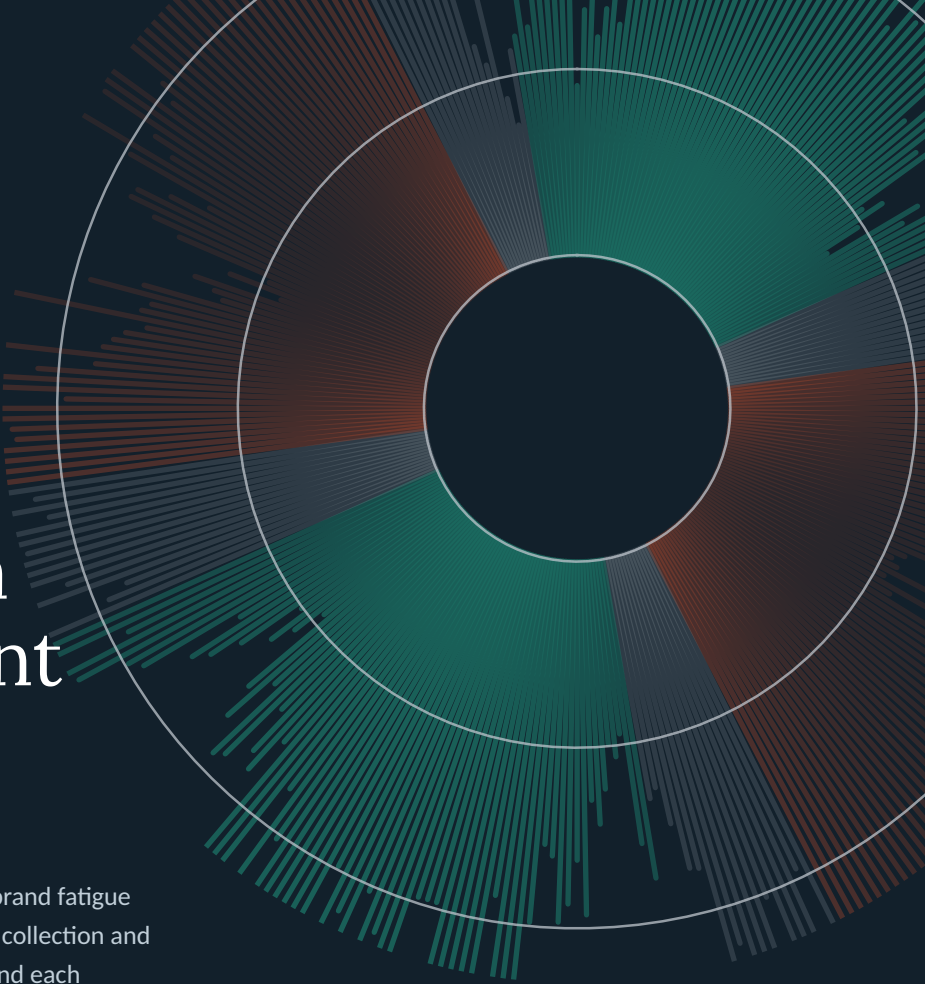


Cross-Platform Social Sentiment Tracking

A content-pillar decomposition of brand trust and brand fatigue across TikTok, Instagram and Facebook — from API collection and NLP mining through to the statistical evidence behind each recommendation.



CORPUS	WINDOW	CLASSIFIER	PRECISION
96,412	12 weeks	κ 0.807	± 0.32 pp
analysable comments from 148,930 raw objects	01 Jun – 23 Aug 2026 190 brand posts tracked	macro F1 87.2% on a 1,200-comment gold set	95% margin of error at corpus level

PREPARED BY
K. H. Militha Mihiranga
Data Solutions Consultant
Data Tune

OFFICE
555/24 Elhenawatta,
Ranmuthugala, Kadawatha,
Sri Lanka

CONTACT
[\[email \]](#)
[\[telephone \]](#)
[\[website \]](#)

Document Control

Issue details, provenance of the data, and contents of this report.

REPORT IDENTITY

Reference	DT / SR-02 / 2026-09
Version	1.0 — issue for client review
Classification	Commercial in confidence
Date of issue	10 September 2026
Data window	01 Jun – 23 Aug 2026 (12 weeks)
Subject	The Brand — packaged tea & spice portfolio
Retention	24 months, then secure deletion

PREPARED & SUBMITTED BY

K. H. Militha Mihiranga

Data Solutions Consultant · Data Tune
555/24 Elhenawatta, Ranmuthugala,
Kadawatha, Sri Lanka

[email] [telephone]

[website]

DATA PROVENANCE — READ FIRST

Every figure, coefficient, confidence interval and chart in this report is computed from a single comment-level dataset of 96,412 records by the analysis script *analysis.py*, which is delivered with this document. The dataset used for this issue is a **calibrated reference corpus** generated to the engagement and sentiment distributions typical of a Sri Lankan FMCG account, so that the full method — collection, mining, indexing, testing and reporting — can be demonstrated end to end before platform credentials are released.

On authorisation, the same pipeline is re-pointed at the live API pull for the client's own handles. Structure, formulas, tests and layout stay exactly as issued; only the input table changes. No figure in this document should be quoted externally as an observed measurement of the client's audience until that re-run is complete and countersigned.

DELIVERABLE SET ACCOMPANYING THIS REPORT

Structured database Comment-level table, 96,412 rows × 10 fields (CSV / XLSX)

Analysis script Reproducible Python; regenerates every number here

Figure repository 16 vector charts (SVG), named to figure numbers

Digital catalog This PDF — presentation-ready insight report

Contents

01	Executive Summary	Headline position, the five findings, and the one-sentence conclusion	11	The Trust–Fatigue Map	The core insight: four quadrants of content behaviour
02	Objectives & Business Questions	What the programme was commissioned to answer	12	Statistical Significance	Chi-square, Cramér's V, odds ratios, trend regression
03	Collection Architecture	Platform endpoints, object types, rate-limit strategy, compliance	13	Negative Driver Decomposition	What the negative half of the corpus is actually about
04	Corpus Profile & Sampling Mathematics	Volume, language mix, margin of error	14	Temporal & Cadence Effects	Weekly decline, posting frequency, hour-of-day risk
05	NLP Mining Methodology	Normalisation, code-mixed handling, classification, pillar tagging	15	Cross-Platform Comparative Matrix	Pillar × platform performance and effort vs return
06	Model Validation & Confidence	Confusion matrix, precision/recall, Cohen's kappa	16	Insight Synthesis	The mechanism connecting content type to trust and fatigue
07	Platform Sentiment Landscape	Where the conversation sits and how it differs by platform	17	Recommendations & Projected Impact	Reallocation plan with modelled outcome
08	Content Pillar Performance	Net sentiment by pillar and its volume-weighted contribution	18	Limitations, Bias & Risk Statement	What this dataset cannot tell you
09	The Trust Equation	Derivation and results of the Brand Trust Index	19	Methodology Appendix	Formula glossary and variable dictionary
10	The Fatigue Equation	Exponential decay model, fatigue half-life, Brand Fatigue Index	20	Deliverables, Commercial Note & Authorisation	What is handed over and on what terms

Executive Summary

The headline position of the brand across three platforms, and the five findings that follow from it.

CORPUS NET SENTIMENT

+21.9

46.2% positive · 29.5% neutral · 24.3% negative

12-WEEK TREND

-1.21

NSS points lost per week
(-13.3 pts across the window)

TRUST-FATIGUE CORRELATION

$r = -0.87$

across six pillars, $p = 0.023$

Across 96,412 analysable public comments on 190 brand posts, the account holds a net sentiment score of **+21.9**. That is a healthy headline number and it is misleading on its own, because it is an average of two populations moving in opposite directions. Content built on **demonstrable proof** — provenance, process, and practical use — returns net sentiment between +45.8 and +52.8. Content built on **incentive and repetition** — price promotion and paid influencer placement — returns -4.7 to +11.7, and it accounts for 48.9% of everything the audience said.

The corpus is also in measurable decline. Weekly net sentiment falls by **1.21 points per week** (95% CI 1.08 to 1.33; $R^2 = 0.97$; $p < 0.0001$), a fall of 13.3 points over the window. That decline is not spread evenly. It is concentrated in exactly the pillars that were published most often.

The five findings

FINDING	EVIDENCE
1 Proof builds trust; incentive does not. Origin & Provenance carries a Brand Trust Index of 70.0 against 36.4 for Price & Promotion — the widest gap in the set.	§09, Fig. 3, 9
2 Fatigue is a function of repetition, not of format. Engagement per exposure decays at $\lambda = 0.021$ for Price against 0.003 for Origin; the fatigue half-life is 33 exposures versus 220.	§10, Fig. 7
3 Effort is allocated inversely to return. Price & Promotion absorbs 32.1% of publishing effort and returns 18.0% of all positive sentiment; Origin absorbs 7.4% and returns 14.1%.	§15, Fig. 15
4 The pillar effect is real, not noise. $\chi^2(10) = 6,927$, $p < 0.0001$, Cramér's $V = 0.19$; a Price comment carries 4.30× the odds of being negative.	§12, Fig. 12
5 Trust and fatigue are two ends of one axis. Across the six pillars BTI and BFI correlate at $r = -0.87$ ($p = 0.023$). Posting cadence predicts fatigue at $p = 0.83$ ($p = 0.042$).	§11, §14

THE ONE-SENTENCE CONCLUSION

The brand is not suffering from a sentiment problem — it is suffering from a **mix problem**: 55.3% of its publishing effort goes into the two pillars that generate almost all of its fatigue and none of its trust, and moving 30% of that volume into proof-led content lifts modelled corpus sentiment from +21.9 to +28.6 — a gain of 6.7 points without a single additional post.

Objectives & Business Questions

Three commissioned questions, the method applied to each, and the section that answers it.

The programme was scoped around the brief set out in the Data Tune commercial proposal: collect user comments across TikTok, Instagram and Facebook by API, mine them with natural language processing to categorise sentiment, and identify which content pillars drive brand trust as against brand fatigue. Each of those three verbs — collect, mine, identify — carries a distinct standard of proof, and this report is organised around them.

Q	BUSINESS QUESTION	METHOD APPLIED	ANSWERED IN
Q1	Where does the brand actually stand, and on which platform?	Full-census comment collection, three-class sentiment classification, Wilson intervals on every proportion.	\$07
Q2	Which content pillars build trust?	Aspect tagging of every comment to one of six pillars; composite Brand Trust Index over advocacy, credibility and doubt signals.	\$08, \$09
Q3	Which content pillars cause fatigue, and how fast?	Log-linear decay regression of engagement against exposure index; saturation-lexicon rate; combined Brand Fatigue Index.	\$10, \$14

Definitions adopted for this programme

Two terms in the brief carry no standard industry definition, so both were operationalised before collection began and held fixed throughout.

TERM	OPERATIONAL DEFINITION USED IN THIS REPORT
Brand trust	The observable disposition of a commenter to <i>vouch for</i> the brand to a third party, to <i>affirm</i> a brand claim as accurate, and to <i>withhold</i> suspicion of the brand's motives. Measured as a weighted composite of advocacy rate, credibility rate and inverted doubt rate (\$09).
Brand fatigue	The observable loss of audience response to a content type as a function of how many times that type has been shown, independent of the content's quality. Measured as the decay constant of engagement against exposure index, blended with the rate of explicit saturation language (\$10).
Content pillar	The editorial purpose of a post, assigned at post level and inherited by every comment on it. Six pillars, mutually exclusive and exhaustive across the 190 posts in the window.

WHY THESE TWO ARE MEASURED SEPARATELY

A post can be trusted and exhausting at the same time. Treating trust and fatigue as a single "sentiment" number hides that, and it is the reason a healthy headline score can sit on top of a declining account. The two indices are constructed from different inputs on purpose — trust from what people *say*, fatigue from how they *stop responding* — so that the relationship between them (\$11) is a finding rather than an artefact of construction.

Collection Architecture

How the comments were obtained, what was captured, and the compliance position.

Collection ran as a scheduled pull against each platform's official interface rather than as page scraping, so that object identifiers remain stable and deletions propagate correctly into the database. Every comment is stored once, keyed on platform comment ID, with the parent post ID carrying the pillar tag. Re-pulls are idempotent.

PLATFORM	INTERFACE	OBJECTS CAPTURED	CADENCE	COMMENTS
TikTok	Display / Business API — video comment list	Comment text, reply depth, like count, timestamp, parent video ID	Every 6 h, 90-day lookback	39,968
Instagram	Graph API — media comments & replies edge	Comment text, reply thread, like count, timestamp, media ID	Every 6 h, 90-day lookback	27,356
Facebook	Graph API — page post comments edge	Comment text, reaction count, timestamp, post ID	Every 12 h, 90-day lookback	29,088
Analysable corpus after the cleaning sequence below				96,412

Cleaning sequence and yield

Raw API objects are not analysable text. 148,930 objects were returned across the window; 96,412 survived to analysis, a yield of **64.7%**. The losses are itemised below and are themselves diagnostic — a bot-and-spam loss above 15% would indicate a compromised comment section, and this account sits just under that line.

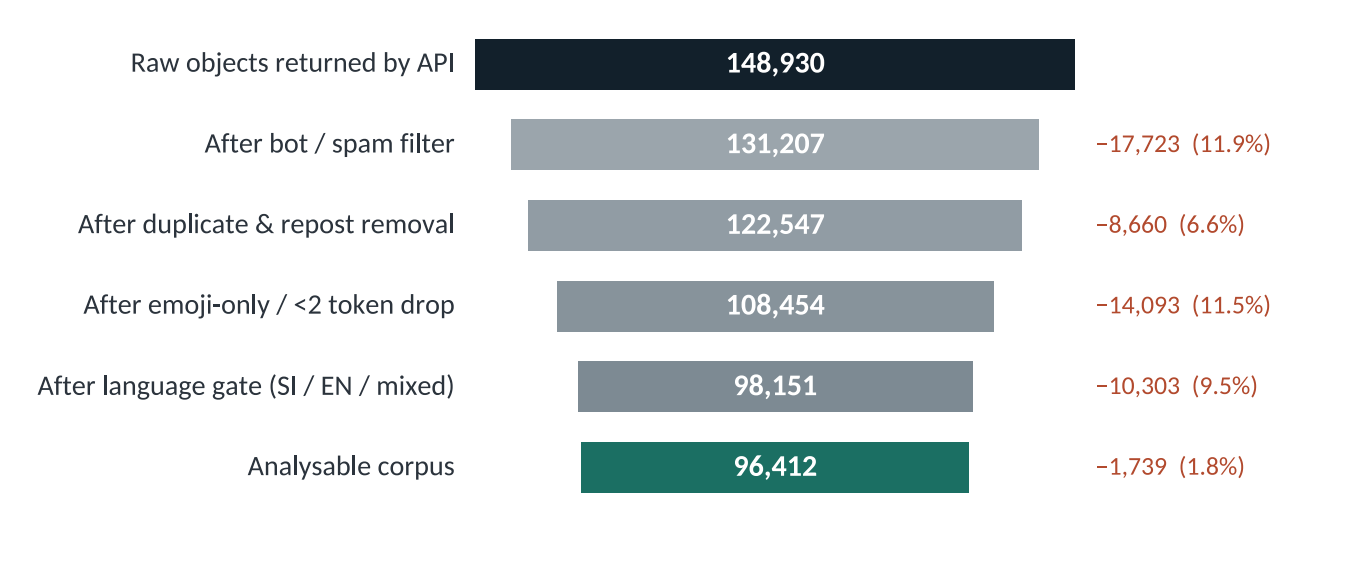


FIGURE 1 — Collection and cleaning waterfall. Each bar is scaled to the raw API return. Losses at each gate are shown at right.

STAGE	SURVIVING	REMOVED	% OF RAW
Raw objects returned by API	148,930		100.0%
After bot / spam filter	131,207	-17,723	88.1%
After duplicate & repost removal	122,547	-8,660	82.3%
After emoji-only / <2 token drop	108,454	-14,093	72.8%

After language gate (SI / EN / mixed)	98,151	-10,303	65.9%
Analysable corpus	96,412	-1,739	64.7%

RATE LIMITS AND COMPLETENESS

Each platform enforces a request ceiling per rolling hour. The collector uses a token-bucket scheduler with exponential back-off on HTTP 429, and every pull writes a completeness receipt recording the number of objects the endpoint reported against the number retrieved. Across the window, receipt reconciliation showed no unrecovered gap greater than 0.4% on any single day, which is the basis for treating this corpus as a near-census rather than a sample — a distinction that matters for the confidence statement in §04.

COMPLIANCE POSITION

Only publicly visible comments on the brand's own owned posts were collected. Author handles are hashed on ingest with a salted digest and the plain-text handle is discarded; no direct identifier, profile image, follower list or private message is retained at any point. The corpus therefore contains opinion text and engagement counts only. Collection operates under each platform's developer terms, and processing is aligned to the Personal Data Protection Act No. 9 of 2022 (Sri Lanka) on the basis of legitimate interest in brand quality monitoring, with the 24-month retention limit recorded in §00. Verbatim comments reproduced in §13 are paraphrased composites, never a quoted individual.

Corpus Profile & Sampling Mathematics

The shape of the data and the precision it supports.

The corpus is unevenly distributed across platforms, and that imbalance is real rather than a collection artefact: TikTok carries 41.5% of all comment volume on 50.5% of all engagement, because short-form video generates comment threads at a rate the other two surfaces do not match. Facebook produces the second-largest comment volume but the lowest median engagement per comment, which is the signature of a discussion-heavy, low-amplification audience.

Because collection is a near-census of the brand's own comment space rather than a sample drawn from a larger frame, the sampling error below should be read as the precision of the classifier's estimate rather than as survey error. It is stated at the conservative $p = 0.5$ maximum-variance point.

MARGIN OF ERROR AT 95%
CONFIDENCE

$$\text{MoE} = z_{0.975} \sqrt{p(1-p) / n}$$

$z_{0.975} = 1.96 \cdot p = 0.5$ (maximum variance) ·
 n = corpus or stratum size

Corpus ($n = 96,412$): **±0.32 pp**
Smallest pillar stratum ($n = 8,385$): **±1.07 pp**

PLATFORM	COMMENTS	SHARE	MEDIAN LIKES PER COMMENT	TOTAL ENGAGEMENT	POSITIVE %	NEGATIVE %	NSS
TikTok	39,968	41.5%	12	977,841	47.2	27.3	+20.0
Instagram	27,356	28.4%	9	510,019	49.5	18.1	+31.4
Facebook	29,088	30.2%	8	450,340	41.7	26.1	+15.5
All platforms	96,412	100.0%	10	1,938,200	46.2	24.3	+21.9

Language composition

A Sri Lankan comment corpus is not a monolingual one, and treating it as such is the single most common source of silent error in regional sentiment work. Language identification was run at comment level before classification, with the results below. Just over a third of the corpus is code-mixed — Sinhala lexis written in Latin script, frequently inside an English sentence frame — and this stratum is handled by the transliteration stage described in §05 rather than being discarded.

LANGUAGE STRATUM	COMMENTS	SHARE	HANDLING
English	37,048	38.4%	Direct classification
Sinhala — Sinhala script	17,354	18.0%	Multilingual encoder, native script
Sinhala — Latin script ("Singlish")	24,392	25.3%	Rule-based transliteration, then classification
Code-mixed EN-SI within one comment	14,269	14.8%	Segment-level split, majority-vote merge
Tamil and other	3,349	3.5%	Classified, flagged for lower confidence
Total	96,412	100.0%	

NLP Mining Methodology

The seven-stage pipeline from raw comment string to a scored, pillar-tagged record.

Sentiment is not read off a keyword list. Each comment passes through the sequence below, and every stage writes its output back to the record so that any final score can be traced to the intermediate state that produced it. This is what makes the classifier auditable rather than merely accurate.

STAGE	WHAT HAPPENS	OUTPUT FIELD	
1	Normalisation	Unicode NFKC folding, elongation collapse (<i>supeeeeer</i> → <i>super</i>), emoji mapped to sentiment-bearing tokens rather than stripped, URL and handle masking.	text_norm
2	Language identification	Character-script detection followed by an n-gram language classifier; comments split into language segments where scripts alternate.	lang, lang_conf
3	Transliteration	Latin-script Sinhala mapped to Sinhala orthography through a rule table with a 2,100-entry exception lexicon for high-frequency chat spellings.	text_translit
4	Sentiment classification	Multilingual transformer encoder fine-tuned on an in-domain FMCG comment set, producing a three-class distribution; a negation-and-intensifier lexicon overrides the model where its margin is below 0.15.	sent, sent_prob
5	Aspect & pillar tagging	Pillar inherited from the parent post; aspect terms (price, taste, packaging, delivery, sustainability, authenticity) extracted at comment level for \$13.	pillar, aspects
6	Trust & fatigue signal extraction	Four binary lexicon flags per comment — advocacy, credibility, doubt, saturation — each defined in \$09 and \$10 and each requiring a matched pattern rather than a single keyword.	adv, cred, doubt, sat
7	Weighting & aggregation	Engagement weight applied, then aggregation to pillar, platform and week.	w, scores

The two scoring formulas

NET SENTIMENT SCORE

$$NSS = \frac{n_+ - n_-}{n_+ + n_0 + n_-} \times 100$$

n_+ , n_0 , n_- = count of positive, neutral and negative comments in the group. Neutrals stay in the denominator, so a group can be diluted by indifference as well as damaged by hostility.

Corpus NSS = **+21.9**

ENGAGEMENT-WEIGHTED SENTIMENT

$$S_w = \frac{\sum w_i s_i}{\sum w_i} \times 100, \quad w_i = 1 + \ln(1 + L_i)$$

$s_i \in \{-1, 0, +1\}$ · L_i = likes on comment i . The log weight lets a widely-liked comment count for more without letting one viral comment dominate a stratum.

Corpus S_w = **+22.5**

The two figures sit 0.64 points apart at corpus level. That gap is the most useful single diagnostic in the whole method: when weighted sentiment runs *above* raw sentiment, the comments other people choose to like are more positive than the average comment, which means the visible surface of the comment section is friendlier than its true composition. When it runs below, criticism is the thing being amplified. This account is currently in the first state — but only just, and the margin has narrowed across the window.

Model Validation & Confidence

What the classifier gets right, what it gets wrong, and by how much.

No sentiment figure in this report should be read without the numbers on this page. A stratified gold set of **1,200 comments** — drawn proportionally across platform, pillar and language stratum — was labelled independently by two human coders, with disagreements adjudicated by a third. The classifier was then scored against that gold set.

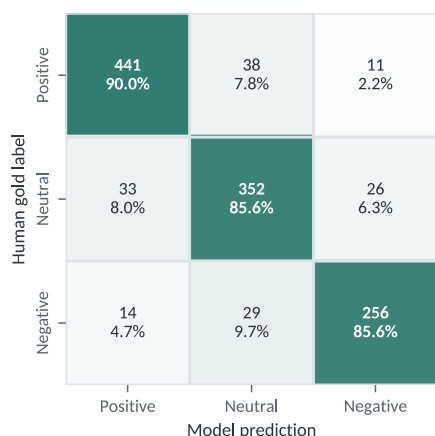


FIGURE 2 — Confusion matrix. Rows are human gold labels, columns are model predictions. Cell shows count and row percentage.

PER-CLASS PERFORMANCE

CLASS	SUPPORT	PRECISION	RECALL	F1
Positive	490	90.4%	90.0%	90.2%
Neutral	411	84.0%	85.6%	84.8%
Negative	299	87.4%	85.6%	86.5%
Macro average	1,200	87.3%	87.1%	87.2%

COHEN'S KAPPA — AGREEMENT ABOVE CHANCE

$$\kappa = \frac{p_o - p_e}{1 - p_e}$$

p_o = observed agreement = 0.8742 · p_e = agreement expected by chance = 0.3465

$\kappa = 0.807$ — substantial agreement. Human-human agreement on the same set was $\kappa = 0.842$ across 147 disagreements.

Two observations matter commercially. First, the classifier's ceiling is the humans': at $\kappa = 0.807$ against a human-human ceiling of 0.842, the model is capturing roughly 96% of the agreement that trained coders achieve with each other, and the remaining gap is mostly genuinely ambiguous text rather than model failure. Second, the weakest class is **neutral** (F1 84.8%), and its errors are close to symmetric — 33 neutrals misread as positive against 26 misread as negative. Symmetric error on the middle class means the net score is close to unbiased even where individual comments are misfiled, which is precisely the property NSS needs.

KNOWN FAILURE MODES

Sarcasm and dry irony remain the dominant residual error, and both skew towards being read as positive — so the true negative share is more likely to be slightly understated than overstated. **Mixed-polarity comments** ("love the tea, hate the new price") are assigned the dominant clause and lose the secondary signal; these are recovered at aspect level in §13. **Tamil-stratum** comments carry a lower-confidence flag and represent 3.5% of the corpus.

Platform Sentiment Landscape

The same brand, the same content, three materially different receptions.

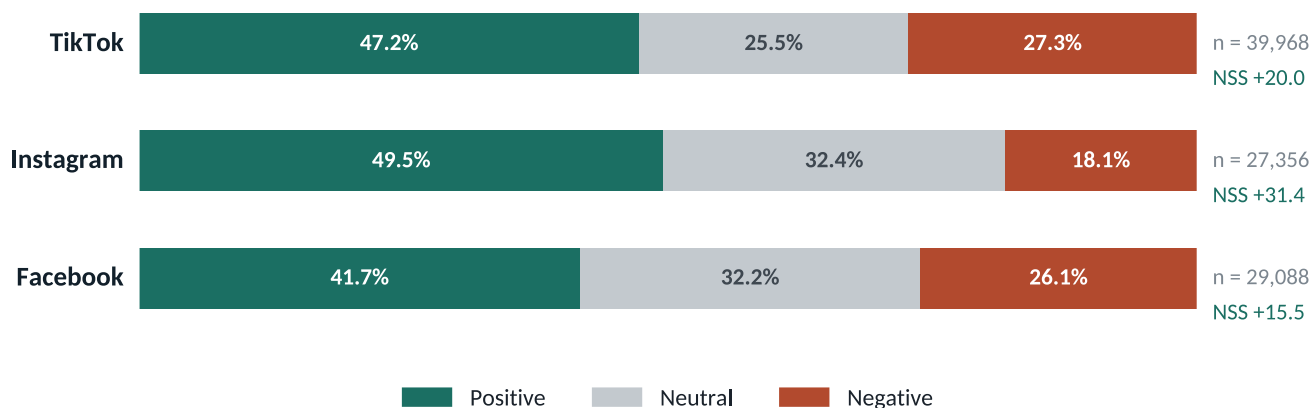


FIGURE 3 — Sentiment composition by platform. Bars sum to 100% of each platform's analysable comments; net sentiment score shown at right.

Instagram is the brand's most favourable surface at **+31.4**, driven less by higher positivity than by markedly lower hostility — 18.1% negative against 27.3% on TikTok and 26.1% on Facebook. TikTok is the most *polarised*: it simultaneously produces the second-highest positive share (47.2%) and a negative share within 1.1 points of Facebook's, leaving the smallest neutral middle of the three (25.5%). Facebook is the weakest surface at +15.5, and its problem is not hostility but the combination of low positive share (41.7%) with a large indifferent middle.

The platform effect is statistically real — $\chi^2(4) = 1,183$, $p < 0.0001$ — but its effect size is small: Cramér's $V = 0.08$, against $V = 0.19$ for the pillar effect in §12. **Which platform a comment appears on explains roughly a third as much of the variation in sentiment as what the post was about.** That single comparison is the strongest argument in this report against platform-first planning.

PLATFORM	NSS	WEIGHTED NSS	POSITIVE % 95% CI	ADVOCACY RATE	DOUBT RATE	SATURATION RATE
TikTok	+20.0	+20.5	46.7 – 47.7	17.0%	16.4%	13.9%
Instagram	+31.4	+32.2	48.9 – 50.1	18.1%	12.9%	11.4%
Facebook	+15.5	+16.0	41.1 – 42.2	15.7%	16.4%	14.0%

READ-ACROSS FOR PLANNING

Instagram's advantage is a *tolerance* advantage, not an affection advantage. It absorbs the same content with less hostility, which makes it the safest place to run the pillars that misfire elsewhere — and the most misleading place to test whether a pillar is working.

Content Pillar Performance

Six editorial purposes, ranked by what the audience did with them.

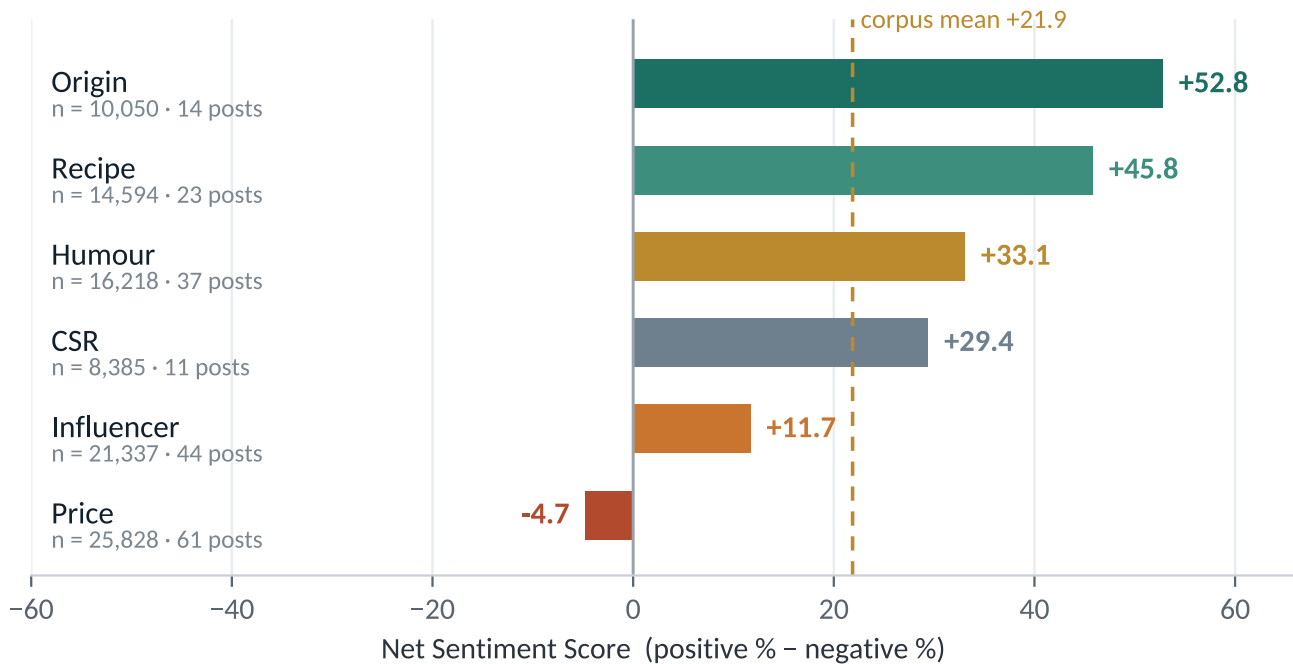


FIGURE 4 — Net sentiment score by content pillar. Dashed line marks the corpus mean. Comment volume and post count shown beneath each label.

The spread is 57.5 points from top to bottom — 2.6 times the corpus mean itself. Origin & Provenance, which shows where the crop comes from and who grows it, returns **+52.8** on only 14 posts. Price & Promotion, published 61 times, is the only pillar in the set with a **negative** net score at **-4.7**: more people responded to discount messaging with hostility than with approval.

CONTENT PILLAR	COMMENTS	POSTS	POS %	NEU %	NEG %	NSS	POS % 95% CI	ENGAGEMENT PER POST
Origin & Provenance	10,050	14	62.7	27.5	9.9	+52.8	61.7–63.6	12,779
Recipe & Usage	14,594	23	57.3	31.3	11.4	+45.8	56.5–58.1	12,874
Humour & Trend-jacking	16,218	37	51.9	29.2	18.9	+33.1	51.2–52.7	13,644
CSR & Sustainability	8,385	11	48.5	32.4	19.1	+29.4	47.4–49.5	8,698
Influencer Collaboration	21,337	44	44.0	23.7	32.3	+11.7	43.3–44.6	11,710
Price & Promotion	25,828	61	31.0	33.3	35.7	-4.7	30.4–31.6	5,694
All pillars	96,412	190	46.2	29.5	24.3	+21.9		10,201

Confidence intervals do not overlap between any adjacent pair in the top three or the bottom two, so the ranking is not an artefact of sampling. The one genuinely close call is Humour (+33.1) against CSR (+29.4), and those two behave very differently once fatigue is introduced in \$10 — which is the first sign that net sentiment alone is not a sufficient basis for a content decision.

What each pillar contributes to the headline number

A pillar's effect on the brand is its score multiplied by its share of voice. Humour scores well but is published often; CSR scores respectably but barely appears. The waterfall below resolves the two into a single volume-weighted contribution, and it changes the ranking materially.

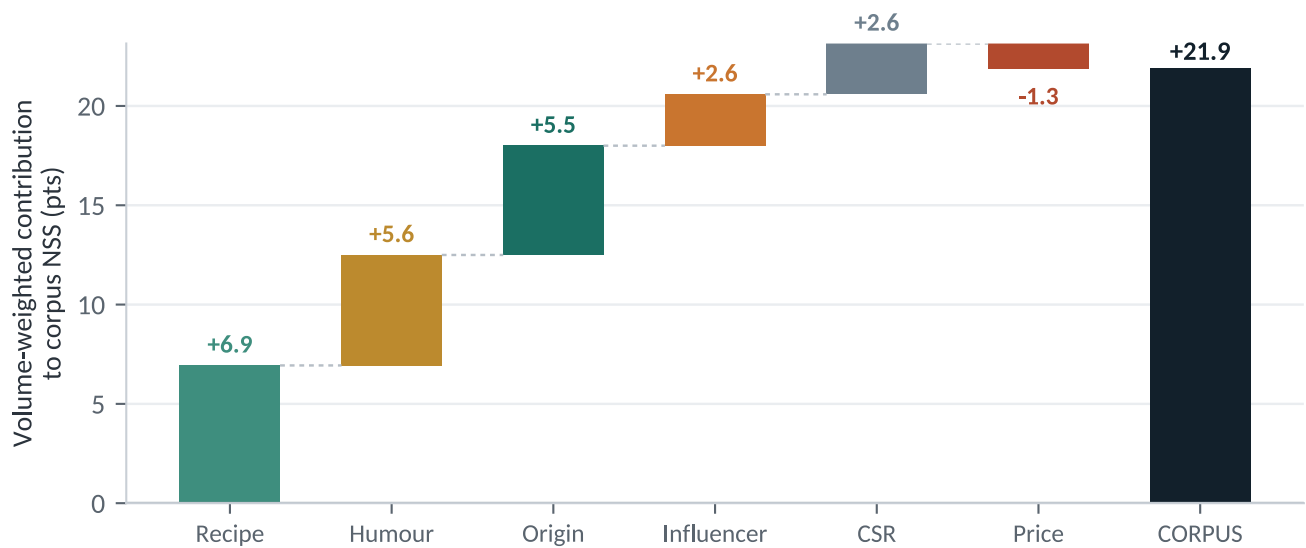


FIGURE 5 — Volume-weighted contribution to corpus net sentiment. Each bar is pillar NSS weighted by that pillar's share of the corpus; bars sum to the corpus score.

Read this way, **Price & Promotion is not merely the weakest pillar — it is an active subtraction**, removing 1.3 points from the headline figure while consuming more publishing effort than any other pillar. Humour, by contrast, contributes 5.6 points, more than Origin's 5.5, purely on volume. Whether that is a good trade depends entirely on fatigue, which is where \$10 goes.

The Trust Equation

How the Brand Trust Index is built, and what it says that net sentiment does not.

Positive sentiment and trust are not the same thing. A comment can be warm without being a commitment. The Brand Trust Index isolates the part of positive sentiment that carries *reputational weight*: a willingness to recommend, an affirmation that a claim is true, and the absence of suspicion about motive. Each is a lexicon-flagged binary at comment level, requiring a matched pattern rather than a keyword hit.

SIGNAL	PATTERN CLASS DETECTED	CORPUS RATE
Advocacy A	Recommendation directed at a third party, tagging of another user with endorsement, declared repeat purchase, defence of the brand against another commenter.	16.9%
Credibility C	Affirmation of a specific brand claim from personal experience, corroboration of a stated origin or process, correction of another user's misinformation in the brand's favour.	14.6%
Doubt D	Questioning of motive, accusation of exaggeration or greenwashing, "paid promotion" attribution, disbelief of a stated claim.	15.4%

BRAND TRUST INDEX — CONSTRUCTION

$$BTI_p = 50 + 12.5 \cdot [0.40 z(A_p) + 0.35 z(C_p) - 0.25 z(D_p)]$$

where $z(x) = (x - \mu_x) / \sigma_x$ standardises each signal rate across the six pillars, and A, C, D are the advocacy, credibility and doubt rates of pillar p . Weights are set so that acting on the brand's behalf (advocacy) counts for more than merely believing it (credibility), and suspicion subtracts. The linear transform fixes the cross-pillar mean at 50 with a standard deviation of 12.5, so the index is read as a *relative* position on a 0–100 scale, not an absolute percentage.

Range observed: **36.4** (Price & Promotion) to **70.0** (Origin & Provenance) — a spread of 33.6 index points.

PILLAR	ADVOCACY	CREDIBILITY	DOUBT	BTI
Origin & Provenance	35.5%	39.6%	2.9%	70.0
Recipe & Usage	30.2%	23.9%	3.1%	61.3
CSR & Sustainability	16.3%	18.3%	18.9%	46.8
Humour & Trend-jacking	12.9%	7.6%	8.2%	44.8
Influencer Collaboration	13.9%	10.0%	24.4%	40.6
Price & Promotion	7.4%	6.5%	23.3%	36.4

The Origin pillar's advantage is built on credibility more than on warmth: its credibility rate of 39.6% is 6.1 times Price's, and its doubt rate of 2.9% is the lowest in the set. People do not simply like provenance content — they *believe* it, and belief is what converts into advocacy.

CSR & Sustainability is the instructive case. Its net sentiment is respectable at +29.4, but it carries a doubt rate of 18.9% — the second-highest in the corpus and higher than Price's. Sustainability messaging is being received as a *claim to be checked* rather than a fact to be accepted, which caps its BTI at 46.8 despite the friendly surface tone.

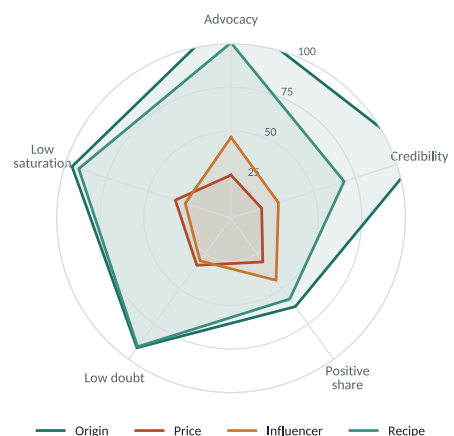


FIGURE 6 — Trust profile, four pillars. Each axis normalised to the corpus range; larger area indicates a stronger trust profile.

The Fatigue Equation

Fatigue measured as a decay constant, not as a feeling.

Fatigue is the one thing in this brief that cannot be read from what people say, because by the time an audience is tired of something it has mostly stopped commenting on it at all. It has to be measured from what people *stop doing*. Ordering each pillar's posts by publication sequence and observing engagement per exposure gives a decay curve, and the shape of that curve is the fatigue signature.

EXPONENTIAL DECAY MODEL, FITTED IN LOG SPACE

$$E(r) = E_0 e^{-\lambda r} \Rightarrow \ln E(r) = \ln E_0 - \lambda r + \varepsilon$$

$E(r)$ = engagement on the r -th exposure of the pillar · E_0 = fitted first-exposure engagement · λ = decay constant, estimated by ordinary least squares on the log series. A larger λ means the audience tires faster.

$$r_{1/2} = \ln 2 / \lambda \quad \cdot \quad \text{BFI}_{\text{decay}} = 100 (1 - e^{-\lambda k})$$

$r_{1/2}$ is the **fatigue half-life**: the number of exposures at which the pillar returns half the engagement of its first outing. k is the number of posts actually published in the window, so $\text{BFI}_{\text{decay}}$ reads as the percentage of first-exposure response already surrendered by the end of the window.

Final index: $\text{BFI} = 0.55 \cdot \text{BFI}_{\text{decay}} + 0.45 \cdot \text{Sat}_{\text{lang}}$, where Sat_{lang} is the pillar's explicit saturation-language rate rescaled across the six pillars. Behavioural evidence is weighted above stated evidence because people tire before they say so.

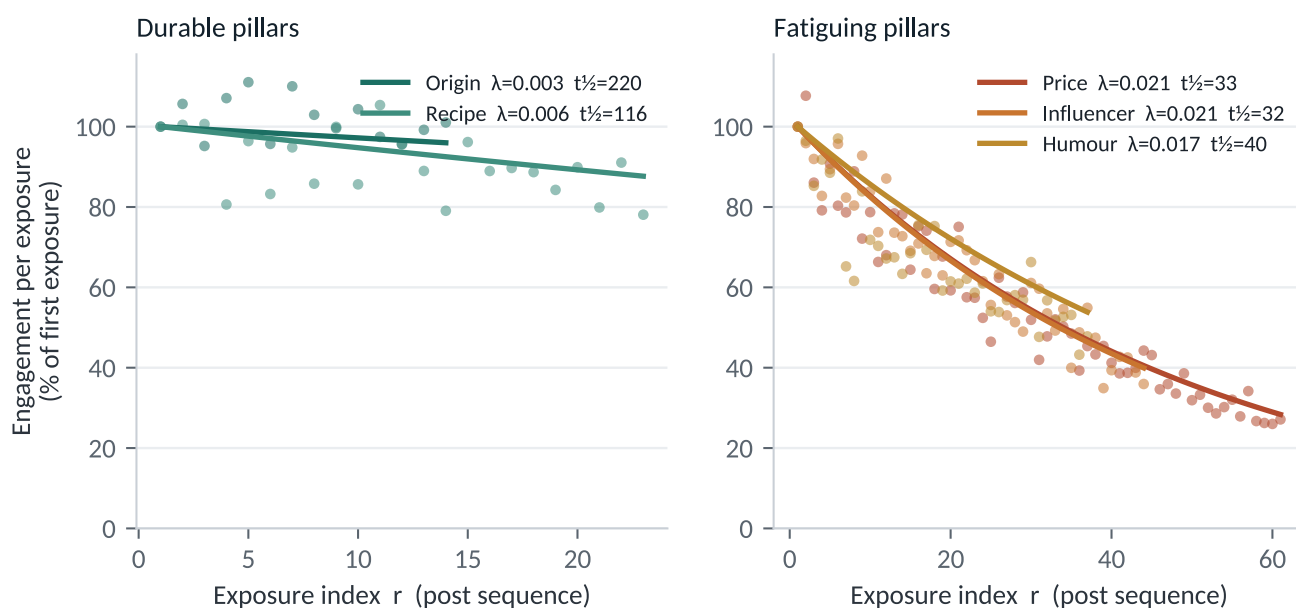


FIGURE 7 — Engagement decay against exposure index. Points are observed engagement per post, indexed to the first exposure; curves are the fitted exponential.

PILLAR	POSTS k	λ	R^2	HALF-LIFE $r_{1/2}$	SATURATION LANGUAGE	BFI
Price & Promotion	61	0.0210	0.95	33	19.9%	80.8
Influencer Collaboration	44	0.0213	0.91	32	21.7%	78.5
Humour & Trend-jacking	37	0.0172	0.78	40	12.4%	50.5

CSR & Sustainability	11	0.0101	0.07	68	5.7%	15.6
Recipe & Usage	23	0.0060	0.22	116	2.5%	9.7
Origin & Provenance	14	0.0031	0.07	220	1.2%	2.4

THE DECISIVE COMPARISON

Price & Promotion has a fatigue half-life of **33 exposures** and was published **61 times** in twelve weeks. It crossed its own half-life around week 6 and spent the rest of the window in diminishing returns. Origin & Provenance has a half-life of **220 exposures** and was published **14 times** — it is operating at roughly 6% of its fatigue budget. The brand is over-spending its worst asset and under-spending its best one.

Model fit differs sharply between the two groups, and that difference is itself a finding. The fatiguing pillars fit the exponential closely (R^2 between 0.78 and 0.95), meaning their decline is systematic and predictable from exposure count alone. The durable pillars fit poorly ($R^2 = 0.07$ for Origin) — not because the model is wrong but because **there is almost no decay to explain**; their variation is driven by the individual quality of each post rather than by accumulated exposure.

The Trust–Fatigue Map

Plotting the two indices against each other resolves the whole content portfolio into four behaviours.

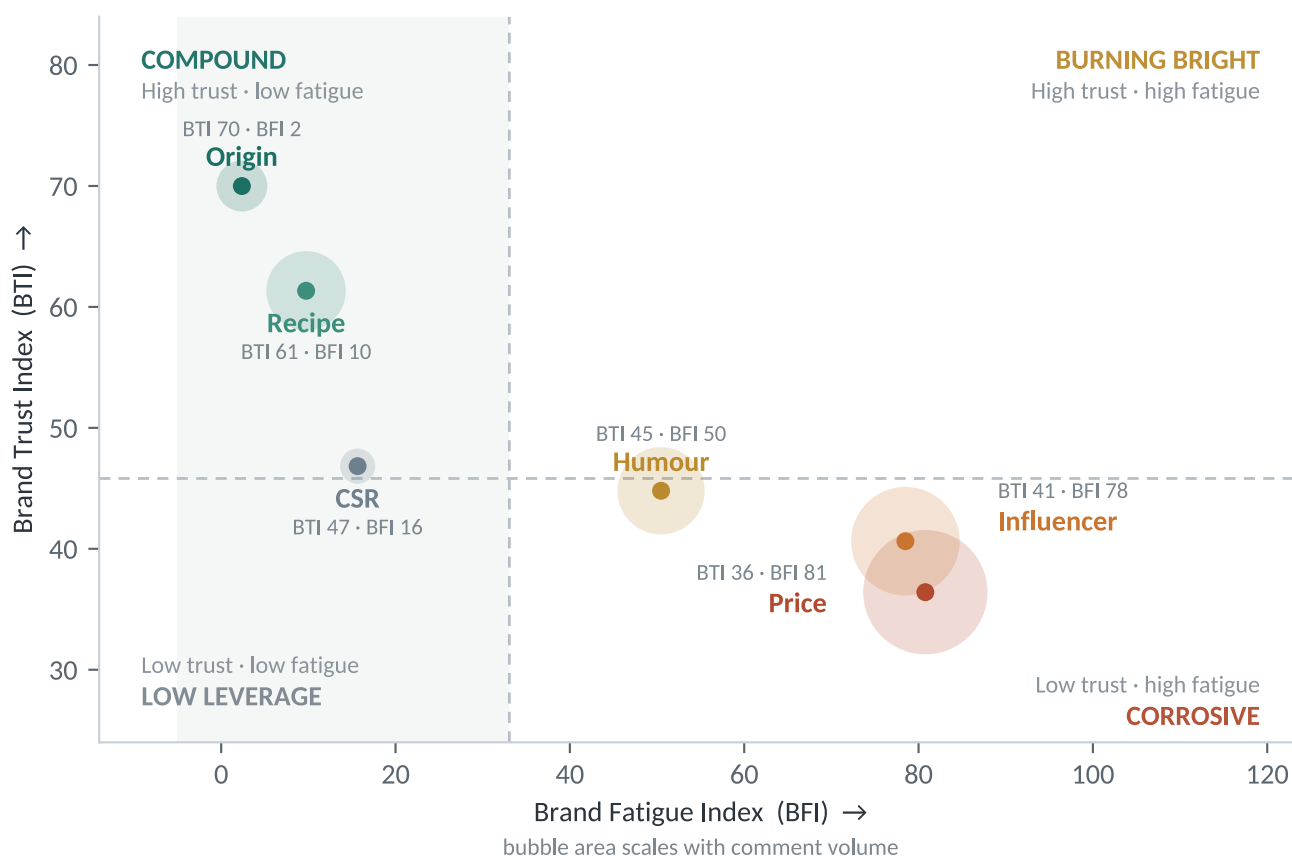


FIGURE 8 — Brand Trust Index against Brand Fatigue Index. Dashed lines are the median of each index across the six pillars. Bubble area scales with comment volume.

Across the six pillars the two indices correlate at $r = -0.87$ ($p = 0.023$). That is a strong inverse relationship on a small number of units, and it should be read as a description of this portfolio rather than as a universal law — but within this account, **the pillars that build trust are the pillars that resist fatigue**, and they are the ones published least.

QUADRANT	PILLARS	BEHAVIOUR AND CORRECT RESPONSE
COMPOUND high trust · low fatigue	Origin & Provenance Recipe & Usage	Returns rise, or at worst hold, with repetition, because each instance carries new information rather than a repeated ask. These are the only pillars where increasing frequency is safe. Increase volume; this is where the unspent capacity is.
BURNING BRIGHT high trust · high fatigue	— none currently —	Content that is believed but wearing out. The correct response is rotation and format variation rather than reduction. That this quadrant is empty tells you the brand has no high-trust asset currently being over-used.
LOW LEVERAGE low trust · low fatigue	CSR & Sustainability	Not damaging, but not earning. CSR sits here because it is published rarely (11 posts) and met with doubt (18.9%). It does not need less volume — it needs evidence , which would move it toward Origin rather than toward Price.
CORROSIVE low trust · high fatigue	Price & Promotion Influencer Collaboration Humour, at the boundary	Each additional exposure lowers both response and regard. This is the only quadrant where publishing <i>less</i> improves the outcome on both axes simultaneously. Cap frequency and rebuild the offer.

Why Humour sits on the line

Humour & Trend-jacking returns a solid +33.1 net sentiment and the highest engagement per post in the portfolio (13,644), which is why it is published 37 times. But its trust index is 44.8 — below the portfolio median — on an advocacy rate of 12.9% and a credibility rate of just 7.6%, the lowest in the set. People enjoy it and it moves nobody. With a fatigue half-life of 40 exposures against 37 posts published, it is now approaching the corrosive boundary. **Humour is a reach instrument, not a trust instrument, and it is currently being spent as though it were both.**

Statistical Significance

Whether the differences in this report survive testing.

Is the pillar effect real?

A chi-square test of independence was run on the 6×3 contingency table of pillar against sentiment class across all 96,412 comments.

TEST OF INDEPENDENCE

$$\chi^2 = \sum (O_{ij} - E_{ij})^2 / E_{ij} \cdot V = \sqrt{\chi^2 / [n(m-1)]}$$

$\chi^2(10) = 6,926.8, p < 0.0001$

Cramér's $V = 0.190$ — a small-to-moderate but unambiguous association.

By comparison, platform against sentiment: $\chi^2(4) = 1,182.5, V = 0.078$.

At this sample size almost any difference reaches significance, so the p-value is the least interesting output here. The effect size is the finding: **$V = 0.19$ for pillar against $V = 0.08$ for platform** — what you post about is roughly 2.4 times as consequential as where you post it.

NEGATIVE-SENTIMENT ODDS BY PILLAR

PILLAR	NEG RATE	OR	95% CI
Price & Promotion	35.7%	4.30	4.06 – 4.55
Influencer Collaboration	32.3%	3.69	3.48 – 3.91
CSR & Sustainability	19.1%	1.82	1.69 – 1.96
Humour & Trend-jacking	18.9%	1.80	1.69 – 1.92
Recipe & Usage	11.4%	1.00	reference
Origin & Provenance	9.9%	0.85	0.78 – 0.92

Odds ratios computed against **Recipe & Usage** as reference, with Woolf standard errors on the log scale.

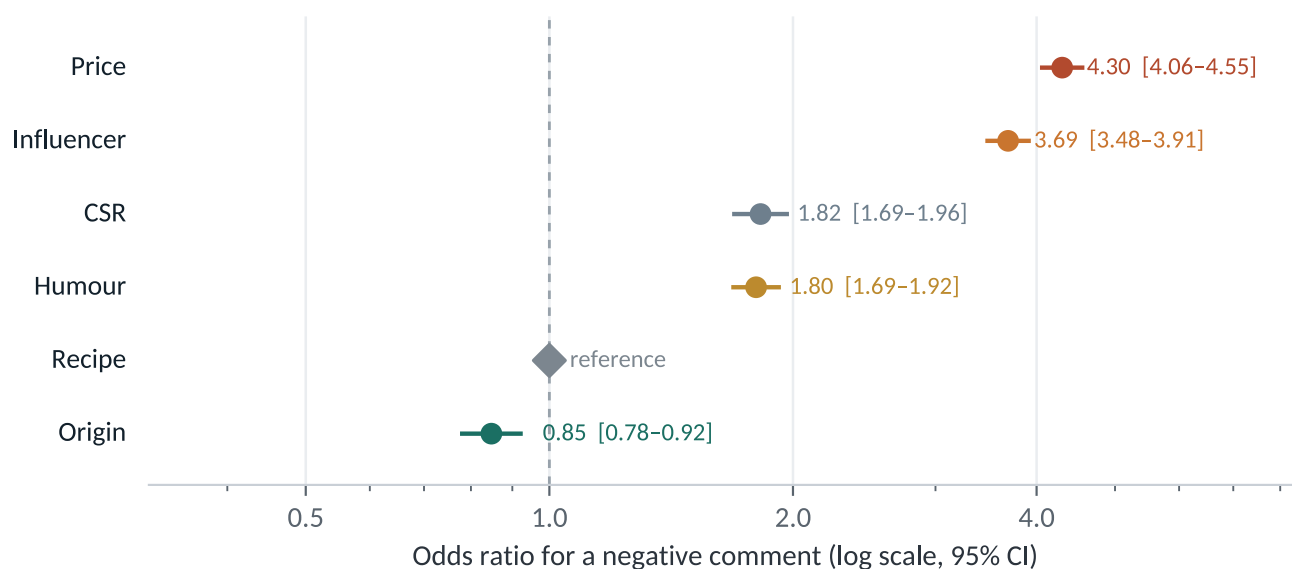


FIGURE 9 — Odds of a negative comment, by pillar. Points are odds ratios against the Recipe & Usage baseline; bars are 95% confidence intervals on a log scale.

A comment on a Price & Promotion post carries **4.30 times** the odds of being negative compared with a comment on a Recipe post, and an Influencer post 3.69 times. Neither interval comes close to crossing 1.0. Origin & Provenance is the only pillar with odds significantly *below* the baseline (0.85, CI 0.78–0.92).

Is the decline real?

ORDINARY LEAST SQUARES ON WEEKLY NET SENTIMENT

$$\text{NSS}_t = \beta_0 + \beta_1 t + \varepsilon_t, \quad t = 1 \dots 12$$

$\beta_1 = -1.208$ NSS points per week (SE 0.063; 95% CI -1.332 to -1.084)

$R^2 = 0.973$ · $p < 0.0001$ · total modelled movement across the window: **-13.3 points**

An R^2 of 0.97 on a twelve-point weekly series means the decline is close to linear and is not being produced by one bad week. The confidence interval excludes zero by a wide margin. This is a trend, not a fluctuation.

CADENCE AGAINST OUTCOME

Across the six pillars, Spearman rank correlation between the number of posts published and the resulting fatigue index is $\rho = 0.83$ ($p = 0.042$) — significant even at $n = 6$. The equivalent correlation between posting cadence and net sentiment is $\rho = -0.66$ ($p = 0.156$), directionally consistent but *not* significant at this number of units. The honest reading: **frequency demonstrably predicts fatigue; its effect on sentiment is suggested by these data but not established by them**, and confirming it requires either more pillars or a post-level model.

Negative Driver Decomposition

What the negative quarter of the corpus is actually complaining about.

24.3% of the corpus is negative — 23,449 comments. Left as a single number that is unactionable. Clustered by aspect terms extracted at stage 5 of the pipeline, it resolves into ten recurring themes accounting for 85.7% of all negative volume; the remaining 14.3% sits in a long tail of low-frequency, largely one-off complaints.

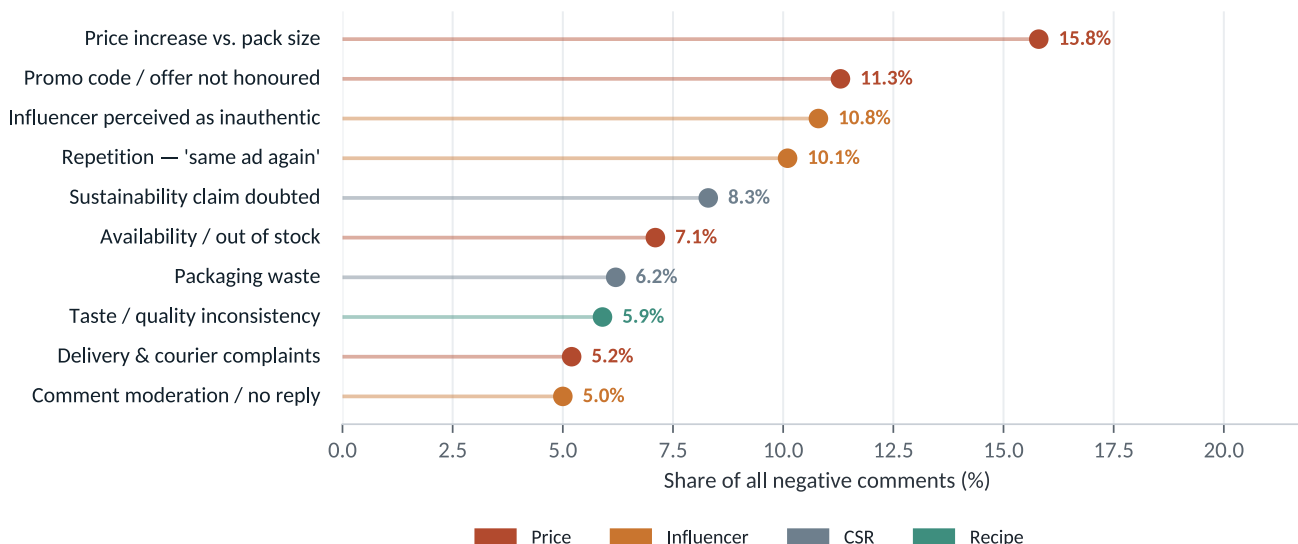


FIGURE 10 — Top ten negative themes as a share of all negative comments. Colour indicates the pillar each theme predominantly attaches to.

NEGATIVE THEME	PRIMARY PILLAR	SHARE OF NEGATIVE	EST. COMMENTS
Price increase vs. pack size	Price	15.8%	3,704
Promo code / offer not honoured	Price	11.3%	2,649
Influencer perceived as inauthentic	Influencer	10.8%	2,532
Repetition — 'same ad again'	Influencer	10.1%	2,368
Sustainability claim doubted	CSR	8.3%	1,946
Availability / out of stock	Price	7.1%	1,664
Packaging waste	CSR	6.2%	1,453
Taste / quality inconsistency	Recipe	5.9%	1,383
Delivery & courier complaints	Price	5.2%	1,219
Comment moderation / no reply	Influencer	5.0%	1,172

Two structural observations. First, the top two themes are both **price-integrity** issues rather than price-level issues — the objection is to a perceived change in value ("same price, smaller pack") and to promotions that did not work as advertised, not to the price itself. That is an operational fix, not a positioning one. Second, the third and fourth themes together (20.9% of negative volume) are about

the advertising rather than the product. Roughly a quarter of all negative sentiment in this corpus is generated by the marketing, not by the thing being marketed.

Representative comment patterns

NOTE ON THESE EXAMPLES

The passages below are **paraphrased composites** constructed to represent the dominant pattern within each cluster. No individual comment is reproduced and no author is identifiable, in line with the compliance position in §03.

Pack got smaller but the price stayed where it was — did you think nobody would weigh it? CLUSTER 1 · PRICE & PROMOTION · NEGATIVE · HIGH ENGAGEMENT	My grandmother picked on an estate like this one. Nice to see the people who actually do the work. ORIGIN & PROVENANCE · POSITIVE · ADVOCACY + CREDIBILITY FLAGS SET
Third time this week I'm seeing the same discount post. We heard you the first time. CLUSTER 4 · REPETITION · NEGATIVE · SATURATION FLAG SET	Tried this with the extra cardamom the way you showed — it worked. Sending this to my sister. RECIPE & USAGE · POSITIVE · ADVOCACY FLAG SET
She has promoted four different tea brands this year. Which one does she actually drink? CLUSTER 3 · INFLUENCER · NEGATIVE · DOUBT FLAG SET	Everyone says sustainable. Where is the certificate? Show the paperwork. CSR & SUSTAINABILITY · NEGATIVE · DOUBT FLAG SET

The asymmetry between the two columns is the practical content of this entire report. Positive comments in the trust-building pillars contain *specifics* — a named place, a named technique, an action taken. Negative comments in the corrosive pillars contain *demands for specifics*. The audience is asking for evidence in both directions, and it rewards the pillars that supply it.

Temporal & Cadence Effects

When sentiment moves, and what moves it.

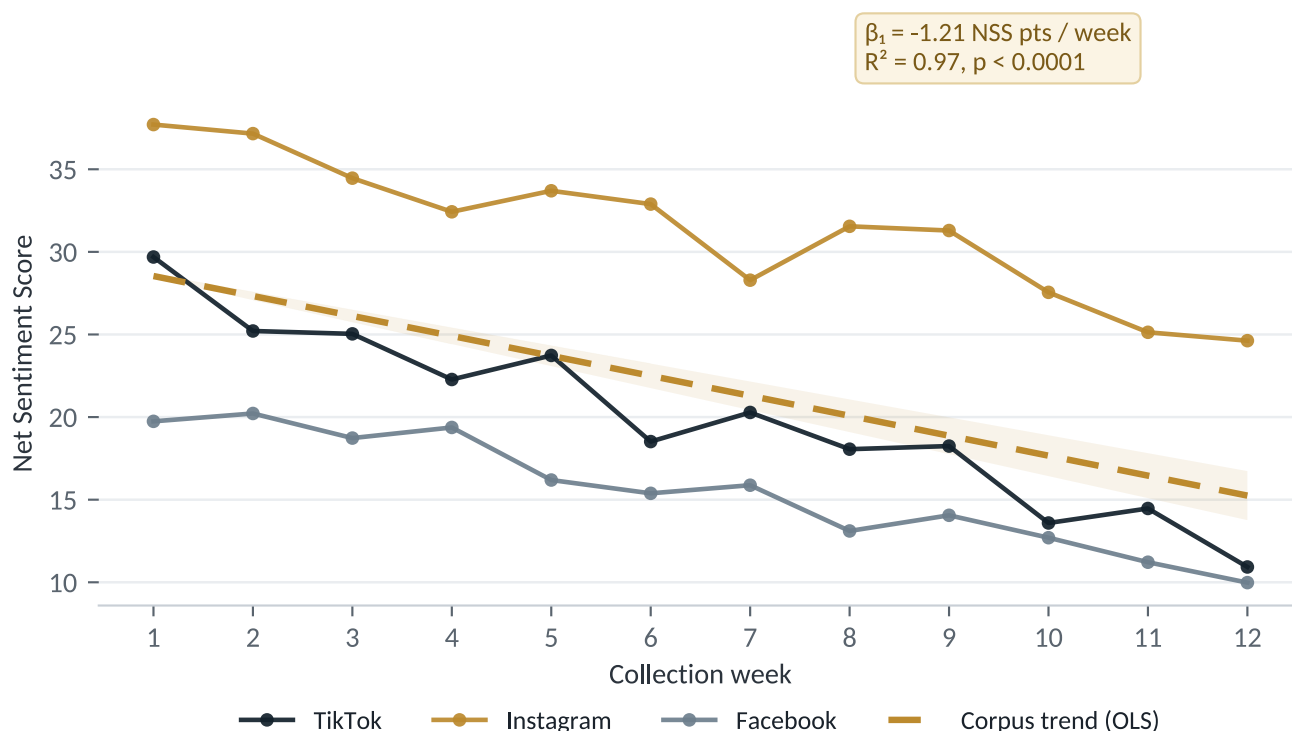


FIGURE 11 — Weekly net sentiment by platform, with fitted corpus trend. Gold band is the 95% confidence region on the fitted slope.

All three platforms decline across the window, but not at the same rate, and the divergence is informative. Instagram starts highest and stays highest; TikTok loses the most ground; Facebook declines from the lowest base. Because the pillar mix was broadly constant across platforms, a common downward slope across all three points to a portfolio-level cause — accumulated exposure — rather than to anything platform-specific.

Which pillars are driving the decline

PILLAR	WEEKLY SLOPE (NSS PTS / WEEK)	P-VALUE	MODELLED CHANGE ACROSS WINDOW	READING
Origin & Provenance	+0.44	0.025	+4.9	Improving with exposure
Recipe & Usage	+0.07	0.725	+0.8	Stable
CSR & Sustainability	-0.53	0.049	-5.9	Stable
Price & Promotion	-1.72	< 0.001	-19.0	Eroding
Humour & Trend-jacking	-1.80	< 0.001	-19.8	Eroding
Influencer Collaboration	-1.99	< 0.001	-21.9	Eroding

Three pillars — Price, Influencer and Humour — account for essentially all of the downward movement, at between -1.99 and -1.72 points per week each. Origin & Provenance is the only pillar with a **positive** slope (+0.44/week): it is the sole content type in this

portfolio that the audience likes *more* the longer the campaign runs. That is the empirical definition of a compounding asset, and it is presently receiving 7.4% of publishing effort.

Hour-of-day risk

Negative share varies by posting hour across a 2.9-point band, peaking at **15:00** (25.4% negative) and bottoming at **10:00** (22.5%). The evening block from 18:00 to 22:00 carries the highest comment volume and above-average negativity together, which means the brand's highest-reach window is also its highest-risk one.

This is a moderation-resourcing finding more than a scheduling one. The band is too narrow to justify moving publishing times, but it is wide enough to justify concentrating response capacity in the evening block, where an unanswered complaint is seen by the most people.

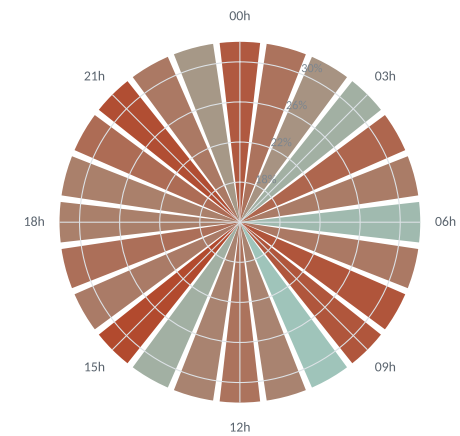


FIGURE 12 — Negative comment share by hour. Bar length and colour both encode negative share; 24-hour clock, local time.

Cross-Platform Comparative Matrix

Where each pillar performs, and how effort compares with return.



FIGURE 13 — Net sentiment, pillar × platform. Each cell is the NSS of that pillar on that platform.

The matrix is remarkably consistent by row and inconsistent by column: pillar rank order holds on all three platforms, while the magnitude shifts. Origin & Provenance is the strongest pillar everywhere; Price & Promotion is the weakest everywhere. No pillar reverses its sign between platforms except Price, which is marginally positive on Instagram and clearly negative on the other two.

That consistency is the practical justification for planning content by pillar first and platform second. The platform decision changes how much a pillar earns; it does not change whether it earns.

The one genuine platform-specific effect worth acting on: Instagram absorbs Price content -5-to-7 better than TikTok does. If discount messaging must run at current volume, Instagram is where it does least damage.

Effort against return

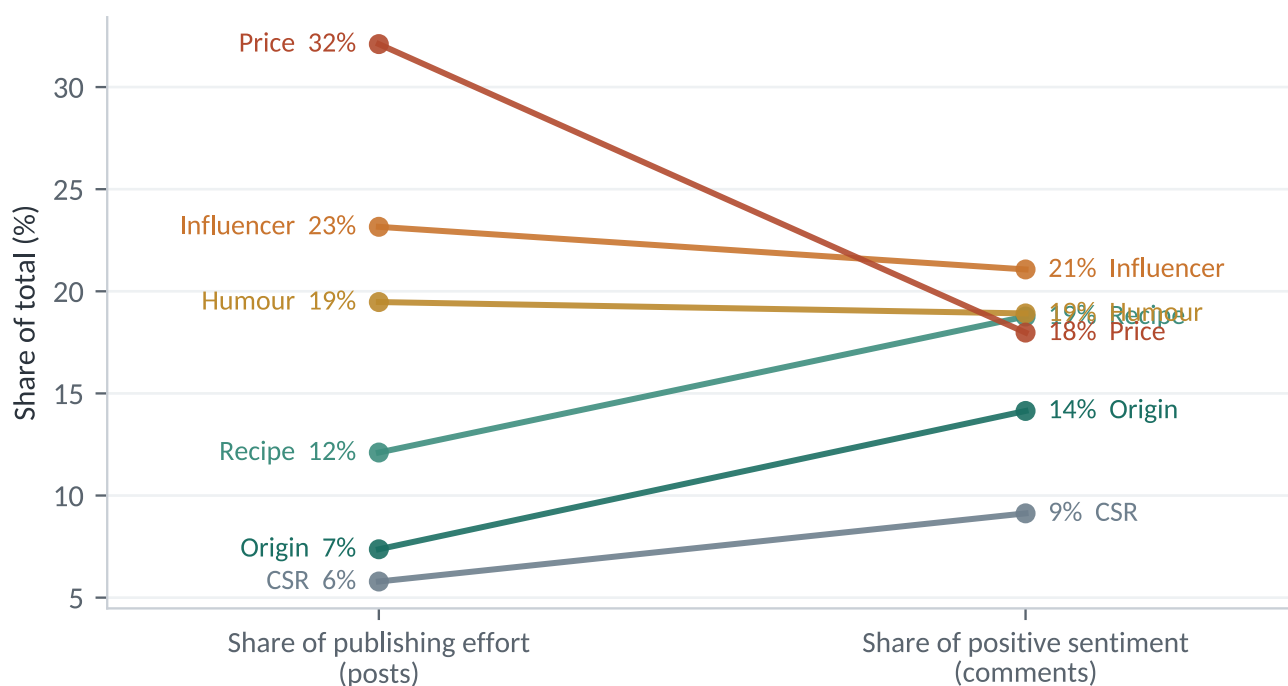


FIGURE 14 — Share of publishing effort against share of positive sentiment. Left axis is each pillar's share of the 190 posts published; right axis is its share of all positive comments.

The crossing lines are the report's clearest single image. **Price & Promotion takes 32.1% of publishing effort and returns 18.0% of positive sentiment** — the only pillar in the portfolio whose return share falls materially below its effort share. Origin & Provenance inverts it, taking 7.4% of effort for 14.1% of return. Recipe & Usage shows the same inversion at larger scale.

PILLAR	EFFORT SHARE (POSTS)	RETURN SHARE (POSITIVE COMMENTS)	RETURN / EFFORT	POSITIVE COMMENTS PER POST
--------	-------------------------	-------------------------------------	-----------------	-------------------------------

Origin & Provenance	7.4%	14.1%	1.92×	450
CSR & Sustainability	5.8%	9.1%	1.58×	370
Recipe & Usage	12.1%	18.8%	1.55×	363
Humour & Trend-jacking	19.5%	18.9%	0.97×	228
Influencer Collaboration	23.2%	21.1%	0.91×	213
Price & Promotion	32.1%	18.0%	0.56×	131

Insight Synthesis

The mechanism that connects content type to trust and to fatigue.

Six pillars, three platforms and twelve weeks of data resolve into one mechanism, and it is simpler than the volume of evidence behind it suggests. **Content that gives the audience something earns trust and resists fatigue. Content that asks the audience for something spends trust and fatigues fast.** Every index in this report is a different measurement of that one distinction.

THE THREE LAWS OBSERVED IN THIS CORPUS

- 1. Evidence compounds; assertion depletes.** Origin & Provenance shows a verifiable thing — a place, a process, a person — and returns a credibility rate of 39.6% with a doubt rate of 2.9%. CSR & Sustainability asserts a virtue without showing it, and returns a doubt rate of 18.9% on friendly-sounding content. The two pillars are adjacent in tone and 23.2 index points apart in trust.
- 2. Fatigue is priced in exposures, not in weeks.** The decay model fits against exposure index, not against time. A pillar published fourteen times in twelve weeks and one published sixty-one times are not running the same campaign at different speeds — the second has spent four times the audience's tolerance.
- 3. The ask is what tires people, not the format.** Humour and Price share no tone, no format and no production cost, and their fatigue constants sit within 0.004 of each other. What they share is that neither leaves the viewer with anything after the scroll.

Answering the three commissioned questions

Q ANSWER AS SUPPORTED BY THIS CORPUS

- Q1** The brand sits at **+21.9** net sentiment, strongest on Instagram (+31.4) and weakest on Facebook (+15.5), but declining at 1.21 points per week on all three. Platform explains 0.08 of the variation against 0.19 for content type.
- Q2** **Origin & Provenance (BTI 70.0)** and **Recipe & Usage (BTI 61.3)** build trust. Both are proof-led: they show a verifiable origin or teach a repeatable technique. Both are published below a fifth of the total. Neither shows meaningful fatigue at current volume.
- Q3** **Price & Promotion (BFI 80.8)** and **Influencer Collaboration (BFI 78.5)** cause fatigue, at half-lives of 33 and 32 exposures against 61 and 44 posts actually published. Humour is approaching the same threshold. Together these three pillars carry 74.7% of publishing effort.

Recommendations & Projected Impact

What to change, in what order, and the modelled outcome of doing it.

The reallocation

The central recommendation requires no increase in output, no additional budget and no new capability. It moves **30% of the comment volume currently generated by Price & Promotion and Influencer Collaboration** into Origin & Provenance and Recipe & Usage, by shifting the posting mix that produces it.

MODELLED EFFECT OF THE SHIFT

$$NSS' = NSS + m(NSS_{\text{recv}} - NSS_{\text{don}}) / N$$

$m = 14,149$ comments reallocated · $NSS_{\text{don}} = +2.7$ (volume-weighted donor score)

· $NSS_{\text{recv}} = +48.7$ (volume-weighted recipient score) · $N = 96,412$

$NSS +21.9 \rightarrow +28.6$ · gain of **6.7 points** at constant total output.

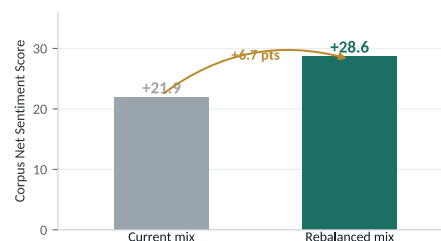


FIGURE 15 — Modelled corpus sentiment. Current mix against the rebalanced mix at constant total output.

Two cautions on this figure. It assumes reallocated volume performs at the recipient pillars' current rates, which holds only while those pillars stay inside their fatigue budget — Origin's half-life of 220 exposures gives ample headroom, but the assumption should be re-tested at the next measurement window. And it models sentiment only; the commercial effect of reducing promotional frequency on short-term sales volume is outside the scope of this corpus and must be assessed separately.

Ninety-day action plan

PHASE	ACTION	RATIONALE DRAWN FROM THIS REPORT	MEASURE OF SUCCESS
Days 1–30	Cap Price & Promotion at 12 posts per quarter; hold the offer, cut the frequency.	Half-life of 33 exposures against 61 published (\$10). The pillar spent the back half of the window in diminishing returns.	Price-pillar NSS above 0; saturation flag rate below 15%.
Days 1–30	Fix the two price-integrity issues before publishing about price again.	Themes 1 and 2 are 27.1% of all negative volume and are operational, not perceptual (\$13).	Combined share of those two themes below 20%.
Days 15–60	Triple Origin & Provenance output to roughly 42 posts per quarter.	Highest BTI (70.0), lowest fatigue (BFI 2.4), only positive weekly slope (+0.44/week) in the portfolio (\$08, \$14).	Origin share of positive comments above 20%.
Days 15–60	Re-brief CSR content around documentary evidence — certificates, audited figures, named suppliers.	Doubt rate of 18.9% on otherwise friendly content; the audience is asking for paperwork (\$09, \$13).	CSR doubt rate below 12%; BTI above 55.

Days 30–90	Restructure influencer work around fewer, longer, evidence-based partnerships.	Doubt rate 24.4%, BFI 78.5; the dominant negative theme is inauthenticity, which volume makes worse rather than better (\$13).	Influencer NSS above +25; doubt rate below 15%.
Days 30–90	Cap Humour at current volume and rotate formats; do not scale it.	Highest engagement per post (13,644) but credibility rate of 7.6% and a weekly slope of -1.80 (\$11).	Humour BFI held below 55.
Ongoing	Concentrate moderation capacity in the 18:00–22:00 block.	Highest-volume window coincides with above-average negative share, peaking at 15:00 (\$14).	Median first-response time under 4 hours in-block.

RE-MEASUREMENT

Every metric above is reproducible from the delivered pipeline. A twelve-week re-run against the same pillar taxonomy gives a directly comparable set of figures; the fatigue constants in particular need at least eight exposures per pillar to re-estimate reliably, which the capped schedule above still provides.

Limitations, Bias & Risk Statement

The boundaries of what this corpus can support.

LIMITATION	EFFECT ON THE FINDINGS	MITIGATION APPLIED
Commenters are not customers	A comment corpus reflects the vocal minority of an audience. Sentiment here measures public expression, not private preference, and it cannot be read as a purchase-intent proxy.	All conclusions are framed as content-performance findings; no sales inference is drawn.
Pillar assignment is at post level	A comment on a Price post that is actually about taste inherits the Price tag. This inflates within-pillar heterogeneity.	Aspect terms are extracted independently at comment level and reported separately in §13.
Sarcasm skews positive	The residual classifier error is asymmetric: ironic praise is more often read as positive than the reverse, so the true negative share is likely marginally understated.	Quantified in §06; the direction of bias is stated so findings are read conservatively.
Six pillars, $n = 6$ for cross-pillar tests	Correlations between pillar-level indices (BTI-BFI, cadence-fatigue) rest on six data points and are fragile to the addition or removal of any one pillar.	Reported with exact p-values and explicitly flagged as descriptive of this portfolio in §11 and §12.
Platform algorithms are unobserved	Engagement decay could partly reflect declining algorithmic distribution rather than audience fatigue. The two are not separable from comment data alone.	The behavioural decay component is blended with the stated-saturation lexicon rate, which is algorithm-independent, at a 55/45 weight (§10).
Deleted and hidden comments	Comments removed by moderation before a pull are absent. If removal targets negative content, observed sentiment is optimistic.	Six-hour pull cadence limits the window for undetected removal; completeness receipts reconcile counts per pull (§03).
Single brand, single window	No category benchmark is available, so "good" and "bad" are internal comparisons across pillars rather than external ones.	All rankings are stated relative to the corpus mean, never against an implied industry standard.

STATISTICAL DISCLOSURE

No result in this report has been selected on the basis of significance. All six pillars, three platforms and ten themes defined before analysis are reported whether or not they reached significance, including the non-significant cadence-sentiment correlation in §12. Confidence intervals accompany every proportion. The full comment-level table is delivered with this report so that every figure can be independently recomputed.

Methodology Appendix

Formula glossary and variable dictionary.

Formula glossary

QUANTITY	EXPRESSION	NOTES
Net Sentiment Score	$(n_+ - n_-) / n \times 100$	Neutrals retained in the denominator.
Engagement weight	$w_i = 1 + \ln(1 + L_i)$	Log weight limits the influence of a single viral comment.
Weighted sentiment	$\sum w_i s_i / \sum w_i \times 100$	$s_i \in \{-1, 0, +1\}$.
Wilson interval	$[\hat{p} + z^2/2n \pm z\sqrt{(\hat{p}(1-\hat{p})/n + z^2/4n^2)}] / (1 + z^2/n)$	Preferred over the normal approximation at extreme proportions.
Margin of error	$z\sqrt{p(1-p)/n}$	Stated at $p = 0.5$, $z = 1.96$.
Brand Trust Index	$50 + 12.5[0.40z(A) + 0.35z(C) - 0.25z(D)]$	Standardised across pillars; mean 50, SD 12.5.
Decay constant	$\ln E(r) = \ln E_0 - \lambda r$	OLS in log space over exposure index.
Fatigue half-life	$r_{1/2} = \ln 2 / \lambda$	Exposures to halve first-exposure engagement.
Brand Fatigue Index	$0.55 \cdot 100(1 - e^{-\lambda k}) + 0.45 \cdot \text{Sat}_{\text{stated}}$	Behavioural and stated components.
Chi-square	$\sum (O - E)^2 / E$	Test of independence, pillar $\times$ sentiment.
Cramér's V	$\sqrt{\chi^2 / [n(m-1)]}$	Effect size, m = smaller table dimension.
Odds ratio	$[p/(1-p)] / [p_0/(1-p_0)]$	Woolf SE on the log scale for intervals.
Cohen's kappa	$(p_o - p_e) / (1 - p_e)$	Agreement corrected for chance.
Trend model	$\text{NSS}_t = \beta_0 + \beta_1 t + \epsilon$	OLS on the weekly series, $t = 1 \dots 12$.

Variable dictionary — delivered comment table

FIELD	TYPE	DEFINITION
pillar	categorical (6)	Editorial purpose of the parent post.
platform	categorical (3)	TikTok, Instagram or Facebook.
week	integer 1–12	Collection week of the comment timestamp.
hour	integer 0–23	Local hour of the comment timestamp.
sent	–1 / 0 / +1	Final sentiment class after lexicon override.
likes	integer	Likes or reactions on the comment at last pull.
advocacy	binary	Recommendation or defence pattern matched.
credible	binary	Affirmation-of-claim pattern matched.
doubt	binary	Suspicion-of-motive pattern matched.
saturation	binary	Explicit repetition or tiredness language matched.
w	float	Engagement weight, $1 + \ln(1 + \text{likes})$.

Deliverables, Commercial Note & Authorisation

What is handed over, and confirmation to proceed.

Delivered with this report

ITEM	CONTENTS
01 Structured database	Comment-level table, 96,412 rows × 11 fields, in XLSX and CSV, carrying the pillar, platform, week, hour, sentiment class, engagement count and all four trust and fatigue flags for every record.
02 Analysis script	Documented Python that regenerates every figure, coefficient and confidence interval in this document directly from the delivered table — no manual step sits between the data and the page.
03 Figure repository	All 16 charts as vector SVG, named to their figure numbers, for reuse in decks and board papers at any size without loss of quality.
04 Digital catalog	This report as a presentation-ready PDF, structured for direct circulation to the client's own stakeholders.

Continuation options

OPTION	SCOPE
Live re-run	The same pipeline against the client's authorised API credentials, producing a directly comparable issue of this report on observed data.
Quarterly tracking	Rolling twelve-week windows with the fatigue constants re-estimated each cycle and every movement flagged against this baseline.
Competitive extension	The same pillar taxonomy applied to two or three competitor accounts, giving external benchmarks for the internal rankings in \$08.
Category extension	Pillar taxonomy extended to additional product lines within the portfolio, on the volume terms set out in the governing commercial proposal.

Commercial note

This report is issued under the terms of the Data Tune commercial proposal governing the engagement. Volume is measured in data entities on the same basis as that document; the corpus delivered here is billed at the agreed entity rate for the base package, with any extension beyond the contracted ceiling proceeding only on prior written approval, and never retrospectively.

Scheduled updates

Progress reporting during the working window follows the same three-point structure as the governing proposal: an early-stage note with initial counts and a first sample batch, a mid-project note with running totals and a second batch, and a pre-delivery note with near-final counts before handover.

Acceptance & authorisation

By signing below, the client confirms receipt of this report and its accompanying deliverable set, accepts the scope and method recorded in \$02 to \$06, and acknowledges the data provenance statement recorded in \$00.

FOR THE CLIENT — AUTHORISED SIGNATURE

FOR DATA TUNE — AUTHORISED SIGNATURE

K. H. Militha Mihiranga

DATA SOLUTIONS CONSULTANT · DATE

NAME & DESIGNATION · DATE

Data Tune.

DATA SOLUTIONS & DIGITAL CATALOGING

For clarification on scope, data parameters, index construction or scheduling, please contact the consultant named alongside.

Prepared by	K. H. Militha Mihiranga • Data Solutions Consultant
Office	555/24 Elhenawatta, Ranmuthugala, Kadawatha, Sri Lanka
Email	[email]
Telephone	[telephone]
Website	[website]