

CCaaS, WEM, and AI customer service landscape as of 8/26/26

Twelve platforms across three markets, scored on one matrix, with nine buying situations and a call on each

David Soden

Independent analyst and practitioner, enterprise automation and applied AI
davidsoden.com/market-assessment

About this report

Every substantive claim in this report carries an evidence classification and, where applicable, a numbered source. Facts are separated from vendor claims, third-party estimates, the author's assessments, and untested hypotheses. Forward-looking sections use scenarios with observable tripwires rather than point forecasts. Stated assumptions are listed before the analysis begins.

PUBLISHED FOR PUBLIC READING | SHARE FREELY, UNMODIFIED, WITH ATTRIBUTION

© 2026 David Soden, davidsoden.com/market-assessment. All rights reserved. You may share this file complete and unmodified, and quote short passages with attribution. Republishing under another name, excerpting into a commercial deliverable, resale, and use as machine learning training data require written consent. Full terms appear under Rights and Permitted Use on the final page.

Market Assessment reports are published and commissioned at davidsoden.com/market-assessment. This document is analysis, not investment, legal, or procurement advice.

Scope and evidentiary standard

Every price, product boundary and corporate structure in this document is a claim about 26 August 2026 and nothing else. That date is printed on the matrix headings for a reason. Two of the ten vendors on the original shortlist became one company nine months ago, a third bought its conversational AI vendor outright eleven months ago, and one of the four vendors that used to publish a full list price now publishes two of five. A reader picking this up in six months should treat the commercial cells as history and the architectural cells as still broadly true.

The subject is three markets that share one buying committee: contact centre as a service, workforce engagement management, and AI applied to customer service. The vendor set is the nine names handed over as a shortlist, plus three the shortlist omitted and which compete for the same budget. Cloud primitives that a buyer can assemble themselves are included on purpose, because leaving them out is how a suite quote goes unchallenged.

The test applied to every claimed difference: a capability both vendors have is not a differentiator, and a capability only one has is not a differentiator either unless a buyer's outcome changes because of it. Roughly two thirds of what vendors market as differentiation fails that test and is not in this document. What survives is mostly commercial structure, telephony ownership, model provenance and exit cost, none of which appear on a feature comparison sheet.

On sources, the order of preference is: published pricing pages, including archived snapshots from the Internet Archive where the question is whether a price moved; regulatory text; product and API documentation; company filings and earnings releases; peer-reviewed and preprint research. Vendor marketing is used, labelled as a vendor claim, and named alongside the product it is selling. Third-party comparison sites, best-of listicles and vendor-published competitor roundups are used only to establish that a complaint recurs, never to quantify anything.

No third-party market share or growth figure appears here. The number that circulates most widely in this category, a CCaaS revenue split placing NICE at roughly 22 percent, Genesys at

roughly 20 percent and Five9 at roughly 13 percent, appears on several commercial comparison sites with no named research firm, no stated denominator and no method. It may well be a repackaging of a real syndicated estimate. It is not traceable to one, so it is named here and then not used. Analyst quadrant placements are treated the same way: they are positioning that vendors pay to participate in and brief into, and their inputs are not published, so a placement is evidence about a vendor's analyst relations function rather than about its product.

One entry on the source shortlist, "Fear," could not be resolved to any vendor in these three markets. It appears to be a transcription error. It is flagged here as unresolved and no vendor has been guessed from it.

Label	Definition	How to treat it
FACT	Verifiable in the cited primary source: a filing, press release, official documentation, or standards body announcement.	Reliable as stated. Check the source if the exact wording matters.
VENDOR CLAIM	Reported by a vendor about its own business or product. Self-measured and not audited by any third party.	The vendor said it. Directionally useful, not a measurement. Definitions of the metric are usually undisclosed.
THIRD-PARTY ESTIMATE	A research firm forecast or survey result. Methodology is proprietary and generally not reproducible.	Use for direction and order of magnitude only. Do not build a model on it.
ASSESSMENT	The author's reasoned judgment from the cited evidence. Falsifiable, and the reasoning is shown so it can be argued with.	This is analysis, not reporting. Disagree with it where your evidence differs.
HYPOTHESIS	A proposition offered for testing. Not currently supported by sufficient evidence to assert.	Treat as a research question, not a conclusion. Each one names what would confirm or refute it.
SCENARIO	A structured future state with a stated subjective probability. The probability is the author's judgment, not a calculated figure.	Use the leading indicators attached to each, not the probability. The indicators are observable; the probability is not.

Where this view comes from

The author spent almost five years at Concentrix, the second largest contact centre operator in the world, where he built and led the Low-Code / No-Code Center of Excellence. In that role he consulted with a large number of client companies and their decision makers, assessing what each one needed, what it expected, and what outcome it was buying toward, then making platform recommendations against those things.

On decision-making role, stated plainly: he did not hold signing authority on any of these platform decisions. The role was advisory. That is a different vantage point from a buyer's and this document does not blur the two.

Direct working experience covers NICE, Genesys, RingCentral, Salesforce, Microsoft, ServiceNow, AWS and Google Cloud. That is the list. No familiarity rating, deployment count or usage claim beyond it is asserted anywhere in this document.

What this vantage point is good for: seeing how these platforms behave once deployed at scale across many different client estates, at once, under real volume, with the operator carrying the service level. That is a rarer view than one buyer's view of one deployment, and it is the reason the assessments in this document weight operational behaviour and exit cost more heavily than feature coverage.

What it is not: a procurement view. The author has not signed one of these contracts, holds none of the discount sheets, and has not sat on the buy side of a renewal negotiation with any vendor named here. Every price in this document is list. Real enterprise deals settle well below list, the discount varies by vendor, quarter and seat count, and nothing in this document tells you what you will actually pay. It tells you what the vendor has decided to publish, which is a different and separately useful thing.

Assumptions this analysis rests on

1. Published list prices are a signal about vendor strategy even though almost nobody pays them. The argument in the commercial section rests on what vendors chose to publish and when they stopped, not on transaction prices. If it turns out that a vendor's published page is decorative and bears no relation to its quoted structure, the pricing-movement finding weakens considerably, though the observation that disclosure was withdrawn survives.
2. The capability floor across the named suites has converged far enough that routing, recording, forecasting and basic conversational self-service are no longer selection criteria at enterprise scale. This is the load-bearing assumption behind every scenario call. If a buyer's evaluation finds a genuine capability gap on a core function at one of the top four suites, the scenario calls in this document should be discarded for that buyer and the evaluation should proceed on capability.
3. Vendor-reported AI revenue figures are directionally honest about scale even though their definitions differ and none are separately audited. The finding that AI is roughly a tenth to a sixth of revenue at the disclosing vendors depends on this. If any of these figures turns out to include bundled seat revenue attributed to AI, the finding overstates AI penetration and the argument that the meter is where the money is going gets stronger, not weaker.
4. Switching cost is real, large, and roughly proportional to integration depth and years live. No vendor publishes it and no independent study measures it, so the exit-cost column is the author's assessment from operator-side experience of migrations, not a measured figure. If a buyer has run a like-for-like migration and found it cheap, that buyer's evidence beats this document's.

5. EU AI Act Article 50 will be enforced against deployers of customer-facing conversational systems rather than being treated as a dead letter. The regulatory section and one of the four scenarios depend on this. If enforcement never materialises, the contractual asks in the contract-language section remain sensible hygiene but stop being urgent.

The call

ASSESSMENT Treat this as three purchases, not one. The workforce specialists are now a single company, the suites bundle workforce management to defend the seat, and the AI sits on a meter the seat contract does not govern. On the evidence I have, the deciding variable at renewal is no longer capability, which has converged, but who controls the meter and what it costs to leave. I would buy the suite for global voice and complex routing, the specialist where the problem is workforce cost, and the cloud primitives only with engineers on staff. What would move me: a major vendor publishing an unmetered agentic tier at a fixed seat price, which would collapse the meter argument entirely.

My judgment on the evidence assembled here, nothing more. There is data I can't see and data I may have missed. New evidence can move me, including to the opposite position, and I reserve the right to move.

Executive summary

Anyone handed these ten names as one comparable list has a category error in front of them before evaluation starts. Verint and Calabrio are workforce specialists that never sold a routing engine. NICE and Genesys are suites that bundle workforce management to defend a seat. Five9 and Talkdesk are contact centre platforms without workforce heritage. RingCentral arrives from business telephony. Salesforce and Microsoft attach contact centre capability to a seat the buyer already owns and pays for. Naming that is not a preliminary to the answer, it is part of the answer, and the matrix that follows is built so a buyer can read across a row rather than down a column.

FINDING 1 Two of the ten shortlisted vendors are now one company, and the buyer-facing consequence is a roadmap collision rather than a merger announcement.

FACT Thoma Bravo agreed to acquire Verint on 25 August 2025 at \$20.50 per share, an enterprise value of about \$2 billion, closed on 26 November 2025, and combined Verint with its existing portfolio company Calabrio. On 18 February 2026 the combined organisation announced it would operate under the single name Verint, with Calabrio solutions continuing inside the Verint CX Automation Platform, and on 24 February 2026 Dave Rhodes, formerly Calabrio's chief executive, was named chief executive of the combined firm. Running Verint against Calabrio on a shortlist in August 2026 is comparing a company to itself under two brands, with one of the two product lines eventually losing the roadmap argument. ^{[25] [26] [27] [28]}

FINDING 2 After three years of AI-first positioning across the whole category, AI is roughly a tenth to a sixth of revenue at every vendor that discloses anything, and no two of them are counting the same thing.

VENDOR CLAIM NICE reported AI and self-service annual recurring revenue of \$362 million for the second quarter of 2026, which it described as 15 percent of cloud revenue. Genesys

reported cloud annual recurring revenue of nearly \$2.6 billion for the fiscal year ended 31 January 2026 and AI annual recurring revenue above \$250 million, roughly a tenth. Five9 management put AI revenue at an annual run rate above \$150 million on the second quarter 2026 earnings call, against a full-year revenue guide of about \$1.27 billion. RingCentral reported that 13 percent of annual recurring revenue comes from customers using at least one paid AI product, which counts the whole subscription of any customer who buys any AI and is therefore not comparable to the other three. The number that matters commercially is not the percentage, it is that these four figures are presented side by side in the market as though they measure the same thing. ^{[19] [20] [21] [22] [23]}

FINDING 3 Five9 withdrew published list pricing from exactly the three tiers that carry workforce management and AI, and did it between May 2024 and May 2025.

FACT Archived snapshots of Five9's own pricing page show all five tiers carrying a published monthly rate in June 2023 (Digital \$149, Core \$149, Premium \$169, Optimum \$199, Ultimate \$229) and again in May 2024, at materially higher numbers (Digital \$175, Core \$175, Premium \$235, Optimum \$290, Ultimate \$325). By May 2025 the top three tiers read "Contact Sales" and only Digital and Core carried a price, both \$119. As of 26 August 2026 the tier names have changed to Digital, Core, Plus, Pro and Enterprise, the two published rates are \$119 and \$159, and the three highest tiers remain quote-only. The top published rate rising then falling is a disclosure event, not a price cut. ^{[5] [6] [7] [8]}

FINDING 4 Genesys held its seat prices almost flat through the AI cycle and moved the money onto a meter instead.

FACT In May 2023 Genesys published CX 1 at \$75, CX 2 Digital at \$110 and CX 3 Digital plus workforce engagement at \$150, alongside a flat AI Experience add-on at \$40 per agent per month. On 26 August 2026 the same three tiers publish at \$75, \$115 and \$155, a top-tier CX 4 has appeared at \$240, and the AI Experience add-on no longer exists as a flat seat charge. AI capability is now metered in AI Experience tokens at \$1.00 each, with an agentic virtual agent interaction consuming 1.2 tokens and AI-based scoring consuming 0.05 tokens per invocation. Three years of AI investment moved the seat price by five dollars and created a variable line item with no ceiling. ^{[1] [2] [3] [4]}

FINDING 5 Genesys' own reporting documentation defines a contained interaction as one that never reached an agent, which means a customer who gives up inside the bot scores as a success.

FACT The Genesys self-service statistics report defines the metric "Contained in Self-Service" as "the total number of interactions that entered the Designer application in Self-Service and were concluded without entering Assisted-Service." Nothing in that definition requires the customer's problem to have been solved, or the session to have ended in anything other than abandonment. Genesys' own marketing has since acknowledged the weakness, writing that "a high containment rate might look like a win, but that might not be the case if your customers didn't get what they needed." No vendor in this set publishes a containment figure alongside a definition of what counts as contained and an independent measurement of it. ^{[30] [31]}

FINDING 6 On published list meters, the cost of one AI-handled interaction spans roughly twenty-five to one across this vendor set, and the spread is a purchasing decision rather than a technical one.

ASSESSMENT Salesforce publishes Agentforce at \$2.00 per conversation, or 20 Flex Credits (\$0.10) per action at \$500 per 100,000 credits. Genesys prices an agentic virtual agent interaction at 1.2 AI Experience tokens, \$1.20 at the published \$1.00 token rate. NICE's Ultimate Suite carries \$0.25 per session on top of \$249 per agent per month. Google prices a generative voice agent at \$0.002 per second, about \$0.36 for a three-minute call. Amazon Connect prices end-customer self-service at \$0.0080 per voice minute on top of \$0.018 per voice minute, under ten cents for the same call. The arithmetic is this author's, on published rates; the underlying rates are the vendors' own. [\[1\]](#) [\[2\]](#) [\[9\]](#) [\[13\]](#) [\[16\]](#) [\[17\]](#)

FINDING 7 Microsoft's contact centre AI draws from a credit pool shared with the rest of the tenant, which puts the burn rate outside the contact centre budget owner's control.

FACT Microsoft's Copilot Credits Guide of June 2026 states that Copilot Credits are the common currency across Microsoft 365 Copilot experiences including Copilot Cowork, Copilot Studio, Dynamics 365 agents, Power Platform workloads and Work IQ APIs, that they are "pooled at the tenant level," and that an organisation's total cost is the sum of credits consumed across all supported experiences. Pay-as-you-go is \$0.01 per credit. The pre-purchase plan runs from 300,000 credits at a 5 percent discount to 300,000,000 at 20 percent, and unused credits expire at the end of the annual term. A contact centre director buying Dynamics 365 Contact Center at \$110 per user per month is buying a seat whose AI consumption competes with every other Copilot workload in the company. [\[14\]](#) [\[15\]](#) [\[42\]](#)

FINDING 8 The regulatory obligation most likely to affect a live deployment took effect three weeks before this document's date and is almost never priced at selection.

FACT EU AI Act Article 50(1), (2) and (5) became applicable on 2 August 2026. Providers of AI systems that interact directly with natural persons, which includes customer-service chatbots, voice bots and AI agents, must ensure the person is informed they are interacting with an AI system unless that is obvious to a reasonably well-informed observer, and the information must be given clearly and distinguishably at the latest at the time of first interaction. The obligation applies to systems already on the market, not only to new ones. Penalties reach €15 million or 3 percent of worldwide annual turnover, whichever is higher. Nothing in the published pricing of any vendor in this set reflects the cost of meeting it, and the compliance duty falls on the deployer, not the platform. [\[33\]](#) [\[34\]](#)

FINDING 9 The claim that automated quality management scores every interaction is true; the claim that those scores mean what a human reviewer's scores mean is unevicenced, and that gap is worth more money than most of the platform comparison.

ASSESSMENT Traditional quality management samples a low single-digit percentage of interactions, so full-coverage automated scoring is a genuine change in kind. What no vendor in this set publishes is a chance-corrected agreement statistic between its automated scores and trained human reviewers on the same interactions. The nearest adjacent evidence is the largest systematic evaluation of language models used as judges published to date, covering 21 judges from nine providers across roughly 541,000 individual judgments, which found

mean Cohen's kappa against human labels of 0.451 on one benchmark and 0.593 on another, and found that high test-retest reliability coexisted with severe position bias. Moderate agreement across all interactions may still beat excellent agreement across two percent of them. That is an empirical question a buyer can settle in a fortnight and nobody has published the answer. ^[35]

FINDING 10 Salesforce and ServiceNow each own a piece of Genesys, which changes how a Salesforce-versus-Genesys evaluation should be read.

FACT On 31 July 2025 Genesys announced a \$1.5 billion investment, \$750 million each from Salesforce and ServiceNow, with proceeds used to repurchase shares from existing equity holders. The announcement names two existing integrations: CX Cloud, combining Genesys Cloud with Salesforce Service Cloud, and Unified Experience, combining Genesys Cloud with ServiceNow Customer Service Management. A buyer being told by a Salesforce account team that Service Cloud Voice removes the need for a separate contact centre platform is being told that by a shareholder in the separate contact centre platform. ^[24]

Ten names, three markets

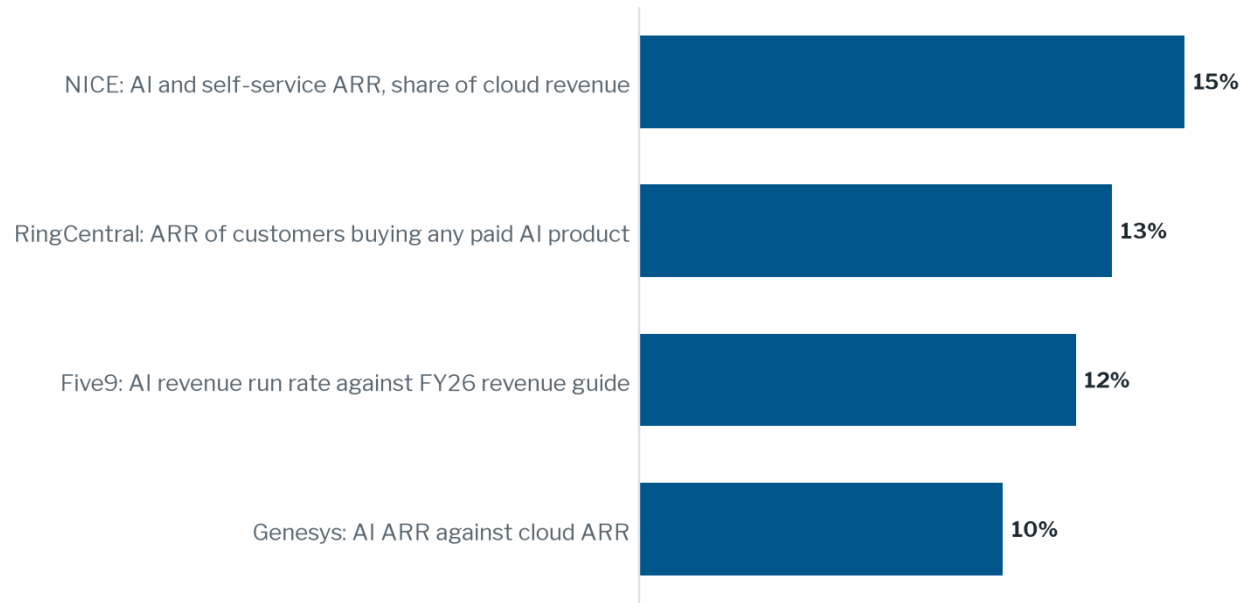
ASSESSMENT The shortlist is not a shortlist. It is three shortlists that have been stapled together, and the staple is the buying committee rather than anything about the products. Verint and Calabrio sell against a labour cost line: forecasting, scheduling, adherence, recording, evaluation. NICE and Genesys sell a routing engine and give away enough workforce management to make buying a specialist look like duplication. Five9 and Talkdesk sell a routing engine and acquired or bolted on workforce management later, which shows. RingCentral sells a telephony seat and attaches a contact centre to it. Salesforce and Microsoft sell a system of record and attach a contact centre to that. Amazon and Google sell components and let you attach whatever you like.

ASSESSMENT The consequence for evaluation is specific and not merely taxonomic. A scoring matrix that gives every vendor a workforce management score will rank Verint above Five9 and then be unable to explain why Five9 won the deal, because the deal was decided on global voice coverage that Verint does not sell. A matrix that scores routing will rank Genesys above Calabrio and be unable to explain why the operation with a 41 percent attrition problem bought Calabrio, because that operation was not buying routing. The fix is not a better weighting scheme. The fix is deciding which of the three purchases you are actually making, and running the other two as separate evaluations with separate business cases.

ASSESSMENT Three names the shortlist omitted belong in the comparison because they take money from the same budget line. Amazon Connect is not a niche developer product; it is the telephony and routing substrate underneath a large number of enterprise contact centres, and it prices as meters with no seat fee at all. Google's Customer Engagement Suite competes directly for the AI self-service budget at prices roughly an order of magnitude below the suites. And Microsoft's Dynamics 365 Contact Center is a different purchase from the Teams-attached route that most Microsoft-standardised organisations actually end up on, which pairs a Teams telephony seat with a partner-supplied contact centre and produces a materially different

support and integration position. Leaving all three out of an evaluation is how a suite quote goes unchallenged, and the omission tells you the market frames itself around vendors that sell seats.

FIGURE 1 Four published "AI share" figures that do not measure the same thing



VENDOR CLAIM Each vendor's own most recent disclosure. The RingCentral figure counts the entire subscription of any customer who buys any AI product and is not comparable to the other three; the Five9 figure is management commentary on an earnings call rather than a line in the release; the Genesys figure is unaudited and the company is private. None is independently verified. The comparison is offered to show the spread of definitions, not to rank the vendors. [\[19\]](#) [\[20\]](#) [\[21\]](#) [\[22\]](#) [\[23\]](#)

The comparison matrix, as of 26 August 2026

Every cell below is a claim about 26 August 2026. Where a cell rests on the vendor's own documentation or pricing page, that is stated in the cell or the paragraph beneath the table. Where a cell is the author's assessment from operator-side experience rather than a citable source, it is marked as such and should be weighted accordingly. Read across a row to make a decision; reading down a column produces a feature score, which is the thing this document argues against.

Routing, voice and who can change a rule

Platform	Routing customisation without professional services	Telephony model	Who changes a routing rule in practice
NICE CXone Mpower	Deep. Studio scripting is genuinely powerful and genuinely a programming language.	Owns and resells. Carrier-grade voice plus bring-your-own-carrier.	IT or a trained platform specialist. Not an operations supervisor.
Genesys Cloud CX	Deep. Architect flows plus Predictive Routing. Data	Owns Genesys Cloud Voice; also supports bring-your-own-carrier and cloud carrier.	A trained admin. Operations can change queue

Platform	Routing customisation without professional services	Telephony model	Who changes a routing rule in practice
	actions reach external systems.		membership, not routing logic.
Five9	Moderate. Studio 7 covers common patterns; unusual logic goes to services.	Owens and resells. Strongest in North America, thinner outside it.	Admin. Simpler than NICE or Genesys, which is the point.
Talkdesk	Moderate to good. Studio is approachable; Navigator adds intent-based routing.	Resells with a global carrier mesh. Coverage claimed broadly, verify per country.	Admin, with a lower skill floor than the two suites.
RingCentral RingCX	Shallow to moderate. Deliberately constrained for time-to-value.	Owens the network. This is RingCentral's actual asset.	Admin, and the shallow model means fewer people are needed.
Verint (CX Automation)	Not applicable. Verint does not sell a primary routing engine.	None. Sits behind whoever owns the voice path.	Not applicable.
Calabrio ONE	Not applicable. Workforce platform, now inside Verint.	None.	Not applicable.
Salesforce Service Cloud	Omni-Channel routing is capable for cases and digital; voice routing is Amazon Connect underneath.	Resells. Service Cloud Voice runs on Amazon Connect or a partner BYOT.	A Salesforce admin, which is usually a different team from contact centre operations.
Dynamics 365 Contact Center	Good for digital and case routing; unified routing is real. Voice is Azure Communication Services.	Resells via Azure Communication Services, priced separately from the seat.	A Power Platform or Dynamics admin. Rarely the contact centre.
Teams-attached contact centre	Depends entirely on the partner product bolted on. Ranges from shallow to deep.	Microsoft Teams Phone plus partner. Two vendors in the voice path.	Partner admin, and escalations cross a vendor boundary.
Amazon Connect	Unlimited, because contact flows plus Lambda is code. No ceiling and no guardrails.	Owens. AWS is the carrier. Regional coverage narrower than the suites.	An engineer. This is a feature or a disqualifier depending on your staffing.
Google CCAI / CES	Unlimited in the same sense. Conversational Agents plus your own orchestration.	Partner or bring your own. Google does not own the voice path end to end.	An engineer.

ASSESSMENT The routing column is where most evaluations waste their time, because all four suites clear the bar for almost every real requirement. The column that decides deployments is the last one. An operation that needs to move traffic at nine in the morning because a competitor's outage has doubled its volume, and whose only person capable of changing a routing rule sits in a different function with a two-week change queue, has bought a slow contact centre regardless of what the routing engine can express. This is the single most consistent post-deployment complaint across client estates in the author's experience, it is invisible on every

feature comparison, and it is not a product defect. It is a consequence of how deep the customisation model goes.

Omnichannel, reliability, reporting and residency

Platform	Single queue across channels	Deployment and data residency	Real-time reporting
NICE CXone Mpower	Genuinely single queue. Voice and digital share routing state.	Multiple AWS regions plus NICE-operated processing centres; FedRAMP voice points of presence in the US.	Strong. Dashboards refresh fast enough to run a floor from.
Genesys Cloud CX	Genuinely single queue, and the cleanest of the four suites at it.	21 AWS regions as published, plus an EU sovereign region in Brandenburg on the AWS European Sovereign Cloud.	Strong. Analytics API is open enough to build your own.
Five9	Mostly single queue. Digital arrived later and occasionally shows its seams.	Fewer regions than NICE or Genesys. Verify residency per country before signing.	Adequate. Historically the weakest of the four suites at real time.
Talkdesk	Single queue for the channels it supports. Fewer channels than the suites.	Multiple regions claimed; verify the specific one you need is generally available.	Good, and Talkdesk's reporting is better than its market position suggests.
RingCentral RingCX	Voice-first with digital attached. Honest about being voice-first.	Follows RingCentral's own footprint, which is broad for voice.	Adequate for the operation size it targets.
Verint (CX Automation)	Not applicable; ingests interactions from the routing platform.	Cloud and on-premises. On-premises is a real option here, which matters for some regulated buyers.	Not applicable to routing; workforce dashboards are strong.
Calabrio ONE	Not applicable.	Cloud, on AWS.	Workforce dashboards only.
Salesforce Service Cloud	Cases and digital share a queue. Voice joins it through the Amazon Connect integration.	Salesforce Hyperforce regions; the voice leg's residency is Amazon Connect's, not Salesforce's.	Good for case metrics, weaker for voice-native real-time.
Dynamics 365 Contact Center	Unified routing across channels is real. Voice residency follows Azure Communication Services.	Azure regions; the AI leg follows the tenant's Copilot configuration, not the contact centre's.	Good, through Dynamics and Power BI rather than a purpose-built wallboard.
Teams-attached contact centre	Depends on the partner. Frequently two queues wearing one skin.	Two vendors, two residency positions. Reconcile them explicitly.	Partner-supplied. Quality varies more than any other row here.
Amazon Connect	Single queue if you build it that way. Nothing stops you building it badly.	Published as available in a subset of AWS regions, materially fewer than AWS overall.	You build it. Contact Lens plus your own pipeline. Excellent or absent.

Platform	Single queue across channels	Deployment and data residency	Real-time reporting
Google CCAI / CES	You assemble it. There is no default single queue.	Google Cloud regions; strong residency controls, weak turnkey contact centre.	You build it.

ASSESSMENT On published reliability and outage history, this matrix reports an absence rather than a ranking, and the absence is the finding. No vendor in this set publishes measured availability against an independent audit. Every vendor operates a status page that it controls, on which it decides what counts as an incident and when the incident started and stopped. The third-party aggregators that appear in search results scrape those same pages, so their outage histories are re-reported vendor declarations wearing the clothes of independent measurement. A buyer who wants a reliability comparison has exactly one honest route: require, contractually, the last twenty-four months of incident records at the specific region and service the buyer will occupy, and make the vendor's answer a warranty rather than a slide.

FACT What is checkable is the contractual remedy, and it is worth reading before it is needed. Genesys publishes its service level agreement openly, which is more than most of this set does. The commitment is to "use commercially reasonable best efforts to provide 100% uptime," and the remedy is a credit against subscription fees: 10 percent if uptime falls below 99.99 percent in a month, 30 percent below 99.0 percent, 100 percent below 97 percent. The customer must request the credit within thirty days and have raised a support case validating the impact, and credits apply only to annual prepay or annual month-to-month contracts. Failures traced to the customer's own network or directly contracted telecoms are excluded. ^[32]

ASSESSMENT Work the arithmetic on that remedy. A thousand-seat operation on Genesys Cloud CX 3 at the published \$155 pays \$155,000 a month at list. A month containing a single outage long enough to breach 99.99 percent, which is about four and a half minutes, returns a credit of \$15,500. A month containing an outage long enough to breach 99.0 percent, roughly seven hours, returns \$46,500. Seven hours of a thousand-seat operation being unable to take calls costs far more than \$46,500 in abandoned contacts, agent idle time and service level penalties on the operation's own downstream commitments. This is not a criticism of Genesys, whose disclosure is better than its peers'. It is the shape of every service level agreement in this category, and the point is that the standard remedy is not designed to make a buyer whole and should not be treated as risk transfer.

The workforce layer: what is commoditised and what still separates

ASSESSMENT Recording, adherence and basic quality management are commoditised. Every platform in this set that sells workforce management records interactions, tracks schedule adherence and provides an evaluation form builder, and the differences between them do not change a buyer's outcome. Anyone still scoring vendors on recording is scoring a solved problem. Long-interval forecasting for a stable voice queue is close to commoditised too; the published algorithms are old, well understood and broadly equivalent.

ASSESSMENT Four things still separate vendors and each of them changes an outcome. First, multi-skill and multi-channel scheduling under real constraint, which is where Verint and Calabrio earned their positions and where the suite-bundled modules remain visibly thinner. Second, intraday reforecasting that actually reruns rather than merely alerting. Third, speech and interaction analytics that produce categories an operations manager can act on rather than a sentiment score, which is a data-science quality difference and not a feature difference. Fourth, coaching workflow that closes the loop from an evaluation to a scheduled coaching session to a re-evaluation, which sounds like process software because it is, and which is where the labour saving actually sits.

Platform	Where its workforce management came from	Forecasting and scheduling under constraint	Why a buyer would still buy a specialist alongside it
NICE CXone Mpower	Native, and NICE's oldest business. This is not a bolt-on.	Strong. Competitive with the specialists on most operations.	Rarely worth it. NICE's own workforce module is the reason to buy NICE.
Genesys Cloud CX	Native, built into the platform, sold as a tier or an add-on.	Good, and improving. Thinner than NICE on complex multi-skill.	Above roughly 2,000 seats with heavy multi-skill, sometimes.
Five9	Acquired and integrated. Better than it was, still visibly later than the routing.	Adequate for a stable voice-heavy queue. Strained by complex skill mixes.	Often, and this is the most common specialist attach in this set.
Talkdesk	Built later than the routing engine. Present in the Elite tier.	Adequate for mid-market. Not a reason to choose Talkdesk.	Frequently, above a few hundred seats.
RingCentral RingCX	RingWEM, recent and positioned at the operation size RingCX targets.	Basic. Honest about it.	Almost always, if the operation is large enough to have a real workforce problem.
Verint (CX Automation)	Native and the company's core competence for two decades.	Best in this set alongside NICE, particularly under constraint.	This is the specialist.
Calabrio ONE	Native. Now a product line inside Verint.	Strong, and historically easier to administer than Verint's own.	This is also the specialist, which is now the problem.
Salesforce Service Cloud	None of its own. Workforce Engagement is thin and rarely the reason anyone buys.	Weak. Not a workforce platform.	Effectively always, for any operation where labour cost is the problem.
Dynamics 365 Contact Center	Limited native capability; Microsoft leans on partners and on Viva for adjacent needs.	Weak.	Effectively always.
Teams-attached contact centre	Partner-supplied, or absent.	Varies from strong to nonexistent depending on the partner.	Depends entirely on which partner.

Platform	Where its workforce management came from	Forecasting and scheduling under constraint	Why a buyer would still buy a specialist alongside it
Amazon Connect	Forecasting, capacity planning and scheduling, priced at \$27 per agent per month.	Functional and improving. Not a match for NICE or Verint under constraint.	Above a few hundred seats with complex skills, usually.
Google CCAI / CES	None. Google does not sell workforce management.	Not applicable.	Always, because there is nothing to compare.

ASSESSMENT That last column is the honest answer to "should I buy the specialist or take what the suite bundles," and the crossover is smaller than vendors on either side will tell you. Below roughly 300 seats, a suite's bundled workforce module is almost always enough and the specialist's licence, implementation and second admin skill set are not recoverable. Between 300 and 1,500 seats it depends almost entirely on skill complexity rather than headcount: a 600-seat operation with four skills does not need a specialist, and a 400-seat operation with thirty skills, three languages and split shifts probably does. Above 1,500 seats with genuine multi-skill complexity, the specialist usually pays for itself on schedule efficiency alone, and the argument stops being about software and becomes an argument about whether you have the workforce management analysts to run it. Buying Verint without staffing it is the most expensive way to not solve a scheduling problem.

The AI layer: where it runs, who owns the model, and what breaks at renewal

Platform	Where the AI runs and whose model	Agentic resolution into back-end systems	Renewal exposure if the model provider changes terms
NICE CXone Mpower	Proprietary Enlighten models trained on NICE's own interaction data, plus Cognigy acquired outright in September 2025, plus AWS Bedrock, Q and Nova integration.	Yes, through Cognigy and Mpower Agents.	Lowest in this set. NICE owns the models that matter and bought its conversational vendor rather than partnering with it.
Genesys Cloud CX	Native models plus AI Studio compatible with proprietary, open-source and Amazon Bedrock frontier models including OpenAI, Anthropic and Google.	Yes, agentic virtual agents metered at 1.2 tokens per interaction.	Moderate. Model choice is an advantage until the frontier model you standardised on reprices, and the token rate is Genesys' to change.
Five9	Own AI layer plus partner models; the AI portfolio is sold as Five9 Genius and priced into the upper tiers.	Yes, Voice AI Agents launched in the second quarter of 2026.	Moderate, and unusually opaque because the three tiers that carry it are quote-only.

Platform	Where the AI runs and whose model	Agentic resolution into back-end systems	Renewal exposure if the model provider changes terms
Talkdesk	Own layer plus partner models. Autopilot, Copilot and Navigator.	Partial. Automation exists; deep back-end action is thinner.	Moderate, with the additional exposure that Talkdesk is the smallest independent balance sheet in this set.
RingCentral RingCX	Partner models, with an OpenAI relationship RingCentral cites publicly in its own results.	Emerging. AIR Pro and AVA added agentic capability during 2026.	Highest single-provider exposure in this set. One model provider's terms change is RingCentral's problem and then yours.
Verint (CX Automation)	Own models for behavioural scoring and analytics, plus bots for specific tasks.	Limited by design. Verint is not in the routing path.	Low, because the models do narrow tasks on data you already have.
Calabrio ONE	Own analytics models, now converging with Verint's.	Limited.	Low on models. High on roadmap, for a different reason.
Salesforce Service Cloud	Agentforce on the Atlas reasoning engine, with multiple underlying model providers.	Yes, and this is Salesforce's genuine strength: the agent acts inside the system of record.	Moderate. Salesforce absorbs model cost into Flex Credits, which insulates you from model repricing but hands Salesforce the pricing lever.
Dynamics 365 Contact Center	Copilot, on Microsoft's model portfolio, metered in tenant-pooled Copilot Credits.	Yes, through Copilot Studio agents and Power Platform connectors.	Low on model risk, high on meter risk. Microsoft controls both the models and the credit conversion rate.
Teams-attached contact centre	Whatever the partner uses, plus whatever Microsoft supplies. Two AI stories in one deployment.	Fragmented across two vendors.	Highest structural risk, because neither vendor owns the whole path.
Amazon Connect	Amazon Q in Connect and Contact Lens, on AWS models, with Bedrock available for anything custom.	Yes, if you build it. No ceiling.	Low on lock-in, because you can swap the model. High on effort, because you own the swap.
Google CCAI / CES	Gemini, on Google's own infrastructure. First-party throughout.	Yes, if you build it.	Low. Google owns the model and the serving layer, so there is no third party to reprice.

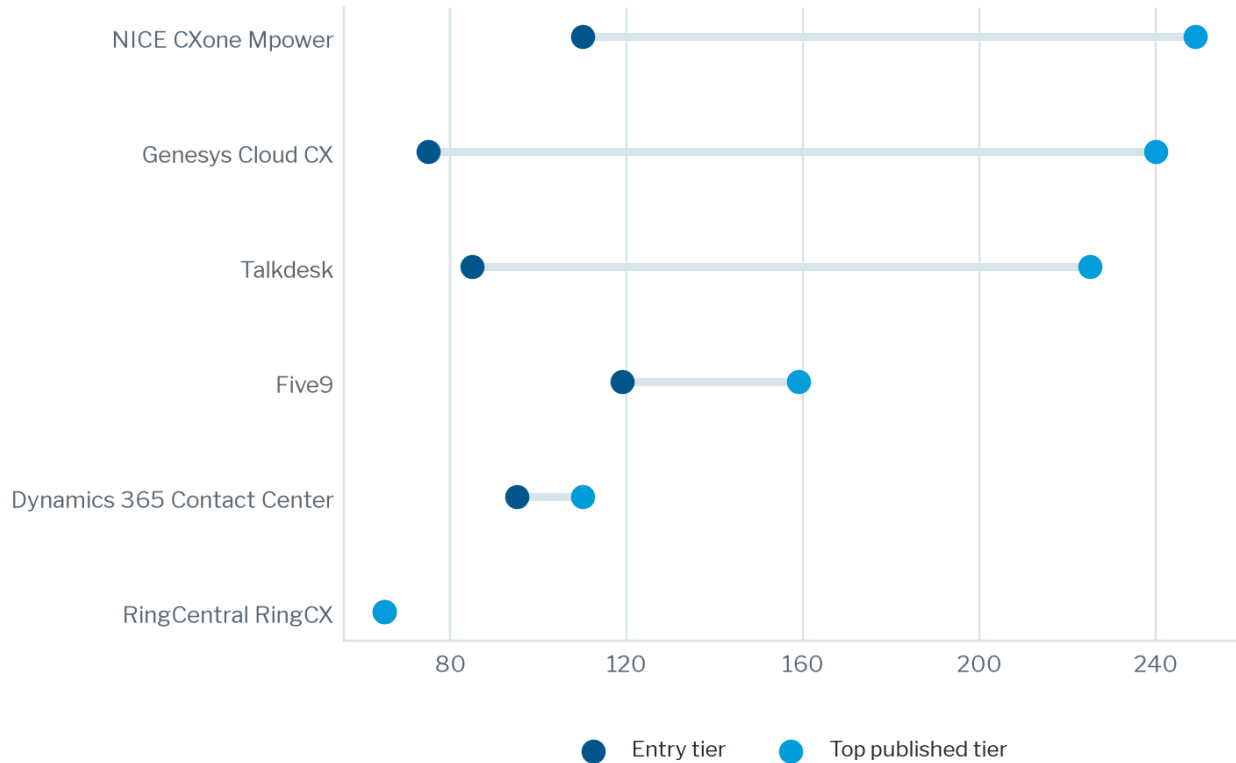
ASSESSMENT The model provenance column is the one that will look most different in twelve months, and it is the one buyers currently under-weight most. The question to ask a vendor is not which model it uses. It is: if your model provider raises its price by 40 percent mid-term, does my rate change, and where in the contract does it say so. Three answers are possible. NICE and Google own the model, so the answer is a commercial decision rather than a pass-through.

Salesforce and Microsoft absorb the model into a credit unit whose conversion rate they control, which protects you from the provider and exposes you to the platform. RingCentral, Talkdesk and Five9 depend materially on a third party, and the honest answer is that the pass-through terms are not published. Genesys sits in between by design, letting you bring your own model, which converts a pricing risk into an engineering responsibility. ^[40] ^[41]

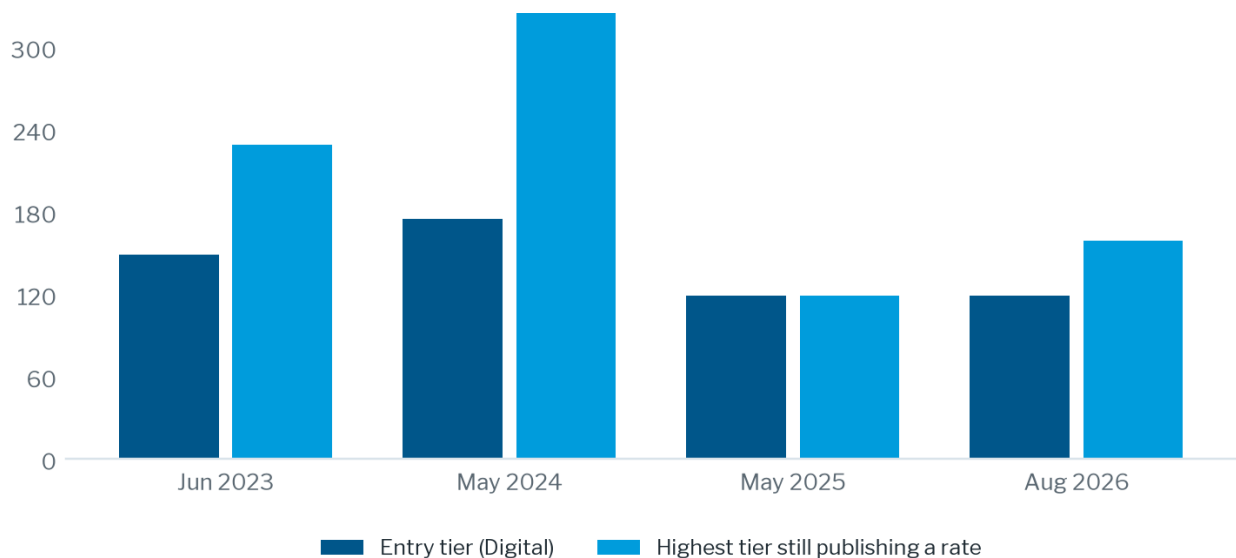
Commercial: what is published, what is metered, and where the AI sits

Platform	Published seat list price, 26 Aug 2026	What is metered separately	Is AI in the seat	Published or quote-only
NICE CXone Mpower	\$110 to \$249 per agent per month across five suites	\$0.25 per session on Ultimate Suite; Enlighten capabilities by package	Partly. Tier-dependent, with a per-session meter at the top	Published, and unchanged since at least September 2024
Genesys Cloud CX	\$75 / \$115 / \$155 / \$240 per user per month, billed annually	AI Experience tokens at \$1.00; voice per minute; text-to-speech per million characters; SMS per segment	No. Token allocations are included, consumption above them is metered	Published, including a separate pricing hub with token conversion rates
Five9	\$119 Digital, \$159 Core. Plus, Pro and Enterprise are Contact Sales	Not published for the tiers that carry AI	Unknown from published material	Partly published. The three AI-bearing tiers went quote-only between May 2024 and May 2025
Talkdesk	\$85 / \$105 / \$165 / \$225 per user per month	AI capabilities present in tiers; separate metering not published	Tier-dependent, not itemised	Published, and unchanged since at least June 2025
RingCentral RingCX	\$65 per agent per month at general availability, including unlimited domestic minutes	AI products sold as paid add-ons; rates not published	Partly. AI summaries included; agentic products priced separately	Partly published. The contact centre pricing page renders no rate to a plain fetch
Verint (CX Automation)	None published	Not published	Not published	Quote-only
Calabrio ONE	None published	Not published	Not published	Quote-only
Salesforce Service Cloud	Edition prices are rendered client-side and were not retrievable from a Salesforce primary source on this date	Agentforce Flex Credits at \$500 per 100,000, 20 credits (\$0.10) per action; or \$2.00 per conversation	No. Agentforce sits on top of a Service Cloud seat	Partly published, and harder to verify than any other vendor here
Dynamics 365 Contact Center	\$110 per user per month paid yearly;	Copilot Credits at \$0.01, pooled tenant-wide; voice via Azure	No. Credits are sold separately and shared	Published, including a 12-page credits guide

Platform	Published seat list price, 26 Aug 2026	What is metered separately	Is AI in the seat	Published or quote-only
	Digital-only and Voice-only at \$95	Communication Services, priced separately	with the rest of the tenant	
Teams-attached contact centre	Teams Phone seat plus partner seat. Two bills.	Both vendors meter independently	No	Partly, and never in one place
Amazon Connect	No seat fee at all	Everything: \$0.018 per voice minute, \$0.004 per chat message, \$12.00 per agent evaluated per month, \$27 per agent per month for forecasting and scheduling, \$0.12 per case created, and about twenty more meters	Metered: \$0.0080 per voice minute for self-service, \$0.0080 per minute for real-time agent assistance	Fully published, to four decimal places
Google CCAI / CES	No contact centre seat. Components only	\$0.007 per chat request on Flows, \$0.012 on generative Playbooks; \$0.001 and \$0.002 per second of voice respectively	Metered, and by a wide margin the cheapest generative self-service on published rates	Fully published

FIGURE 2 Published per-seat list prices, entry tier to top published tier, 26 August 2026

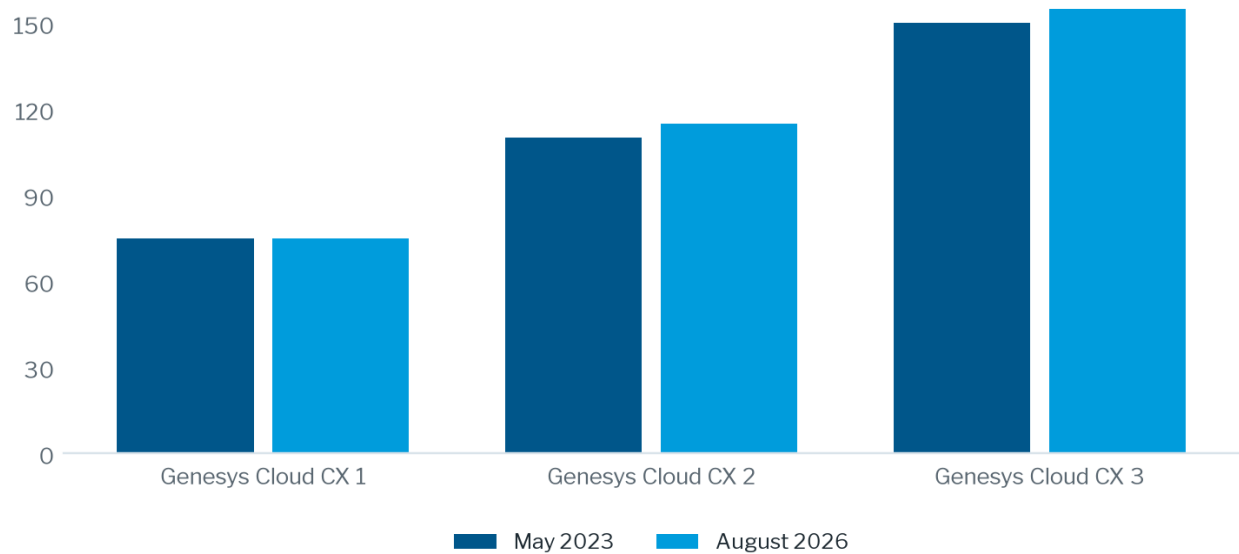
FACT Each vendor's own pricing page on 26 August 2026, or its own general availability release in RingCentral's case. Five9's top published tier is \$159 only because the three tiers above it are quote-only, so the bar understates its real range. NICE's top tier additionally carries \$0.25 per session. Verint, Calabrio and Salesforce Service Cloud editions are absent because no per-seat rate was retrievable from a vendor primary source on this date. Amazon Connect and Google are absent because they charge no seat fee. List prices only; nobody at scale pays these. [\[1\]](#) [\[5\]](#) [\[9\]](#) [\[11\]](#) [\[14\]](#) [\[18\]](#)

FIGURE 3 Five9's published list price, and the year the top tiers stopped publishing one

FACT Internet Archive snapshots of five9.com/products/pricing, plus the live page. In June 2023 and May 2024 all five tiers carried a published rate. From May 2025 the top three, which carry workforce engagement and AI, read "Contact Sales." The

2025 and 2026 upper bars therefore measure disclosure, not price: the actual top tier is unknown and is almost certainly above \$325. Reading this chart as a price cut is the error it exists to prevent. ^{[5] [6] [7] [8]}

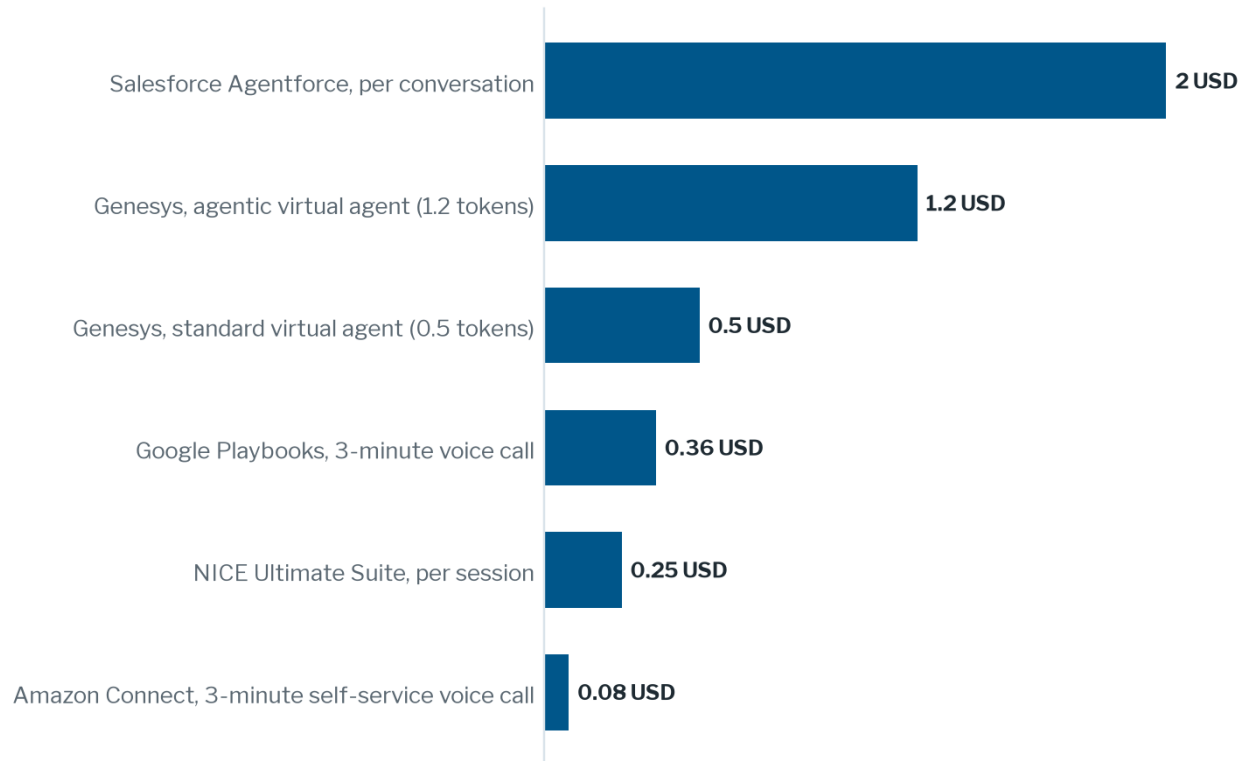
FIGURE 4 Genesys held its seat prices nearly flat and moved AI onto a meter instead



FACT Genesys' own pricing page, May 2023 snapshot against the live page. Over the same three years Genesys retired the flat AI Experience add-on, which was \$40 per agent per month in both the 2023 and 2024 snapshots, and replaced it with AI Experience tokens at \$1.00 each, of which an agentic virtual agent interaction consumes 1.2. A fourth tier, CX 4, appeared at \$240. The seat moved five dollars; the AI moved off the seat entirely. ^{[1] [2] [3] [4]}

What AI costs when you read the meters

ASSESSMENT Nobody sells against this number because nobody publishes it in a comparable form, but it is the number that decides whether an AI deflection business case survives contact with a finance team. Take one customer interaction handled entirely by AI, no human, and cost it on each vendor's published list meters. The vendors have chosen units that make the comparison awkward on purpose, so the arithmetic below normalises them and shows its working.

FIGURE 5 One AI-handled interaction, costed on published list meters

ASSESSMENT The arithmetic is the author's; the rates are the vendors' own published list rates on 26 August 2026. Google: 180 seconds at \$0.002 per second on a generative Playbooks voice agent, excluding telephony. Amazon Connect: 3 minutes of end-customer self-service at \$0.0080 per minute plus 3 minutes of voice at \$0.018, excluding the DID and any Contact Lens analysis. NICE and Salesforce figures sit on top of a per-seat fee that the Google and Amazon figures do not carry, so the comparison is of marginal cost per interaction, not of total cost of ownership. Microsoft is absent because Copilot Credit consumption varies by task and Microsoft publishes no per-interaction rate. [\[1\]](#) [\[2\]](#) [\[9\]](#) [\[13\]](#) [\[16\]](#) [\[17\]](#)

ASSESSMENT Two things follow. First, a deflection business case built on "AI handles the contact for pennies" is true on Amazon and Google published rates and false on Salesforce and Genesys published rates, and the difference is roughly twenty-five to one at the extremes. An operation deflecting two million contacts a year is looking at about \$160,000 on Amazon Connect's meters and about \$4 million on Salesforce's per-conversation rate. Second, and less obviously, the vendors at the expensive end are not overcharging for compute. They are charging for the orchestration, the integration into the system of record, and the fact that you did not have to build any of it. Whether that is worth twenty-five times the marginal cost depends entirely on whether you have engineers, which is the question the assemble-versus-buy section returns to.

ASSESSMENT The third thing that follows is a contracting instruction rather than an analysis. Every one of these meters converts a fixed cost into a variable one, and the variable moves with contact volume, which is the one thing a contact centre cannot control. A seat licence is a bad deal in a downturn and a good one in a surge. A per-interaction meter is the reverse. An operation whose volume doubles in a bad quarter has just doubled its AI bill in the quarter it can least afford it, and no published contract in this set caps that. Ask for the cap. It is the single most valuable thing in the contract-language section and it costs nothing to request.

What moved when AI arrived

FACT Checked against archived snapshots of each vendor's own pricing page, the seat did not absorb the AI. NICE's five published suite prices, \$110, \$135, \$169, \$209 and \$249, are identical in a September 2024 snapshot and on the live page today, a period covering the launch of CXone Mpower and the \$955 million acquisition of Cognigy. Talkdesk's four published rates, \$85, \$105, \$165 and \$225, are unchanged since at least June 2025. Genesys moved CX 2 and CX 3 by five dollars each in three years and added a new top tier. Five9 raised every published tier sharply between June 2023 and May 2024, then withdrew publication from the three tiers that carry AI.

[3] [4] [5] [6] [7] [8] [9] [10] [11] [12] [29]

ASSESSMENT The pattern is consistent and it is a strategy rather than a coincidence. Published seat rates are the number that appears in a renewal conversation and in a procurement benchmark, and raising them invites a fight with an incumbent customer. Meters do not appear in that conversation. So the industry put AI on a meter, held the seat, and moved the growth into a line item that renews automatically with consumption rather than annually with negotiation. NICE's per-session charge on Ultimate Suite, Genesys' tokens, Salesforce's Flex Credits and Microsoft's Copilot Credits are four expressions of the same decision. Five9's choice to stop publishing at all is the fifth, and it is the most revealing, because withdrawing a price is what a vendor does when the number would raise a question it would rather answer in a sales conversation.

ASSESSMENT For a buyer, the practical consequence is that year-on-year seat price benchmarking has stopped being a useful control. An operation that congratulates itself on holding its seat rate flat at renewal while its token, credit or session consumption trebles has lost the negotiation and recorded a win. The control that replaces it is a blended cost per handled contact, calculated across seats, meters and telephony together, tracked monthly, and put in front of whoever signs. That number is computable from any of these vendors' billing data and almost nobody computes it.

Exit cost, the differentiator nobody markets

ASSESSMENT Switching cost is a differentiator and no vendor markets it, for the obvious reason. It is also the dimension on which this vendor set differs most, and the one buyers assess least, usually because it feels like a problem for a future team. It is not a single number. It decomposes into four things that behave differently, and the fourth is the one that surprises people.

Platform	Integration depth once live	Historical interaction data portability	Routing logic portability	Realistic elapsed time to leave, 1,000 seats
NICE CXone Mpower	Deep. Studio scripts, custom integrations, Enlighten models trained on your data	Recordings and metadata exportable; the trained model behaviour is not	Low. Studio does not translate to anything	12 to 18 months

Platform	Integration depth once live	Historical interaction data portability	Routing logic portability	Realistic elapsed time to leave, 1,000 seats
Genesys Cloud CX	Deep. Architect flows, data actions, custom analytics pipelines	Good APIs, and Genesys is better than most at bulk export	Low. Architect flows are Genesys-shaped	12 to 18 months
Five9	Moderate. Fewer integration surfaces means fewer to unpick	Adequate. Export is a project, not an API call	Low to moderate	9 to 12 months
Talkdesk	Moderate	Adequate	Low to moderate	9 to 12 months
RingCentral RingCX	Shallow by design, which is the strongest exit position among the seat vendors	Adequate	Moderate, because there is less logic to move	6 to 9 months
Verint (CX Automation)	Deep on data, shallow on the call path	Recordings and evaluations exportable; scoring model history is not	Not applicable	9 to 12 months, and it does not interrupt the contact centre
Calabrio ONE	Same shape as Verint	Same	Not applicable	9 to 12 months
Salesforce Service Cloud	The deepest in this set, because leaving means leaving the system of record	Case data is portable. The rest of the org is the problem	Omni-Channel config does not travel	18 to 36 months, and frequently never
Dynamics 365 Contact Center	Deep, for the same reason, plus Power Platform sprawl	Dataverse export is real but the surrounding automation is not	Unified routing config does not travel	18 to 36 months
Teams-attached contact centre	Two shallow integrations rather than one deep one	Partner-dependent	Partner-dependent	6 to 12 months, the lowest exit cost in this set
Amazon Connect	As deep as you built it, which is usually very deep	Total. It is your S3 bucket and your data model	Zero portability of contact flows; total portability of the Lambda logic behind them	12 to 24 months, and the variance is about your own code
Google CCAI / CES	As deep as you built it	Total	Same shape as Amazon	12 to 24 months

ASSESSMENT The elapsed-time column is the author's assessment from operator-side experience of migrations across multiple client estates, not a measured figure, and it should be read as a planning range rather than a number. What it captures is that migration duration is driven by parallel-running and requalification, not by data movement. Moving five years of recordings is a week of engineering. Rebuilding two hundred routing rules, revalidating them against the

compliance evidence that justified them, retraining fourteen hundred agents, and running both platforms in parallel long enough to prove the new one holds service level is the twelve months. Any vendor migration plan that does not budget for a parallel-run period is a plan to discover the problem in production.

ASSESSMENT The fourth component, the one buyers miss, is the trained model. An Enlighten behavioural model or a Contact Lens category set that has been tuned against three years of your interactions is an asset you cannot export and cannot rebuild quickly at the new vendor, because rebuilding it requires the interactions to have accumulated on the new platform first. This did not exist as an exit cost five years ago. It now adds something on the order of six to twelve months of degraded analytical capability on top of the migration itself, and it grows with every year the incumbent stays. It is the strongest lock-in mechanism any vendor in this set has ever had, none of them advertises it, and it is the reason the exit-cost column will look worse in every future edition of this document than it does in this one.

Nine buying situations, and the call on each

Each of these names a platform, the business reason, the trade-off being accepted, and the condition that would change the answer. Where the honest answer is that it depends, the thing it depends on is named rather than hedged.

1. Large enterprise, global voice, complex routing

ASSESSMENT Genesys Cloud CX. The business reason is single-queue omnichannel that is genuinely single-queue plus 21 published AWS regions including an EU sovereign region, which is the combination that survives a routing requirement written by someone who understands the business. NICE is the close second and wins where workforce management complexity outweighs routing complexity. The trade-off accepted with Genesys is a token meter with no published cap and an administration model that puts routing changes behind a trained specialist. What changes the answer: if the operation's genuine problem is scheduling twelve thousand agents across thirty skills rather than routing them, buy NICE, because NICE's workforce management is its oldest and best business and Genesys' is not.

2. Mid-market, a few hundred seats, no in-house platform engineering

ASSESSMENT Five9 or Talkdesk, and between them it is close enough that price should decide. The business reason is that both are engineered for an operation without a platform team: the customisation ceiling is lower and that is the feature, because a lower ceiling means a shorter implementation and an administrator rather than an engineer. Talkdesk publishes its full price ladder and Five9 does not, which is worth something at this size where negotiating leverage is thin. The trade-off is that both workforce management modules are adequate rather than good, and both will strain if skill complexity grows. What changes the answer: if the operation is already Salesforce-standardised, skip to situation three; if the seat count is heading past a thousand within the contract term, buy the suite now rather than migrating twice.

3. A Salesforce-standardised organisation

ASSESSMENT Service Cloud with Amazon Connect underneath for voice, not Service Cloud alone and not a separate suite alongside. The business reason is that agentic resolution inside the system of record is the one thing Salesforce does that nobody else in this set does as well, and every alternative reintroduces the two-system problem that Salesforce standardisation existed to remove. The trade-off is severe and should be stated plainly: voice routing capability is materially below Genesys or NICE, the Agentforce meter at \$2.00 per conversation or \$0.10 per action is the most expensive published AI rate in this set, and exit cost is the highest here because leaving means leaving the system of record. What changes the answer: complex global voice routing. If the routing requirement is genuinely hard, run Genesys with the CX Cloud integration into Service Cloud and accept two platforms, which is what the \$750 million Salesforce put into Genesys in July 2025 suggests Salesforce itself expects for that segment.

4. A Microsoft-standardised organisation

ASSESSMENT This is the one where the honest answer depends, and here is what on. Dynamics 365 Contact Center at \$110 per user per month is the right call if the organisation already runs Dynamics as its service system of record and has Power Platform governance in place. It is the wrong call if "Microsoft-standardised" means Teams and Microsoft 365 but a different CRM, in which case you are buying a system of record you did not want in order to get a contact centre. The Teams-attached route with a partner contact centre is the right call for the second group and carries the lowest exit cost in this whole set, at the price of two vendors in the voice path and escalations that cross a support boundary. The trade-off in both cases is the Copilot Credit pool: contact centre AI consumption is metered at \$0.01 per credit against a pool shared with every other Copilot workload in the tenant, which means the contact centre's costs are visible to, and competing with, the whole company. What changes the answer: whether your Microsoft 365 administrator will grant the contact centre a ring-fenced spend policy. If yes, Dynamics is straightforward. If no, do not put a customer-facing service on a budget you do not control.

5. A heavily outsourced or BPO-delivered operation, multiple client programmes on shared infrastructure

ASSESSMENT Genesys, with NICE a genuine second, and the reason is neither routing nor AI. It is multi-tenancy discipline: the ability to partition configuration, reporting, recording and data residency by client programme without deploying a separate instance per client, and to prove that partition to a client's auditor. This is the requirement that separates suites from everything else in this set, and it is the reason the assemble-it-yourself options are usually wrong here despite being cheaper. The trade-off is that per-seat licensing across a portfolio of client programmes with different volumes is a poor commercial fit for a business whose own revenue is per-hour or per-transaction, and the mismatch is absorbed by the operator. What changes the answer: a client contract that mandates its own platform, which happens often enough in this segment that a BPO's real answer is usually "one strategic platform plus the ability to run client-mandated ones," and the platform decision then becomes a decision about which one you are best at operating.

6. A regulated operation with recording, consent and data residency constraints

ASSESSMENT Genesys on the AWS European Sovereign Cloud region in Brandenburg for an EU-constrained operation; NICE where the constraint is US federal, given its FedRAMP voice points of presence; Verint where the constraint requires on-premises recording, which is the one place in this set where an on-premises option is still a live and defensible answer rather than a legacy one. The business reason is that residency is a contractual and architectural question rather than a feature question, and the vendors that publish a region list and a sovereignty position are the ones you can hold to it. The trade-off is cost and, for the sovereign region specifically, a service catalogue that lags the commercial regions. What changes the answer, and it now changes it for everyone: EU AI Act Article 50 became applicable on 2 August 2026 and applies to systems already deployed, so any conversational AI in the customer path needs a disclosure design and an evidence trail regardless of which platform it runs on. Choose the vendor that will warrant its part of that in the contract, not the one that documents it in a whitepaper.

7. The real purchase is workforce cost, and the contact centre platform is incidental

ASSESSMENT Verint, and this is the situation where the answer changed in November 2025. The business reason is unchanged: for an operation whose problem is that it employs too many people for the volume it handles, or schedules them badly, a specialist workforce platform pays for itself on schedule efficiency in a way no bundled module does, and the contact centre platform underneath is genuinely incidental. What changed is that Verint and Calabrio are now the same company under the Verint name, so the historic move of playing them against each other on price is gone, and one of the two product lines will eventually lose the roadmap argument. The trade-off is that the specialist only pays out if you staff it: Verint without workforce management analysts is an expensive way to not solve a scheduling problem. What changes the answer: seat count and skill complexity. Below roughly 300 seats, or above 300 with fewer than about ten skills, take the suite's bundled module and spend the difference on the analysts.

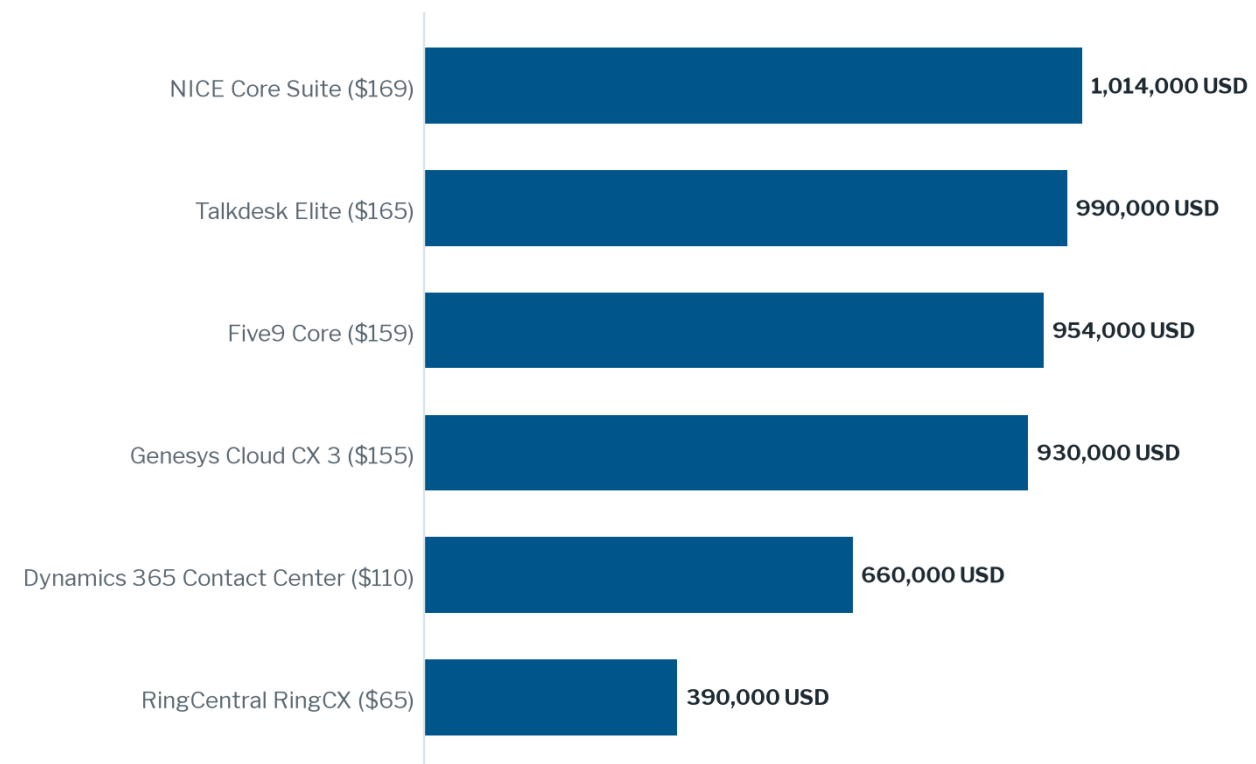
8. Buying primarily for AI deflection

ASSESSMENT Do not buy a contact centre platform for this. Buy the AI layer separately, meter it, and keep the routing platform you have for as long as the contract runs. The business reason is arithmetic: on published rates the marginal cost of an AI-handled interaction ranges from about eight cents to two dollars across this set, the platform you already own is not the constraint on deflection, and rebuilding your routing estate in order to change your bot is the most expensive route to a bot that exists. Google's Conversational Agents at \$0.012 per generative chat request, or Amazon Connect's self-service meters, will out-economise any suite-bundled AI by an order of magnitude if you have the engineering to integrate them. The trade-off is integration effort and a second vendor in the customer path. What changes the answer: no engineering capacity at all, in which case buy the AI from your incumbent suite, negotiate a consumption cap into the contract, and accept that you are paying for the integration you cannot build. What should also change the answer, and rarely does: run the containment definition audit in the next section before you sign anything, because a deflection business case built on the vendor's containment definition is a business case built on abandoned calls.

9. Real engineering capacity, weighing Amazon Connect or Google against buying a suite

ASSESSMENT Amazon Connect, if and only if you can name the four engineers by name and they are not currently fully committed. The business reason is that no seat fee plus published meters to four decimal places produces a cost base roughly half to a third of a suite at equivalent scale, total portability of your own interaction data, and no ceiling on routing logic. Google is the better answer where the AI is the point and the telephony is somebody else's problem, because Google owns the model and the serving layer and prices generative self-service far below anyone else here. The trade-off is that you are now a software team with a contact centre attached: on-call rotation, upgrade regression testing, and the reporting layer that Contact Lens does not give you for free. What changes the answer, and it is the only thing that does: staff turnover. The failure mode is not technical, it is that the two engineers who understood the contact flows leave within eighteen months and the operation discovers it owns a system nobody can safely change. A suite is, among other things, a hedge against your own attrition. Price that hedge honestly and the assemble-it-yourself case is weaker than its spreadsheet suggests.

FIGURE 6 What 500 seats costs per year on published list seat rates alone



ASSESSMENT Seat rate times 500 times 12, at published list, on 26 August 2026. Excludes telephony minutes, AI meters, workforce management add-ons where they are separate, professional services, and any discount. Nobody pays these numbers. The chart is here to establish the order of magnitude the meters sit on top of, and to make the point that the seat spread across the four suites is under ten percent while the AI meter spread across the same vendors is roughly twenty-five to one. [\[1\]](#) [\[5\]](#) [\[9\]](#) [\[11\]](#) [\[14\]](#)

The head-to-heads

NICE against Genesys at enterprise scale

ASSESSMENT The case for NICE: workforce management is NICE's oldest business rather than a module, Enlighten models are trained on NICE's own interaction corpus rather than licensed, the September 2025 acquisition of Cognigy for \$955 million bought the conversational layer outright instead of partnering for it, and NICE has held its published list prices flat since at least September 2024 through the entire Mpower cycle. For an operation where the labour line dominates and the routing requirement is merely hard rather than exotic, NICE is the better economics. The case for Genesys: the cleanest single-queue omnichannel implementation in this set, 21 published regions including an EU sovereign region, model flexibility including bring-your-own, and a published pricing hub that discloses token conversion rates down to 0.05 tokens per AI scoring invocation, which is a level of commercial transparency nobody else here offers. For an operation where routing complexity or residency is the binding constraint, Genesys is the better fit. The tiebreaker is not a feature. It is whether the AI meter or the workforce line is the bigger number in your five-year model, and both vendors will happily let you not do that arithmetic. [\[1\]](#) [\[2\]](#) [\[9\]](#) [\[10\]](#) [\[29\]](#) [\[37\]](#) [\[38\]](#)

Five9 and Talkdesk against both

ASSESSMENT Neither wins on capability at enterprise scale and neither is trying to. Both win on implementation time, administrative simplicity and a lower total delivery cost for an operation that does not have a platform team, and that is a real and defensible position that describes most of the market by count. Where they lose is global voice coverage, workforce management under complexity, and multi-tenant partitioning. Five9's specific weakness in August 2026 is commercial opacity: the three tiers that carry its AI and workforce capability have been quote-only for over a year, and a mid-market buyer negotiating without a published anchor is negotiating badly. Talkdesk's specific weakness is the balance sheet: the last disclosed funding round closed in 2021, headcount has fallen materially, and a buyer signing a five-year term with a private company that has not raised since 2021 should price change-of-control risk into the contract rather than assume it away. Both of these are honest reasons to prefer the other one, and neither is a product criticism. [\[5\]](#) [\[8\]](#) [\[11\]](#) [\[43\]](#)

RingCentral and the telephony-attached route

ASSESSMENT Attaching a contact centre to a telephony seat you already own is a saving for an operation under roughly two hundred seats whose contacts are overwhelmingly voice and whose routing needs are simple, and a downgrade for anyone else. RingCX at \$65 per agent per month including unlimited domestic minutes is genuinely cheap and the network underneath it is RingCentral's real asset rather than a resold one, which matters more than most feature comparisons acknowledge. The downgrade shows up in three places: workforce management is basic, the customisation ceiling is low, and the AI portfolio depends on a single external model provider more heavily than anything else in this set, which RingCentral itself cites publicly. The condition that flips it: skill complexity. The moment an operation needs more than about ten skills or genuine multi-channel queueing, the saving evaporates into workarounds and the operation

ends up buying a second platform anyway. What is genuinely underrated here is the exit position, which is the best of any seat vendor in this set precisely because there is less to unpick. [\[22\]](#) [\[44\]](#)

Salesforce against Microsoft, and what happens to voice

ASSESSMENT Both attach a contact centre to a seat the buyer already pays for, and both do the same thing to voice: they resell it. Salesforce Service Cloud Voice runs on Amazon Connect underneath, which means the voice quality and residency position are Amazon's and the buyer now has three vendors in one call path. Microsoft's Dynamics 365 Contact Center runs voice on Azure Communication Services, priced separately from the \$110 seat, which means the voice bill is a different line item from the contact centre bill and frequently a different budget. Routing is where they diverge more usefully than voice does. Salesforce's Omni-Channel is genuinely strong for cases and digital and weak for voice-native scenarios; Microsoft's unified routing is more balanced across channels but less deep than either suite. On AI, Salesforce charges the most per interaction of anyone in this set on published rates and delivers the deepest agentic action inside the record; Microsoft charges from a tenant-pooled credit balance the contact centre does not control. Pick Salesforce if the resolution has to happen inside the record and the voice requirement is modest. Pick Microsoft if the organisation genuinely runs on Dynamics and the Copilot budget can be ring-fenced. Pick neither if global voice is the hard part, and be aware when Salesforce tells you otherwise that Salesforce owns \$750 million of Genesys. [\[13\]](#) [\[14\]](#) [\[15\]](#) [\[16\]](#) [\[24\]](#)

Verint and Calabrio standalone against workforce management bundled into a suite

FACT This head-to-head no longer exists in the form the question assumes. Thoma Bravo completed its acquisition of Verint on 26 November 2025 and combined it with Calabrio, the combined organisation announced on 18 February 2026 that it would operate under the single Verint name with Calabrio solutions continuing inside the Verint CX Automation Platform, and on 24 February 2026 Calabrio's former chief executive was named chief executive of the combined firm. A shortlist that still runs Verint against Calabrio is comparing two product lines owned by the same private equity sponsor and managed by the same executive team. [\[25\]](#) [\[26\]](#) [\[27\]](#) [\[28\]](#)

ASSESSMENT The question that replaces it is specialist against bundled, and the crossover is roughly 300 seats with meaningful skill complexity, or roughly 1,500 seats regardless. Below that, the suite's bundled module wins on total cost because the specialist's licence, its implementation and the second administrative skill set are not recoverable from the schedule efficiency it produces. Above it, the specialist wins on schedule efficiency alone, provided the operation employs workforce management analysts who can use it. The new risk a buyer should now price, and could not have priced a year ago, is roadmap consolidation: two overlapping workforce platforms under one owner will converge, and a customer on the line that loses will face a migration inside its own vendor. Ask for a written product lifecycle commitment for whichever line you buy, with a minimum support horizon and a stated migration entitlement if it is sunset. That ask is cheap today and impossible after the convergence is announced.

Assemble it yourself against buying a suite

ASSESSMENT On published rates, assembling wins the cost argument decisively and loses the risk argument quietly. Amazon Connect charges no seat fee and publishes every meter, and at 500 seats a workload that costs roughly \$930,000 a year in Genesys seat licences alone costs a fraction of that in Connect meters before telephony. Google prices generative self-service at \$0.012 per chat request against Salesforce's \$2.00 per conversation. Those gaps are too large to be explained by margin, and they are real. What the spreadsheet does not carry is that you have bought a software product to operate rather than a service to consume: on-call, regression testing every time AWS changes a service, a reporting layer you build, and a single point of human failure in whoever wrote the contact flows. The suites are not selling routing. They are selling the absence of that, plus an escalation path with a name on it at three in the morning. The honest decision rule: assemble if you already run production software with an on-call rotation and can commit named engineers for the life of the platform, not the life of the project. Buy the suite otherwise, and stop treating the price gap as evidence you are being overcharged.

Containment, and what it actually counts

FACT The most repeated number in this category is containment, and its definition is almost never disclosed alongside it. Where it is disclosed, it is worse than buyers assume. Genesys' self-service statistics report defines the metric "Contained in Self-Service" as "the total number of interactions that entered the Designer application in Self-Service and were concluded without entering Assisted-Service." That definition is satisfied by a customer who solved their problem, by a customer who hung up in frustration, by a customer whose session timed out, and by a customer who gave up and used a different channel five minutes later. All four score identically.

[30]

VENDOR CLAIM Genesys' own marketing has since conceded the point, writing that "a high containment rate might look like a win, but that might not be the case if your customers didn't get what they needed," and that "a channel can appear successful on paper, showing high usage and low transfers, but it might still deliver poor experiences." That post proposes moving toward journey-level measurement rather than a replacement metric. A vendor publishing a critique of the metric its own product reports is unusual and worth crediting, and it does not change the fact that the product still reports the metric that way. [31]

ASSESSMENT No published containment figure in this category is measured rather than vendor-reported, and none that this author could locate is accompanied by both a definition and an independent audit. That is the finding, and it should be stated as an absence rather than filled with a number. The practical consequence is that a deflection business case built on a vendor-quoted containment rate is a business case built on an unknown mixture of resolutions and abandonments, and the ratio between the two is exactly what determines whether the deployment saves money or moves cost downstream into repeat contacts.

ASSESSMENT The audit that settles it takes about two weeks and any buyer can run it before signing. Take a month of self-service sessions the vendor scores as contained. Join them, by customer identifier, to every subsequent contact on any channel within the following seven days.

The proportion that reappear is your apparent-containment rate, and true containment is the vendor's figure minus that. Then split the remainder by how the session ended: completed a defined success action, timed out, or disconnected mid-flow. In this author's operator-side experience the gap between a vendor-quoted containment figure and a repeat-contact-adjusted one is routinely large enough to change the sign on a business case, though the specific magnitude varies enough by estate that no single number should be quoted for it. Require this join to be possible before signing, because on some architectures the identifiers do not survive the handoff and the audit becomes impossible after go-live.

ASSESSMENT Deflection and resolution deserve separating from containment and from each other, because vendors use all three interchangeably. Deflection is a channel-shift measure: the contact happened, it just happened somewhere cheaper. Containment is a no-agent measure and says nothing about outcome. Resolution is an outcome measure and is the only one of the three a chief financial officer should accept in a business case. A vendor quoting a containment rate as a resolution rate is not lying, it is using the industry's own vocabulary, and the buyer's job is to insist on the third definition and make the vendor's number testable against it.

AI-scored quality management, and the labour case behind it

ASSESSMENT This is the largest financial finding available in this category and it is buried under a feature claim. Traditional quality management samples a low single-digit percentage of interactions, typically two to five per interaction per agent per month, because a human reviewer takes twenty to forty minutes to score one call properly. Automated scoring claims complete coverage. If automated scores are usable, the labour case for a large quality assurance function collapses, and quality assurance headcount in a large operation is a bigger number than most of the platform differences in this document. Amazon Connect prices the capability at \$12.00 per agent evaluated per month, published, which is roughly an order of magnitude below the loaded cost of the analyst it replaces. ^[13]

ASSESSMENT The question is not coverage, which is real, but validity, which is unevidenced. No vendor in this set publishes a chance-corrected agreement statistic between its automated scores and trained human reviewers on the same interactions. Not a correlation, not an accuracy percentage, and specifically not a Cohen's kappa, which is the statistic that corrects for the agreement you would get by chance and which is standard practice in every other field that uses human raters. The absence is uniform across all twelve platforms here and it is not an oversight. Publishing a kappa invites comparison, and a kappa in the range the adjacent literature suggests would be an awkward number to put on a slide.

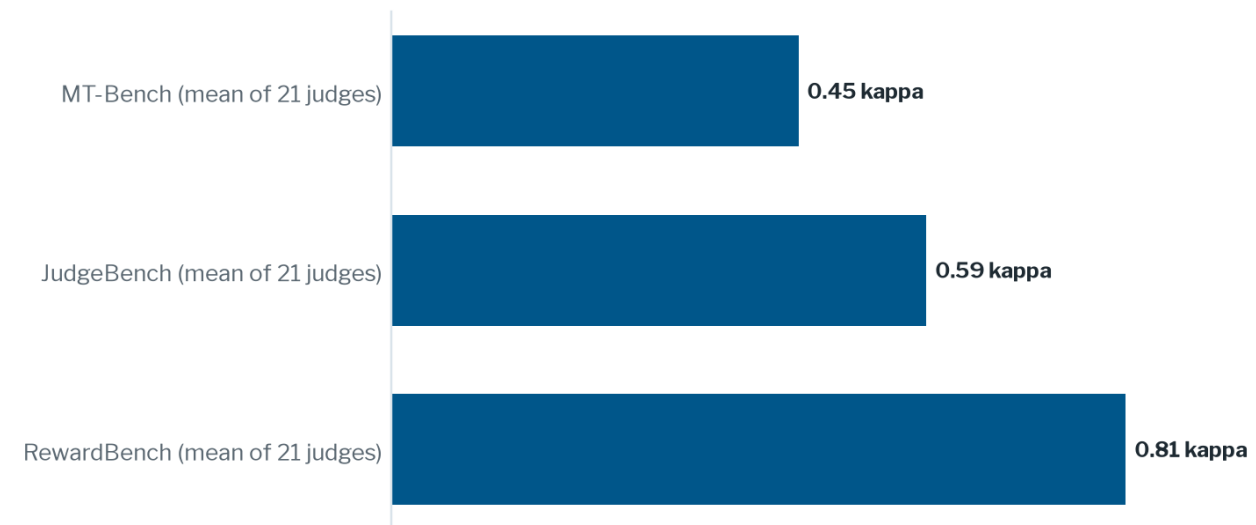
FACT The nearest adjacent evidence is the largest systematic evaluation of language models used as judges published to date, covering 21 judges from nine providers across three benchmarks, 118 runs and roughly 541,000 individual judgments. Mean Cohen's kappa against human labels was 0.451 on MT-Bench (range 0.376 to 0.511), 0.593 on JudgeBench and 0.814 on RewardBench. Test-retest reliability was very high, from 0.889 to 0.992 with a cohort mean of 0.944, while position bias ranged up to 0.192, and two judges achieved above 0.95 test-retest while exhibiting above 0.10 position bias. The paper names this the consistency-bias paradox:

reproducibility masking systematic invalidity. It also reports what it calls kappa deflation, that raw exact-match agreement overstated chance-corrected agreement by 33.8 to 41.3 percentage points across all 21 models on one benchmark. ^[35]

ASSESSMENT Read across to contact centre quality management carefully, because the tasks are not the same and the benchmarks are not conversations. Two things transfer. First, an automated scorer that is highly self-consistent is not thereby valid, and self-consistency is exactly what a vendor demo shows you. Second, any accuracy figure a vendor quotes that is not chance-corrected is likely to overstate agreement by tens of percentage points, and the honest ask is for the kappa. Whether moderate agreement across every interaction beats high agreement across two percent of them is a genuine empirical question with a plausible answer of yes, particularly for coaching, where a directionally correct signal across a whole team is more useful than a precise signal about four calls. It is much less plausibly yes for anything that touches employment consequences, where a scorer with measurable position bias and moderate chance-corrected agreement is a grievance waiting to be filed.

HYPOTHESIS Offered for testing rather than asserted: automated quality scores are usable for coaching direction and unusable for individual performance management, and the operations that get value from AI-scored quality management will be the ones that keep a small human calibration function rather than the ones that eliminate quality assurance entirely. This is confirmed if a buyer's own calibration exercise, scoring a stratified sample of 200 interactions with both the automated scorer and two trained human reviewers, yields a human-to-machine kappa materially below the human-to-human kappa on the same set. It is refuted if the two kappas are comparable, in which case the labour case is real and larger than anything else in this document. That exercise costs about three analyst-weeks and no vendor in this set will run it for you unprompted.

FIGURE 7 Chance-corrected agreement between language-model judges and human labels



FACT Norman et al., arXiv:2606.19544, June 2026, covering 21 judges from nine providers over 118 runs and roughly 541,000 judgments. These benchmarks are not contact centre interactions and the read-across is indicative, not direct. Conventional interpretation places 0.41 to 0.60 at moderate agreement and 0.61 to 0.80 at substantial. No contact centre

platform vendor publishes an equivalent statistic for its own automated quality scoring, which is the point of including the chart. ^[35]

What the three layers do to each other

ASSESSMENT AI is not a fourth product sitting beside contact centre platform and workforce management. It changes what both of them are, and it does so in ways that make a working deployment read as a failing one on the dashboard the operation already has. This is the single most common cause of an AI programme being killed in year two in this author's operator-side experience, and it is a measurement problem rather than a technology problem.

ASSESSMENT Start with self-service. A working deployment removes the simple contacts first, because simple contacts are what automation handles. What reaches a human afterwards is, by construction, the residue: the harder, longer, more emotionally charged half. Average handle time rises. First contact resolution falls. Customer satisfaction on agent-handled contacts often falls too, because the population of contacts changed. Every one of those movements is the correct signature of a deployment that is working, and every one of them reads as failure against a target set on the pre-deployment mix. An operation that holds its agents to last year's average handle time after deploying self-service has built an incentive to fail.

ASSESSMENT Now agent assist. Real-time guidance changes handle time in both directions at once, shortening the search-and-navigate portion and lengthening the read-and-adapt portion, and the net varies by contact type. Whatever the net, it breaks the forecasting model, because workforce management forecasts are fitted on a historical handle-time distribution that no longer describes the work. The practical consequence is that the first eight to twelve weeks after an agent-assist rollout produce systematically wrong schedules, and the workforce management team is usually blamed for it. The fix is not a better forecast. It is knowing in advance that the distribution shifted, freezing the historical fit, and refitting on post-deployment data with a deliberately shortened lookback.

ASSESSMENT Post-interaction summarisation is the quiet one. It removes after-call work, which is real saving, and it also removes the agent's own act of writing the disposition, which was the mechanism by which agents encoded judgment into the record. Summaries are more consistent and less informative about the things a summary model was not told to look for. Six months in, an operation frequently finds its interaction data has become more uniform and less useful for the analytics that justified the platform. Nobody notices, because the metric that would have caught it, the diversity of disposition codes in use, is not on any dashboard.

Metric	Why it stops meaning what it meant	What to measure instead
Average handle time	Self-service removes the short contacts, so the human average rises when the deployment works	Cost per resolved customer issue across all channels, human and automated
First contact resolution	The residue reaching agents is harder, so it falls even as total resolution rises	Seven-day repeat contact rate for the same customer and intent, all channels

Metric	Why it stops meaning what it meant	What to measure instead
Containment rate	Counts abandonment as success by definition on at least one major platform	Verified resolution rate: contained sessions minus those followed by a contact within seven days
Quality score	Automated scores are self-consistent without being validated against human review	The kappa between your automated scorer and two trained reviewers on the same sample, rerun quarterly
Schedule adherence	Agent assist and summarisation change the shape of handle time, so the fit is stale	Forecast error against post-deployment actuals, with the lookback window explicitly shortened
Occupancy	Rises mechanically when simple contacts leave, and is read as an efficiency gain	Occupancy split by contact complexity band, so the mix shift is visible rather than hidden
Customer satisfaction on agent contacts	The contact population changed, so the score is not comparable year on year	Satisfaction measured end to end across the whole journey, including the automated leg
Cost per contact	Splits into two populations with different costs and stops being one number	Blended cost per resolved issue: seats, AI meters and telephony over verified resolutions

ASSESSMENT The last row of that table is the one worth acting on immediately, and it does not require an AI deployment to justify. A blended cost per resolved issue, computed from seat licences plus every meter plus telephony, divided by verified resolutions rather than by contacts handled, is the only number in this category that stays comparable across a platform change, an AI deployment and a volume swing. It is computable today from any of these vendors' billing exports. Almost no operation computes it, which is why almost no operation can tell whether its AI programme saved money.

Who decides, and where deployments stall

ASSESSMENT The committee that renews a contact centre platform and the committee that buys an AI add-on are different committees, and the vendors know it. A platform renewal is run by contact centre operations with procurement, IT and the incumbent account team, it is anchored on the existing seat rate, and its success criterion is continuity at an acceptable price. An AI purchase is increasingly sold above and around that group: to a chief customer officer, a transformation function, or a chief information officer with an AI mandate and a budget that did not come from the contact centre's cost centre. The AI is often bought before operations has been consulted about whether it can be run.

ASSESSMENT That split is where deployments stall, and the stall has a recognisable shape. The AI is procured on a consumption meter against a business case owned by a transformation function. Delivery lands with contact centre operations, who did not write the business case and whose own targets, still set on the pre-deployment metric mix, now move against them. Nobody owns the meter: the transformation budget funded the pilot, the contact centre cost centre

inherits the run rate, and the first quarter in which consumption exceeds the pilot assumption produces a budget argument rather than an optimisation. Microsoft's tenant-pooled Copilot Credits make this structurally worse than the alternatives, because the pool is shared with workloads the contact centre neither controls nor sees.

ASSESSMENT Three things prevent it and all three are free. Name a single owner of the AI meter with authority over both the spend and the configuration that drives it, before the pilot rather than after. Reset the operational targets that the deployment will move, in writing, before go-live, using the replacement metrics in the previous section. And put the AI on the same renewal date as the seat contract, so that the two purchases arrive at the same negotiating table at the same time. Vendors will resist the third one specifically, and their resistance is the best evidence that it matters.

Churn and migration, and the measurement that does not exist

ASSESSMENT Who leaves whom, and why, is the question buyers most want answered and the one with the weakest evidence base in this entire document. What circulates is vendor win claims and competitor-alternatives pages published by the competitors. RingCentral publishes a page on Five9 alternatives. Verint publishes a page on NICE CXone alternatives. These establish that a vendor believes a rival is vulnerable and that a complaint recurs often enough to build a landing page around. They quantify nothing, and using them as evidence of churn direction would be citing a vendor's marketing department about its rival's retention.

ASSESSMENT The absence of independent measurement is itself the finding and should be reported as one. No independent body publishes migration flows between these platforms with a stated method. The nearest thing to a real signal is the retention disclosure each public vendor makes in its own results, and those are not comparable to each other either. NICE reported cloud net revenue retention of 106 percent in the second quarter of 2026. Genesys claims cloud net revenue retention above 120 percent for twelve consecutive quarters, unaudited, as a private company preparing for a listing. Both numbers are net of expansion, so a vendor can lose customers and still report above 100 percent by selling more to the ones it keeps, which in a period when everyone is attaching AI meters to existing seats is exactly what you would expect to happen. Neither number tells you the gross logo churn rate, and neither vendor publishes it.

[19] [23]

ASSESSMENT What a buyer can do instead is bounded but real. Ask the vendor, in writing and as a contractual representation rather than a slide, for the number of customers above your seat count in your region who left in the last twenty-four months, and for three references among customers who joined from the platform you are leaving. Both asks are refusable, and a refusal is information. The pattern in this author's operator-side experience is that migrations away from a suite are driven by commercial disputes and by administrative burden far more often than by capability gaps, and that the capability reason offered in the post-mortem is usually the one the losing account team could most comfortably report upward.

The regulatory layer, arriving now

FACT EU AI Act Article 50(1), (2) and (5) became applicable on 2 August 2026, three weeks before this document's date. Providers of AI systems intended to interact directly with natural persons, which covers customer-service chatbots, voice bots and AI agents, must ensure the natural person is informed that they are interacting with an AI system, unless that is obvious from the point of view of a reasonably well-informed, observant and circumspect person. The information must be provided in a clear and distinguishable manner at the latest at the time of the first interaction or exposure, and must conform to accessibility requirements. The obligations apply to systems already on the market, not only to those placed on it after that date. A limited transitional period to 2 December 2026 applies only to the marking and detection obligation for generative systems already deployed. Penalties reach €15 million or 3 percent of worldwide annual turnover, whichever is higher. ^{[33] [34] [45]}

ASSESSMENT Nothing in the published pricing of any vendor in this set reflects the cost of complying with that, and none of the nine pricing pages examined for this document mentions it. That is not a vendor failing, because the duty in most contact centre deployments falls on the deployer rather than the platform. It is a selection failing, because a buyer choosing a platform in August 2026 is choosing the party whose contractual commitments determine whether compliance is straightforward or a bespoke engineering project. The distinction that matters is between what a vendor will warrant and what it merely documents. A whitepaper describing how to configure a disclosure prompt is documentation. A contractual commitment to provide, maintain and evidence AI-interaction disclosure across every channel the platform supports, with change notification, is a warranty. Ask which one you are being offered and get the answer in the agreement.

ASSESSMENT Recording and consent remain a jurisdiction-by-jurisdiction problem that predates the AI Act and has not been simplified by it. Two-party consent states in the US, national variations across the EU under the ePrivacy regime and member-state law, and sector rules in financial services all still apply, and the platform's job is to make per-jurisdiction consent configuration possible rather than to know your obligations. What has changed is that AI processing of a recording is a separate question from the recording itself. A call lawfully recorded under a customer's consent to recording has not thereby been consented to for behavioural scoring, sentiment inference or model training, and several vendors in this set train or tune on customer interaction data by default with an opt-out rather than an opt-in. Find that clause before signing, because it is usually in the data processing addendum rather than the main agreement.

ASSESSMENT On residency, the vendors that publish a region list are the ones you can hold to a position, and the list is shorter than the marketing implies. Genesys publishes 21 AWS regions plus a European sovereign region in Brandenburg on the AWS European Sovereign Cloud. Amazon Connect is published as available in a subset of AWS regions materially smaller than AWS overall. NICE operates a mixture of AWS regions and its own processing centres and publishes FedRAMP voice points of presence for US federal work. The operative question is not how many regions a vendor has but whether the specific service you need, including the AI services, is generally available in the specific region you need, because the sovereign and

government regions routinely lag the commercial ones on service catalogue. Verify per service and per region, in writing, and treat the general availability date as a contractual milestone rather than a roadmap. ^{[13] [37] [38] [39]}

Agent experience, absent from every vendor comparison

VENDOR CLAIM Verint's State of Agent Experience 2026 surveyed more than 1,000 contact centre agents at organisations with at least 300 agents across five industries, 59 percent of them in centres with more than 1,000 agents, collected online between 18 November and 9 December 2025. It found that 31 percent planned to leave their role within six months, that the figure was 46 percent among agents aged 18 to 34 against 8 percent among those 45 and over, and that the most cited frustrations were unrealistic performance expectations at 47 percent and lack of schedule flexibility at 45 percent. This is a vendor-commissioned survey and Verint sells the products that address both of the top two complaints, which is exactly why the method disclosure matters and why it is worth crediting: sample size, seat-size threshold, industry count, field dates and age breakdown are all published, which is more than most vendor research in this category discloses. ^[36]

FIGURE 8 Agent intent to leave and its stated drivers, from a vendor-commissioned survey with a disclosed method



VENDOR CLAIM Verint, State of Agent Experience 2026. More than 1,000 agents at organisations with 300+ agents across five industries, 59 percent in centres above 1,000 agents, fielded online 18 November to 9 December 2025. Self-reported intent, not observed departures, and commissioned by a vendor that sells against both of the top two drivers. Included because the method is disclosed and because no independent equivalent exists, not because it is neutral. ^[36]

ASSESSMENT Two of those five bars name things a platform decision actually touches. Schedule flexibility is a workforce management capability: self-scheduling, shift bidding and intraday swap

approval are product features, they differ materially between the specialists and the bundled modules, and they are usually scored as nice-to-have in evaluations. Unrealistic performance expectations is a metrics problem, and the preceding section on what AI does to established metrics explains exactly how an operation generates unrealistic expectations without meaning to: by holding agents to a handle-time target set on a contact mix that self-service has since removed. An operation that deploys AI and does not reset its targets has, mechanically, manufactured the top complaint in that survey.

ASSESSMENT The reason agent experience is absent from vendor comparisons is that it is not a product attribute. It is an outcome of the interaction between the product, the targets and the management. That does not make it irrelevant to a platform decision. At 40 percent annual attrition, a 1,500-seat operation replaces 600 agents a year, and at any plausible cost of hire and ramp that is a larger number than the difference between any two vendors in this document. A platform decision that improves schedule flexibility and makes the metric reset possible is worth more than one that wins on routing sophistication, and no evaluation scoring sheet in this author's experience has ever weighted it that way.

Contract language: what to require, what to strike

The following are asks a buyer can make of any vendor in this set. They are ordered by how much money they protect, not by how likely a vendor is to agree. Several will be refused, and a refusal is itself information worth having before signing rather than after.

Area	Require	Strike or renegotiate
AI consumption	An annual cap on total metered AI spend, with overage at the same unit rate rather than a punitive one, and the right to disable metered features at the account level	Any clause letting the vendor change the unit-to-currency conversion rate for tokens, credits or sessions during the term without your consent
Meter definition	A written definition of the billable unit, with a worked example: what counts as a session, a conversation, an action, an interaction	Definitions that appear only in an online help article the vendor can edit unilaterally
Renewal alignment	The AI agreement co-terminating with the seat agreement, so both arrive at one negotiating table	Auto-renewing consumption schedules on a different anniversary from the platform term
Containment claims	The vendor's containment or deflection figure restated as a warranty against your own definition, measured on your traffic during a defined proving period	Any performance figure that survives into the contract only as a recital or an appendix marked non-binding
Quality scoring	The right to run a calibration exercise against human reviewers, with the vendor supplying raw per-item scores rather than aggregates	Clauses making automated scores contractually authoritative for any employment-related purpose
EU AI Act Article 50	A warranty that the platform provides and maintains AI-interaction disclosure across	Compliance language that points to documentation rather than committing the vendor to anything

Area	Require	Strike or renegotiate
	every channel, with notification before any change to disclosure behaviour	
Data and model rights	Explicit opt-in, not opt-out, before your interaction data trains or tunes any model, including the vendor's own	Default training rights buried in the data processing addendum rather than the main agreement
Exit	A written exit assistance obligation: bulk export of recordings, transcripts, metadata and evaluations in a documented format, at a capped fee, for a defined period after termination	Export priced as professional services at the vendor's discretion at the time you ask
Reliability	The last 24 months of incident records for your specific region and service, supplied as a representation rather than a slide	Service credits framed as the sole and exclusive remedy where an outage causes downstream contractual loss
Product lifecycle	A minimum support horizon for the specific product line you are buying, with a migration entitlement if it is sunset or merged	Silence on lifecycle, which is the default and which matters most where two overlapping product lines now sit under one owner
Change of control	Notification and a termination right on a change of control that materially changes the roadmap for your product line	Assignment clauses permitting transfer without notice, which is standard and which is what a sponsor-owned vendor most wants to keep
Telephony	Voice rates fixed for the term where the vendor owns the network, or pass-through with audit rights where it resells	Voice priced as a separate agreement with a different term, which is how the second bill becomes a surprise

ASSESSMENT If only three of those survive the negotiation, make them the AI consumption cap, the meter definition and the exit assistance obligation. The first is the only protection against a variable cost that moves with the one thing a contact centre cannot control. The second is what makes the first enforceable, because a cap on an undefined unit is not a cap. The third is worth more in year four than everything else in the agreement combined, and it costs nothing to obtain in year zero because the vendor is not yet holding your data. Every one of these is easier to get before signature than at any point afterwards, which is obvious and is nonetheless the reason most operations do not have them.

Four scenarios to the end of 2027, and what to watch

No point forecast appears here, because the category's own rate of change in the eighteen months covered by the evidence above rules one out: two of the shortlisted vendors merged, one bought its AI vendor for \$955 million, one withdrew most of its published pricing, and a regulation took effect. The probabilities below are subjective and are the least reliable content in this document. The tripwires are the most useful, because each is observable without private information.

SCENARIO A **The meter wins** 40 percent

AI consumption pricing becomes the primary growth line across every suite vendor, and the seat rate settles into a floor that is negotiated once and then ignored. Published seat prices stay flat or drift down while metered revenue compounds.

Consequence. Buyers lose year-on-year seat benchmarking as a cost control and do not replace it, so total cost per contact rises through a line item nobody owns. Procurement leverage moves decisively to the vendor, because consumption renews continuously and contracts renew annually.

Earliest visible sign. NICE extends per-session pricing below the Ultimate Suite tier, or Five9 restores published pricing with an explicit AI meter line rather than restoring the old bundled tiers.

SCENARIO B The suite reabsorbs AI 30 percent

One major vendor decides metered AI is costing it deals and publishes a tier that includes agentic handling at a fixed seat price, forcing the others to follow within two quarters.

Consequence. The meter argument in this document collapses and seat benchmarking becomes useful again. Buyers mid-term on consumption contracts are structurally disadvantaged against new customers, and renewal timing becomes the most valuable thing in the agreement.

Earliest visible sign. Any of NICE, Genesys, Five9 or Talkdesk publishing a per-seat tier that states unmetered agentic interaction handling. Watch the pricing pages rather than the announcements; the page changes first.

SCENARIO C Consolidation continues 20 percent

Another name in this set changes owner. The candidates on structural grounds are the private vendors with no disclosed raise since 2021 and the sponsor-owned assets whose holding periods are maturing.

Consequence. A buyer signing a five-year term with an independent vendor inherits a roadmap decision made by a new owner in year two. Product lifecycle and change-of-control language, currently treated as boilerplate, becomes the most valuable clause in the agreement.

Earliest visible sign. A change-of-control filing or announcement for Talkdesk or Five9, or a Genesys public listing that resolves its ownership question in the other direction.

SCENARIO D Regulatory drag 10 percent

Article 50 enforcement plus a recording or consent action against a customer-facing AI deployment materially changes how conversational AI is deployed in the EU, adding a compliance workstream to every project.

Consequence. Deployment timelines in the EU extend by a quarter or more, vendors that will warrant disclosure behaviour gain a real commercial advantage over those that only document it, and some EU-constrained operations defer agentic deployment past the contract term they bought it in.

Earliest visible sign. The first penalty or formal proceeding under Article 50 against a customer-facing conversational system, or the first vendor in this set publishing a contractual Article 50 warranty as a competitive feature.

Implications

For heads of customer experience and contact centre operations

ASSESSMENT Your immediate exposure is not the platform, it is the measurement. If AI is deployed or about to be, the metrics your floor is managed on will move against you while the deployment works, and the section above on what the three layers do to each other lists which ones and why. Reset those targets in writing before go-live, not after the first quarter of bad-looking numbers, because after is when it reads as excuse-making. Compute a blended cost per resolved issue from your billing data this quarter regardless of whether you are buying anything, because it is the only number that stays comparable through a platform change and you will need the pre-deployment baseline. And run the containment audit described above on your existing self-service before anyone quotes you a deflection figure for a new one.

For VPs of customer service evaluating or renewing a platform

ASSESSMENT Decide which of the three purchases you are making before you build a scoring matrix, because a single matrix across all twelve of these vendors will produce a ranking that cannot explain your own decision. If routing and global voice are the binding constraint, the shortlist is Genesys and NICE and the rest is noise. If labour cost is the binding constraint, the shortlist is Verint and possibly your incumbent suite's bundled module, and the routing platform is incidental. If AI deflection is the binding constraint, do not replace your platform at all. On renewal specifically: your incumbent's seat rate is probably not where the money is any more, so ask for twenty-four months of your own consumption history before you open the negotiation, and negotiate the meter.

For procurement and vendor management

ASSESSMENT Seat-rate benchmarking has stopped being sufficient and in this category is now close to a decoy. Four of these vendors have moved their growth onto meters, one has withdrawn publication from the tiers that carry the capability, and a year-on-year seat comparison will show a flat or improving picture while total spend rises. The three clauses that protect the most money are the AI consumption cap, the written definition of the billable unit, and the exit assistance obligation with a capped fee, in that order, and all three are cheap before signature and unobtainable after. Treat vendor containment, deflection and quality-accuracy figures as unwarranted marketing unless they are restated as contractual representations measured on your own traffic; every one of them can be, and vendor resistance to doing so is the most reliable signal in the process. Note also that Verint and Calabrio can no longer be played against each other, and that Salesforce holds an equity position in Genesys.

For IT and integration owners who inherit the deployment

ASSESSMENT Two things in this document land on you and neither appears in the business case. First, the answer to "who can change a routing rule" determines your ticket queue for the next five years, and choosing a deep-customisation suite means operations will route every change request through your team. Price that in headcount before the decision, not after. Second, the trained-model component of exit cost is now the largest lock-in mechanism in this category and it accrues silently: an Enlighten or Contact Lens model tuned on three years of interactions cannot be exported and cannot be rebuilt at a new vendor until the new vendor has accumulated the interactions. If assemble-it-yourself is on the table, the honest question is not whether your team can build it but whether you can commit named engineers for the life of the platform rather than the life of the project, because the failure mode is staff turnover rather than technology.

For CIOs approving the spend

ASSESSMENT Three structural facts should shape the approval and none of them is a product question. The AI in this category is priced on meters that renew continuously against contracts that renew annually, which means the growth in this spend line will not surface at the moment you are set up to review it; require the AI agreement to co-terminate with the platform agreement. If the estate is Microsoft, contact centre AI consumption draws from a Copilot Credit pool shared tenant-wide at \$0.01 per credit with unused prepaid credits expiring annually, so the contact centre's costs compete with every other Copilot workload and a ring-fenced spend policy is a prerequisite rather than a refinement. And the compliance duty under EU AI Act Article 50, applicable since 2 August 2026 with penalties up to €15 million or 3 percent of worldwide turnover, sits with your organisation as deployer rather than with the platform vendor, so the question to put to the vendor is what it will warrant rather than what it documents. On the vendor set itself: two of the ten are now one company, one holds equity from a competitor you may also be buying from, and at least one has not disclosed a funding round since 2021.

What this document cannot answer

This list is deliberately long and specific. It is the most useful section here for anyone deciding whether to commission primary research, because every item on it is answerable with fieldwork and none is answerable from public sources.

ASSESSMENT What anyone actually pays. Every price in this document is list. The real question, what a 500-seat, a 2,000-seat and a 10,000-seat buyer pays each of these vendors after discount, how the discount differs between the seat and the meter, and whether AI consumption is discounted at all, requires access to executed agreements. A structured collection of twenty to thirty anonymised contracts across the four suites would answer it and would be worth more to a buyer than everything in this document.

ASSESSMENT Whether AI-scored quality management agrees with human reviewers. No vendor publishes a chance-corrected agreement statistic. The study that settles it is a calibration exercise across three or four platforms: a stratified sample of 200 to 400 interactions per platform, scored by the automated scorer and by two trained human reviewers blind to each

other, reporting human-to-human and human-to-machine kappa separately, broken down by evaluation criterion. It is roughly a three-analyst-month project per platform and its finding would change the labour case for quality assurance functions across the industry.

ASSESSMENT True containment against apparent containment, at scale. The seven-day repeat-contact join described above can be run inside any single operation. What nobody has done is run it across a set of operations on different platforms with a common definition, which is the only way to establish whether the gap between quoted and true containment is a platform property, a configuration property or a customer-base property. That requires data access agreements with a dozen operations and is the single highest-value piece of fieldwork available in this category.

ASSESSMENT Actual measured availability. Every uptime figure in circulation is a vendor declaration on a vendor-controlled status page, or a third party re-reporting one. Independent measurement would require synthetic monitoring from multiple geographies against each platform continuously for twelve months or more, which is expensive, technically straightforward, and has apparently never been published.

ASSESSMENT Real migration cost and duration. The exit-cost table above is this author's assessment from operator-side experience, not a measured figure. A study of completed migrations, capturing elapsed time from decision to decommission, professional services spend, internal effort in full-time-equivalents, parallel-run duration and service level impact during transition, across ten to fifteen migrations, would replace the most judgment-heavy table in this document with evidence.

ASSESSMENT Gross logo churn by vendor. Public disclosures give net revenue retention, which nets expansion against loss and is therefore uninformative about whether customers leave. Gross logo churn, by seat band and region, exists inside every one of these vendors and is disclosed by none of them.

ASSESSMENT What agent assist does to handle time, measured rather than claimed. Vendors publish handle-time reductions. Nobody publishes a controlled comparison with a matched holdout group over a period long enough to separate the tool's effect from the novelty effect and from the contact-mix shift caused by concurrent self-service deployment. A randomised rollout within a single large operation would answer it cleanly and cheaply.

ASSESSMENT Whether metered AI pricing is cheaper or more expensive than bundled pricing at realistic volumes. This document establishes the published unit rates. It cannot establish the realised annual cost, because that depends on consumption patterns that only appear in billing data. Twelve months of anonymised billing from a handful of operations on each pricing model would settle a question that currently gets answered by whichever vendor is presenting.

ASSESSMENT How EU AI Act Article 50 is actually being enforced. The text is clear, the applicability date has passed, and there is as yet no enforcement history to reason from. Whether supervisory authorities treat a pre-conversation disclosure banner as sufficient, and how they treat voice channels where a visual disclosure is impossible, will only be knowable from the first decisions.

ASSESSMENT The Salesforce and ServiceNow positions in Genesys. The \$1.5 billion investment is disclosed; what is not disclosed is whether it carries governance rights, product roadmap influence, or any restriction on competitive behaviour. That is a matter for the share purchase agreement and for whatever appears in a Genesys registration statement if the company lists.

ASSESSMENT Which of Verint's and Calabrio's overlapping product lines survives, and on what timetable. Both companies' customers need this answer to plan and neither has published one. It is knowable only from the combined firm and only when it chooses to say.

Half-life of this assessment

ASSESSMENT Decays within six months: every published price in this document, and specifically the Five9 disclosure position, which could reverse with a single page edit; the AI revenue percentages, which all four vendors restate quarterly; the Genesys token conversion rates, which are published in a help article the vendor can edit unilaterally; the Verint and Calabrio product roadmap position, which will move as soon as the combined firm announces convergence; and Talkdesk's independence.

ASSESSMENT Decays within twelve months: the AI unit-cost comparison, because the whole industry is repricing AI and the twenty-five-to-one spread is unlikely to survive competition; the model provenance table, which will change with every OEM renegotiation and every acquisition; the region and residency lists, which grow steadily; and the scenario probabilities, which should be rescored against the tripwires rather than carried forward.

ASSESSMENT Decays within twenty-four months: the workforce management crossover thresholds, as suite-bundled modules continue to close on the specialists; the containment-definition finding, if regulatory or competitive pressure forces disclosure of measurement definitions; and the observation that no vendor publishes a quality-scoring kappa, which is one competitor's marketing decision away from being false.

ASSESSMENT Architectural, and unlikely to change on this timescale: the three-markets category structure and the category error it produces on a combined shortlist; the observation that the party who can change a routing rule determines operational agility; that trained-model lock-in accrues silently and grows with tenure; that switching cost is dominated by parallel-running and requalification rather than data movement; that net revenue retention cannot answer a question about churn; and that AI changes what the established contact centre metrics mean rather than adding a metric beside them.

Limitations

Stated interest. The author spent almost five years at Concentrix, which is a services buyer of several vendors named in this document and which holds commercial relationships with them that the author was not party to and cannot describe. The author holds no equity in, and receives no compensation from, any vendor named here, and this document was not commissioned, reviewed or influenced by any of them. Readers should nonetheless weigh the Concentrix

relationship when reading assessments of the vendors an operator of that kind buys from, and the operator-side perspective it produced is disclosed in the front matter rather than left implicit.

Vantage point. This is an operator-side and advisory view of how these platforms behave once deployed at scale across many client estates. It is not a procurement view. The author has not signed one of these contracts and holds none of the discount sheets. Every price here is list, and the gap between list and transacted is large, vendor-specific and undisclosed.

Evidence base. The strongest material in this document is the archived pricing evidence, which is checkable by anyone with a browser, and the regulatory text, which is definitive. The weakest is the exit-cost table and the workforce management crossover thresholds, both of which are the author's assessment from experience rather than measured figures, and both of which are labelled as such. Between those extremes sit the vendor financial disclosures, which are self-reported, defined differently by each vendor, and in Genesys' case unaudited and issued by a private company.

Verification failures worth naming. Salesforce's Service Cloud edition prices are rendered client-side and could not be retrieved from a Salesforce primary source on 26 August 2026, so no Salesforce seat price appears in the matrix; the widely circulated figures for those editions are omitted rather than substituted. RingCentral's contact centre pricing page likewise returned no rate to a plain fetch, so the RingCX figure here is taken from RingCentral's own general availability release of November 2023 and should be treated as potentially stale. Five9's three highest tiers have no published price at all. Verint and Calabrio publish none. The Five9 AI revenue figure is management commentary from an earnings call rather than a line in the results release, and its definition is not disclosed.

Coverage. Twelve platforms is not the whole market. Avaya, Cisco Webex Contact Center, 8x8, Vonage, Zoom Contact Center, Zendesk, Sprinklr, NICE's own Financial Crime segment, and the entire population of regional and vertical specialists are outside scope. Their absence reflects the shortlist this document was built against, not a judgment that they are uncompetitive, and for several buyer profiles one of them would be the right answer.

Method. No vendor was briefed, interviewed or given a right of reply, and no non-public information was used. Where a vendor's position is characterised, it is characterised from that vendor's own published material, cited. Vendors who believe a cell in the matrix is wrong should be able to point at a published source that says so, and if they can, this document is wrong and should be corrected.

Sources

-
- [1] Genesys. "Genesys Cloud CX Pricing." Vendor pricing page, retrieved 26 August 2026. *Vendor primary* genesys.com
 - [2] Genesys. "Genesys Cloud pricing hub." Genesys Cloud Resource Center, retrieved 26 August 2026. AI Experience token rates, workforce engagement add-on rates, voice, SMS and text-to-speech rates. *Vendor primary* help.genesys.cloud
 - [3] Genesys. "Genesys Cloud CX Pricing." Internet Archive snapshot, 30 May 2023. *Vendor primary (archived)* web.archive.org

- [4] Genesys. "Genesys Cloud CX Pricing." Internet Archive snapshot, 13 June 2024. *Vendor primary (archived)* web.archive.org
- [5] Five9. "Pricing." Vendor pricing page, retrieved 26 August 2026. *Vendor primary* five9.com
- [6] Five9. "Pricing." Internet Archive snapshot, 2 June 2023. Digital \$149, Core \$149, Premium \$169, Optimum \$199, Ultimate \$229. *Vendor primary (archived)* web.archive.org
- [7] Five9. "Pricing." Internet Archive snapshot, 21 May 2024. Digital \$175, Core \$175, Premium \$235, Optimum \$290, Ultimate \$325. *Vendor primary (archived)* web.archive.org
- [8] Five9. "Pricing." Internet Archive snapshot, 24 May 2025. Digital \$119, Core \$119, three highest tiers "Contact Sales." *Vendor primary (archived)* web.archive.org
- [9] NiCE. "Pricing." Vendor pricing page, retrieved 26 August 2026. Omnichannel Suite \$110, Essential \$135, Core \$169, Complete \$209, Ultimate \$249 plus \$0.25 per session. *Vendor primary* nice.com
- [10] NiCE. "Pricing." Internet Archive snapshot, 8 September 2024. Digital Agent \$71, Voice Agent \$94, Omnichannel \$110, Essential \$135, Core \$169, Complete \$209, Ultimate \$249. *Vendor primary (archived)* web.archive.org
- [11] Talkdesk. "Contact Center Software Pricing." Vendor pricing page, retrieved 26 August 2026. *Vendor primary* talkdesk.com
- [12] Talkdesk. "Contact Center Software Pricing." Internet Archive snapshot, 4 June 2025. *Vendor primary (archived)* web.archive.org
- [13] Amazon Web Services. "Amazon Connect Customer Pricing Appendix." Retrieved 26 August 2026. *Vendor primary* aws.amazon.com
- [14] Microsoft. "Dynamics 365 Contact Center pricing." Retrieved 26 August 2026. *Vendor primary* microsoft.com
- [15] Microsoft. "Copilot Credits Guide." June 2026, 12 pages. Pay-as-you-go \$0.01 per credit; tenant-level pooling across Copilot Cowork, Copilot Studio, Dynamics 365 agents, Power Platform and Work IQ APIs; pre-purchase tiers 300,000 to 300,000,000 credits. *Vendor primary* microsoft.com
- [16] Salesforce. "Salesforce Introduces New Flexible Agentforce Pricing to Accelerate the Digital Labor Revolution." Press release, 15 May 2025. Flex Credits at \$500 per 100,000; 20 credits (\$0.10) per action. *Vendor primary* salesforce.com
- [17] Google Cloud. "Conversational Agents pricing." Retrieved 26 August 2026. Flows chat \$0.007 per request, Playbooks chat \$0.012 per request; Flows voice \$0.001 per second, Playbooks voice \$0.002 per second. *Vendor primary* cloud.google.com
- [18] Calabrio. "Pricing." Retrieved 26 August 2026. No list price published; the page now serves Verint navigation following the merger. *Vendor primary* calabrio.com
- [19] NiCE. "NiCE Exceeds Revenue Guidance Range, Reporting 8% Year-Over-Year Revenue Growth in Second Quarter 2026." 5 August 2026. *Vendor primary* nice.com
- [20] Five9. "Five9 Announces Second Quarter 2026 Financial Results." Revenue \$312.4 million, up 10 percent; subscription revenue up 14 percent; adjusted EBITDA \$70.1 million, 22.4 percent of revenue against 24.0 percent a year earlier. *Vendor primary* investors.five9.com
- [21] Five9. Second quarter 2026 earnings call, 5 August 2026. Management commentary placing AI revenue above a \$150 million annual run rate, growing 78 percent. Definition of AI revenue not disclosed. *Vendor primary (earnings call)* fool.com
- [22] RingCentral. "RingCentral Announces Second Quarter 2026 Financial Results." 23 July 2026. Revenue \$657 million, up 5.9 percent; "13% of ARR is now from customers utilizing a native paid AI product, doubling year-over-year." *Vendor primary* ringcentral.com
- [23] Genesys. "Genesys Reports Record Fourth Quarter as Organizations Accelerate the Adoption of AI-Powered Experience Orchestration." 26 March 2026. Fiscal year ended 31 January 2026. Company-reported and unaudited; Genesys is privately held. *Vendor primary (unaudited)* genesys.com
- [24] Genesys. "Genesys Announces \$1.5 Billion Investment by Salesforce and ServiceNow." 31 July 2025. \$750 million from each; proceeds used to repurchase shares from existing equity holders. *Vendor primary* genesys.com
- [25] Verint. "Verint Agrees to Be Acquired by Thoma Bravo for \$2 Billion." 25 August 2025. \$20.50 per share in cash. *Vendor primary* verint.com
- [26] Thoma Bravo. "Thoma Bravo Completes Acquisition of Verint, a Leader in AI-Driven Customer Experience Automation." Transaction completed 26 November 2025; Verint combined with portfolio company Calabrio. *Vendor primary* thomabravo.com
- [27] Verint. "Verint Announces Corporate Name for Combined Verint-Calabrio Organization." 18 February 2026. Single corporate name Verint; Calabrio solutions continue inside the Verint CX Automation Platform. *Vendor primary* verint.com

- [28] Verint. "Verint Names Dave Rhodes as Chief Executive Officer." 24 February 2026. Rhodes was previously chief executive of Calabrio. *Vendor primary* [verint.com](https://www.verint.com)
- [29] NiCE. "NiCE Closes Acquisition of Cognigy." 8 September 2025. \$955 million. *Vendor primary* [nice.com](https://www.nice.com)
- [30] Genesys. "Self-Service Statistics Report." Genesys product documentation, retrieved 26 August 2026. "Contained in Self-Service: The total number of interactions that entered the Designer application in Self-Service and were concluded without entering Assisted-Service." *Vendor primary (product documentation)* all.docs.genesys.com
- [31] Genesys. "Beyond containment: Measure digital success, not just effort." Vendor blog, retrieved 26 August 2026. *Vendor claim* [genesys.com](https://www.genesys.com)
- [32] Genesys. "Genesys Cloud Service Level Agreement." Genesys Cloud Resource Center, retrieved 26 August 2026. Credit thresholds of 10 percent below 99.99 percent, 30 percent below 99.0 percent, 100 percent below 97 percent; 30-day claim window. *Vendor primary* help.genesys.cloud
- [33] European Union. Regulation (EU) 2024/1689 (Artificial Intelligence Act), Article 50, transparency obligations for providers and deployers of certain AI systems. Article 50(1), (2) and (5) applicable from 2 August 2026. *Standards body primary (regulatory text)* artificialintelligenceact.eu
- [34] European Commission. "Guidelines on transparency obligations for providers and deployers of certain AI systems." Shaping Europe's digital future, 2026. *Standards body primary* digital-strategy.ec.europa.eu
- [35] Norman et al. "Reliability without Validity: A Systematic, Large-Scale Evaluation of LLM-as-a-Judge Models Across Agreement, Consistency, and Bias." arXiv:2606.19544, June 2026. 21 judges, nine providers, three benchmarks, 118 runs, approximately 541,000 judgments. *Preprint research* arxiv.org
- [36] Verint. "Nearly One-Third of Contact Center Agents Plan to Quit as Agent Experience Falls Short." 14 April 2026. State of Agent Experience 2026, more than 1,000 agents at organisations with 300+ agents across five industries, fielded 18 November to 9 December 2025. *Vendor claim (survey, method disclosed)* [verint.com](https://www.verint.com)
- [37] Genesys. "Cloud Platform Global Availability Capabilities." Retrieved 26 August 2026. *Vendor primary* [genesys.com](https://www.genesys.com)
- [38] Genesys. "Genesys to Offer Experience Orchestration Services on AWS European Sovereign Cloud." Sovereign region in Brandenburg, Germany. *Vendor primary* [genesys.com](https://www.genesys.com)
- [39] Amazon Web Services. "Resilience in Connect Customer." Amazon Connect Administrator Guide, retrieved 26 August 2026. Region availability and Global Resiliency scope. *Vendor primary (product documentation)* docs.aws.amazon.com
- [40] Genesys. "Genesys Deepens Strategic Collaboration with AWS to Accelerate Global AI Innovation." 22 July 2026. AI Studio compatibility with proprietary, open-source and Amazon Bedrock frontier models. *Vendor primary* [genesys.com](https://www.genesys.com)
- [41] NiCE. "NiCE Launches CXone Mpower Agents: Enterprise-Grade Agentic AI Agents Built for CX to Deliver Automated Fulfillment." 17 June 2025. Describes Mpower Agents as powered by NiCE's proprietary CX AI models. *Vendor claim* [businesswire.com](https://www.businesswire.com)
- [42] Microsoft Learn. "Standard harness licensing, Microsoft Copilot Studio." Updated 3 August 2026. Copilot Credits as common currency; monthly capacity enforcement; unused credits do not carry over. *Vendor primary (product documentation)* learn.microsoft.com
- [43] Zeff, Maxwell. "Talkdesk conducts third round of layoffs in less than 14 months." TechCrunch, 26 September 2023. *Independent press* [techcrunch.com](https://www.techcrunch.com)
- [44] RingCentral. "RingCentral Announces General Availability of RingCX, a Natively Built AI-First Contact Center." 14 November 2023. \$65 per agent per month including unlimited domestic inbound and outbound minutes. *Vendor primary* [ringcentral.com](https://www.ringcentral.com)
- [45] Cooley LLP. "EU AI Act: Transparency Obligations Take Effect 2 August 2026." 3 August 2026. Practitioner analysis of Article 50 scope, the December 2026 transitional period and penalty exposure. *Independent press* [cooley.com](https://www.cooley.com)

Rights and permitted use

© 2026 David Soden, davidson.com/market-assessment. All rights reserved.

This document, including its structure, analysis, findings, evidence classification, and framing, is the original work of David Soden, published at davidson.com/market-assessment.

It is published for public reading. You are free to share the complete, unmodified file, to link to it, and to quote short passages in commentary, reporting, or discussion, provided the author and the site above are credited.

The following require prior written consent: republishing this work in whole or in substantial part under another name, masthead, or brand; excerpting or modifying it in a way that alters the analysis or removes the evidence classifications; incorporating any portion of it into a paid, commissioned, or client-facing deliverable; resale or distribution behind a paywall; and use of this document as training data for machine learning models.

Commissioned Market Assessment reports, commercial licensing, and briefings on this material are available.

This document is analysis, not investment, legal, or procurement advice. It was not commissioned, reviewed, or funded by any company named in it, and the author holds no position in any of them.

Market Assessment reports, commissioned research, and briefings: davidsoden.com/market-assessment