

**Barriers to Effective Service Delivery in Contact Centres and the Evidence for
Artificial Intelligence Support: A Review of Four Sources**

Ruth-Ann Bravo

Independent Research

Montreal, Quebec, Canada

August 2026

Author Note

This review was conducted independently and was not funded or commissioned. The author has six years of experience in bilingual contact centre operations, which informed the selection of sources and the interpretation of practitioner accounts. No proprietary or employer-held information was used at any stage. All sources are publicly available.

Correspondence concerning this review may be addressed to Ruth-Ann Bravo, Montreal, Quebec, Canada.

Abstract

This review examined two questions. First, what makes work difficult for people employed as contact centre agents. Second, whether artificial intelligence (AI) tools address those difficulties. Four sources were used: 45 practitioner accounts posted to a public online forum in 2025; a survey of 2,100 unionized United States contact centre workers conducted in 2017; a survey of 2,891 United States and 385 Canadian contact centre workers conducted in 2022 and 2023; and a field study of an AI assistant deployed to 5,172 customer support agents. Analysis of the forum accounts produced thirteen categories of reported difficulty. Organizational conditions were referenced most frequently and are not amenable to a technical solution. The category most amenable to such a solution concerned the inability to locate accurate information during customer interactions, referenced by seven respondents. One respondent, who had held responsibility for maintaining a knowledge base, described the mechanism producing contradictory documentation: uncontrolled editing access distributed across fifteen staff without change notification. Survey evidence corroborated the pattern at scale, with 52% of respondents in 2017 reporting that the information held in their systems was inadequate and 68% reporting contacts at least weekly that their training had not prepared them to answer. When AI tools were assessed in aggregate, workers reported predominantly negative effects. When tool types were assessed separately, AI used to retrieve product and customer information was associated with higher job satisfaction, while AI used for workplace monitoring was associated with lower satisfaction. A field study of an assistant of the former type found a 15% increase in issues resolved per hour, with the largest gains among less experienced agents. Findings indicate that the category of AI deployed matters more than the quantity.

Keywords: contact centres, customer service, knowledge management, artificial intelligence, worker well-being, qualitative content analysis

Barriers to Effective Service Delivery in Contact Centres and the Evidence for Artificial Intelligence Support: A Review of Four Sources

Contact centres are among the first workplaces in which artificial intelligence has been deployed at scale. Decisions about which tools to adopt are typically informed by operational metrics such as call duration and customer satisfaction scores. The people performing the work are rarely consulted about what actually impedes them.

This review reverses that order. It begins with what agents report, tests whether those reports are supported by larger studies, and then examines whether AI tools address the difficulties identified. Four questions guided the review:

1. What barriers do contact centre agents report when responding to an inquiry not administered by their employer?
2. Are those reported barriers supported by large-sample survey evidence?
3. Which barriers could plausibly be addressed by a software tool, and which reflect organizational rather than technical conditions?
4. What measured effect does AI assistance have on the barriers identified as technically addressable?

Definition of Terms

The following terms recur throughout this review. They are defined here because they are specific to contact centre operations and to research methodology, and are not assumed to be familiar.

Agent. The employee who handles customer contacts by telephone or written chat. Also referred to as a representative or customer service representative.

Average handle time (AHT). The mean duration of an agent's customer interactions. Most contact centres set a target for this measure and evaluate agents against it.

Customer satisfaction (CSAT). A score derived from surveys completed by customers following an interaction. Agents are commonly evaluated on this measure, including in cases where dissatisfaction originates in company policy rather than agent conduct.

First contact resolution. Whether a customer's issue was resolved during the initial interaction, without requiring a transfer or a subsequent contact.

Knowledge base. The internal repository of policies, procedures, and product information that agents consult during interactions.

Closed key time. Paid working time during which an agent is not required to handle customer contacts. It is used for documentation and administrative tasks.

Interactive voice response (IVR). The automated menu system a customer navigates before reaching an agent.

Quality assurance (QA). A function that reviews recorded interactions and scores agents against defined criteria.

Agent assistance tool. Software that analyses an interaction as it occurs and supplies the agent with relevant information or suggested responses.

Correlational finding. A finding that two variables occur together, which does not establish that one causes the other. Confounding factors may explain the association.

Grey literature. Research published as an institutional report rather than in a peer-reviewed journal. Such work may be methodologically sound but has not undergone peer review.

Method

Four sources were selected. They differ substantially in method, population, and period. This was intentional: each source compensates for a limitation in the others. Practitioner accounts supply detail but not representativeness. Surveys supply scale but not causal evidence. The field study supplies measured effects but within a single organization.

Source 1: Practitioner Accounts From a Public Forum

Data were drawn from a discussion thread posted on August 18, 2025 to r/callcentres, a public online community for contact centre workers. The thread posed four open-ended questions concerning sources of delay during interactions, information that could not be located quickly, difficulties with existing systems, and a single change the respondent would make to their role.

The complete thread was reviewed. Contributions from the thread author were excluded, as these consisted of follow-up questions rather than accounts of working conditions. Two contributions expressing agreement with another respondent but containing no substantive content were also excluded. The resulting sample comprised 45 respondents.

Where a respondent posted more than once because the thread author had requested elaboration, the contributions were treated as a single case, with the follow-up recorded as additional detail. Respondents self-selected by choosing to reply. No demographic, geographic, or employment information was collected and none was verifiable. Response length ranged from a single phrase to several hundred words, and four respondents supplied only a phrase.

Responses were analysed using inductive thematic coding. An initial pass generated candidate categories from the response text without reference to a predetermined framework. A second pass consolidated overlapping categories and assigned final codes. Responses were

multi-coded, so category frequencies exceed the sample size and should be read as the number of respondents referencing a theme rather than as a proportional distribution. Coding was performed by a single analyst without an inter-rater reliability measure.

Source 2: Survey of Unionized United States Contact Centre Workers

Findings were drawn from Doellgast and O'Brady (2020), a report prepared for the Communications Workers of America and directed through the ILR School at Cornell University. The survey was administered by email in December 2017 to members of union locals. It yielded 2,199 usable responses, of which 2,100 were matched to employers and collective agreements.

Respondents were experienced. Mean age was 44 years, mean tenure with the current employer was 9.6 years, and mean total contact centre experience was 12.3 years. Customer service was the largest job title group at 48% of respondents.

Source 3: Survey of AI Use in United States and Canadian Contact Centres

Findings were drawn from Doellgast et al. (2023), a report issued by the ILR School at Cornell University. The lead researchers were the same as those in Source 2, and a number of survey items were identical, permitting direct comparison between 2017 and 2023.

The researchers surveyed 2,891 United States and 385 Canadian contact centre workers between December 2022 and January 2023, and conducted case studies in the United States, Canada, Germany, and Norway. This source is grey literature. It is an institutional research report and has not undergone peer review.

It is the only source in this review containing Canadian data, and the only one that assessed categories of AI separately rather than in aggregate.

Source 4: Field Study of an AI Assistant Deployment

Findings were drawn from Brynjolfsson et al. (2025), published in the Quarterly Journal of Economics. The authors examined the staggered introduction of a generative AI conversational assistant across 5,172 customer support agents at a large firm. Because the deployment reached different groups at different times, agents with access could be compared against agents without access at the same point in time. This design provides stronger evidence than a simple comparison of outcomes before and after adoption. Agents retained discretion over whether to act on the assistant's recommendations. The setting was written chat rather than voice.

Results

Findings From the Forum Analysis

Table 1 presents the categories generated from the 45 responses, the number of respondents referencing each, and an assessment of whether the category could plausibly be addressed by a software tool.

Table 1

Categories of Reported Difficulty in Forum Responses (N = 45)

Category	Respondents referencing	Addressable by tooling
Management practices, compensation, and monitoring	19	No
Performance metric design	12	Partially
Customer conduct and capability	12	No
Number and reliability of systems	9	Partially
Locating accurate information	7	Yes
Inconsistent information from elsewhere in the organization	3	Partially
Training adequacy	3	Yes

Category	Respondents referencing	Addressable by tooling
Script and policy rigidity	3	Partially
Emotional labour and blame absorption	3	Indirectly
Worker perspectives on technology adoption	3	Yes
Post-contact documentation	1	Yes
Verification and authentication friction	1	Partially
Colleague non-adoption of available tools	1	Indirectly

Note. Responses were multi-coded. Column totals exceed N = 45. The final three categories were referenced by a single respondent each and are reported for their substantive content rather than their frequency.

The largest category concerned organizational conditions: supervision quality, compensation, break scheduling, and monitoring intensity. Nineteen respondents referenced conditions of this kind, including one who reported being reprimanded for sneezing and another who described being timed to the minute for start, break, and finish times. No software addresses these conditions, and they are reported here for proportional context rather than because they inform tool design.

Customer conduct was referenced with equal frequency to metric design. Accounts included condescension, disputes over the agent's expertise, harassment, and complaints regarding colleagues' accents. Two respondents described customers who were unable to perform the task they had called for assistance with, such as entering a web address.

Locating Accurate Information

Seven respondents referenced difficulty locating information. The substance of these accounts is more instructive than their frequency, and four are reported in detail.

The first respondent stated that their organization maintained an adequate knowledge base and that the ease with which search terms returned relevant results, rather than content

coverage, constituted the operative limitation. A second described documentation distributed across multiple pages in which sources contradicted one another, alongside frequent policy change. A third described receiving case notes written in internal terminology by an upstream department, with no accompanying explanation, while being expected to convey case status to the customer. A fourth described operating across seven systems spanning nine regional configurations, each with distinct workflows, and reported that most working time was spent verifying procedure rather than executing it.

One respondent supplied an account of causal mechanism. Having previously held responsibility for maintaining knowledge base content, they described a manager distributing editing access to agents as a gesture of recognition. With approximately fifteen individuals making changes without notification to the content owner, contradictory information emerged almost immediately. Content that was structurally constrained could be controlled; free text could not.

This account is significant because it identifies contradictory documentation as the product of an access control and change notification failure rather than of poor writing. It indicates that contradiction in an operational knowledge base is generated continuously rather than arising as an occasional defect.

The Workaround Sequence

One respondent described the procedure followed when an answer could not be located. The customer is placed on hold. Colleagues seated nearby are consulted. A floor walker or team leader is approached. A message is posted to a team chat channel and a response awaited. If none of these produce an answer, the contact is transferred to another line.

Two features of this account warrant attention. First, the same respondent stated that their organizational resources were adequate and identified nothing they would change about their tools, locating the failure in access under time pressure rather than in content quality. Second, the procedure consumes the working time of three or four additional staff in addition to the agent, and terminates in a transfer. None of this appears in standard metric sets, which record only the duration of the original agent's contact.

Worker Perspectives on Technology Adoption

Three respondents addressed how technology should be introduced rather than what was wrong with it. Two proposed AI applications unprompted: one suggested an AI tool that would guide customers through password resets step by step to enable self-service, and a second proposed AI-guided browsing or the automated dispatch of a direct link.

A third respondent proposed a governance principle, stating that management should use a tool themselves before mandating its use across an agent population. This corresponds to the voluntary adoption condition documented in the case study evidence reported below, arrived at independently by a practitioner.

These accounts indicate that the negative aggregate assessment of AI reported in the survey evidence should not be read as generalised opposition to the technology. Respondents proposed AI applications where those applications addressed their workload, while objecting to applications directed at their own conduct.

A Documented AI Deployment Failure

One respondent described an automated note-taking deployment in detail. Transcript accuracy was poor, particularly for names, pronouns, and spellings, requiring manual correction. Because the transcript was generated only after the interaction concluded, correction occurred

during after-call work status, which supervisors subsequently identified as excessive. The tool performed its nominal task while increasing total workload and exposing agents to criticism under the existing performance regime.

Reported Job Satisfaction Within the Sample

Two respondents stated that they were satisfied with their employment. One opened their account by describing their position as a good one that they enjoyed, and proceeded to describe operating across seven systems and nine market configurations. A second described their workplace as generally good while criticising break time tracking. A third reported that a system replacement had substantially improved a previously poor experience.

These accounts are relevant to the interpretation of the sample. The information access barrier was reported by respondents who expressed satisfaction with their employer as well as by those who did not, which indicates that the barrier is not solely an artifact of generalised dissatisfaction.

Findings From the 2017 Survey

The conditions described in the forum accounts appear in the survey data with associated frequencies.

Systems and Technology

Doellgast and O'Brady (2020) reported that 86% of respondents considered the technology used to process contacts too slow, 80% reported frequent system failures, 52% considered the information held in their systems inadequate, and 59% reported difficulty arising from the number of separate applications in use. Respondents reported an average of 3.6 distinct technology problems from a list of six.

Training

Across the sample, 68% reported receiving contacts at least weekly that their training had not prepared them to answer. Only 13% received updated product and service information daily. In-person training was rated very or extremely effective by 51% of respondents, compared with 15% for computer-based training, although computer-based methods accounted for approximately 60% of training hours delivered.

Discretion

A quarter of respondents were required to follow scripts verbatim, and a further 39% could modify them only slightly. Approximately two-thirds therefore operated under substantial script constraint. In customer-related decisions such as adjusting charges or issuing credits, 34% reported little or no discretion and 14% reported significant or complete discretion.

Attribution of Responsibility

Reported customer incivility centred on matters originating outside the agent's control. Eighty percent of respondents reported that customers frequently held them responsible for events beyond their control, 82% encountered customer frustration regarding transfers between departments, and 81% addressed errors made by external vendors at least several times weekly.

Monitoring

Voice recording was near-universal at 98%. Respondents reported an average of three concurrent monitoring methods. Seventy percent considered monitoring information to be applied for disciplinary rather than developmental purposes, and more than 70% considered performance metrics unreasonable, the interval between contacts insufficient, and scheduling inflexible.

Association With Organizational Cost

The survey data associate these conditions with outcomes carrying measurable cost. Respondents reporting very high stress reported an average of 3.9 technology problems, compared with 2.3 among those reporting low stress. Among high-stress respondents, 59% frequently received contacts their training had not covered, compared with 36% of low-stress respondents. Stress was in turn associated with absence, at a mean of 25.2 days annually among very-high-stress respondents against 12.5 days among low-stress respondents, and with turnover intention, with 44% of the very-high-stress group planning to seek other employment within a year against 8% of the low-stress group.

Findings From the 2023 AI Survey

Experience with AI tools was near-universal by 2023. Only 12% of United States respondents and 20% of Canadian respondents reported no experience with any of the eight tool categories assessed.

Assessed in aggregate, worker evaluation of AI was negative. More United States respondents disagreed than agreed that AI tools increased working speed, made work easier, made work more interesting, or improved customer service. Approximately half agreed that AI increased stress.

The pattern changes substantially when tool categories are assessed separately. Table 2 presents this comparison.

Table 2
United States Worker Evaluation of AI Tools by Category

Tool category	Improves customer service	Makes work easier	Increases stress
Retrieval of product and customer information	53%	40%	24%

Tool category	Improves customer service	Makes work easier	Increases stress
Automated feedback on tone and script adherence	19%	12%	55%
Automated workspace monitoring	13%	8%	67%

Note. Data from Doellgast et al. (2023). Percentages indicate respondents agreeing or strongly agreeing.

The researchers additionally conducted a regression analysis. The total number of AI tools in use was significantly associated with lower job satisfaction. However, two specific applications, retrieval of product information and support for training, were each associated with higher satisfaction, while workspace monitoring was associated with lower satisfaction.

This finding is central to the present review. The question of whether AI benefits contact centre workers is insufficiently specified. Some categories are associated with benefit and others with harm, and aggregate measures conceal both.

Work Intensity and Recovery Time

Respondents reporting higher AI intensity also reported spending more of their working time in customer contact. In the United States sample this rose from 69% of working time among respondents reporting no AI to 88% among those reporting high-intensity AI, and in Canada from 56% to 83%. Closed key time declined correspondingly, from a weekly mean of 108 minutes to 48 minutes in the United States and from 110 minutes to 16 minutes in Canada.

Well-Being

In the Canadian sample, 57% of respondents reporting no AI indicated job satisfaction, compared with 37% among those reporting high-intensity AI. Emotional exhaustion showed the inverse pattern, reported a few times weekly or daily by 38% of the no-AI group and 63% of the high-intensity group. Differences in the United States sample followed the same direction but were smaller.

Change Between 2017 and 2023

Monitoring intensified over the period. Keystroke tracking rose from 34% to 58%, chat monitoring from 28% to 69%, and online activity monitoring from 68% to 85%. The mean number of concurrent monitoring methods rose from three to five.

Customer conduct improved over the same period. Respondents reporting that customers frequently held them responsible for matters beyond their control declined from approximately 80% in 2017 to 51% in 2023, and frustration regarding transfers declined from 82% to 67%.

Case Evidence

The report documents an agent assistance tool deployed at Deutsche Telekom that supplied suggested responses, removing the requirement to search several databases manually. Agents were permitted to decide whether to use it. Some declined, finding it distracting. A majority adopted it, and first contact resolution more than doubled. Two conditions accompanied the deployment: employment protection was guaranteed, and workers were involved in decisions about implementation.

Findings From the Deployment Study

Brynjolfsson et al. (2025) reported that access to the AI assistant increased issues resolved per hour by 15% on average. This figure conceals substantial variation. Less experienced and lower-performing agents improved by approximately 30% or more, and improved in output quality as well as speed. The most experienced and highest-performing agents recorded small speed gains accompanied by small declines in quality.

The authors characterise the mechanism as the diffusion of practices associated with higher-performing agents to those earlier in their tenure. Agents with two months of experience performed comparably to agents with six months of experience operating without assistance.

The study additionally reported improved customer sentiment and increased employee retention. During periods when the software was unavailable, agents who had previously had access performed above their own pre-deployment baseline, with the largest residual gains among those who had used the tool most extensively. This indicates durable learning rather than dependency.

Convergence Across Sources

Table 3 maps each barrier identified in the forum accounts against the corresponding survey and deployment evidence.

Table 3
Correspondence of Reported Barriers Across Four Sources

Barrier	Forum, 2025	Survey, 2017	AI survey, 2023
Information present but not retrievable	7 respondents; mechanism identified	52% report system information inadequate	Retrieval AI is the only category rated positively
Training does not cover live demands	Agents improvising during system fault	68% receive such contacts weekly	Training AI associated with higher satisfaction
Agent absorbs blame for upstream error	Website, vendor, and self-service errors reaching the agent	80% frequently blamed for matters outside control	68% of high-AI United States workers, against 42% with no AI
Interface fragmentation	9 respondents; seven systems across nine regions	59% report too many applications	Not directly addressed
Monitoring experienced as disciplinary	19 respondents; largest category	70% report monitoring used for discipline	Monitoring AI is the least favourably rated category

Note. Forum figures indicate the number of respondents referencing the theme out of 45. Sources differ in country, union status, sampling method, and period.

Discussion

The Agent as Point of Final Reconciliation

A single structural pattern recurs across the categories amenable to intervention. The agent occupies the position of final reconciliation for inconsistencies generated elsewhere in the organization.

Contradictory documentation, region-specific workflow variation, public-facing content misaligned with internal policy, untranslated upstream terminology, and self-service failures invisible to the agent are described as distinct problems. Each requires the same activity: real-time resolution, by the agent, under time pressure, of a discrepancy the agent did not create and cannot correct at source.

That 80% of surveyed workers were frequently held responsible by customers for matters outside their control indicates that this reconciliation burden is characteristic of the role rather than incidental to it.

The activity is largely unmeasured. Respondents described substantial time allocated to procedure verification and system troubleshooting, neither of which appears in standard metric sets, while both contribute to the handle time against which agents are evaluated. The workaround sequence described in the forum accounts extends this further: the consultation of colleagues, supervisors, and team chat channels consumes the working time of several staff for a single unresolved query, and none of that time is attributed to the information failure that produced it.

The Generation of Contradictory Documentation

The survey sources establish that agents consider system information inadequate. The forum account from a former knowledge base owner supplies a mechanism for how this condition arises.

Editing access was distributed without corresponding change notification. Multiple contributors modified overlapping content independently, and contradictions emerged as a routine consequence rather than as an occasional error. Where content was structurally constrained, the owner retained control; where it was free text, they did not.

Two implications follow. First, interventions that address contradiction by improving writing quality address a symptom rather than a cause. Second, any tool operating over an organizational knowledge base should treat contradiction as a persistent structural condition rather than as an exception requiring occasional handling. A system that assumes internal consistency in its source material assumes a condition that the operational evidence indicates is not met.

Reconciling the Two AI Studies

Two sources appear to reach opposing conclusions. Brynjolfsson et al. (2025) reported a 15% productivity gain and increased retention. Doellgast et al. (2023) reported that AI intensity was associated with lower satisfaction and higher absence.

The apparent conflict resolves on inspection of what each measured. The deployment study assessed a single agent assistance tool that agents were free to disregard. The survey measured intensity across eight categories including workspace cameras, automated coaching, script enforcement, and scheduling algorithms. Aggregating beneficial and detrimental applications produces a negative mean.

Critically, when the survey researchers disaggregated by category, the retrieval application aligned with the deployment study's findings. The two sources are therefore consistent, and together support a more specific proposition than either supports alone: the category of AI deployed is more consequential than the quantity.

Implications for Tool Design

Six implications follow from the convergent findings.

Retrieval, not content coverage, is the operative constraint. Organizations continue to expand documentation while agents continue to report an inability to locate information within the interaction. Expanding a repository does not improve access to it.

Tools must accommodate contradiction as a normal condition. A system that returns a single confident answer where sources disagree reproduces the existing failure at greater speed. Conflicting sources should be surfaced as conflicting.

Tools must be capable of declining to answer. The documented note-taking failure illustrates the consequence of a system that produces incorrect output in preference to no output.

Evaluation should measure net workload rather than task completion. The note-taking tool completed its assigned task and increased total work. Task-level accuracy is an insufficient evaluation criterion.

Adoption should be optional. The Deutsche Telekom case indicates that voluntary adoption produced majority uptake and substantial performance improvement. The least favourably rated tool categories are those which agents cannot decline.

Retrieval and monitoring functions should remain separate. Automated feedback and workspace monitoring are the most negatively evaluated categories. Combining performance

data generation with a retrieval function would compromise the component associated with benefit.

A seventh consideration concerns evaluation design. Because effects varied by agent experience level in the deployment study, with quality declining slightly among the most capable agents, an aggregate measure is insufficient. Evaluation should be disaggregated by tenure and performance level.

Limitations

The forum sample was self-selected. Participation in a community of this kind is plausibly motivated in part by dissatisfaction, and some negative skew should be assumed. That assumption requires qualification, however. Three respondents expressed satisfaction with their employment or reported improvement following a system replacement, and the information access barrier appeared in their accounts as well as in more critical ones. The sample is therefore not uniformly composed of dissatisfied workers, although it cannot be treated as representative.

Employer, sector, geography, and tenure were unknown for all respondents. Four respondents supplied only a single phrase, which was sufficient for category assignment but not for interpretation. Coding was conducted by a single analyst without a reliability check, and category boundaries were permeable, with several responses plausibly assignable to more than one category. Three categories were referenced by a single respondent each; these are reported for their content and should not be read as indicating prevalence.

The 2017 survey population was restricted to union members at United States employers. Unionized contact centres may differ systematically from non-union and offshore operations in monitoring intensity, tenure, and training investment. Mean tenure of 9.6 years substantially

exceeds typical contact centre norms, so the findings may understate conditions affecting newer agents. The survey predates the general availability of generative AI tools.

The 2023 survey included only 385 Canadian respondents. Once disaggregated by AI intensity, several Canadian subgroups comprise approximately twenty individuals. Canadian figures should be read as indicating direction rather than magnitude. This source is also grey literature and has not undergone peer review.

Both surveys are correlational and cross-sectional. Workplaces deploying AI intensively may have differed in monitoring culture and work organization prior to adoption. Causal direction cannot be established from these data, and the original authors do not claim otherwise.

The deployment study examined one firm, one product category, one AI system, and written chat rather than voice interaction. Generalization to voice channels and to regulated sectors such as financial services is not established. Agents in that study retained discretion over adherence, so the findings may not transfer to deployments in which use is mandatory.

Finally, this review identifies which barriers are commonly reported. It does not establish their operational magnitude or cost within any particular organization.

Directions for Further Research

Semi-structured interviews with practitioners recruited outside a self-selecting forum context would permit assessment of whether the identified patterns hold in a non-self-selected sample.

Additional Canadian data, and data from non-unionized workplaces, would address the geographic and labour relations limitations of the survey sources.

Replication of the deployment study design in a voice-based and regulated setting would address the most consequential generalization gap in the intervention evidence.

Direct measurement of information retrieval time, currently absent from standard contact centre metric sets, would establish the operational magnitude of the barrier that all four sources implicate as central.

References

- Brynjolfsson, E., Li, D., & Raymond, L. (2025). Generative AI at work. *The Quarterly Journal of Economics*, 140(2), 889–942. <https://doi.org/10.1093/qje/qjae044>
- Doellgast, V., & O'Brady, S. (2020). *Making call center jobs better: The relationship between management practices and worker stress. A report for the CWA*. ILR School, Cornell University, and DeGroote School of Business, McMaster University.
- Doellgast, V., O'Brady, S., Kim, J., & Walters, D. (2023). *AI in contact centers: Artificial intelligence and algorithmic management in frontline service workplaces*. ILR School, Cornell University.
<https://ecommons.cornell.edu/items/2c04c957-d672-400e-9aed-adb5f6ace640>
- r/callcentres. (2025, August 18). *Contact centre folks: What makes your job harder?* [Online forum post]. Reddit.
https://www.reddit.com/r/callcentres/comments/1mu4v99/contact_centre_folks_what_makes_your_job_harder/

Appendix A

Coding Categories Applied to Forum Responses

Management practices, compensation, and monitoring. Supervision quality, scheduling and break control, compensation, recognition, monitoring practices, and process change implemented without consultation.

Performance metric design. Handle time targets, attribution of customer satisfaction scores, revision of targets, and quality assurance sampling and criteria.

Number and reliability of systems. Count of concurrent applications, software reliability, interface efficiency, and technical support arrangements.

Locating accurate information. Retrieval speed, contradiction between sources, gaps in documentation coverage, and terminology opacity.

Customer conduct and capability. Customer behaviour during interactions and technical capability affecting resolution.

Inconsistent information from elsewhere in the organization. Public-facing content conflicting with internal policy, third-party misinformation, and self-service failures invisible to the agent.

Training adequacy. Duration and sufficiency of onboarding relative to system and process complexity.

Script rigidity. Verbatim script requirements, constraints on resolution authority, and escalation requirements for routine exceptions.

Absorbing customer reaction to organizational error. Managing customer response to organizational failure and exposure to hostile conduct.

Post-contact documentation. After-call work volume, note generation, and correction of automated output.

Worker perspectives on technology adoption. Proposals for AI applications, and principles concerning how tools should be introduced or tested before mandatory deployment.

Verification and authentication friction. Customer resistance to identity verification procedures required before information can be released.

Colleague non-adoption of available tools. Accounts of colleagues declining to use available systems or procedures, with consequences transferring to subsequent agents.