 1–3 replies, 1.0 total if there are 4 or more
- Each reply gets pool / replyCount

If engagement weighting is also on: **weight = h × multiplier**.

Negativity cap

If enough comments score **below 3**, the maximum the score can reach is lowered. The displayed score cannot exceed that maximum.

Siftsy Score = min(raw, max)

Very-negative share	Maximum
< 2%	10.0 (no cap)
2%+	9.5
5%+	9.0
10%+	8.5
20%+	8.0

Overall score uses the pooled very-negative share of Sentiment readings – each comment counted once for Content and once for Topic (readings below 3). Content uses opposed (Content below 3). Topic uses hostile (Topic below 3). Pooling keeps the overall cap consistent with the dimension caps: 6% opposed and 0% hostile reads 3% very negative overall.

