

*Re-writing economics with AI*

# Beyond the 2% Sample

*Agent Performance 100% evaluated*



# The flawed mechanics of Agent Evaluation



For decades, enterprise contact centers and customer operations have evaluated human agent performance through a dangerously narrow lens.

Bound by the constraints of manual quality assurance (QA), **organisations routinely sample barely 1% to 3% of agent interactions.**

Supervisors score these isolated encounters using rigid, static spreadsheets, attempting to grade everything from technical accuracy and protocol compliance to empathy and soft skills.

This traditional framework creates **a profound governance blind spot.**

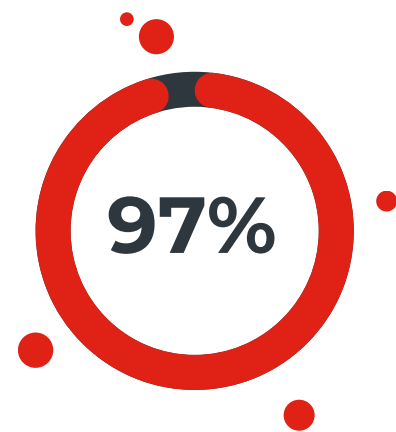




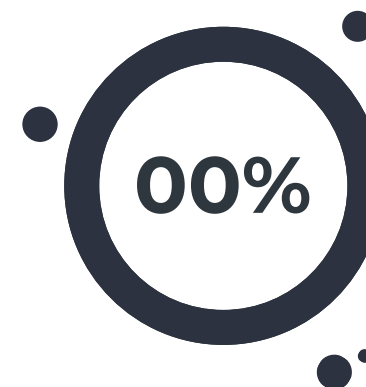
# The silent crisis in Modern Quality Assurance

When **97% of an agent's monthly work goes entirely unreviewed**, performance metrics distorted by statistical noise.

Exceptional agents are penalised for isolated anomalies, chronic compliance slips go undetected, and coaching sessions degenerate into subjective debates between supervisors and frustrated staff.



**MANUAL AGENT  
WORK EVALUATION**



**AUTOMATED  
WORK SCORING**



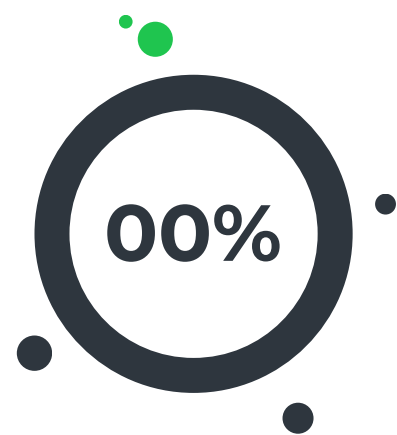
**EVALUATOR  
SUBJECTIVITY**



# Enter AI Autoscoring in Modern Quality Assurance

The introduction of **AI Agent Autoscoring completely redefines this dynamic.**

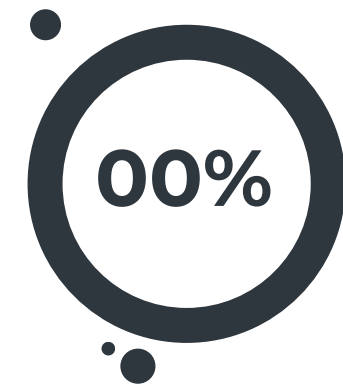
By deploying asynchronous AI language models to evaluate agents' work and performance across managed interactions based on customisable evaluation forms, enterprises can finally **measure, score, and coach 100% of agent output** with mathematical consistency.



**MANUAL AGENT  
WORK EVALUATION**



**AUTOMATED  
WORK SCORING**



**EVALUATOR  
SUBJECTIVITY**



# Core business problems in traditional agent QA

To understand why modern enterprises are abandoning manual scorecard reviews, we must dissect the structural failures plaguing legacy agent evaluation models.



## **The statical illusion** of agent performance

Manual QA reviews barely 3% of agent work, distorting monthly scores with statistical noise based on a tiny fraction of actual interactions.



## **Inter-rater reliability** and supervisor bias

Human evaluations are inherently subjective, with different supervisors applying scorecards according to their own personal biases rather than objective standards.



## **The high cost** of admin overhead

Adopting manual scorekeeping consumes massive supervisory time on calls and spreadsheets instead of coaching agents.

# AI Autoscoring

## the architecture



Modern autoscoring shifts the paradigm **from manual sampling to continuous, automated agent evaluation**. By structuring agent performance around configurable evaluation areas, specific intents, and weighted scoring models, enterprises can automate grading across voice, chat, and asynchronous channels.



### The Core Principle

human work

Autoscoring evaluates the agent's work. Not the customer's behavior. Every interaction handled by an agent is automatically analyzed against customised criteria.

| Overall Score | Getting Started | Product/Service Knowledge | Interpersonal skills |
|---------------|-----------------|---------------------------|----------------------|
| 71.56         | 96.00           | 52.00                     | 86.56                |
| 17.43         | 14.29           | 7.14                      | 33.29                |

An example of our AI Autoscoring results for two different agents.



# AI Autoscoring

## transforming Management & Coaching

Moving from a 3% sample to 100% automated agent scoring fundamentally **alters how enterprises develop talent and drive operational efficiency:**



### **Objective, unbiased** performance baseline

AI evaluates every interaction against identical standards, eliminating personal bias and replacing subjective critiques with reliable, data-backed insights across hundreds of interactions.



### **Surgical data-driven** agent coaching

Autoscoring isolates exact behavioral drivers - such as technical knowledge gaps or tone issues - allowing supervisors to conduct laser-focused coaching sessions targeting specific problem areas.



### **Quality governance** across all queues

Running silently in the background across 100% of agent work, autoscoring immediately flags compliance risks, policy deviations, and training gaps before they cause systemic degradation.



# Elevating Human Workforce through Precision Intelligence

Contact center success ultimately hinges on frontline human performance.

Enterprises have attempted to manage human performance using outdated, analog sampling methods that breed frustration and obscure reality.

By implementing AI Agent Autoscoring, organisations **bridge the gap between operational expectation and actual execution.**

Eliminating subjective bias, **evaluating 100% of agent work**, and anchoring coaching in granular intent metrics allows enterprise leaders to transform quality assurance from a punitive administrative chore into a powerful engine for agent mastery and operational excellence.





**BECLLOUD**

Solutions for innovation

# Find out more

**[becLOUDsolutions.com](https://becLOUDsolutions.com)**

[marketing@becLOUDsolutions.com](mailto:marketing@becLOUDsolutions.com)

