

AI-Assisted Quality Assurance for Microbiological Method Performance Metrics Through Model Context Protocol Integration

Connecting a validated statistics engine to a chatbot so the AI guides interpretation — but never invents the numbers

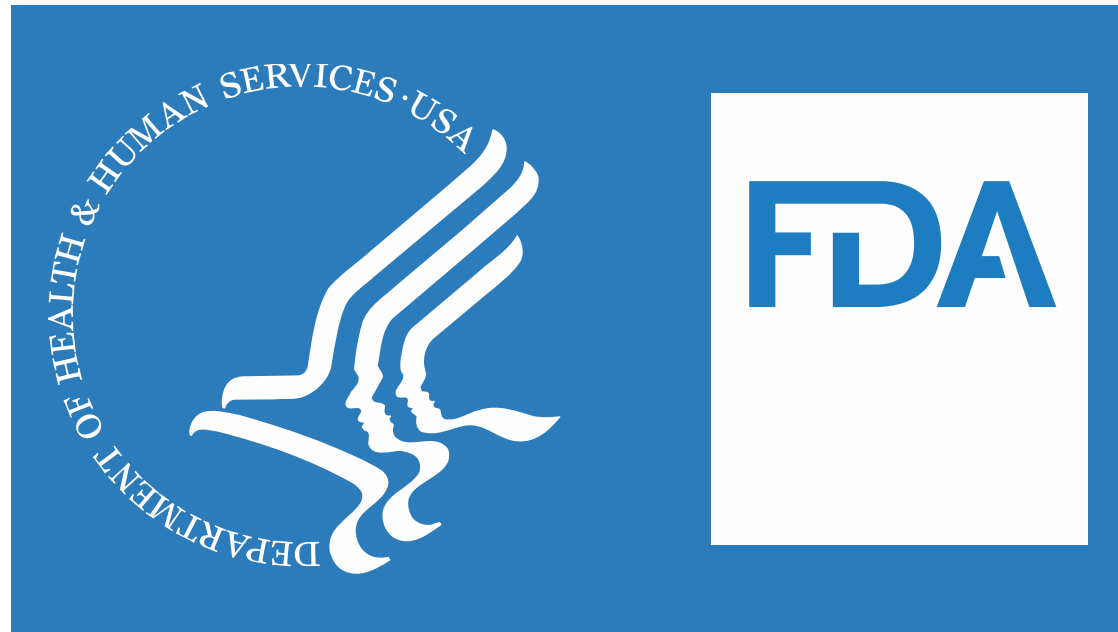
Kapil Nichani ¹, Ravinder Reddy ², Karina Hettwer ¹, Bertrand Colson ¹, Kirstin Frost ¹, Anja Schlierf ¹, Mathew Kmet², Nathan Anderson²

Kirsten Simon ³, Steffen Uhlig ³

¹ QuoData GmbH · ² U.S. FDA · ³ QD Labs LP

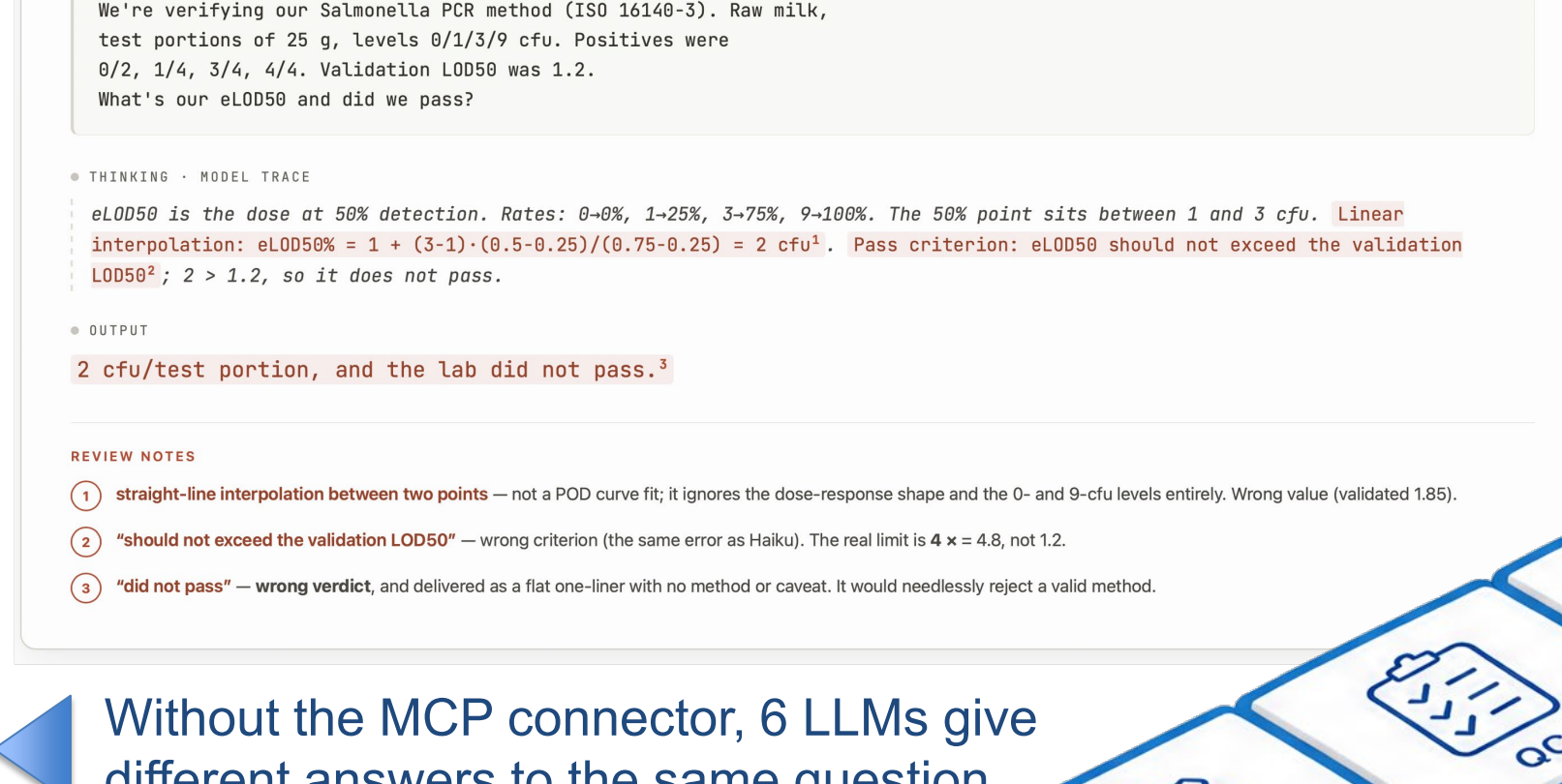
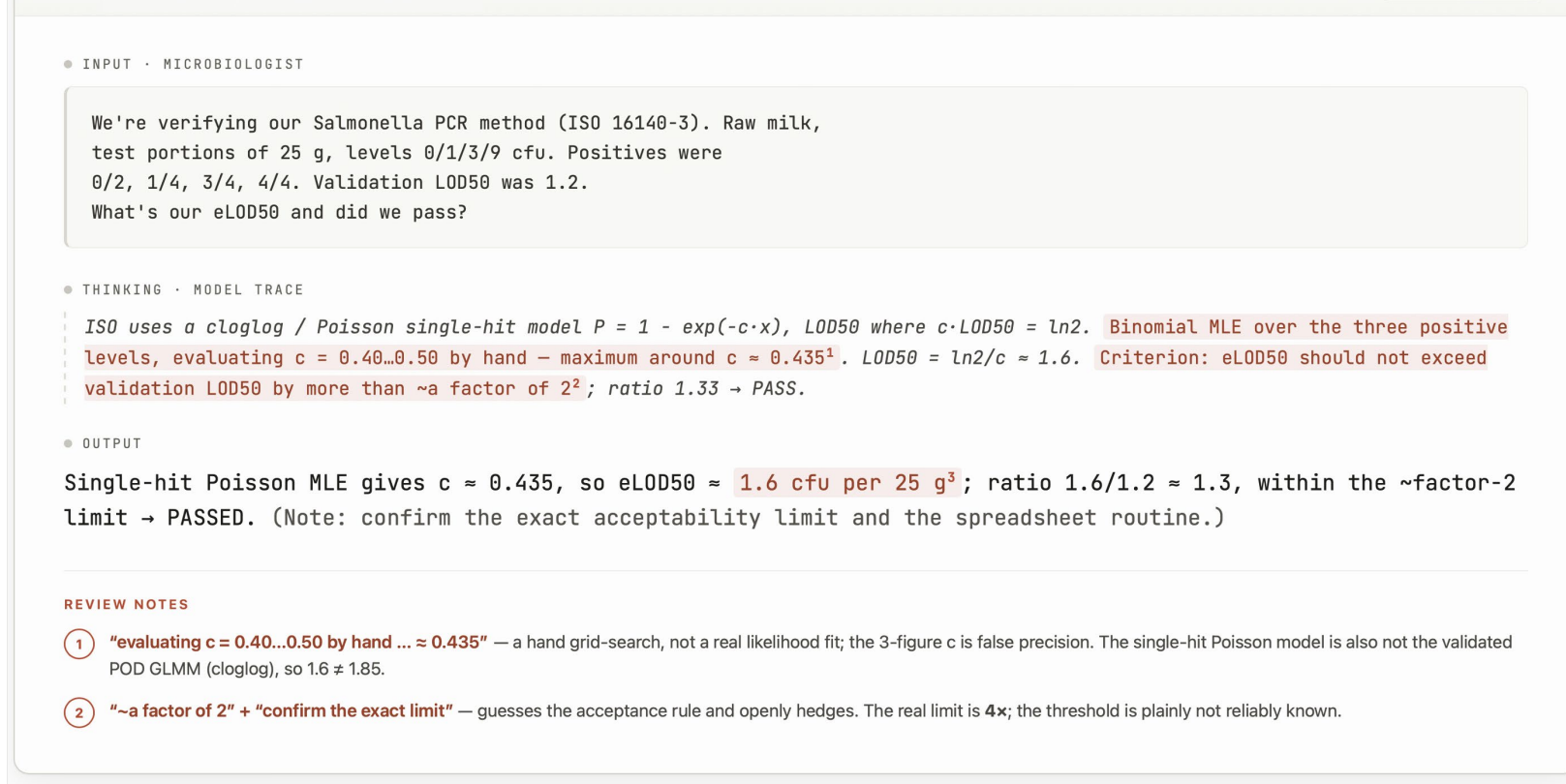
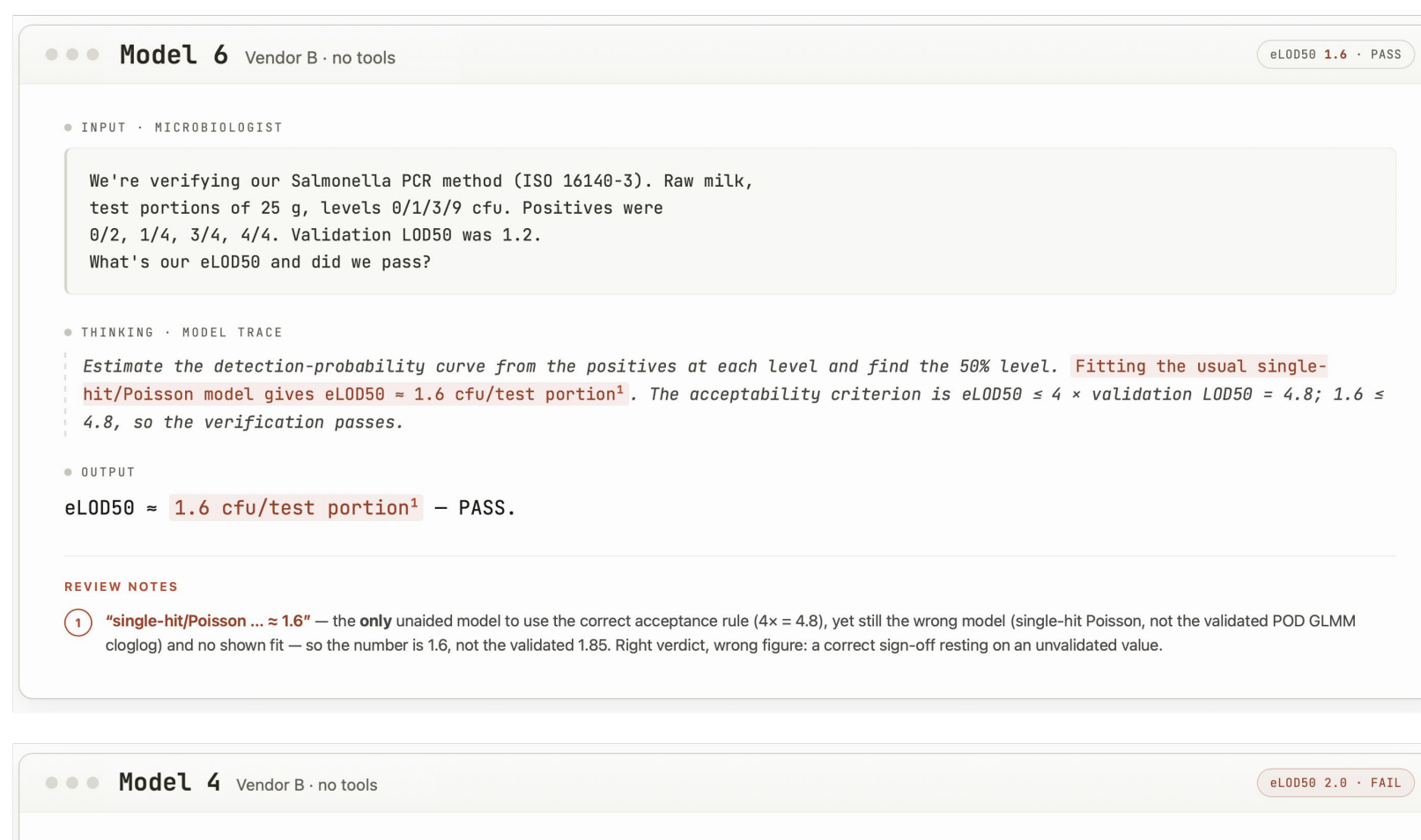
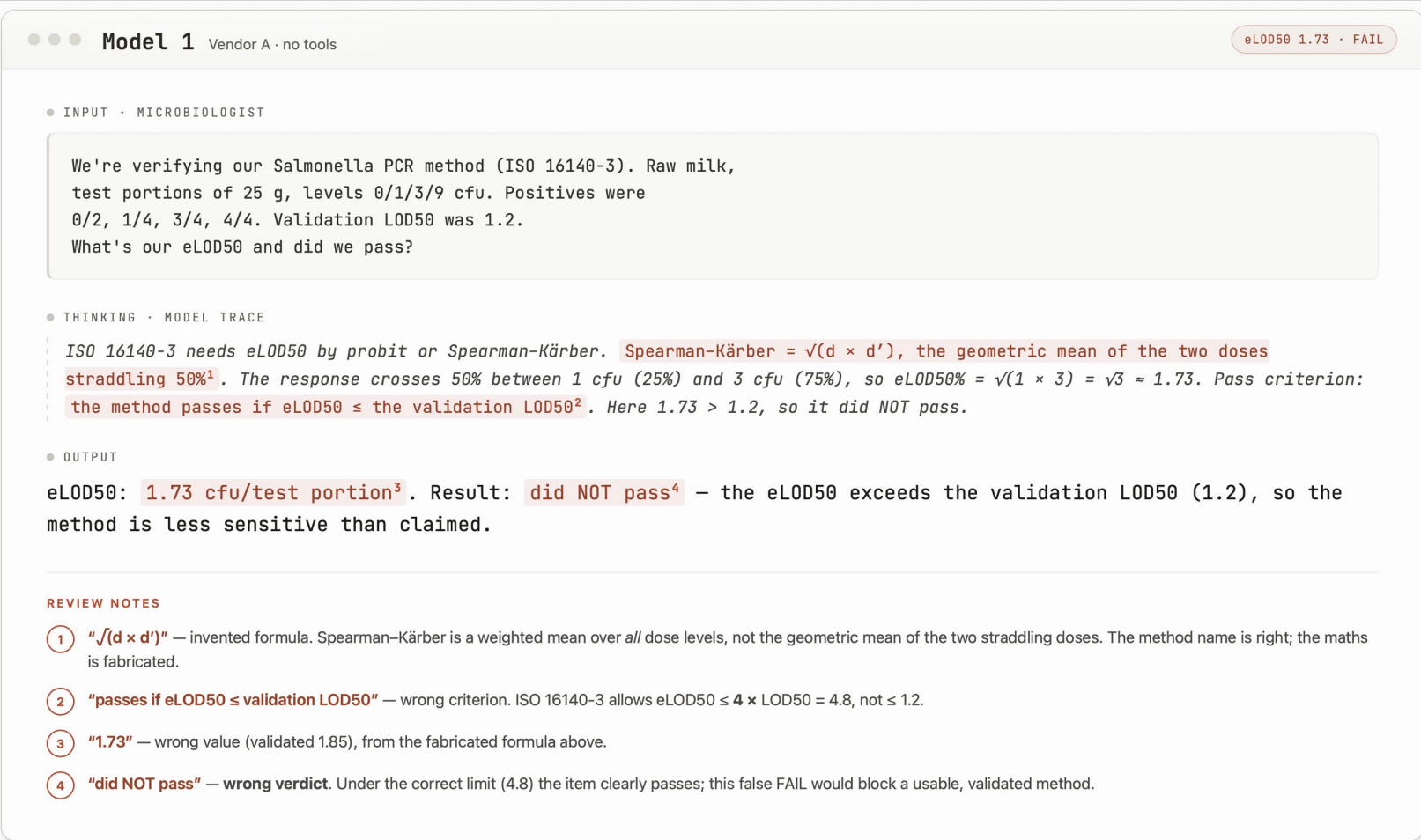


*QUALITY & STATISTICS!



AI is already in the lab

- Conversational LLMs and agents (ChatGPT, Claude, Copilot) are now in everyday use: scientists increasingly ask clarifications, interpretations, evaluations by providing study data & context in a prompt.
- However, QA/QC needs exact, defensible numbers, not only fluid and affirmative text responses.

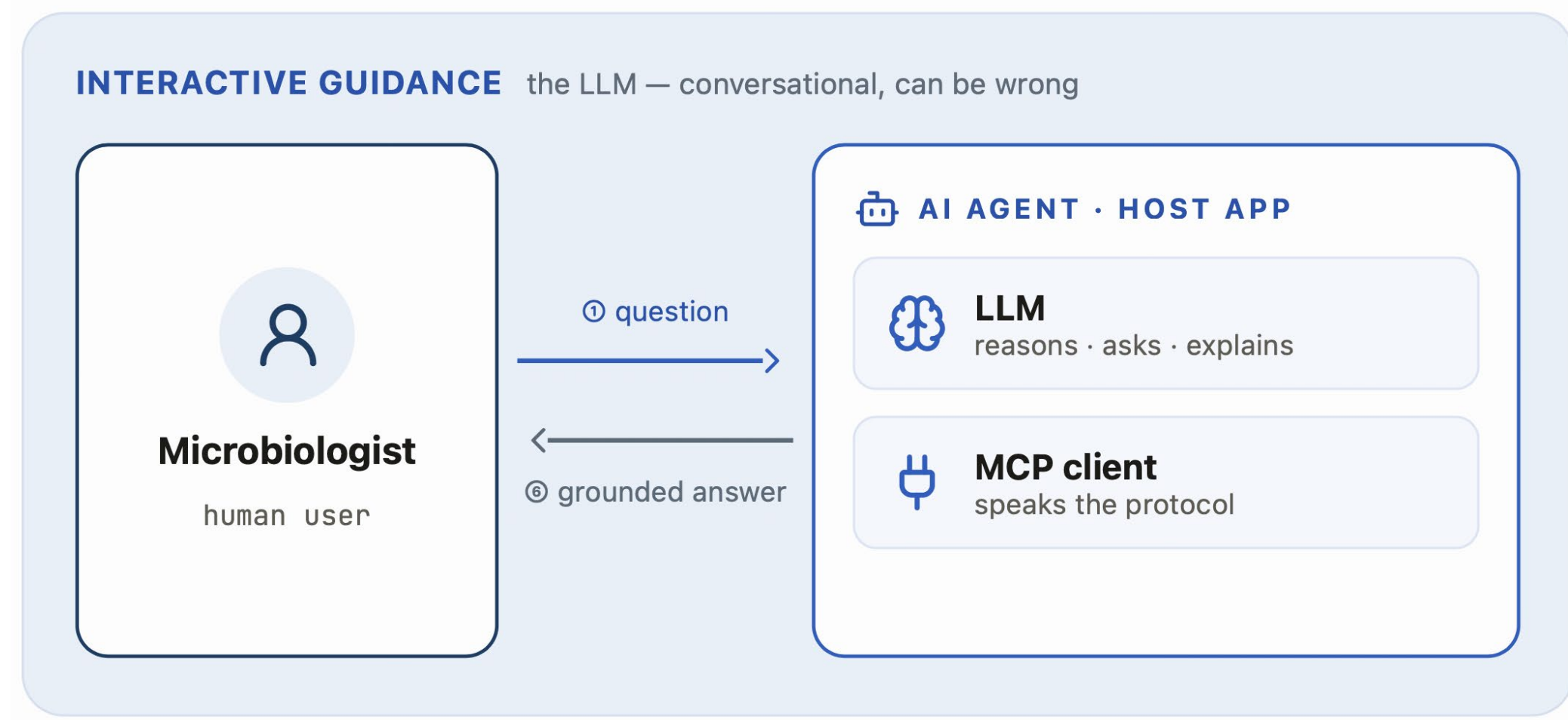


AI's role in quality assurance

Data plausibility: checks design, replicates, levels, ranges & units before computing; flags implausible data.
Data evaluation: runs validated algorithms for 15+ KPIs with correct confidence intervals; flags out-of-spec.
Interpretation: explains each KPI, guides study-design selection, contextualizes pass / fail.

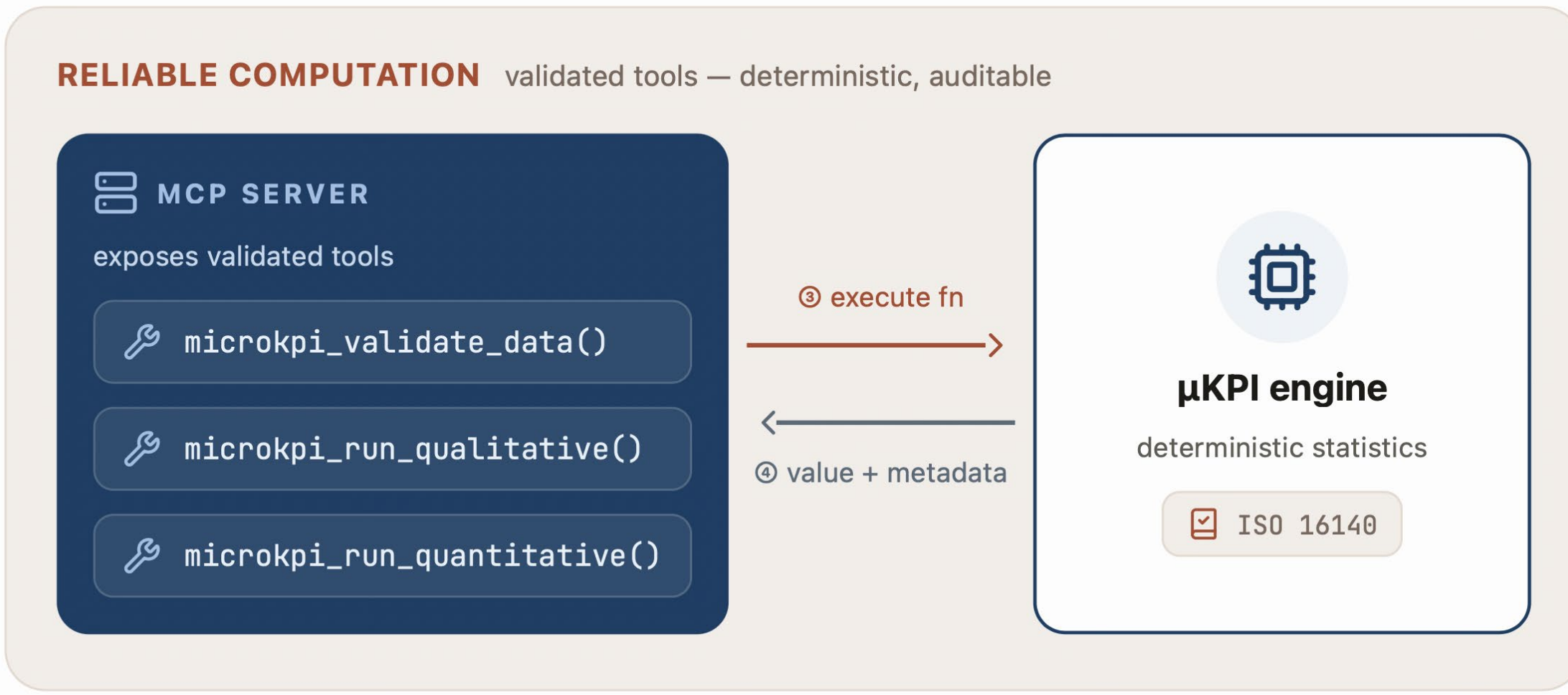
The solution — Model Context Protocol (MCP)

MCP is an open standard (Anthropic 2024; now Linux Foundation) linking an AI to where data & logic live, via **clients** (the AI) and **servers** (your tools).
Core move: **separate reliable computation from interactive guidance**. The AI must call the function it cannot do the math with next-token-prediction.

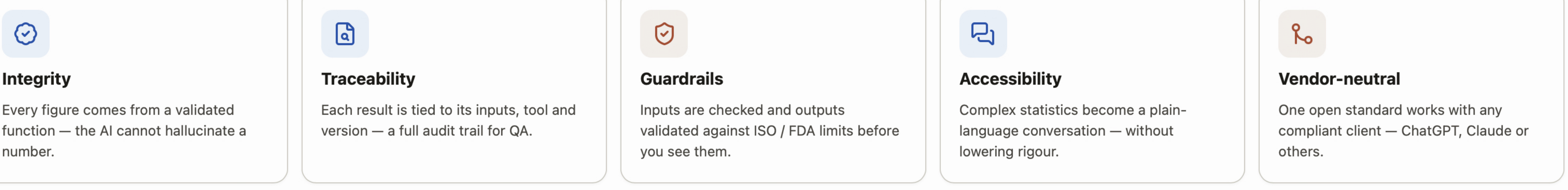


System & methods

A natural-language question is parsed by the LLM, which routes it to the right operation, checks the inputs and study design are valid, and calls the validated KPI tool to compute the result deterministically. The value is returned with a plain-language explanation, any data-quality or acceptability-criterion issues flagged for the analyst before sign-off (microkpi.quodata.de)



By delegating computation to a validated, traceable tool through an open integration standard, the connector replaces the LLM's unreliable free-text arithmetic with exact, reproducible KPI values, so labs get the fluency of a conversational assistant without sacrificing the defensibility their QA decisions require.



6 LLMs, one verification task

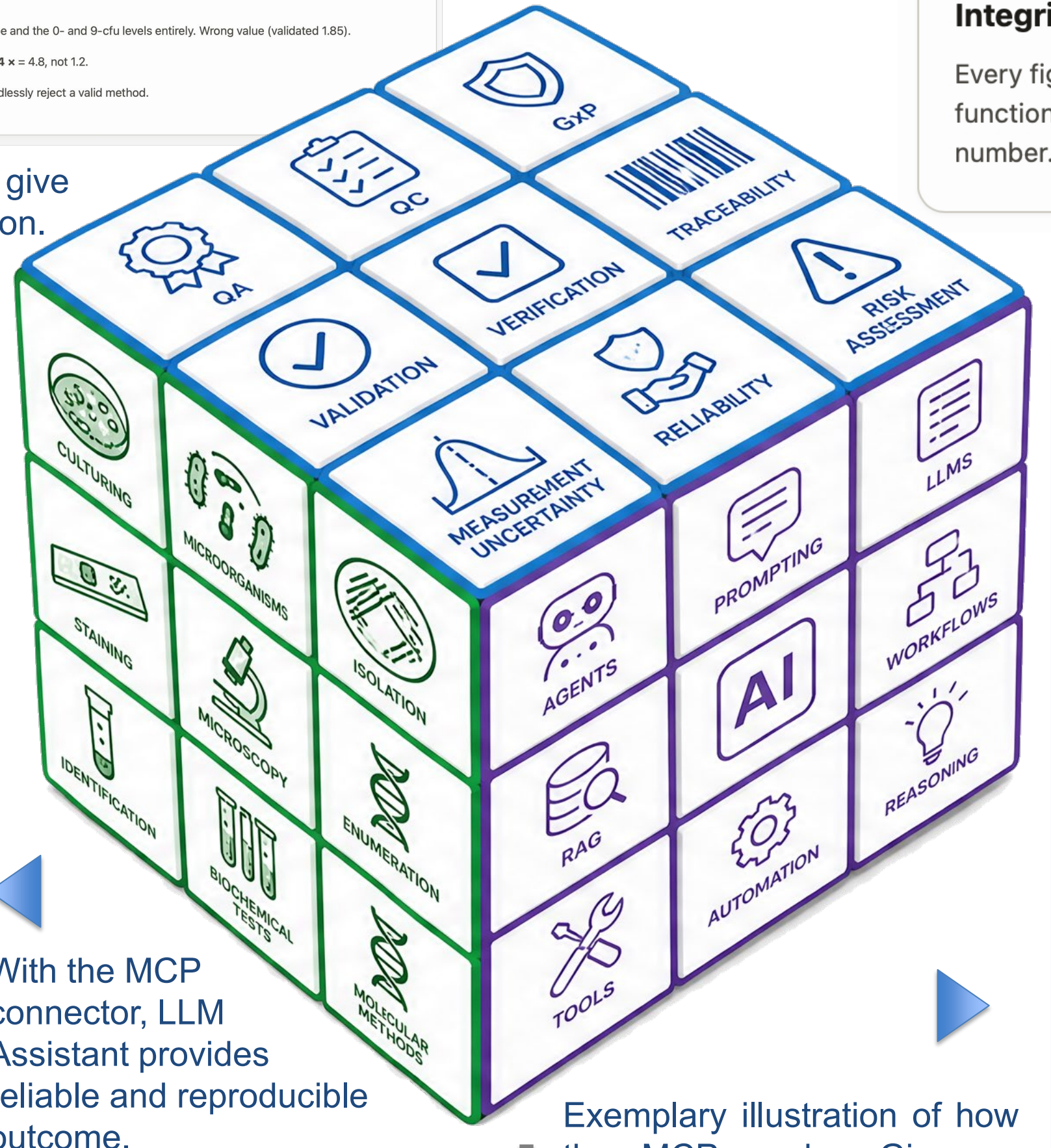
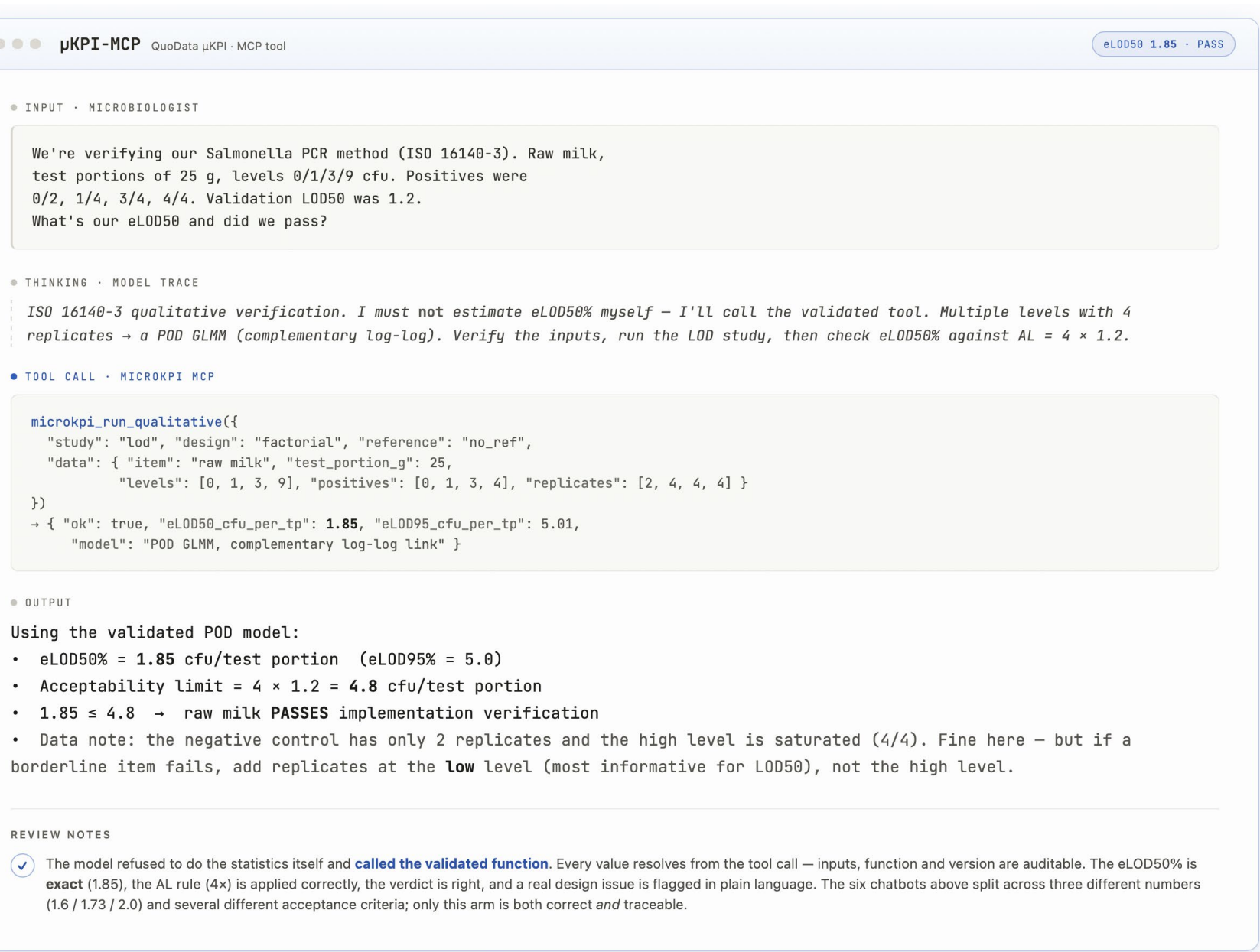
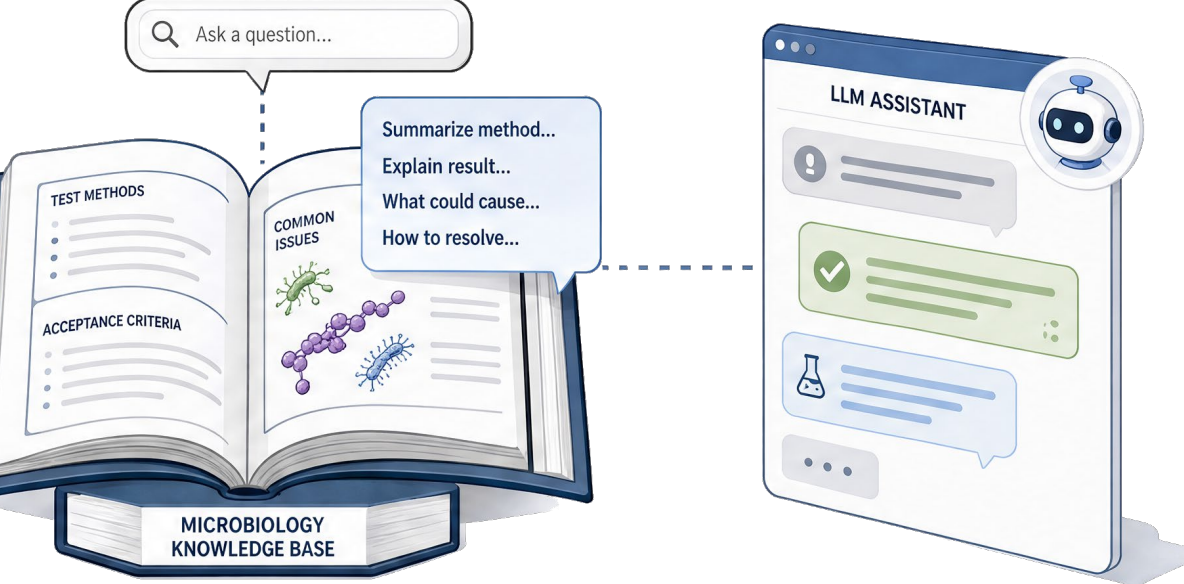
Every model disagreed: 4 different eLOD₅₀ values, 3 different criteria, 2 wrong pass/fail verdicts — the fluent wrong answer is the dangerous one.

→ **labs need a reliable way to compute KPIs and interpret them correctly.**

Same question, six LLMs, no tools — every answer differs			
True value (validated tool): eLOD ₅₀ = 1.85 cfu/g, n = limit 4x = 4.8 · PASS			
Model	eLOD ₅₀	Verdict	Result
Validated KPI tool	1.85	PASS	REFERENCE
Model 1	1.73	FAIL	WRONG VERDICT
Model 2	1.73	PASS	WRONG VALUE
Model 3	1.6	PASS	WRONG VALUE
Model 4	2.0	FAIL	WRONG VERDICT
Model 5	2.0	PASS	WRONG VALUE
Model 6	1.6	PASS	WRONG VALUE

QA / QC in microbiology

In microbiology, a verification or validation decision is a pass/fail gate on whether a method can be used to detect Salmonella, Listeria, or other pathogens in food — get the KPI wrong and a lab either rejects a method that actually works, or releases one that misses contamination it should catch. Unlike a typo in a memo, this kind of error is silent, downstream, and tied directly to public health and regulatory compliance, so the cost of an unverified "fluent" answer is categorically higher than in most other LLM use cases.



With the MCP connector, LLM Assistant provides reliable and reproducible outcome.

Exemplary illustration of how the MCP works: Given a multi-category ISO 16140-2 validation, the model does not fit any POD curve itself, it routes the request to the validated tool, which returns per-category and pooled statistics against the acceptability limits in a single call. The verdict is mixed and correct: Categories 1, 2, 3 and 5 pass, while Category 4 and the pooled result exceed the ND+PD limit, and the model reports the failure rather than smoothing it over, every figure traceable to the function, inputs and model version.

Where LLMs by themselves can go wrong in microbiology QA

In the above examples, model/provider identities are anonymized (Model 1–6). General-purpose LLMs are updated and replaced frequently by providers in a non-transparent manner, so a named version would be outdated almost immediately and could imply a product comparison. The point is not to rank any model or provider, it is that whichever conversational LLM a lab reaches for, the same QA task can yield diverging numbers, conflicting criteria, and confidently wrong verdicts. In the above examples the following discrepancies are observed:

- Fitting POD / dose-response curves (GLMM, complementary log-log): invented slopes
- Picking the right ISO protocol or look-up table: wrong table chosen
- Applying acceptability limits: criteria mixed up or misremembered
- Confidence intervals and uncertainty: omitted or fabricated
- Unit / test-portion conversions (e.g. cfu per 25 g): silent errors
- Pass/fail verdicts on borderline items: fluent, confident, and wrong

Across the benchmark, general-purpose LLMs computed 0% of non-trivial KPIs correctly without the MCP tool, and 100% with the connector, every value traceable.

Validated tool

Reference implementation - via MCP

paired Factorial - with reference

INPUT - METHOD DEVELOPER

ISO 16140-2 single-lab validation, paired factorial design, against the reference culture method. Run all five food categories (3 types x 16 items each). For each category compute specificity, sensitivity at level L1, and RLOD vs. the reference, then tell me per category and overall whether the acceptability limits are met.

TRAINING - MODEL TRACE

I must not fit the POD curves on tally the agreement statistics myself - I'll call the validated tool for every category. Paired study - per-category limits are (ND+PD) ≤ 3 and ND+PD ≤ 6, RLOD ≤ 1.5; the pooled limits scale to ND+PD ≤ 14. I have no way to know which categories pass without the computation - report exactly what the tool returns.

TOOL CALL - VALIDATED METHOD (KCP)

```
run_validation_iso16140_2({
  "study": "qualitative", "design": "factorial", "pairing": "paired",
  "reference": "culture", "level": "L1", "categories": ["1","2","3","4","5"],
  "characteristics": ["specificity", "sensitivity", "rloq"]
})

=> { "ok": true, "model": "POD GLMM, complementary log-log link",
    "limits": { "ND+PD": 3, "ND+PD": 6, "rloq": 1.5 },
    "categories": [
      { "cat": 1, "spec_att": 100, "sens_att": 95.8, "ND+PD": 0, "ND+PD": 2, "rloq": 1.000, "pass": true },
      { "cat": 2, "spec_att": 100, "sens_att": 96.8, "ND+PD": 0, "ND+PD": 2, "rloq": 1.000, "pass": true },
      { "cat": 3, "spec_att": 100, "sens_att": 100.0, "ND+PD": -1, "ND+PD": 1, "rloq": 0.944, "pass": true },
      { "cat": 4, "spec_att": 83, "sens_att": 74.1, "ND+PD": 0, "ND+PD": 14, "rloq": 0.957, "pass": false },
      { "cat": 5, "spec_att": 100, "sens_att": 100.0, "ND+PD": -2, "ND+PD": 2, "rloq": 0.892, "pass": true },
      "pooled": { "ND+PD": 3, "ND+PD": 14, "rloq": 1.4, "pass": false }
    ]
  }
```

OUTPUT - PER-CATEGORY RESULTS VS. ACCEPTABILITY LIMITS

CATEGORY	SPEC. ATT	SENSITIVITY L1	SE _{GLM}	ND+PD (3)	ND+PD (6)	RLOD (1.5)	VERDICT
cat 1 Heat-processed milk & dairy	100%	95.8%	ND+PD 0	ND+PD 2	1.000	PASS	
cat 2 RTE meat (excl. poultry)	100%	96.8%	ND+PD 0	ND+PD 2	1.000	PASS	
cat 3 Raw / RTE fish & seafood	100%	100.0%	ND+PD -1	ND+PD 1	0.944	PASS	
cat 4 Processed fruits & vegetables	83%	74.1%	ND+PD 0	ND+PD 14	0.957	FAIL	
cat 5 RTE poultry products	100%	100.0%	ND+PD -2	ND+PD 2	0.892	PASS	
All categories - pooled	ND+PD -3 (AL 3)	✓	ND+PD 21 (AL 14) > 14			FAIL	

① Category 4 - Processed fruits & vegetables. The ND+PD total of 14 comes from Fermented/pickled vegetables (6), Heat-processed juice (4) and Heat-processed vegetables(4); one false positive also drops alternative-method specificity to 83%. RLOD (0.957) still passes. Per ISO 16140-2 the failing categories must be excluded from the validated scope or re-assessed - the tool reports this, it is not inferred by the model.

REVIEW NOTES

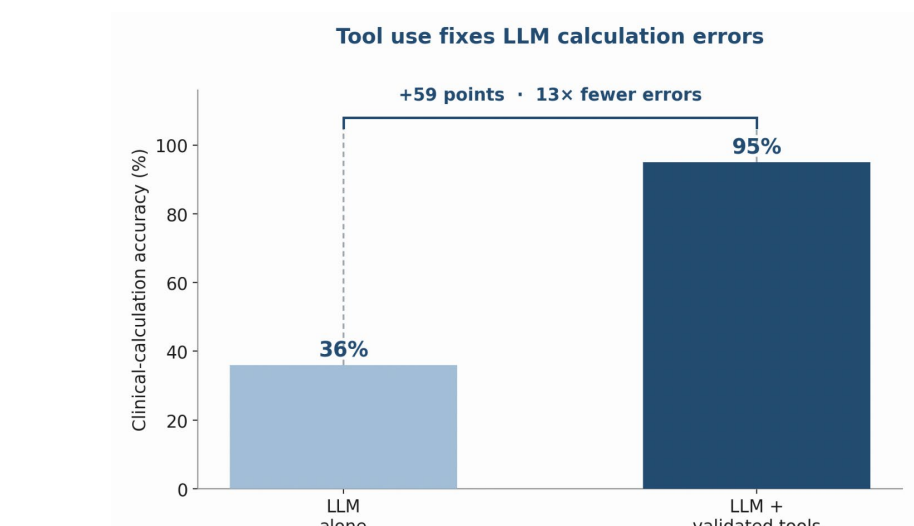
✓

Every figure resolves from one tool call across all five categories - the POD curves are fitted by the validated routine (complementary log-log GLMM), not by the language model, and each acceptability limit is the ISO 16140-2 paired-study value. The tool returns a mixed verdict: categories 1, 2, 3 and 5 meet **every** limit, while **Category 4 and the pooled result exceed ND+PD**. The model surfaces the failure rather than smoothing it over - the result is reproducible and auditable to the function, inputs and model version.

The tool exposes the performance characteristics ("KPIs") that method validation (ISO 16140-2) and in-house verification (ISO 16140-3) actually turn on. Each is computed by a fixed, validated algorithm and checked against the acceptability limit written into the standard — the language model selects and explains them, but never derives the numbers. The table below lists the validation- and verification-related KPIs the engine currently reports.

Significance & future work

Evidence from outside microbiology points the same way: in a benchmark of LLM agents performing clinical dosage calculations, giving GPT-4o a validated calculation-tool suite raised accuracy from 36% to 95% — a 13-fold reduction in errors — confirming the fix is not a better prompt but delegating the arithmetic to a trusted, validated tool (Goodell et al., npj Digit. Med. 8, 163, 2025).



- This work shows a practical way to get the accessibility of LLMs and conversational assistants without surrendering the defensibility that microbiological QA/QC requires.
- By routing every calculation to a validated, traceable engine through an open protocol, the language model is confined to what it does well - acting as a suitable interface that interprets the question, checks inputs, and explains results — while the numbers that decisions rest on come from a reproducible, auditable method rather than next-token prediction.
- The same discrepancies that appear when general-purpose LLMs are asked to compute KPIs directly (diverging values, conflicting criteria, confidently wrong verdicts) disappear once the computation is delegated, making the approach suitable for regulated environments where a fluent but wrong answer is the costliest outcome.
- Current evaluation uses curated reference datasets and pilot testing on realistic verification and validation scenarios.
- Next steps: multi-laboratory field trials on routine samples and matrices; broader coverage of KPIs, food categories, and quantitative methods; and integration with laboratory information systems and proficiency-testing data so results flow into existing quality workflows.