



Reading the Simulator

A complete reference: every derived figure, the exact formula behind it, the statistics used to judge it, and how to read every dashboard — written for someone who has never worked in a contact centre

Simulated centre	169 agents · 8,396 contacts/day · 6 queues
Evidence window	2026-06-17 to 2026-08-15 (60 days)
Live demonstration	demo.abldigitech.com
Generated	2026-08-16

Every figure in this document was queried from the simulator's own warehouse when it was generated. Nothing is illustrative.

Table of contents

How to use this guide	4
Part 1 • What a contact centre is	5
The life of one call	5
Part 2 • The two grains — read this before any average	6
Where the numbers physically come from	6
Part 3 • Volume and outcome measures	7
Offered	7
Answered	7
Abandoned	7
Short abandoned	7
Abandon rate	7
Transfer rate	7
Part 4 • Time measures	8
Talk time	8
Hold time	8
After-call work (ACW)	8
Handling time (one touch)	8
Average handling time — AHT	8
Average speed of answer — ASA	8
Part 5 • Service level — the one everyone quotes	9
Service level	9
Service level for a day or a queue group	9
Part 6 • Agent state, occupancy and shrinkage	10
Productive time	10
Occupancy	10
Shrinkage	10
Login hours	10

Adherence	10
Part 7 • Resolution, repeat contact and the banking view	11
First-contact resolution — FCR	11
Repeat contact rate	11
Call-back rate by outcome	11
Complexity band	11
Part 8 • Like-for-like — how agents are actually compared	12
Expected value for an agent	12
Index	12
Part 9 • The statistics, and how an outlier is identified	13
Describing a distribution	13
Robust z-score	13
The three gates — all must pass	13
Persistence — pattern versus bad fortnight	13
Comparing teams	13
Kruskal-Wallis H	14
Welch's t-test	14
Cohen's d	14
Chi-squared	14
Team scorecard composite	14
Peer groups — the strictest comparison	14
Part 10 • Workforce planning	16
Forecast	16
WAPE — weighted absolute percentage error	16
Bias	16
Erlang C	16
Service level from Erlang C	16
Part 11 • Reading each dashboard	17
How this centre is staffed, and why it varies by day	17
Part 12 • Cautions, and a closing checklist	18

How to use this guide

This is a reference, not a narrative. Part 1 explains what a contact centre is and how one call moves through it. Part 2 says where every number comes from. Parts 3 to 6 give the exact formula for every derived figure the simulator produces, including the ones that are easy to misread. Part 7 explains the statistics behind the analytics pages, and Part 8 walks through each dashboard.

Formulae are written as they are actually computed, taken from the queries that produce them, not restated from a textbook. Where two reasonable definitions exist — and for handling time, two do — both are given, with the measured difference on this centre's own data.

The one idea to take away

A contact-centre figure almost never means what it appears to on its own. It has to be read against the work that person, team or hour was given, and at the right grain. This guide keeps returning to that, because it is where the mistakes are.

Part 1 · What a contact centre is

A contact centre is a group of people answering telephone calls, and a system deciding who answers which call. Everything in these dashboards describes one of three things: the **calls** arriving, the **people** available to take them, and the **gap** between the two.

Almost every argument in contact-centre reporting comes from that third thing. If more calls arrive than there are people free, callers wait; if they wait too long they hang up. Staff too generously and nobody waits, but most of the staff sit idle and expensive. The discipline is finding the point between those two failures, hour by hour.

The life of one call

Stage	What happens
1. Arrival	The customer dials. The call is now offered .
2. IVR	The automated menu answers. Some callers get what they need and never reach a person (self-service); some hang up.
3. Queue	The call joins the queue for the skill it needs — cards, banking or priority, in the customer's language.
4. Wait	Time passes. If the customer hangs up while waiting, the call is abandoned .
5. Answer	An agent with the right skill becomes free and takes it. The wait is that call's speed of answer .
6. Talk	The conversation. The agent may place the customer on hold , or transfer them to someone else.
7. Wrap-up	The call ends but the agent is not free: notes, updating the account. This is after-call work .
8. Outcome	Was the problem actually solved? If not the customer rings back — and that is a new call at step 1.

On this centre about **8,396 contacts** arrive a day, of which roughly **6,661** reach a queue — the rest are handled or abandoned in the IVR.

Part 2 · The two grains — read this before any average

A call handled by one agent produces one record. A call **transferred** produces one record *per agent who handled it*. So there are two honest ways to count, and they give different answers:

Grain	Records	Average handling time	Answers the question
Per agent touch	408,494	346.1 s	What one agent spent on their part of a call
Per whole call	382,981	369.2 s	What the centre spent on the customer's issue

The two differ by **23.1 seconds** — exactly the transfer volume, which runs at **6.41%** of touches on this centre. Neither number is wrong. They answer different questions, and quoting one while meaning the other is the most common way a contact-centre report misleads.

Which one the dashboards use

Every handling-time figure on these dashboards — the agent pages, the interval reports and the daily summary alike — is computed **per agent touch**, because they are reporting on what agents did. A whole-call figure is the right choice for asking what a customer's issue cost the business; it is not what the agent tables show.

Records with no handling time at all — a call that rang an agent who did not pick up — are excluded from every handling average. Including them would drag the mean toward zero and reward not answering.

Where the numbers physically come from

Table	One row means
INTERACTION_FACT	One row per contact, end to end. A transferred call is still one row here.
INTERACTION_RESOURCE_FACT	One row per agent involvement. This is the touch grain, and the source of every handling-time figure.
MEDIATION_SEGMENT_FACT	One row per queueing episode — how long the call waited and what became of it. Source of service level, speed of answer and abandons.
SM_RES_STATE_FACT	One row per span of time an agent spent in a state (ready, busy, hold, wrap-up, not ready). Source of occupancy and shrinkage.
SM_RES_SESSION_FACT	One row per login session.
CORE_BANKING_FACT	What the bank did about the call: the reason, the action, and whether it was resolved.
CALLBACK_PENDING	Call-backs an unresolved contact has scheduled for a later day.

Part 3 · Volume and outcome measures

Offered

Formula	count of queue episodes in the period
Where	one row per call that reached a queue
Watch out	Not the same as answered. A centre can look busy while serving badly — the difference is the people who gave up.

Answered

Formula	count where DELIVERED_FLAG = 1
Where	an agent picked the call up
Watch out	Read beside offered, never alone: a rising count can mean better service or simply more demand.

Abandoned

Formula	count where ABANDONED_FLAG = 1 and the result is 'AbandonedWhileQueued'
Where	the customer hung up while waiting in the queue
Watch out	Deliberately excludes very short abandons — see below.

Short abandoned

Formula	count where the result is 'ShortAbandoned'
Where	hung up almost immediately, typically a misdial
Watch out	Counted separately and kept OUT of the service-level denominator. Including them punishes the centre for wrong numbers.

Abandon rate

Formula	abandoned / offered
Watch out	This centre averages 4.3%. On a small interval this is violently unstable — one abandon in a three-call overnight hour reads as 33%.

Transfer rate

Formula	sum(TRANSFER_INIT_FLAG) / count of touches
Where	share of an agent's calls passed to someone else
Watch out	Runs at 6.41% here. High transfer rates can mean poor routing rather than a reluctant agent — check what work they were sent before concluding anything.

Part 4 · Time measures

Talk time

Formula `sum(TALK_DURATION)`

Where time in conversation

Hold time

Formula `sum(HOLD_DURATION)`

Where customer parked, agent still on the call

After-call work (ACW)

Formula `sum(ACW_DURATION)`

Where wrap-up after the customer has gone

Watch out A distinct habit from slow talking, which is why the dashboards score it separately.

Handling time (one touch)

Formula `TALK_DURATION + HOLD_DURATION + ACW_DURATION`

Where everything one call costs one agent, from answer to free again

Average handling time — AHT

Formula `sum(handling time) / count(touches)`

Where over whatever slice is being shown: agent, skill, interval or day

Watch out A VOLUME-WEIGHTED mean, not the average of interval averages. Averaging percentages or averaging averages would let a three-call overnight interval count as much as the midday peak.

Average speed of answer — ASA

Formula average queue wait over ANSWERED calls only

Where abandoned calls are excluded from this average

Watch out An average hides its tail: half the callers can be answered instantly while a minority wait minutes, and ASA still looks fine. Read it beside service level, never instead of it.

Part 5 · Service level — the one everyone quotes

Service level is the share of calls answered within a target time, written as '85 in 40' — 85% answered within 40 seconds. It is the number a contact centre is usually judged on, and its exact denominator matters more than most people realise.

Service level

Formula answered within threshold ÷ (answered + abandoned-while-queued)

Where the threshold is per queue: 85% in 15s, 85% in 40s

Watch out Short abandons are NOT in the denominator. Callers who hung up instantly were never really waiting for service, and including them makes a centre's score move with its misdial rate.

Service level for a day or a queue group

Formula $\text{sum}(\text{answered} \times \text{interval service level}) \div \text{sum}(\text{answered})$

Where weighted by how many calls each interval actually answered

Watch out NOT the mean of the interval percentages. A quiet 02:00 interval would otherwise count as much as the midday peak and flatter the day.

This centre targets 85% in 15s, 85% in 40s and runs at **89.4%** across the week, with a worst day of 71.9% and a best of 99.2%.

Why the same centre shows very different numbers by day

Demand across a week is not flat — this one swings more than two to one between its busiest and quietest day. A roster that cannot flex with it will over-serve the quiet days and starve the busy ones, and the service level will swing far more than the staffing does.

Part 6 · Agent state, occupancy and shrinkage

An agent is always in exactly one state, and the time spent in each is recorded as spans. Every productivity figure is built from those spans.

State	Meaning
Ready	Logged in, available, waiting for a call
Busy	On a call
Hold	On a call, customer parked
AfterCallWork	Wrapping up the last call
NotReady	Logged in but unavailable — break, lunch, training, meeting, system problem

Productive time

Formula	Busy + Hold + AfterCallWork
Where	time genuinely spent on customer work

Occupancy

Formula	$\text{productive} \div (\text{productive} + \text{Ready})$
Where	of the time an agent was AVAILABLE, the share actually worked
Watch out	NotReady time is not in the denominator — an agent on a break is not idle capacity. Occupancy above roughly 85% sustained burns people out; well below 70% means you are paying for idle time. This centre runs at 80.6%.

Shrinkage

Formula	$\text{NotReady} \div (\text{productive} + \text{Ready} + \text{NotReady})$
Where	share of paid, logged-in time not available for calls
Watch out	Configured at 26.5% here. Forgetting shrinkage is the classic staffing error: thirty rostered agents are not thirty agents on the phone.

Login hours

Formula	$\text{sum of all state spans} \div 3600$
Where	time logged in at all

Adherence

Formula	$\text{time in the state the schedule called for} \div \text{scheduled time}$
Where	schedule versus what actually happened
Watch out	A centre showing 100% adherence is measuring a model with no punctuality in it, not a well-run operation.

Part 7 · Resolution, repeat contact and the banking view

First-contact resolution — FCR

Formula resolved ÷ answered, on the banking record

Where whether the customer's issue was actually settled on that call

Watch out Usually self-reported by the agent who handled the call, which makes it close to meaningless. It is only trustworthy when something independent confirms it — here, a real repeat contact.

Repeat contact rate

Formula share of contacts where the SAME caller ID appears again within the repeat window (7 days)

Where matched on the caller's number

Watch out Only computable because callers are drawn from a stable customer pool. If every call carried a fresh random number this metric would silently become noise.

Call-back rate by outcome

Formula share of contacts followed by another contact from the same customer within 14 days, split by whether the first was resolved

Where the test of whether FCR means anything

First contact	Calls	Customer called back
Resolved	225,925	38.3%
Unresolved	67,445	71.5%

That gap is the point. It makes first-contact resolution a capacity number rather than a quality slogan: failed resolutions come back as demand, are handled again, and consume the roster twice.

Complexity band

Formula a property of the reason the customer called (1 simple, 3 complex)

Where set per reason, not per agent

Watch out Complex reasons resolve far less often than simple ones BY DESIGN. Comparing an agent on complex work with one on simple work measures the work. Every banking comparison is scored inside the band.

Part 8 · Like-for-like — how agents are actually compared

This is the heart of the analytics, and the part most contact-centre reporting gets wrong. Measured on this centre, **99.4% of the apparent spread in agent handling time was the mix of work each agent was given**, not how well they did it.

So no agent is ever ranked on a raw level. Each is scored against what their own mix of work costs the centre — an **index**, where 1.00 means exactly par for that mix.

Expected value for an agent

Formula $\Sigma(\text{agent's calls in skill} \times \text{centre-wide average for that skill}) \div \text{agent's total calls}$

Where the centre-wide rate is computed over the SAME window, so a busy fortnight moves everyone's reference together

Index

Formula $\text{agent's actual value} \div \text{their expected value}$

Where 1.00 = par for their work mix; 1.20 = 20% above what that mix costs

Watch out An index still needs volume behind it before it means anything — see the gates in Part 9.

Six indices are computed, each answering a different question:

Index	What it measures	Direction
Handle time	Talk + hold + wrap-up per call, against what this agent's skill mix costs	lower is better
Wrap-up	After-call work per call — a distinct habit from slow talking	lower is better
Transfers	Share of calls passed on rather than resolved	lower is better
Short calls	Calls ended under 30s — the call-avoidance signal	lower is better
Hold use	Share of calls placed on hold	lower is better
Occupancy	Busy against available, versus the centre's occupancy in the HOURS this agent worked	higher is better

Occupancy is standardised by HOUR rather than by skill, because occupancy is driven by how busy the centre was when someone was on shift, not by what they were skilled in. A night-shift agent is not lazy; the night is quiet.

The short-call threshold is a calibrated number, not a guess

At 60 seconds it conflated fast agents with call-avoiding ones — a 3.7× separation, and it wrongly flagged a top performer. At 30 seconds the separation is 33×. The evidence sits beside the value in the configuration.

Part 9 · The statistics, and how an outlier is identified

Being at the bottom of a list is not evidence of anything. On a 152-seat centre somebody is always last. These are the tests that separate a real difference from ordinary variation.

Describing a distribution

Every metric is summarised with median, quartiles, 10th and 90th percentiles, standard deviation, skew, and a 95% confidence interval on the mean. The median and the **median absolute deviation** (MAD) — the median of each value's distance from the median — do the work, because both are unmoved by extreme values.

Robust z-score

Formula $(\text{value} - \text{median}) \div (1.4826 \times \text{MAD})$

Where how far out a value is, in robust standard deviations

Watch out A plain z divides by the standard deviation, which the very outliers being hunted have already inflated — so everyone's z shrinks and the detector hides its own targets. The 1.4826 rescales MAD so it is comparable with a standard deviation.

The three gates — all must pass

Gate	Threshold	Why
Volume	at least 150 calls in the window	an index built on 40 calls is mostly sampling error
Tail	in the bottom decile of peers, widening to a quartile below 40 peers	a decile of eight people is not a decile
Robust z	at least 2.0 robust standard deviations from the median	far enough out to be more than ordinary spread

Any one of these alone produces noise. Requiring all three is what makes the flag worth acting on.

Persistence — pattern versus bad fortnight

The window is cut into whole 7-day periods, and an agent must appear in the tail in **75% of them** to be confirmed — 3 of 4 weeks at 30 days, 6 of 8 at 60. A genuine pattern repeats; an unlucky fortnight does not. With fewer than 3 whole periods the test cannot distinguish the two, so it is not applied and the dashboard says so rather than confirming someone on evidence that cannot carry it.

How well the detector actually works

The simulator seeds a known number of genuinely different agents, so the detector can be scored against the truth: **100% recall, 89% precision, and no top performer wrongly flagged**. That is a measured result, not a claim.

Comparing teams

Kruskal-Wallis H

Formula rank-based test that ANY group differs from the others

Where run BEFORE any team is named best

Watch out Sort ten teams by mean and there is always a 'best' one, even when all ten are identical. If H is not significant the ordering is noise and the dashboard says so instead of crowning someone.

Welch's t-test

Formula difference between two groups, without assuming equal variances

Where used for a specific pairwise comparison

Cohen's d

Formula difference in means ÷ pooled standard deviation

Where effect size — how big the gap is, not just whether it is real

Watch out A significant p on 150 agents can still be a trivial gap. This says whether anyone should care.

Chi-squared

Formula test of whether a mix has shifted between periods

Where used for reason-mix drift

Team scorecard composite

Formula $\Sigma(\text{weight} \times z \text{ of the team's mean index, signed so positive is good}) \div \Sigma(\text{weights})$

Where z is taken against the spread BETWEEN teams

Watch out Working in z rather than raw units stops whichever metric happens to have the widest natural spread from dominating the total by accident.

Component	Weight
Aht	25%
Short	20%
Transfer	20%
Occupancy	15%
Acw	10%
Hold	10%

Peer groups — the strictest comparison

Mix adjustment prices an agent's work at centre-wide rates, but it does not ask whether the people they are ranked against ever do that work at all. The peer-group view compares each agent only with others

holding an **identical skill signature**. It is the fairest comparison available and the one to use when a finding is challenged.

Part 10 · Workforce planning

Planning runs forwards: forecast the volume, convert it to a staffing requirement, build a roster, then measure afterwards how well the forecast did.

Forecast

Formula	recent volume level × day-of-week factor, shaped across the day by the interval profile for that weekday
Where	built STRICTLY from days before the target date
Watch out	The target day is filtered out rather than trusted to be absent. A forecast that has seen the answer is not a forecast.

WAPE — weighted absolute percentage error

Formula	$\Sigma \text{forecast} - \text{actual} \div \Sigma \text{actual}$
Where	the headline accuracy number
Watch out	Volume-weighted, so it reflects staffing pain. Plain MAPE averages percentage errors across intervals and lets a quiet interval with a large relative error dominate.

Bias

Formula	$\Sigma (\text{forecast} - \text{actual}) \div \Sigma \text{actual}$
Where	signed, unlike WAPE
Watch out	Persistent over- or under-forecasting is a different disease from imprecision and needs a different fix.

Erlang C

Formula	probability an arriving call has to wait, given traffic and agents
Where	traffic (in erlangs) = calls per hour × handling time ÷ 3600
Watch out	Assumes nobody hangs up and every agent can take every call. Both are false, so it is a starting point that must be checked against what actually happened.

Service level from Erlang C

Formula	$1 - P(\text{wait}) \times e^{-(\text{agents} - \text{traffic}) \times \text{threshold} \div \text{handling time}}$
Where	the analytic prediction the roster is sized against

Why the Erlang target is calibrated rather than trusted

Pooled Erlang is optimistic wherever skill coverage is thin, because it assumes any agent can take any call. On this centre, sizing to an 85% Erlang target produced about 80% in reality. The configured target is therefore set from measured outturn, not from the number the business wants to see.

Part 11 · Reading each dashboard

Page	What it answers
Wallboard	Right now: queues waiting, agents by state, calls in progress. Read the LONGEST wait, not the average — the average hides the customer about to hang up.
Intraday	Today, building through the day on a fixed axis so the shape is comparable with yesterday's.
Intervals	Fifteen minutes at a time, by queue. This is where understaffing actually shows up; a day can look fine while two intervals were unstaffed.
Adherence	Scheduled versus actual. Whether the plan was followed.
Trends	Weeks and months, for direction. Not for judging one day.
Agents	Per-agent detail, indexed against comparable work. Sort by index, never by the raw level.
Performance	Team and individual scorecards, work-mix adjusted, with the raw figure shown alongside so the difference between the two is visible.
Analytics	Outliers, significance and persistence — who is genuinely different rather than briefly unlucky. Start here when asking who needs help.
Workforce	Forecast, required staffing, schedule, and forecast accuracy measured after the fact.
Banking Ops	Why customers called and what the bank did about it, scored inside complexity bands.

How this centre is staffed, and why it varies by day

Day	Calls offered	Agents	Service level	Occupancy
Sunday	4,072	83	85.1%	77.6%
Monday	8,215	153	85.4%	84.7%
Tuesday	7,579	153	93.1%	80.2%
Wednesday	6,841	138	91.0%	80.1%
Thursday	6,783	136	88.8%	81.4%
Friday	7,269	142	90.5%	80.8%
Saturday	5,859	119	89.9%	79.7%

The busiest day carries roughly twice the quietest. A roster with the same number of people every day is therefore wrong twice over — wasteful on the quiet day and short on the busy one. Staffing here is solved per day, which is why the agent count moves with the volume beside it.

The busiest day is deliberately staffed slightly below what the formula asks for. A demonstration in which nothing is ever under pressure would not show how the reporting behaves when a centre is genuinely stretched, which is the condition the reporting exists for.

Part 12 · Cautions, and a closing checklist

Before quoting any figure from any of these pages, ask four things.

- **What is the grain?** Per call or per agent touch — they differ by the transfer volume, 6.41% here, and both are correct.
- **What is the denominator?** Service level excludes short abandons; occupancy excludes not-ready time. Two systems quoting 'service level' can be computing different things.
- **How big is the sample?** Ratios on small samples are noise. An overnight interval may hold three calls, and the dashboards deliberately blank a ratio rather than print one that will be misread.
- **What else differs?** If the things being compared were given different work, the comparison is about the work — not the people.

If you remember one sentence

Never rank anybody on a raw level; rank them on the difference from what their own work should have cost.

The data is simulated. It is realistic and internally consistent, and it is not any real bank's traffic. Every figure quoted in this guide was queried from the simulator's own warehouse on 2026-08-16, over the 60 days from 2026-06-17 to 2026-08-15.