

PROPOSED STANDARD

The Human Fluency Benchmark

Can AI actually know a person? Nobody measures.

Hundreds of public benchmarks score what AI knows. None score whether it can form an accurate model of the specific human it's talking to — validated the way psychology validates any instrument. The Human Fluency Benchmark is a proposed standard for exactly that: AI trait judgments from real conversation history, scored against validated psychometric ground truth, across standardized doses of interaction.

Resonant is convening this effort, not running it. Scientific leadership sits with academic collaborators; studies run under university IRB oversight; governance follows the convener-not-owner model of efforts like MLCommons. Resonant's contribution is the substrate the question requires: consented interaction data joined to validated multi-instrument ground truth.

The benchmark's value depends on its independence — including from us. That separation is maintained on principle and by design.

Researchers in person perception, psychometrics, and machine behavior:

[COPY TODO] — full protocol, scoring rubric, and governance charter.