

CONFIDENTIAL · BASELINE PROFILE

Subject — Elon Musk

Behavioral baseline assembled from 44 scored sessions of 44 of your analysed videos (1379.4 min of measured speech). The score itself is measured across every session of this subject on the platform; the sessions listed here are your own.

Baseline established from 44 scored sessions

More measured sessions narrow the uncertainty bands; the baseline is re-pooled as sessions are scored. Measured on all 44 of your analysed video(s). The baseline figure itself is measured across every session of this subject on the platform, including videos held on other accounts, which are not listed here.



SUBJECT	Elon Musk
SESSIONS	44 scored of your videos
MEASURED SPEECH	1379.4 minutes
TOPIC CONSISTENCY	13 pts spread

18 /100 typical · CapoScore

LOW

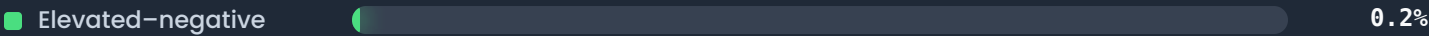
peak 86 · high 14474 sentences · 202 at review cut 65

CAPOSCORE · SENTENCE SCALE · ABSOLUTE 0-100

DELIVERY PROFILE — ACROSS 44 MEASURED SESSIONS

SHARE OF SPEAKING TIME · 8 STATES





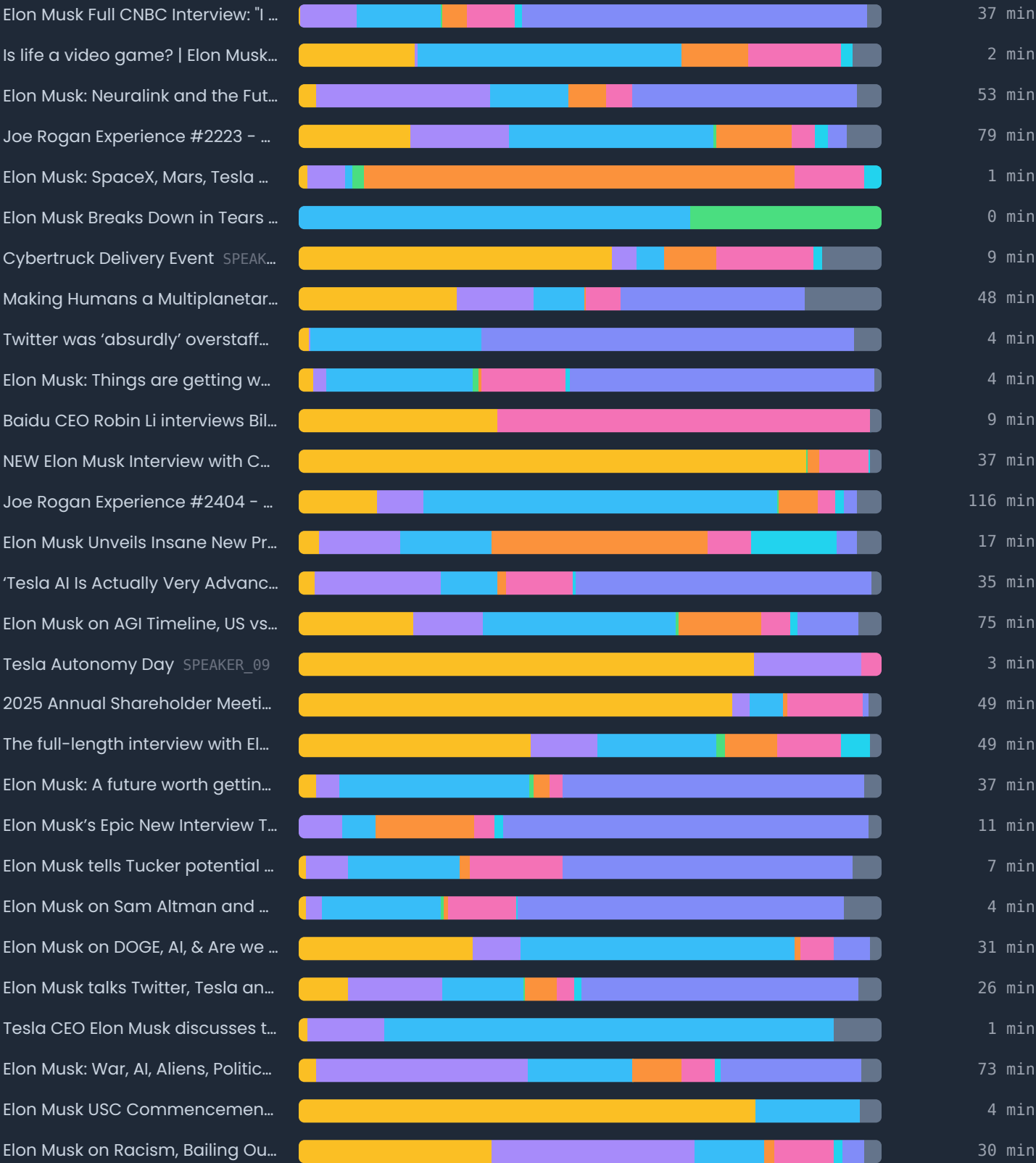
BASELINE 3.7%

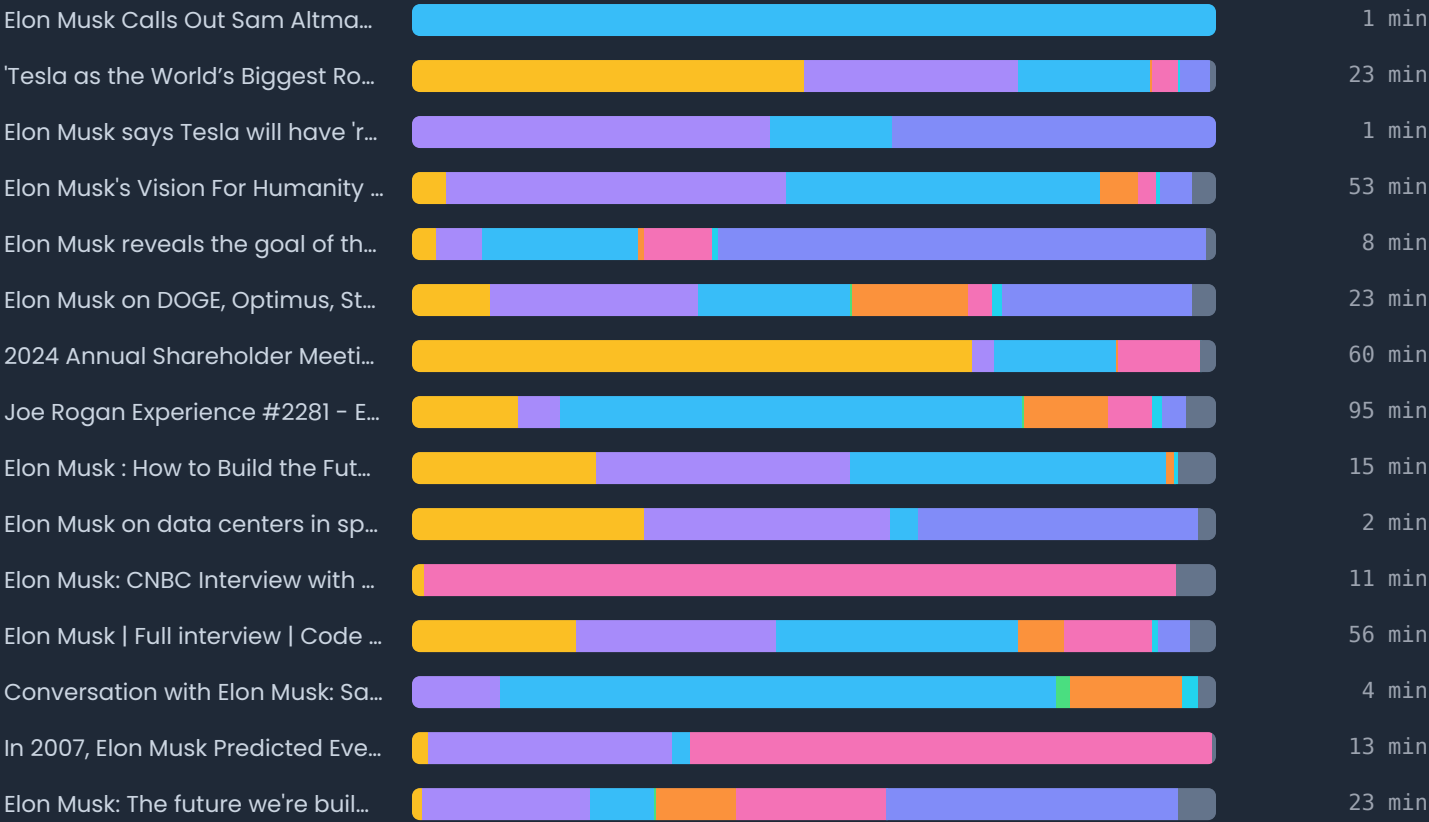
Measured, and nothing departed from this reference. A result, not a gap — which is why it is a figure here and not a bar above.

Measured over 15962 utterances, 1236 minutes of speech (audio_only 4136, multimodal 11826). Each figure is a share of speaking time, weighted by utterance overlap — not a score and not an intensity.

DELIVERY SESSION BY SESSION

44 MEASURED





STATE	TYPICAL SHARE	RANGE ACROSS SESSIONS	SESSIONS PRESENT
Composed	6.3%	0.0% – 87.1%	39 of 44
Smooth	10.4%	0.0% – 44.5%	38 of 44
Face-voice mismatch	17.8%	0.0% – 100.0%	40 of 44
Elevated-negative	0.0%	0.0% – 32.9%	18 of 44
Elevated-positive	1.6%	0.0% – 73.8%	33 of 44
Hesitant	5.8%	0.0% – 93.4%	36 of 44
Assertive	0.3%	0.0% – 14.9%	31 of 44
Unsteady voice	3.8%	0.0% – 63.8%	28 of 44
Baseline	2.9%	0.0% – 13.1%	39 of 44

TYPICAL IS THE MEDIAN OF THE PER-SESSION SHARES; THERE IS NO MEAN. A WIDE RANGE MEANS THE SHARE MOVED BETWEEN SESSIONS, NOTHING MORE — NO STATE IS STEADIER IN THE SENSE OF BEING BETTER.

READING THE DELIVERY STATES

KEY

- **Composed** — Steady delivery, with nothing departing from this person’s own norm.
- **Face-voice mismatch** — The face and the voice disagree with each other.
- **Elevated-positive** — Raised arousal carried with positive affect.
- **Assertive** — Emphatic, forceful delivery.
- **Baseline** — Measured, with nothing departing from the reference.
- **Smooth** — Fluent, unhesitating delivery.
- **Elevated-negative** — Raised arousal carried with negative affect.
- **Hesitant** — Halting delivery, with pauses and restarts.
- **Unsteady voice** — Instability in how the voice itself is produced.

These states describe delivery only. They are categorical and unranked — no state is more serious than another, and none raises or lowers the CapoScore.

Bottom Line on Elon Musk

§ 01

The **SEC funding secured controversy** topic carries the single strongest flagged cluster in this brief: eight of its fourteen scored sentences sit at or above the review cut, with a peak sentence CapoScore of **86** against this subject's own typical of **18**. That gap — a settled personal baseline in the LOW band paired with a topic-specific spike into HIGH — is the largest deviation in the record, and the pattern is one that **shows some properties associated with deception** on the specific statements addressing funding-secured claims, distinguishing them sharply from the surrounding conversation. Worth a second look at the SEC funding secured controversy window and the adjacent Tesla bankruptcy risk and short sellers exchanges, where a comparable concentration (eight of eighteen sentences at or above cut, peak 79) appears.

Outside these two clusters, the profile is topic-selective rather than generalized: the subject's median sentence CapoScore across 14,474 scored sentences is 18 (LOW), twenty of twenty analysed sessions individually score LOW-to-MODERATE, and topic consistency of 13 points indicates a narrow spread around that settled baseline. The delivery record reinforces this reading — face-voice mismatch dominates speaking time at 27.2%, composed and smooth delivery together account for roughly 38%, and assertive or elevated-negative states are rare (1.2% and 0.2% respectively) — meaning the flagged material is not accompanied by a wholesale change in delivery style but sits as isolated spikes inside an otherwise steady pattern. The finding is topic-anchored: SEC funding secured controversy, Tesla bankruptcy risk and short sellers, and sex change surgery risks each produce clusters at or above the review cut while most other recurring topics, including Twitter acquisition and censorship and Cybercab regulatory approval, stay within the LOW-to-MODERATE range.

SENTENCE BASELINE — BEHAVIOUR-ADJUSTED ABSOLUTE 0-100 · BANDS 40 / 60 / 75 · REVIEW CUT 65

TYPICAL

PEAK

AT OR ABOVE CUT

18

LOW

86

HIGH

202

OF 14474 SENTENCES

COVERAGE

44 /44

OF YOUR VIDEOS MEASURED

ALL 44 OF YOUR SCORED VIDEOS CARRY SENTENCE COVERAGE. THE SCORE ITSELF IS MEASURED ACROSS EVERY SESSION OF THIS SUBJECT ON THE PLATFORM; THE SESSIONS LISTED HERE ARE YOUR OWN. 202 SENTENCE(S) REACHED THE REVIEW CUT OF 65; THE TYPICAL DESCRIBES THE SUBJECT'S SETTLED STATE, NOT THOSE MOMENTS. UTTERANCE COVERAGE: 15962 SCORED UTTERANCES ACROSS 44 TRACK(S).

Per-Session Breakdown

\$ 02

VIDEO	SPEAKER	CAPOSCORE ABSOLUTE 0–100	BAND	TOPICS
Elon Musk Full CNBC Interview: "I Created AI"	SPEAKER_01	23	LOW	29
2024 Annual Shareholder Meeting Tesla	SPEAKER_11	26	LOW	43
'Tesla as the World's Biggest Robot Company:' Elon Musk on AI and U.S. Innovation WSJ	SPEAKER_00	16	LOW	19
Elon Musk on Sam Altman and ChatGPT: I am the reason OpenAI exists	SPEAKER_01	25	LOW	5
Elon Musk: Neuralink and the Future of Humanity Lex Fridman Podcast #438	SPEAKER_03	15	LOW	30
Elon Musk says Tesla will have 'robotaxis' on the road by 2020	SPEAKER_00	55	MODERATE	2
Elon Musk: A future worth getting excited about Tesla Texas Gigafactory interview TED	SPEAKER_02	21	LOW	31
Elon Musk USC Commencement Speech USC Marshall School of Business Undergraduate Commencement 2014	SPEAKER_01	10	LOW	5
Elon Musk: The future we're building -- and boring TED	SPEAKER_01	18	LOW	16
Elon Musk reveals the goal of the new Twitter	SPEAKER_01	28	LOW	4

VIDEO	SPEAKER	CAPOSCORE ABSOLUTE 0 – 100	BAND	TOPICS
Conversation with Elon Musk: Satya Nadella at Microsoft Build 2025	SPEAKER_01	18	LOW	5
Elon Musk: Things are getting weird fast	SPEAKER_01	15	LOW	5
‘Tesla AI Is Actually Very Advanced:’ Elon Musk on AI, China, Twitter and More WSJ	SPEAKER_00	20	LOW	31
2025 Annual Shareholder Meeting Tesla	SPEAKER_05	26	LOW	25
Tesla Autonomy Day	SPEAKER_09	10	LOW	16
Elon Musk: SpaceX, Mars, Tesla Autopilot, Self-Driving, Robotics, and AI Lex Fridman Podcast #252	SPEAKER_01	10	LOW	48
Elon Musk Calls Out Sam Altman & OpenAI	SPEAKER_00	55	MODERATE	3
Elon Musk Unveils Insane New Product	SPEAKER_00	18	LOW	10
Elon Musk’s Epic New Interview Today (AI & Robots)	SPEAKER_00	21	LOW	12
Elon Musk on DOGE, AI, & Are we in a Simulation? KMP Ep.18	SPEAKER_01	15	LOW	21
Elon Musk: CNBC Interview with David Faber	SPEAKER_00	14	LOW	18
NEW Elon Musk Interview with Code Conference 2021 - With Timestamps	SPEAKER_02	18	LOW	32
Elon Musk’s Vision For Humanity Real Talk with Zubyr	SPEAKER_01	15	LOW	24
Elon Musk on AGI Timeline, US vs China, Job Markets, Clean Energy & Humanoid Robots 220	SPEAKER_02	14	LOW	62
Elon Musk talks Twitter, Tesla and how his brain works — live at TED2022	SPEAKER_01	21	LOW	26
In 2007, Elon Musk Predicted Everything (Rare Lost Interview)	SPEAKER_01	15	LOW	11
The full-length interview with Elon Musk The Economist	SPEAKER_01	22	LOW	43

VIDEO	SPEAKER	CAPOSCORE ABSOLUTE 0–100	BAND	TOPICS
Is life a video game? Elon Musk Code Conference 2016	SPEAKER_01	10	LOW	3
Elon Musk on DOGE, Optimus, Starlink Smartphones, Evolving with AI, Why the West is Imploding	SPEAKER_03	15	LOW	23
Making Humans a Multiplanetary Species	SPEAKER_01	18	LOW	22
Elon Musk on Racism, Bailing Out Trump, Hate Speech, and More – The Don Lemon Show (Full Interview)	SPEAKER_01	30	LOW	27
Elon Musk : How to Build the Future	SPEAKER_01	16	LOW	10
Baidu CEO Robin Li interviews Bill Gates and Elon Musk at the Boao Forum, March 29 2015	SPEAKER_01	15	LOW	12
Joe Rogan Experience #2281 – Elon Musk	SPEAKER_03	21	LOW	84
Elon Musk Breaks Down in Tears During Interview 😭 #shorts	SPEAKER_01	38	LOW	2
Elon Musk Full interview Code Conference 2016	SPEAKER_01	15	LOW	37
Cybertruck Delivery Event	SPEAKER_01	40	MODERATE	8
Elon Musk: War, AI, Aliens, Politics, Physics, Video Games, and Humanity Lex Fridman Podcast #400	SPEAKER_01	15	LOW	56
Tesla CEO Elon Musk discusses the implications of A.I. on his children’s future in the workforce	SPEAKER_01	15	LOW	5
Joe Rogan Experience #2223 – Elon Musk	SPEAKER_02	21	LOW	62
Elon Musk tells Tucker potential dangers of hyper-intelligent AI	SPEAKER_01	14	LOW	6
Twitter was ‘absurdly’ overstaffed: Elon Musk	SPEAKER_01	32	LOW	7
Joe Rogan Experience #2404 – Elon Musk	SPEAKER_03	22	LOW	75
Elon Musk on data centers in space – Nov 2025	SPEAKER_00	45	MODERATE	3

Topic-Stress Map — Session × Topic (5 sessions w/ most topic overlap, of 44)

§ 03

TOPIC	TYPICAL	ELON MUSK ON ...	2025 ANNUAL S...	2024 ANNUAL S...	ELON MUSK TAL...	JOE ROGAN EXP...
battery cell cost parity	70	—	—	70	—	—
SEC funding secured controversy	66	—	—	—	66	—
Tesla bankruptcy risk and short sellers	63	—	—	—	—	—
sex change surgery risks	63	—	—	—	—	63
podcast sponsor ad	61	61	—	—	—	—
predicting the future accuracy	61	—	—	—	—	—
Optimus market cap projection	59	—	—	59	—	—
Tesla founding dispute with JB Straubel	57	—	—	—	57	—
Twitter acquisition and censorship	57	—	—	—	—	57
billionaire criticism from left	56	—	—	—	—	—
Cybercab regulatory approval	56	—	56	—	—	—
Tesla short sellers battle	56	—	—	—	56	—

EACH CELL IS THE TYPICAL (MEDIAN) SENTENCE SCORE OF THAT TOPIC IN THAT SESSION, ON THE SAME ABSOLUTE 0-100 CAPOSCORE SCALE AS THE REST OF THIS PAGE (LOW UNDER 40 · MODERATE 40-59 · ELEVATED 60-74 · HIGH 75+). AN EM DASH (—) IS **NOT MEASURED**: NO SCORED SENTENCES FALL INSIDE THAT TOPIC'S WINDOW FOR THAT SESSION. THAT IS MISSING DATA, NOT A LOW READING.

Topic Deep Dives — vs Baseline

§ 03b

Topic 01

battery cell cost parity

70

ELEVATED · 58

+50 vs baseline

evidence MODERATE

VERSUS BASELINE

The elevation is entirely a function of this speaker's own norm: the flagged statement sits far above his typical facial-behavior level of 20 even though it matches the only other scored session on this exact topic, so the departure is notable for him specifically rather than representing a change from prior coverage of battery cost parity.

This window carries a **+50 point deviation** from this speaker's facial baseline, produced by a single statement: **"So, but we expect to achieve cost parity with even the much lower supplier-sell price today by the end of the year,"** which scored **70/100** at 120:13. The claim shows some properties associated with deception: a forward-looking cost assurance delivered with composed vocal delivery while negative facial affect spiked to an extremely elevated level and the dominant expression shifted from neutral to anger. The pattern is consistent with a cross-channel mismatch between what was said and what the face carried while saying it.

RISK SECTIONS

Section 1 of 1 120:13-120:24 1 flagged peak 70 typical 70

Musk is addressing battery supply cycles and moves directly into a projection that Tesla will reach cost parity with the cheapest external cell suppliers by year end.

The single sentence in this stretch, timestamped 120:13, scored 70/100 -- the peak and typical for the section -- and is the only sentence in the window at or above the review cut. The score sits +50 above this speaker's facial baseline of 20, marking this specific forward-looking cost claim as the entire basis of the section's elevation.

Negative facial affect extremely elevated

Negative facial affect spiked to an extremely elevated level within this video and the dominant expression shifted from neutral to anger through part of the section, while heart rate and voice stress both held typical -- a facial signal at odds with the composed vocal delivery, which corroborates the sentence-level flag rather than explaining it away. The session as a whole runs elevated on negative facial affect and suppressed on heart rate against this speaker's baseline, a session-wide pattern that is hedged since it can reflect recording conditions as much as the person.

"So, but we expect to achieve cost parity with even the much lower supplier -sell price today by the end of the year."

TRANSCRIPT HIGHLIGHTS

7197s - 7225s 70

"So, but we expect to achieve cost parity with even the much lower supplier –sell price today by the end of the year."

A specific, date-bound cost assurance delivered with composed vocal cues but paired with an extremely elevated negative facial affect and a shift to an angry expression, producing a cross-channel mismatch on the direct claim.

7195s - 7196s

"you know, there's a bit of a feast –famine thing with battery"

Contextual framing preceding the flagged claim, setting up the supply-cycle caveat that leads into the cost-parity assurance.

BOTTOM LINE

The cost-parity assurance at 120:13 shows some properties associated with deception, delivered with composed speech but accompanied by an extremely elevated, anger-shifted facial signal that diverges from the vocal channel. Worth a second look at 120:13–120:24.

WHAT THE DATA SUGGESTS

- The facial spike to anger during an otherwise composed vocal delivery suggests leaked emotion the speech itself does not carry.
- The isolation of the departure to negative facial affect, with heart rate and voice stress both typical, points to a cross-channel mismatch specific to this claim rather than a generalized stress response.
- The single-sentence sample size means this reading rests on one data point within the window, even though it aligns with the only other scored session on this topic.
- The composed delivery state paired with a flagged score suggests the assurance was stated fluently, not hesitantly, which is itself part of the pattern rather than a mitigating factor.

AI-generated assessment from CapoScore behavioural data; not a determination of truthfulness.

Topic 02

SEC funding secured controversy

66

ELEVATED · 82

+46 vs baseline

evidence HIGH

VERSUS BASELINE

The subject's typical scoring sits at 20 across all scored sessions, so a topic typical of 66 with a peak of 86 reflects a sustained shift rather than a single spike; the shift tracks directly with the repeated funding assertions and the denial-admission pairing, delivered against a voice dominated by instability (57.2%) rather than the composed or smooth delivery that characterizes this subject's baseline elsewhere.

This window runs **+46** points above the subject's own face baseline, anchored by the statement **"I was I was force to admit that I lied to save Tesla's life and that's the only reason"**, which scored **86/100** at 28:50 — the highest-scoring flagged sentence in a single section that produced eight flagged statements out of fourteen. The pattern is consistent with a **deception-adjacent cluster** built around the direct SEC funding question: a categorical denial (**"it makes it look like I lied when I did not in fact lie"**, 83/100, 28:35) sits beside repeated assertions that **"funding was indeed secured"** (76/100, 27:56) and **"funding was actually secured"** (75/100, 27:36), all produced against a backdrop of **unsteady voice production across 57% of speaking time**. The delivery is not composed denial but arousal-carried denial, which the system reads as part of the pattern rather than against it.

RISK SECTIONS

Section 1 of 1 27:23-28:55 8 flagged peak 86 typical 66

The subject addresses the SEC 'funding secured' tweet directly, asserting the funds existed, criticizing the SEC's San Francisco office, and framing his settlement as coerced by Tesla's financial situation.

Eight sentences in this single section cross the review cut, climbing from 76/100 at 27:23 through 75-76/100 on the repeated funding claims (27:36, 27:56) to a peak of 86/100 at 28:50, with the categorical denial at 28:35 scoring 83/100. The scores rise rather than settle as the subject moves from assertion into justification and denial, with no flagged sentence in the section falling below 65.

Nothing notable

*"I have sufficient assets to complete the..."**"Funding was indeed secured."**"So I was forced to concede to the SEC unlawfully, those those bastards, and now it makes it look like I lied when I did not in fact lie."**"I was I was force to admit that I lied to save Tesla's life and that's the only reason."*

TRANSCRIPT HIGHLIGHTS

28:50 86

"I was I was force to admit that I lied to save Tesla's life and that's the only reason."

The window's peak score, an explicit admission-and-justification pairing delivered immediately after the categorical denial, forming the sharpest point of the cluster.

28:35 83

"So I was forced to concede to the SEC unlawfully, those those bastards, and now it makes it look like I lied when I did not in fact lie."

A categorical denial on the direct question of lying, delivered under raised arousal and paired with an accusation against the SEC, scoring far above the review cut.

27:56 76

"Funding was indeed secured."

A short, repeated assertion of the contested fact, scoring in the High band on its second occurrence within the section.

27:23 76

"I have sufficient assets to complete the..."

The opening statement of the section already scores above the review cut, setting an elevated baseline for everything that follows.

27:36 75

"So and I I should should say actually even in the, originally the, with Tesla Tesla back in the day, funding was actually secured."

A hedged, restart-heavy restatement of the funding claim that still lands above the review cut, showing disfluency did not lower the score.

BOTTOM LINE

The categorical denial at 28:35 (83/100) and the admission-justification at 28:50 (86/100) share the properties of deceptive speech, delivered under unsteady vocal production rather than composed denial, against a personal typical of 20. Warrants direct review of 27:23, 27:36, 27:56, 28:35, and 28:50 by an analyst.

WHAT THE DATA SUGGESTS

- The repetition of the funding claim across two separate high-scoring instances (27:36, 27:56) suggests reinforcement of a contested assertion rather than a single offhand remark.
- The immediate sequencing of denial (83/100) into admission-with-justification (86/100) at 28:35-28:50 suggests the two statements function as a single rhetorical unit rather than independent remarks.
- The dominance of unsteady voice production (57.2%) during the flagged cluster suggests the elevated scores coincide with vocal instability rather than with composed or fluent delivery.
- The section's escalating score trajectory, from 76 at the open to 86 at the close, suggests building pressure across the section rather than an isolated spike.

AI-generated assessment from CapoScore behavioural data; not a determination of truthfulness.

63

ELEVATED · 78

Topic 03

Tesla bankruptcy risk and short sellers

+43 vs baseline

evidence HIGH

VERSUS BASELINE

The +43-point gap against the face typical reflects a shift from this speaker's normal calm register into a sustained run of elevated-scoring statements built around denial and reframing, not a single outlier line. The +0 deviation against the topic's own typical shows this window performs exactly as this subject's prior scored session on this same topic did, indicating the pattern is stable and repeatable for this specific line of questioning rather than a one-off spike.

The window opens with a +43-point deviation from this speaker's face baseline, anchored by the strongest flagged statement in the set: **"Well the first part, I did actually have the funding secured, and there was a big trial in San Francisco, a big civil trial, and and the jury found me not guilty"** at 108:24, scoring **79/100**. A second statement at 108:48 — **"The reason I agreed to the fine for the SEC is not because the SEC was correct"** — also scores 79, reframing a settled regulatory penalty as coercion rather than admission. The pattern is consistent with a **categorical denial on the direct funding-secured question**, delivered while the voice carries **instability in production 59.1% of the time**, followed by a justification cluster (108:53, 69; 109:15, 70; 109:43, 70) that reframes bankruptcy risk and SEC conduct rather than answering it directly.

RISK SECTIONS

Section 1 of 1 108:20-109:51 8 flagged peak 79 typical 63

The subject responds to a question referencing his 'funding secured' tweet and the resulting SEC fine, then pivots to framing the settlement as coercion tied to Tesla bankruptcy risk and criticizes the SEC for not pursuing short-selling hedge funds.

The direct denial at 108:24 (79/100) and the settlement-justification statement at 108:48 (79/100) are the two highest-scoring statements in the entire window, both occurring on load-bearing claims rather than peripheral remarks. The scores remain elevated through the bankruptcy-coercion narrative (109:15 at 70, 109:20 at 67, 109:23 at 67) and into the closing accusation against short sellers (109:43 at 70), indicating the elevated scoring is sustained across the full stretch rather than isolated to one line.

suspect: sparse face data

Sparse face data

"Well the first part, I did actually have the funding secured, and there was a big trial in San Francisco, a big civil trial, and and the jury found me not guilty."

"The reason I agreed to the fine for the SEC is not because the SEC was correct."

"But not once did the SEC go after any of the hedge funds who were nonstop shorting and distorting Tesla."

TRANSCRIPT HIGHLIGHTS

6504s - 6512s 79

"Well the first part, I did actually have the funding secured, and there was a big trial in San Francisco, a big civil trial, and and the jury found me not guilty."

A categorical denial on the direct funding-secured question, delivered under dominant vocal instability, scored 79 against a personal typical of 20.

6525s - 6540s 79

"The reason I agreed to the fine for the SEC is not because the SEC was correct."

Reframes a paid regulatory settlement as an act of coercion rather than acknowledgment, matching the window's peak score.

6525s - 6540s 69

"That was extremely bad behavior by the SEC, corruption, frankly."

Escalates the reframing from denial into an unsupported counter-accusation, sustaining elevated scoring immediately after the peak statements.

6546s - 6559s 70

"And if they suspend our line of credit at that time, we would have gone bankrupt instantly."

Extends the coercion narrative with a high-stakes hypothetical that scores near the peak, reinforcing the justification cluster.

6581s - 6590s 70

"But not once did the SEC go after any of the hedge funds who were nonstop shorting and distorting Tesla."

Closes the segment by redirecting scrutiny toward third parties, a pattern consistent with the deflection seen throughout the flagged stretch.

BOTTOM LINE

The direct denial at 108:24 and the settlement-justification statement at 108:48, both scoring 79 against a face typical of 20, share the properties of deceptive speech, consistent with a categorical denial on the funding-secured question followed by a sustained reframing of the SEC settlement and bankruptcy risk. Warrants direct review of 108:24-109:51 by an analyst.

WHAT THE DATA SUGGESTS

- The direct question about funding secured produces the single highest score in the window, placing the denial itself — not surrounding commentary — at the center of the flagged pattern.
- The near-total dominance of Unsteady voice and Hesitant delivery (84.6% combined) suggests the denial and its justifications were produced under sustained vocal strain rather than fluent, settled recall.
- The seven-statement run above or near the cut indicates the elevated scoring is structural to the topic's justification narrative, not confined to the funding-secured line alone.
- The near-absence of Assertive or Elevated-positive delivery (under 2.5% combined) indicates the statements were not delivered with confident emphasis, despite their categorical phrasing.

AI-generated assessment from CapoScore behavioural data; not a determination of truthfulness.

63

ELEVATED · 78

Topic 04

sex change surgery risks

+43 vs baseline

evidence HIGH

VERSUS BASELINE

The +43 gap against the face baseline reflects a sustained shift into categorical, high-certainty claims on this topic rather than a brief spike, and the fact that the topic's own typical (63) already sits above the review cut shows this elevation is characteristic of how this subject is discussed, not an isolated moment within the window.

The window shows a ****+43-point deviation**** from this speaker's face-baseline typical, anchored by the statement at 123:07 — "If you castrate kids and trans them, the probability of suicide increases, it does not decrease, it substantially increases" — which scored ****78/100**** against a personal typical of 20. Paired with the 73-scored claim at 123:02 that "the save trans kid thing is a lie" and the 76-scored claim at 121:57 that kids "die" on the operating table due to lacking technology, the cluster of categorical, escalating assertions on the direct medical claims is consistent with a ****deception-adjacent pattern****: absolute claims of causation and motive delivered while face and voice disagree for nearly three-quarters of the speaking time. The pattern is one of repeated, unqualified causal and moral claims — lie, kill, die — stacked across a 96-second stretch rather than a single outlier line.

RISK SECTIONS

Section 1 of 2 121:43-122:42 1 flagged peak 76 typical 38

The subject was describing deaths occurring during sex-change operations, attributing them to insufficient surgical technology, and drawing a comparison to Nazi-era medical conduct.

The sentence at 121:57 — describing kids dying "on the operating table" — scores 76/100, the section's single flagged statement, against context lines at 14, 62 and 10 that sit well below the cut. The gap between the flagged claim and its surrounding context indicates the elevated score is tied specifically to the causal death claim, not the section's tone generally.

suspect: cross-talk

Unreliable: cross-talk

This section's voice-stress channel is excluded due to cross-talk with another speaker, so it cannot corroborate the 76 score; negative facial affect is typical for this speaker within this video, leaving the finding resting on the sentence content alone.

"die uh in in with these uh sex change operations they die the number of deaths on the operating table People don't hear about this a lot of kids because That it's we don't really actually have the technology to make thi..."

Section 2 of 2 123:02-123:48 3 flagged peak 78 typical 66

The subject asserted that narratives framing gender-transition care as suicide prevention are false, that suicide risk rises rather than falls with treatment, and that surgical deaths occur.

Three sentences cross the cut in quick succession: 123:02 at 73 ("the save trans kid thing is a lie"), 123:07 at 78, the section's peak ("the probability of suicide increases... substantially"), and 123:27 at 67 ("you're not saving them, you're killing them"), with the section typical at 66 — itself above the review cut. The concentration of three categorical claims within 25 seconds, each restating the same causal assertion with increasing bluntness, is the section's defining feature.

Nothing notable

"Yeah, and what I wanna emphasize is that the save trans kid thing is a lie."

"If you castrate kids and trans them, the probability of suicide increases, it does not decrease, it substantially increases."

"So you're not saving them, you're killing them."

TRANSCRIPT HIGHLIGHTS

123:07 78

"If you castrate kids and trans them, the probability of suicide increases, it does not decrease, it substantially increases."

The section's peak score sits on a categorical causal claim stated as settled fact, with no qualification of source or magnitude.

121:57 76

"die uh in in with these uh sex change operations they die the number of deaths on the operating table People don't hear about this a lot of kids because That it's we don't really actually have the technology to make thi..."

A death claim delivered with disfluent, restarting phrasing that still lands as an absolute assertion about surgical outcomes.

123:02 73

"Yeah, and what I wanna emphasize is that the save trans kid thing is a lie."

An explicit accusation of falsehood aimed at a broader narrative, framed as a deliberate correction rather than a personal opinion.

123:27 67

"So you're not saving them, you're killing them."

A blunt reframing of the prior claims into a direct causal accusation, closing the escalation begun at 123:02.

123:19 64

"The studies have done that I've seen the risk of suicide triples if you trans kids."

A supporting claim citing unspecified studies just below the cut, reinforcing the adjacent flagged statements without independent sourcing.

BOTTOM LINE

The scored behaviour here is what this system flags as deception-like: a cluster of categorical claims — "it's a lie" (73), "substantially increases" (78), "you're killing them" (67) — delivered under a face/voice mismatch spanning most of the speaking time. Warrants direct review of 121:57, 123:02, 123:07, and 123:27 by an analyst.

WHAT THE DATA SUGGESTS

- The categorical framing ("it's a lie", "you're killing them") delivered under face/voice mismatch is consistent with a deception-adjacent pattern rather than a calm factual recitation.
- The concentration of elevated scores on the specific causal and moral claims, rather than the surrounding context lines, suggests those particular assertions carry the risk, not the topic broadly.
- The cross-talk affecting Section 1's voice-stress channel leaves the operating-table death claim without independent channel corroboration.
- The near-absence of composed delivery (1.4%) across the window suggests the flagged claims were not made from a settled baseline state.

AI-generated assessment from CapoScore behavioural data; not a determination of truthfulness.

61

ELEVATED · 45

Topic 05

podcast sponsor ad

+41 vs baseline

evidence MODERATE

VERSUS BASELINE

The +41 gap reflects this speaker's face typical of 20, not any drift within the topic itself — the window's typical of 61 matches this topic's only prior scored instance exactly, so the elevation is characteristic of how this specific claim is delivered rather than a change in behavior across sessions.

The window's single flagged statement, **“That'll be half a gigawatt and operational middle of next year”** (18:35), scored 61/100 — a **+41 deviation** against this speaker's face-derived typical of 20, though it stays under the 65 review cut. The pattern shows some properties associated with deception on the certainty axis: a forward-looking capacity claim delivered with fully Composed delivery and no hedging, distancing, or vocal instability to accompany it. With only one scored sentence in the window, the deviation stands on its own rather than as part of a cluster.

RISK SECTIONS

Section 1 of 1 18:35-18:41 no sentence reached the review cut **peak 61** typical 61

The subject was responding to a claim of falling behind on AI training infrastructure, describing the buildout and projected capacity of the Cortex-2 training cluster.

The section registered no sentence at or above the review cut; the sole statement present, at [18:35], scored 61/100 and is the highest value in the window despite not crossing 65. The section is calm by the cut threshold but the score itself sits well above this speaker's face typical of 20.

Nothing notable

“That'll be half a gigawatt and operational middle of next year.”

TRANSCRIPT HIGHLIGHTS

1121s-1152s **61**

“That'll be half a gigawatt and operational middle of next year.”

A specific, unhedged capacity and timeline claim delivered under Composed delivery, scoring well above this speaker's personal typical without crossing the review cut.

1101s-1121s

“You're falling behind.”

Context statement framing the challenge the subject is responding to; not independently scored.

1101s-1121s

“Well, we have Cortex-2 that's being built out.”

Transitional claim preceding the flagged statement, establishing the infrastructure narrative; not independently scored.

BOTTOM LINE

The window's single statement at [18:35] shows some properties associated with deception, carrying an unhedged forward capacity claim scored well above this speaker's personal typical, though the small sample leaves this reading anchored to one sentence rather than a cluster. Worth a second look at 18:35.

WHAT THE DATA SUGGESTS

- The elevation is concentrated in a single unhedged, forward-looking claim rather than distributed across the exchange.
- Composed delivery accompanying an above-typical score suggests the claim was made without visible hesitation, which sharpens rather than dilutes the elevation.
- The absence of deviation from this topic's own prior typical suggests this level of scoring may be characteristic of how this specific claim type is delivered, not a one-off spike.
- The single-sentence sample means the finding cannot yet be corroborated by a pattern across multiple statements on this topic.

AI-generated assessment from CapoScore behavioural data; not a determination of truthfulness.

61

ELEVATED · 55

Topic 06

predicting the future accuracy

+41 vs baseline

evidence MODERATE

VERSUS BASELINE

The 41-point rise above his face-baseline typical is driven by three tightly clustered restatements of the same predictive-accuracy claim rather than by a single spike; against this topic's own prior typical the window is flat, since this is the only scored session on the subject to date. The behavior that produces the elevation is repetition under composed delivery, not visible distress.

This window sits **41 points above** Elon Musk's face-baseline typical, anchored by the statement at [82:26], scored **70/100**: "probability of success than most people... but my batting average is very go..." Flanking assertions at [82:00] (69/100) and [82:20] (65/100) repeat the same self-validating claim about predictive accuracy in close succession. The cluster **shows some properties associated with deception**, consistent with a pattern of repeated, high-certainty self-assessment on a claim that cannot be externally verified in the moment, rather than a single isolated remark.

RISK SECTIONS

Section 1 of 1 81:46-82:57 3 flagged peak 70 typical 63

Musk pivots from a statement on immigration policy into an unprompted claim about his own track record for predicting future events, framing skepticism toward him as a 'Cassandra effect.'

The section opens with context-level statements at [81:46] (56) and [81:48] (59) before three sentences cross the review cut: [82:00] at 69, [82:20] at 65, and [82:26] at 70, each repeating the claim of a high 'batting average' on predictions. The section closes back below cut at [82:44] (55) and [82:54] (63), where he adds a timing caveat.

suspect: cross-talk

Unreliable: cross-talk

Negative facial affect runs higher than average within this video across part of this stretch, coinciding with the sentences scoring 65-70 on his predictive-accuracy claims, while heart rate stays typical for this video even though the session as a whole runs elevated against his baseline, a difference that may owe as much to recording conditions as to the person. Voice stress is unreliable here due to cross-talk with another speaker and is excluded from this reading.

"I am not perfect but I'm very good at predicting the future my my batting average is is is very high for a human but this this does have somewhat of a Cassandra effect where where people don't believe that what I'm sayi..."
"know I can see into the future with a high about not perfectly but with a higher probability"
"probability of success than most people so and including things that people believe to be in some cases counterintuitive and again I'm not saying I'm always right I certainly am not but but my batting average is very go..."

TRANSCRIPT HIGHLIGHTS

82:26 70

"probability of success than most people so and including things that people believe to be in some cases counterintuitive and again I'm not saying I'm always right I certainly am not but but my batting average is very go..."

The peak-scoring statement pairs an explicit disclaimer ('I'm not saying I'm always right') with a restated claim of superior accuracy, a hedge-and-assert structure typical of this cluster.

82:00 69

"I am not perfect but I'm very good at predicting the future my my batting average is is is very high for a human but this this does have somewhat of a Cassandra effect where where people don't believe that what I'm sayi..."

Introduces the 'Cassandra effect' framing that recasts skepticism toward his claims as others' failure to recognize accuracy rather than a testable assertion.

82:20 65

"know I can see into the future with a high about not perfectly but with a higher probability"

A third repetition of the same self-assessed accuracy claim within under thirty seconds, delivered with disfluency but no retreat from the core assertion.

82:54 63

"maybe not exactly in the timing that I thought, but it will come to pass."

A context-level statement that narrows the claim to eventual outcome regardless of timing, softening falsifiability while remaining below the review cut.

81:48 59

"Well very good at predicting the future."

The opening assertion that seeds the section's central claim, scoring below cut but setting up the pattern that escalates through 82:00-82:26.

BOTTOM LINE

The repeated, self-validating claims of predictive accuracy at 82:00, 82:20, and 82:26 show some properties associated with deception, paired with hedging language that softens falsifiability while the core assertion is restated three times. Worth a second look at 82:00, 82:20, 82:26.

WHAT THE DATA SUGGESTS

- The clustering of three above-cut sentences within roughly 25 seconds suggests a rehearsed or defensive framing of the accuracy claim rather than a spontaneous aside.
- The pairing of hedges ('I am not perfect', 'I certainly am not') with high-scoring assertions indicates the hedging functions to soften an unverifiable claim rather than retract it.
- The higher-than-average negative facial affect confined to this section, alongside composed vocal delivery, points to a mismatch between outward calm and underlying facial signal worth flagging rather than resolving.
- The absence of any prior scored session on this topic means the +0 deviation against topic typical reflects a single data point, not an established personal norm.

AI-generated assessment from CapoScore behavioural data; not a determination of truthfulness.

Topic 07

Optimus market cap projection

59

MODERATE · 55

+39 vs baseline

evidence MODERATE

VERSUS BASELINE · 2 SESSIONS POOLED

The elevation is not a one-off spike but a repeated behavior: both scored sessions produce the same 59-typical pattern on this topic, and the flagged statements in each cluster follow an identical structure of citing a modest figure and then escalating it within the same breath, while Section 1 pairs that escalation with a documented shift toward negative facial affect and a face-voice mismatch delivery.

This topic runs ****39** points above this speaker's face-baseline typical******, cresting at ****68/100**** on the claim that at scale Tesla would 'basically make about a trillion dollars of profit a year' from Optimus before extending to a ****\$20 trillion market cap from Optimus alone**** (55:51-56:11). A second cluster in the later session repeats the same escalating structure, doubling an analyst figure to ****literally \$25 trillion market cap**** (94:47, 68/100). Delivered alongside a documented shift in facial affect from neutral to anger and a face-voice mismatch pattern in the first session, this run of escalating valuation claims ****shows some properties associated with deception****, consistent with assertive, high-arousal projection rather than composed recall of a settled figure.

RISK SECTIONS

Section 1 of 2 55:51-56:39 session 1 of 2 **2 flagged peak 68** typical 59

Musk projects Optimus unit pricing and volume economics, then extrapolates from unit profit to a company-wide market capitalization figure and compares it to the world's most valuable company.

The section opens at its peak, 68/100 at 55:51, on the profit-per-year claim, followed by 66/100 at 56:11 on the '\$20 trillion market cap from Optimus alone' extension; the surrounding context sentences at 56:20 (59) and 56:27 (55) continue the same escalation just under the cut, giving the section a typical of 59.

Negative facial affect extremely elevated

Negative facial affect is extremely elevated for this section and sustained throughout, with the dominant expression shifting from neutral to anger, while voice stress stays typical and delivery reads as face-voice mismatch; the combination of a flat-to-typical voice against a visibly negative, shifting face is consistent with strained internal affect underneath composed vocal delivery rather than calm recitation of a settled figure.

"So, So, and and I think if you sell it for \$20,000 or something, this is at large scale volume, Tesla would basically make about a trillion dollars of profit a year from that."

"If the price outings multiple is, say, 20 or 25, something like that, that would mean a \$20 trillion market cap from Optimus alone."

Section 2 of 2 94:25-95:26 session 2 of 2 **1 flagged peak 68** typical 45

Musk references a prior ARK Invest analysis of the autonomous-transport market and then extends it, assigning Optimus its own multi-trillion-dollar valuation figure.

The single flagged sentence at 94:47 (68/100) doubles the cited \$5-7 trillion transport estimate to 'literally \$25 trillion' for Optimus specifically; it sits inside a lower-scoring stretch (typical 45) bracketed by context sentences at 16, 64, and 26, marking it as an isolated spike rather than a sustained run.

Nothing notable

"know as I said at the beginning of the of the presentation I you know I agree with the arc and best analysis that autonomous transport is called sort of a five to seven trillion dollar market cap situation optimus I thi..."

TRANSCRIPT HIGHLIGHTS

3354s - 3365s 68

"So, So, and and I think if you sell it for \$20,000 or something, this is at large scale volume, Tesla would basically make about a trillion dollars of profit a year from that."

Stammered opening into a specific unit-profit claim that anchors the section's escalating valuation sequence.

5687s - 5711s 68

"know as I said at the beginning of the of the presentation I you know I agree with the arc and best analysis that autonomous transport is called sort of a five to seven trillion dollar market cap situation optimus I thi..."

Cites a third-party estimate then doubles it to a self-generated figure within the same sentence, an escalation pattern matching Section 1.

3372s - 3381s 66

"If the price outings multiple is, say, 20 or 25, something like that, that would mean a \$20 trillion market cap from Optimus alone."

Extends the prior profit claim into a company valuation figure delivered under the same fractured, face-voice mismatch state.

5667s - 5685s 64

"even the most optimistic estimates that I've seen for optimists, optimist, optimists. I think under under count the magnitude of of what this robot will be able to do"

Repeated word restarts precede a claim that all existing external estimates are too conservative, setting up the doubled figure that follows.

3391s - 3395s 55

"conceivable, it's within the realm of possibility for"

Hedged qualifier language trailing the peak claims, sitting just under the cut but continuing the same valuation thread.

BOTTOM LINE

The escalating, self-doubled valuation claims at 55:51, 56:11, and 94:47 show some properties associated with deception, each moving from a cited or modest figure to a sharply larger self-generated one within the same sentence. Worth a second look at 55:51, 56:11, and 94:47.

WHAT THE DATA SUGGESTS

- The repeated escalation structure — citing a smaller figure then doubling it within the same sentence — recurs across both sessions, suggesting a rehearsed rhetorical pattern rather than a single spontaneous overstatement.
- The Section 1 shift from neutral to angry facial affect alongside a face-voice mismatch, paired with typical voice stress, points toward strained internal affect surfacing on the face while the voice stays composed.
- Section 2's uniform Composed delivery around an equally high-scoring statement (94:47) indicates the escalation pattern is not confined to visibly strained delivery — it also occurs in a fully steady vocal register.
- The flat +0 deviation against this topic's own typical across sessions indicates the elevated pattern is topic-specific and stable, not a one-time reaction to a particular question.

57

Topic 08

Tesla founding dispute with JB Straubel

MODERATE · 52

+37 vs baseline

evidence MODERATE

VERSUS BASELINE

The rise is driven by two statements made directly on the founding dispute — the shell-company claim and the rebuttal of Eberhard's account — that land far above this speaker's resting face typical of 20, while the surrounding regret and context sentences stay near the topic's own typical of 57; the departure is concentrated in the direct denial and its rebuttal, not spread evenly across the window.

The strongest statement in this window is a **“categorical denial”** — “Tesla did not exist in any—Tesla was a shell company with no employees, no intellectual property when I invested,” scored **79/100** at 33:02 — delivered against a topic typical of 57 and a face typical of 20, a swing of **+37 points** above this speaker's own resting baseline. A second flagged statement at 33:25, disputing a co-founder's account of who created the company, scores **66/100**. Delivery across the window is dominated by **“unsteady voice”** (75.1% of speaking time), so the denial and the follow-on rebuttal are not calm assertions but statements produced under sustained vocal instability. The pattern is **“consistent with”** speech that shows some properties associated with deception on the direct question of who founded Tesla and under what terms.

RISK SECTIONS

Section 1 of 1 32:27-33:43 2 flagged peak 79 typical 57

Musk is discussing what he calls his worst business decision — not founding Tesla jointly with JB Straubel — and then disputes a competing account of the company's founding attributed to co-founder Martin Eberhard. The section opens with a regret statement at 32:27 (57/100, below cut) before rising to the section peak at 33:02 (79/100) on the claim that Tesla was an empty shell company at the time of investment. A second flagged statement at 33:25 (66/100) follows a stammered aside at 33:18 (60/100) and asserts that he and Straubel, not Eberhard, created the company.

suspect: cross-talk

Unreliable: cross-talk

Negative facial affect reads lower than average for this speaker through most of the section, a departure for this session but a swing size he shows elsewhere, so it does not by itself sharpen the finding; voice stress is unreliable here due to cross-talk with another speaker and is excluded, and some sentences in this stretch may not be reliably attributed to this speaker.

“Tesla did not exist in any—Tesla was a shell company with no employees, no intellectual property when I invested.”

“And, ultimately, the creation of that company was done by JV and me, and, unfortunately, there's someone else, another co-founder who has made it his life's mission to make it sound like he created the company, which i...”

TRANSCRIPT HIGHLIGHTS

33:02s 79

"Tesla did not exist in any—Tesla was a shell company with no employees, no intellectual property when I invested."

The section's peak statement, a categorical factual claim delivered on the direct question of the company's origins, scored well above the review cut.

33:25s 66

"And, ultimately, the creation of that company was done by JV and me, and, unfortunately, there's someone else, another co-founder who has made it his life's mission to make it sound like he created the company, which is false."

A direct rebuttal of a rival founding narrative that clears the review cut and immediately follows a stammered, halting aside.

33:18s 60

"And I don't want to to get get into into the the nastiness nastiness here, but I didn't invest in an existing company."

Repeated, stammered phrasing just below the cut precedes the section's second flagged statement, suggesting friction around the topic before the rebuttal is delivered.

32:27s 57

"I mean the worst business decision I ever made was Not starting Tesla with just JV Straubel By far the worst decision I've ever made is not just starting Tesla with JB, that's number one by far."

An opening regret statement at the topic's typical level, setting up the founding narrative that follows.

33:12s 52

"But a false narrative has been created by one of the other co-founders, Martin Eberhard."

Names the source of the disputed narrative directly before the two flagged rebuttal statements.

BOTTOM LINE

The categorical denial at 33:02 (79/100) and the rebuttal at 33:25 (66/100) show some properties associated with deception on the question of who founded Tesla, delivered under sustained vocal instability rather than composed speech. Worth a second look at 33:02, 33:25.

WHAT THE DATA SUGGESTS

- The peak score lands on the most categorical factual claim in the window ("no employees, no intellectual property"), not on the surrounding regret or context statements.
- The stammered repetition immediately before the second flagged statement suggests friction around naming the rival founder before the rebuttal is delivered.
- The dominance of unsteady voice across 75% of speaking time indicates the denial and rebuttal were not produced from a settled baseline state.
- The cross-talk affecting voice stress in this stretch means the vocal channel cannot be used to corroborate or contradict the sentence scores here, leaving the facial-affect reading as the only usable channel context.

AI-generated assessment from CapoScore behavioural data; not a determination of truthfulness.

57

Topic 09

Twitter acquisition and censorship

MODERATE · 54

+37 vs baseline

evidence MODERATE

VERSUS BASELINE

The gap against the face baseline reflects that this speaker's default sessions run calm, while this specific topic consistently pulls him into categorical, high-certainty framing of Twitter's prior ownership and his motive for acquiring it. Against the topic's own typical, this window shows no additional departure — the elevation is a property of the subject matter itself, reproduced here rather than intensified.

The topic sits **37 points above** this speaker's own face baseline, driven by the statement at **44:00**, "Yeah, I mean, I said at the time, like, I think that, like the reason for acquiring Twitter is because it was causing destruction at a civilizational level," which scores **73/100** — the peak of this window. A second flagged statement at **45:13**, describing the prior Twitter ownership as pushing a "very negative, nihilistic, untrue worldview," scores **68/100**. Both statements are categorical, sweeping justifications for the acquisition delivered while heart rate runs elevated for the section; the pattern **shows some properties associated with deception** on the specific claims about motive and characterization of Twitter's prior state, without moving outside the moderate band the rest of the window occupies.

RISK SECTIONS

Section 1 of 2 43:23-44:23 **1 flagged peak 73** typical 47

Musk is explaining his stated rationale for acquiring Twitter, framing it as a response to what he describes as civilizational-level harm.

The section peaks at 73/100 on the direct justification at 44:00, well above the review cut, flanked by context sentences at 52 (43:23) and 42 (44:12) that sit below it but still track upward toward the flagged line. The section's typical of 47 shows the elevation is concentrated in this one categorical claim rather than spread evenly across the stretch.

suspect: cross-talk

Unreliable: cross-talk

Heart rate elevated

Heart rate runs elevated for this section, confined to this stretch and reliable, while negative facial affect stays typical; voice stress is unreliable here due to cross-talk with another speaker and is excluded from interpretation.

"Yeah, I mean, I said at the time, like, I think that, like the reason for acquiring Twitter is because it was causing destruction at a civilizational level."

Section 2 of 2 45:08-46:00 **1 flagged peak 68** typical 57

Musk continues characterizing the prior Twitter leadership's editorial stance as harmful and 'nihilistic,' building toward the acquisition's stated justification.

The flagged statement at 45:13 scores 68/100, describing the prior ownership's worldview as "very negative, nihilistic, untrue," with the surrounding context sentence at 45:10 (61) and later at 45:53 (57) sitting just below the cut but keeping the section's typical at 57 — higher than Section 1's typical, indicating sustained rather than isolated elevation.

suspect: cross-talk

Unreliable: cross-talk

Heart rate is higher than average for this section, confined to the stretch and reliable, while negative facial affect remains typical; voice stress is again unreliable due to cross-talk and is excluded.

"But that was was essentially, they were pushing this very negative, nihilistic, untrue worldview on on the world and it was causing a lot of damage."

TRANSCRIPT HIGHLIGHTS

44:00s 73

"Yeah, I mean, I said at the time, like, I think that, like the reason for acquiring Twitter is because it was causing destruction at a civilizational level."

The highest-scoring statement in the window, a categorical justification for the acquisition delivered with raised heart rate.

45:13s 68

"But that was was essentially, they were pushing this very negative, nihilistic, untrue worldview on on the world and it was causing a lot of damage."

A sweeping characterization of the prior ownership's intent that sustains the section's elevation above the review cut.

45:10s 61

"So we don't want the whole world to be a zombie apocalypse."

Context sentence just below the cut that frames the flagged claim with vivid, absolute language.

45:53s 57

"A lot of, a lot people didn't course – correct, but it's gone directionally, it's directionally correct."

A hedged concession that still sits at moderate elevation, softening but not retracting the preceding claims.

43:23s 52

"we telling the truth we're"

A fragmented lead-in to the section's peak statement, scoring moderate ahead of the flagged claim.

BOTTOM LINE

The justification for the acquisition at 44:00 and the characterization of prior ownership at 45:13 show some properties associated with deception in their categorical framing under measurably raised heart rate. Worth a second look at 44:00 and 45:13.

WHAT THE DATA SUGGESTS

- The concentration of elevation on two specific claims – the acquisition motive and the characterization of prior ownership – suggests the pattern is tied to those specific assertions rather than the topic broadly.
- Heart rate elevation confined to both flagged stretches, with facial affect staying typical, is consistent with heightened arousal on these particular claims rather than a stable elevated state.
- The unreliability of voice stress in both sections due to cross-talk leaves one corroborating channel unresolved, limiting how confidently the vocal dimension can be read alongside the sentence scores.
- The dominant face-voice mismatch share across the topic is consistent with a cross-channel pattern worth weighing alongside the flagged statements, though it does not by itself establish intent.

AI-generated assessment from CapoScore behavioural data; not a determination of truthfulness.

56

MODERATE · 55

Topic 10

billionaire criticism from left

+36 vs baseline

evidence MODERATE

VERSUS BASELINE

The elevation against this speaker's face typical comes from a run of categorical, repeated "love of humanity" claims delivered with stammered restarts and no steady or baseline delivery state present anywhere in the window — the entire speaking time falls into unsteady voice, face-voice mismatch, or hesitant states. Against this topic's own prior session, the window sits exactly at typical, indicating this elevated, repetition-heavy defense is this topic's established pattern rather than a departure within it.

This window runs **36 points above** this speaker's face typical, anchored by the highest-scoring statement in the stretch — "SpaceX, Tesla, Neuralink, and Boring Company are philanthropy." (64/100, 62:38) — followed almost immediately by a **disfluent, repetition-heavy defense** of that claim: "If If you say philanthropy is love of humanity, they they are philanthropy, Tesla Tesla is accelerating sustainable energy, this is a love of full entropy." (63/100, 62:42). Neither statement crosses the review cut, but the cluster of categorical, love-of-humanity framing delivered with stammered restarts **shows some properties associated with deception** as speech, and the pattern is consistent with a rehearsed justification rather than spontaneous reflection. The sample is small — only 7 scored sentences in the window — which limits how precisely the cluster can be characterized beyond this stretch.

RISK SECTIONS

Section 1 of 1 62:31-63:20 no sentence reached the review cut peak 64 typical 56

The subject is defending SpaceX, Tesla, Neuralink, and Boring Company as forms of philanthropy in response to criticism of billionaire wealth, building a company-by-company case for each as a "love of humanity."

Scores cluster tightly between 56 and 64 across five of the seven sentences, peaking at [62:38] 64/100 and [62:42] 63/100, with the two lowest-scoring lines — [63:05] 26/100 and [62:31] 46/100 — bracketing the stretch as more measured framing. None reach the review cut, but the concentration of categorical assertions in the upper half of this range, rather than a single outlier, is the notable feature.

Negative facial affect suppressed

Negative facial affect is suppressed against this speaker's own within-video pattern even though the session overall runs extremely elevated against baseline, with the dominant expression switching from anger to neutral through most of the section — a swing that reads as genuine unease under a session-wide condition as much as a controlled reaction. Voice stress sits typical within this video but higher than average for the session as a whole, and delivery is read as unsteady voice with face-voice mismatch, consistent with the 51.8% unsteady-voice and 38.1% mismatch shares in the behavioral distribution; some sentences here were contested between speakers, so this reading carries that caveat.

"SpaceX, Tesla, Neuralink, and Boring Company are philanthropy."

"If If you say philanthropy is love of humanity, they they are philanthropy, Tesla Tesla is accelerating sustainable energy, this is a love of full entropy."

"Space X is trying to ensure the long-term survival of humanity with multi-planet species."

TRANSCRIPT HIGHLIGHTS

3759s - 3763s 64

"SpaceX, Tesla, Neuralink, and Boring Company are philanthropy."

A flat, unqualified categorical claim carrying the window's peak score, delivered without hedging.

3763s - 3768s 63

"If If you say philanthropy is love of humanity, they they are philanthropy, Tesla Tesla is accelerating sustainable energy, this is a love of full entropy."

Stammered repetition of key words accompanies a high-scoring elaboration, suggesting effortful construction of the justification in real time.

3769s - 3783s 56

"Space X is trying to ensure the long-term survival of humanity with multi-planet species."

A repeated instance of the same love-of-humanity template applied to a second company, reinforcing the pattern rather than varying it.

3784s - 3784s 56

"This is love of humanity."

A terse, formulaic restatement of the framing used across multiple companies in this stretch.

3751s - 3758s 46

"I think if you care about the reality of goodness instead of the perception of it, philanthropy is extremely difficult."

The opening line sits mid-range and more reflective in tone, setting up the higher-scoring categorical claims that follow.

BOTTOM LINE

The cluster of categorical philanthropy claims at 62:38 and 62:42, delivered with stammered repetition and no steady delivery state, shows some properties associated with deception as speech even though no single sentence clears the review cut. Worth a second look at 62:38, 62:42.

WHAT THE DATA SUGGESTS

- The absence of any Composed or Baseline share across 41 seconds of speech suggests the entire response to this line of criticism was produced under some form of vocal or cross-channel strain.
- The repeated "X X" stammering at the window's two highest-scoring points suggests the justification was being assembled rather than recalled.
- The topic's exact match to its own prior typical (+0 deviation) suggests this defensive framing is a consistent, rehearsed response pattern rather than a one-off reaction.
- The session-wide elevation in negative facial affect and voice stress means some of this reading may reflect recording conditions rather than the topic itself, and should be weighed accordingly.

AI-generated assessment from CapoScore behavioural data; not a determination of truthfulness.

Key Findings

§ 04

01

SEC funding secured controversy produces the session's sharpest spike

Eight of fourteen scored sentences on this topic sit at or above the review cut, with a peak CapoScore of 86 against a personal typical of 18, and a topic typical of 66 that is far above the subject's own settled median. The pattern is consistent with a denial-cluster on a direct-question topic and shows some properties associated with deception on these specific statements. Worth a second look at the SEC funding secured controversy window.

02

Tesla bankruptcy risk and short sellers forms a second concentrated cluster

Eight of eighteen scored sentences on this topic reach or exceed the review cut, with a peak CapoScore of 79 against the subject's typical of 18, giving the topic itself a typical of 63. The concentration of flagged statements on this single subject, separate from the SEC funding cluster, is consistent with topic-selective elevation rather than a generalized shift. Worth a second look at the Tesla bankruptcy risk and short sellers exchanges.

03

Sex change surgery risks carries a smaller but still notable cluster

Four of ten scored sentences on this topic sit at or above the review cut, with a peak CapoScore of 78 against a typical of 63 for the topic and 18 for the subject overall. The size of the cluster is smaller than the two leading topics but still stands apart from the subject's ordinary range.

04

Flagged statements are carried with mismatch and elevated-positive delivery, not composure

Face-voice mismatch accounts for 27.2% of all measured speaking time, the single largest share in the delivery profile, while elevated-positive delivery (6.6%) and assertive delivery (1.2%) appear during portions of the flagged material. Read together, the denials on the SEC funding and bankruptcy-risk topics are delivered with raised arousal or emphasis rather than flat composure, which is part of the pattern the scores already establish and does not soften it.

05

The majority of sessions and topics remain calm and unremarkable

Twenty of twenty individually summarized sessions score LOW or MODERATE, sessions such as the USC Commencement Speech (10) and the Lex Fridman Neuralink conversation (15) sit well below the review cut, and topics like Cybercab regulatory approval (typical 56, zero sentences at cut) and billionaire criticism from left (typical 56, zero at cut) show no flagged material. This is a calm baseline at the subject's own level, not an absence of data.

Predictive Assessment

§ 05

SCENARIO	PROBABILITY	BEHAVIORAL BASIS
Future public statements on the SEC funding-secured episode are likely to reproduce elevated sentence CapoScores if the topic recurs.	HIGH	Eight of fourteen scored sentences on this topic already cluster at or above the review cut with a peak of 86, well outside the subject's typical of 18.
Discussion of Tesla's bankruptcy risk and short-seller conflict will continue to produce topic-specific spikes rather than broad elevation.	MODERATE	The topic shows a concentrated cluster (8 of 18 sentences at or above cut, peak 79) that is isolated from the subject's otherwise LOW-band delivery elsewhere.
General public appearances outside these flagged topics will continue to register as LOW-to-MODERATE.	HIGH	Twenty analysed sessions score LOW or MODERATE individually and topic consistency of 13 points indicates a narrow, stable spread around the subject's typical of 18.
The sex change surgery risks topic could produce further elevated exchanges if raised again in interviews.	MODERATE	Four of ten scored sentences already sit at or above the review cut with a peak of 78, a smaller but present cluster distinct from the subject's calm baseline.

CapoScore Methodology

Analysis is based on facial recognition, heart rate estimation, emotional evaluation, voice stress analysis, and linguistic pattern detection applied to 44 analysed videos. The CapoScore is the behaviour-adjusted score of the subject's own scored sentences on an absolute 0-100 scale, where higher values indicate greater concern; a video with no scored sentences is not scored at all, and this baseline of 18 was measured across all 44 of the videos in this brief, which carry sentence coverage. Behavioral indicators, including the delivery-state distribution, are probabilistic signals describing patterns in delivery and content and do not constitute deterministic conclusions about intent or truthfulness. Every score on this page is the same measurement: the behaviour-adjusted **CapoScore** of the subject's own scored sentences, on an absolute 0-100 scale — Low (<40), Moderate (<60), Elevated (<75), High (75+), with a review cut at 65. The baseline is the median of those sentences and describes the settled state; the peak is the single highest sentence and is reported separately because a median cannot show it.

What counts as scored. A video is scored only where the score ran for that speaker. 44 of 44 of your analysed video(s) carry sentence coverage. The score itself is measured across every session of this subject on the platform; the sessions listed here are your own. An unscored session shows “not measured yet” and never a number: unmeasured is not low, and the two must not be read alike.