

The Colorado Office of Judicial Performance Evaluation (COJPE) provided the operational context for this proof-of-concept study.

Research Objectives

Benchmark automated model outputs against manual coding to measure accuracy for each classification task.

Data Source

Volume: 9,397 comments across **387 judges**

Methodology

Constructive comments are grouped by judge and performance category. The agent produces two-paragraph summaries per judge (Strengths and Weaknesses) across domains including: communication clarity, fairness, legal reasoning, demeanor, case management, preparation, and timeliness.

Prompt Engineering & Validation Logic

Automated decisions are subject to human review. The workflow is designed so that reviewers can efficiently inspect flagged cases without re-processing the full dataset.

Results

Classification Findings

"Many respondents described the judge as respectful, and patient in this domain, demonstrating a clear commitment to excellence and skill. Few remarks suggested negative impressions within this area of performance, allowing focus to remain upon the positives instead."

Conclusions & Implications

Human oversight: Structured flag outputs allow efficient human review of automated decisions, maintaining accountability without re-doing full processing.