



WHEN THE SURVEY STOPS SPEAKING

*A practitioner case for the Customer
Experience Score in support operations*

Table of Contents

01. The CSAT Response Problem	3
02. The Bigger Problem: Who Actually Responds	5
03. What Operations Actually Needs	6
04. The Customer Experience Score	6
05. The Eight Dimensions	8
06. A Worked Example	9
07. Practitioner Notes	9
08. The Shift Ahead	10
09. References	10
10. About the Author	11

At the end of every quarter, support leadership gathers around QBR dashboards filled with familiar metrics and reassuring numbers.

One slide read:

“CSAT 4.6 / 5. Target met.”

The score appears to confirm that the customer experience is healthy, but beneath it sits a very different reality thousands of support interactions where most customers never responded at all.

In this article, we explore why CSAT is becoming an increasingly incomplete measure of support performance, and why operations teams now need a faster, more complete way to understand customer experience through the Customer Experience Score (CES).

01. The CSAT Response Problem

From my investigation of multiple support engagements to date, the range of CSAT response rates (i.e., post-ticket) from tickets reviewed thus far vary from 2% to 5%. In fact, none of the clients I examined had an outlier client fall within that range; however, according to industry standards, a healthy email-CSAT response would be between 20%-30%, while for SMS pulse, it would be between 40% and 50% of all responses received by support as a whole are nowhere near these ranges.

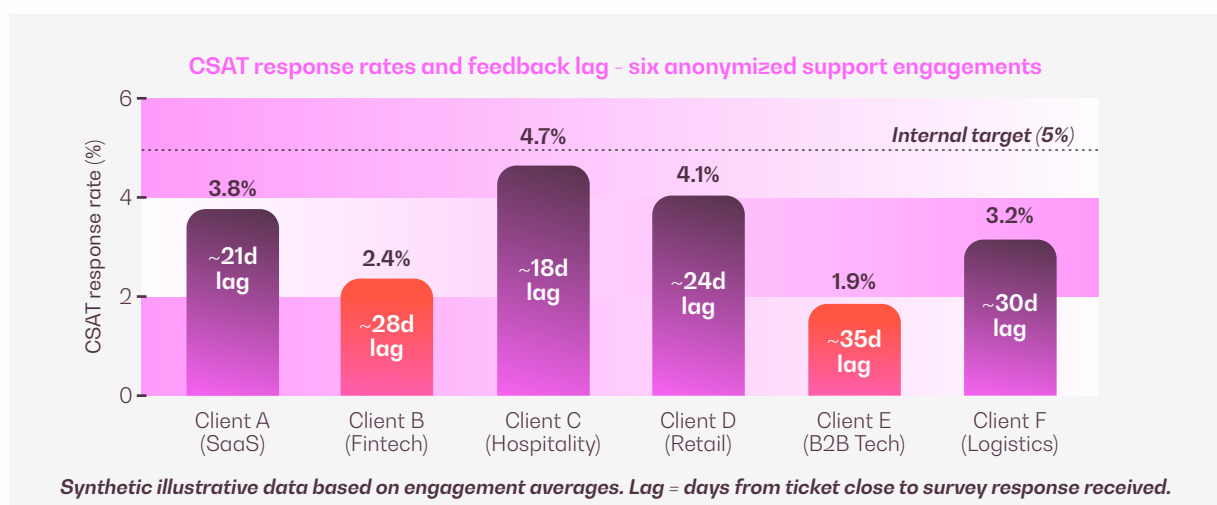


Figure 1 - Response rates and lag across six anonymized engagements. None hit the internal 5% target.

Three things are happening at once.

1

Channel saturation.

In the past 5 years, customers have been getting far more survey invitations than they did in the past. Their inboxes have filled up. Each new survey has to compete for attention with all of the previous surveys that were sent to them. SurveyLab's 2025 report states that Response rates have steadily declined from 2019 (Approximately 1-2% decline per year).

2

Mobile friction.

Long Likert grids and mandatory open-text fields result in abandonment rates above 40% on small screens. Most post-ticket surveys are still designed for desktops.

3

Privacy fatigue.

Customers increasingly question how feedback is used. Even the UK's national Labour Force Survey collapsed to roughly thirteen percent participation. If a statutory survey can't crack double digits, a discretionary support survey is not getting better.

Additionally, there is a delay in receiving feedback because the responses we get occur three to six weeks after the ticket has closed. At that point in time, an agent opened an additional 400 tickets, their team lead has run two retrospectives, and the action that could have been taken based on the original response has been lost. In essence, **we are creating feedback memorials rather than functioning in a feedback loop.**



02. The Bigger Problem: Who Actually Responds

Even if response rates were healthy, the responses themselves would tell a distorted story. A 2019 Swiss study of 717 hospital surveys (Vermeulen et al., BMC Health Services Research) found a J-shaped response curve: the probability of responding was lowest which is around twenty percent for patients in the middle satisfaction band, then climbed sharply at both extremes. Maximum-satisfied patients responded at near seventy percent. Minimum-satisfied responded somewhere between.

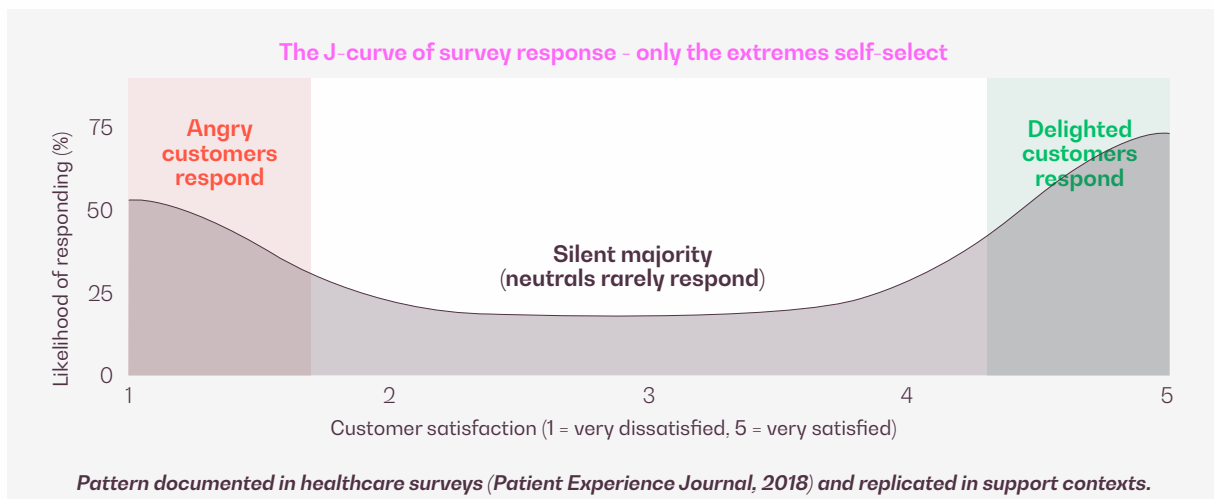


Figure 2 — Customers in the middle rarely respond. Surveys hear the loudest, not the most representative.

Press-Ganey patient satisfaction surveys (Tyser et al., 2016) showed the same: response was patterned by age, gender, payer type, and clinical specialty. Younger patients, men, and Medicaid patients were systematically less likely to respond. The score was not the population's view. It was a slice of the population's view.

Support has the same pattern. The customer who got a clean fifteen-minute resolution shrugs and closes the email. The customer who waited three days through four agents writes paragraphs. The customer in the middle where there is a thirty-minute resolution, mild irritation, no real harm usually deletes the survey. So CSAT becomes, in practice, a barometer of extreme emotion. It tells you who hated you and who loved you. It does not tell you what most of your customers actually experienced.

If **95%** of tickets generate no signal, the 5% that do is not your customer experience.

It is a self-selected slice of it.

03. What Operations Actually Needs

A support operations manager require customer feedback to identify when a support request may escalate, to continue to coach agents while a case is current, and to inform the brand about a decline in customer experience following last week's release. The way that CSAT is currently being used does not provide an operations leader with real-time data about the challenges they are going to face tomorrow. Instead, they are usually asked to manage their team using historical data collected a month prior from one out of twenty customers.

What operations need is a score that is fast, complete, and rooted in observable behaviour:

Fast response and is available within minutes of ticket close, not weeks.

Complete and calculated for every ticket, not just the five percent that respond.

Behavioural is derived from what happened to the customer, not what they remembered later.

That score exists in the data already. Every ticket carries timestamps, agent counts, transfer flags, channel markers, and update logs. The customer's effort is encoded in those signals. We just need to read it.

04. The Customer Experience Score

CES draws directly from Gartner's (formerly CEB's) work on the Customer Effort Score, which Dixon, Toman and DeLisi published in *The Effortless Experience* (2013). Their finding has been replicated repeatedly: the strongest predictor of customer disloyalty is not the absence of delight, it is the presence of effort. Ninety-six percent of customers reporting a high-effort interaction become more disloyal, against nine percent of those reporting low-effort experiences. Gartner's later research positions effort as one-point-eight times more predictive of loyalty than CSAT.

The classic CES is still a survey. It still depends on a response. The Movate CES is a the case-level Customer Experience Score that keeps the philosophy and removes the survey. **We compute customer effort from operational signals that exist on every ticket.**

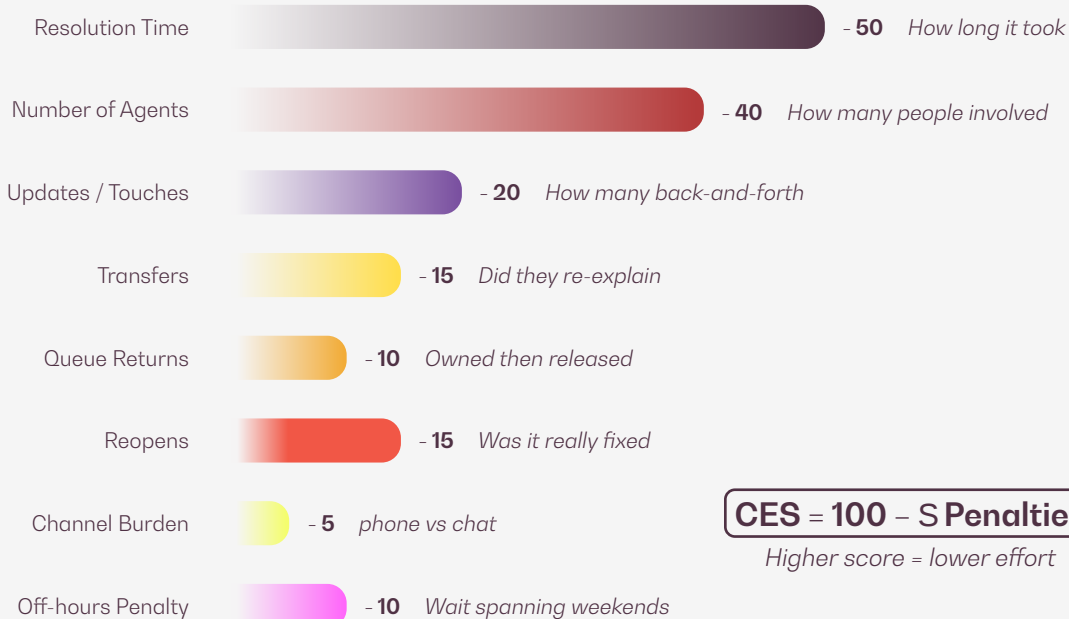
**CES =
100 – S Penalties.**

Every friction event in the case subtracts from a perfect 100.

What remains is the estimated experience quality.

A score of 100 means a single agent resolved the issue in under sixty minutes on first contact. A score near zero means the customer worked hard — multiple agents, follow-ups, transfers, delays. The score is computed for every closed ticket, not for the five percent that bother to respond.

CES Index – eight operational dimensions and their penalty weights



0 10 20 30 40 50
Maximum penalty deducted from 100

Weights are illustrative starting values; calibrate against historic CSAT once sufficient overlap is available.

Figure 3 — Eight dimensions, each with a maximum penalty weight. Resolution time and number of agents carry the heaviest weights.

05. The Eight Dimensions

Eight is a deliberate choice. Six dimensions miss meaningful signals like reopens and off-hours wait. Ten dimensions add complexity without proportionate explanatory power. Eight covers the practical span of effort across most B2B and B2C support contexts.

RESOLUTION TIME

How long the customer waited end-to-end. The single largest weight (max -50) because time-on-issue is the most consistent driver of perceived effort.

NUMBER OF AGENTS

Every additional agent forces the customer to re-explain. Heavy penalty (max -40) for cases handled by four or more.

UPDATES / TOUCHES

The back-and-forth count. Six or more updates suggests the case did not progress smoothly (max -20).

TRANSFERS

Each transfer requires re-explanation and introduces wait at the handoff. Distinct from agent count because it captures the moment of disconnection (max -15).

QUEUE RETURNS

The case was owned, then released back to the queue without resolution. Restarts the customer's wait clock (max -10).

REOPENS

The case was closed, then reopened. The strongest signal that the original "resolution" did not actually resolve (max -15).

CHANNEL BURDEN

The phone requires active engagement throughout; chat allows multi-tasking. Light penalty (max -5) but consistent.

The total maximum penalty across all eight is 165, but real cases rarely accumulate more than 60 to 70. The bands work out to: 85-100 Excellent, 70-84 Acceptable, 50-69 At Risk, 0-49 Poor.

OFF-HOURS PENALTY

The wait that spans weekends or holidays where the customer cannot reasonably expect progress but feels the delay anyway (max -10).

Excellent	Acceptable	At Risk	Poor
85-100	70-84	50-69	0-49

CES Score Bands

06. A Worked Example

A case from a hospitality client. Phone channel. Accounting reports issue. The ticket sat open for nineteen hours, was handled by two agents, returned to queue once, with three customer updates. Starting score 100. Resolution time fell in the under-one-day band: -30. Two agents: -10. Queue return: -10. Phone channel: -5. No transfer, three updates, no reopen has zero on those. Final CES: 45. "Poor." Churn-risk band.

That case never got a CSAT response. The customer closed the email and moved on. But operations now has a flag, dated within an hour of close, with an exact attribution of where the effort came from. The team lead can pull the case, listen to the call, and coach in the same week.

07. Practitioner Notes

Three things from running this in production environments:

First, weights matter less than consistency. A team that runs CES with imperfect weights but reviews every <50 case daily will outperform a team that spent six months tuning weights and never put them into the workflow. Get the version-one calibration done in two weeks, then iterate from operational evidence.

Second, CES exposes process problems that CSAT hides. When CSAT is 4.6, no one investigates. When CES reports that 40% of cases fell into "At Risk" last week, leadership investigates. The same operation. Different visibility.

Third, agents don't resist CES the way they sometimes resist CSAT. CSAT feels personal and usually interpreted as a customer rated "me." CES feels structural where the case had four transfers because routing pushed it there. The conversation moves from blame to fix.

08. The Shift Ahead

CSAT remains in our reports because it captures the customer's voice. That voice is real, but it is partial. We hear the loudest five percent, three weeks late, with a strong pull toward the extremes. No operation function can be run on a signal that thin.

CES fills the gap. It is computed for every ticket, derived from what actually happened, and is available the moment the ticket closes. It is not a replacement for CSAT. It is the daily instrument for a function that has been flying blind on weekly readouts.

Across most engagements, the survey has already stopped speaking. The data has not. CES is how we choose to listen.

09. References

Dixon, M., Toman, N., & DeLisi, R. (2013). *The Effortless Experience: Conquering the New Battleground for Customer Loyalty*. Portfolio / Penguin.

Dixon, M., Freeman, K., & Toman, N. (2010). Stop Trying to Delight Your Customers. *Harvard Business Review*, July–August.

Gartner Research (2024). *How to Measure and Interpret Customer Effort Score (CES)*. Gartner ID G00770907.

Tyser, A.R., Abtahi, A.M., McFadden, M., & Presson, A.P. (2016). Evidence of non-response bias in the Press-Ganey patient satisfaction survey. *BMC Health Services Research*, 16:350.

Vermeulen, B., et al. (2020). Patient satisfaction and survey response in 717 hospital surveys in Switzerland. *BMC Health Services Research*, 20:158.

SurveyLab (June 2025). *Customer Survey Response Rate Benchmarks Review*.

Clootrack (October 2025). *Average Survey Response Rate in 2025: Benchmarks, Drivers, and CX Implications*.

Retently (2026). *CSAT Benchmarks by Industry: 2026 Data*.

10. About the Author



Dr. Kiran Marri

Dr. Kiran Marri leads the company's innovation and digital transformation initiatives. With over 25 years of experience spanning technology, research, and applied innovation, Dr. Marri is recognized for harnessing cutting-edge technologies, including AI and generative AI to solve complex, real-world challenges for clients across industries.

He has authored more than 80 publications in leading conferences and journals, and his work has earned eight award-winning research papers across software engineering, biomedical engineering, analytics, and machine learning. His visionary approach continues to position Movate at the forefront of transformative, AI-driven solutions.

About Movate

Movate is an Applied AI services company that helps enterprises translate AI ambition into measurable business outcomes. With over 12,000 employees across 20 global locations, augmented by AI collaborator agents, it brings together human expertise and technology to help organizations rethink operating models for efficiency and growth. Powered by Mova iO, its intelligent outcomes platform, Movate delivers practical, scalable, and industry-contextualized solutions that address real business challenges.

For more details, please mail us at info@movate.com