

Paper SD06

AI Chatbot in an SCE

Shivoy Pandita, Entimo AG, Berlin, Germany
Marc Jantke, Entimo AG, Berlin, Germany

PHUSE EU Connect 2025, Hamburg

ABSTRACT

As AI adoption evolves, its integration into Statistical Computing Environments (SCEs) presents both opportunities and regulatory challenges in the pharmaceutical industry. *entimICE FastTrack*® introduces a secure, configurable AI assistant, enabling compliant, role-based interactions with a locally deployed large language model (LLM). No clinical or corporate data leaves the environment, ensuring full data privacy and non-training compliance. Users can leverage Retrieval Augmented Generation (RAG) using documents, programs, or datasets, facilitating productivity enhancements such as synthetic data creation or program mock-ups. This AI assistant focuses on secure, domain-specific information retrieval rather than general-purpose chat or coding, which are better handled in IDEs like Posit Workbench or by the more popular Chat GPT and Copilot. The architecture supports full AI feature isolation, user-specific augmentation stores, and future extensions like shared knowledge repositories or public API integrations. This approach makes *entimICE FastTrack*® as a flexible, future-ready solution for AI in regulated environments.

INTRODUCTION

The validation of statistical outputs remains one of the most time-consuming and subjective stages of clinical reporting. Although well-defined processes exist for dataset generation and standard compliance, verifying that final Tables, Listings, and Figures (TLFs) align exactly with their intended mock shells and the Statistical Analysis Plan (SAP) still relies heavily on manual review. This process requires a detailed comparison between documents that are heterogeneous in format, style, and intent, statistical versus visual.

This paper proposes a prototype framework for automating this validation step using AI-assisted workflows embedded within a GxP-compliant Statistical Computing Environment (SCE). The framework is part of a broader AI capability in *entimICE FastTrack*®, where a locally deployed large language model can be accessed either through conversational chat. In this paper we focus on one the AI assistant being used as a validation tool that uses Retrieval Augmented Generation (RAG) and document understanding to extract, compare, and explain differences between TLFs, mock shells, and SAPs. Instead of treating AI as an autonomous reviewer, the system acts as a traceable, explainable assistant that pre-screens potential inconsistencies for human verification.

In every clinical study, TLF validation serves as a critical checkpoint between analysis and submission. While the SAP defines the analytical intent, and the mock shells describe the visual representation, the final TLFs embody the executed outputs of programming activity. Ensuring that these three elements are consistent is essential to scientific accuracy and regulatory compliance.

Yet, this validation is traditionally performed manually. Reviewers meticulously cross-check SAP sections, statistical footnotes, and mock shell layouts against each TLF. They look for discrepancies in statistical methods, variable order, population definitions, and formatting conventions. The process demands deep contextual knowledge and a high level of concentration, especially for large studies with hundreds of outputs and multiple versions. The challenge lies in the heterogeneity of the inputs, SAPs and shells are typically written in natural language or semi-structured tables, while TLFs are machine-generated outputs in formats such as RTF or PDF. Conventional rule-based validation cannot interpret the semantic intent of these documents.

Recent advances in natural language processing (NLP) and large language models (LLMs) open a new avenue for tackling this challenge. Rather than manually reading and comparing documents, AI can be trained to “understand” statistical context, extract key parameters, and identify alignment or deviation between design intent and actual output. Within a controlled, GxP-validated SCE such as *entimICE FastTrack*®, these capabilities can be deployed safely, transparently, and under human supervision.

METHODOLOGY

Overall Architecture

The proposed system integrates three document domains:

1. **SAP (Statistical Analysis Plan):** defines analysis objectives, population criteria, endpoints, and statistical tests.
2. **Mock Shells:** describe the visual and structural format of each planned output, including titles, variables, summary measures, and footnotes.
3. **Final TLFs:** represent the executed outputs of the statistical programming process.

Each artefact is version-controlled and linked to its study metadata within *entimiCE FastTrack*®. The AI layer operates inside the SCE boundary to maintain data privacy and traceability.

To evaluate the feasibility of AI-assisted validation, three representative outputs were selected from a blinded clinical study dataset: **Table 14.1-1 Summary of Demographics of Subjects in the Safety Population**, **Table 14.3.1-1 Frequency of Subjects Experiencing Treatment-Emergent Adverse Events** and **Table_10.1-1 Subject Disposition**. Each of these outputs was compared against its corresponding sections in the **Statistical Analysis Plan (SAP_004)**. The SAP_004 reference shell and specifications were sourced from a publicly available Statistical Analysis Plan posted on *ClinicalTrials.gov* for study NCT03882424. The test tables used in this paper were modelled on the conventions defined in the SAP to ensure realistic validation scenarios.

SAP_004 and three representative TLFs were used to simulate AI-assisted validation.

Statistics	TRS003, 3 mg/kg IV Dose (A)	China-Approved Benicizumab, 3 mg/kg IV Dose (B)
N	27	29
Mean	32.23	36.14
SD	8.1	7.3
Median	34.0	37.0
Min, Max	21.0-52.0	22.0-51.0
n (%)	0 (0.0)	0 (0.0)
n (%)	19 (70.4%)	20 (69.0%)
n (%)	8 (29.6%)	9 (31.0%)
n (%)	27 (100.0%)	29 (100.0%)
n (%)	0 (0.0%)	0 (0.0%)
n (%)	27 (100.0%)	29 (100.0%)
n (%)	27 (100.0%)	29 (100.0%)
n (%)	2 (7.4%)	3 (10.3%)

Category	Statistic	TRS003, 3 mg/kg IV Infusion Dose (A)	China-Approved Benicizumab, 3 mg/kg IV Infusion Dose (B)
Age (years)	N	xx	xx
	Mean	xx.x	xx.x
	SD	xx.x	xx.x
	Median	xx.x	xx.x
Age Groups	Min, Max	xx-xx	xx-xx
	n (%)		
	<18	x (xx.x)	x (xx.x)
	18-40	x (xx.x)	x (xx.x)
Efficacy: Not Hispanic or Latino	n (%)		
	>40	x (xx.x)	x (xx.x)
	n (%)		
	Hispanic or Latino	x (xx.x)	x (xx.x)
Race	n (%)		
	White	x (xx.x)	x (xx.x)
	Black	x (xx.x)	x (xx.x)
	n (%)		

The workflow begins with document ingestion. NLP pipelines segment and parse the SAP and mock shells, converting them into structured metadata. Key analytical elements such as endpoint names, test methods, population identifiers, and summary statistics are extracted into a semantic model.

The TLF outputs are processed through a similar pipeline. Text and tabular components are extracted and converted into comparable structures. Using vector-based similarity analysis and natural-language reasoning, the system performs a three-way alignment:

- **TLF vs. mock shell:** structural and visual conformity.
- **TLF vs. SAP:** analytical and statistical conformity.
- **SAP vs. mock shell:** design-level consistency.

AI Reasoning Layer

At the core of the system is a domain-tuned LLM capable of understanding statistical language and table structures. The model is prompted through structured templates and produces the explanatory response. We selected the **Mistral 7B** open-weight model to allow on-premises deployment, ensuring full control over data access and model behaviour. This was essential for GxP validation and to meet data privacy standards that prohibit sending data to external APIs.

Embedding Structured Clinical Data

Embeddings are numerical representations of text that capture the semantic meaning and contextual relationships between words, sentences, or documents. These vector-based representations enable AI systems to compare and retrieve relevant information based on meaning rather than exact keyword matching.

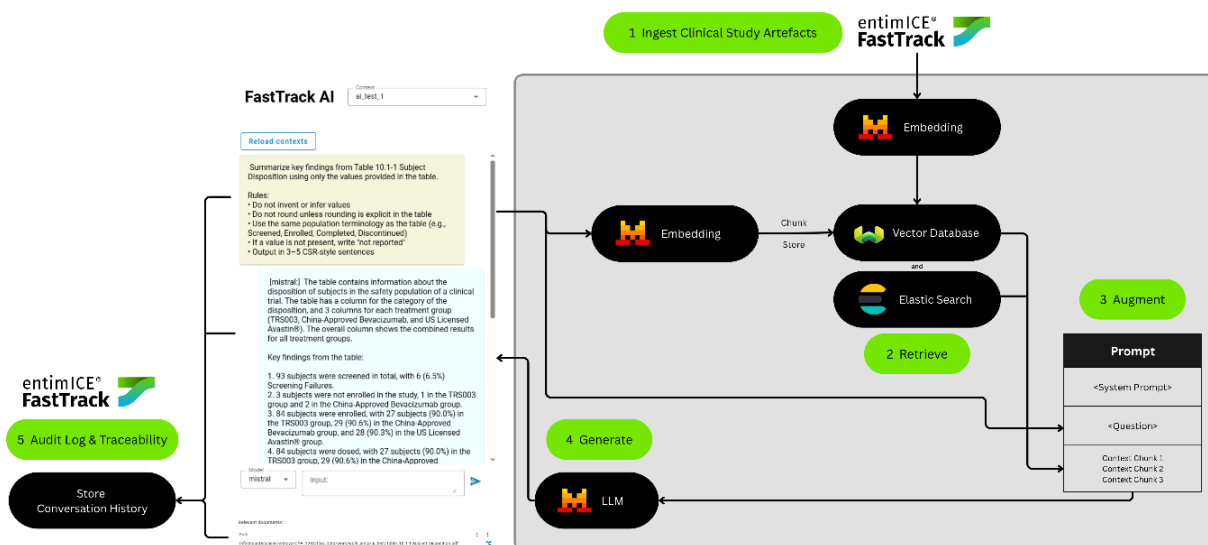
In *entimICE FastTrack*®, study artifacts such as the Statistical Analysis Plan (SAP), mock shells, and TLF outputs are embedded and indexed in a Vector Database (Vector DB). Then tables from SAPs and TLFs are extracted, converted into JSON, preserving column structure, denominators, statistical values, and hierarchical groupings. These structured extractions are then put into a searchable index (Elasticsearch), alongside embeddings to enable hybrid retrieval, combining semantic search with precise value-based and hierarchy-aware comparisons. This design ensures that AI outputs reference authoritative study documentation, reduces hallucination risk, and aligns with GxP expectations for traceability, explainability, and data provenance.

Retrieval Augmented Generation (RAG)

Retrieval Augmented Generation (RAG) combines traditional language models with a retriever module to enhance content generation. It leverages pre-existing knowledge by retrieving relevant information from external sources, allowing the model to generate more informed and contextually rich responses. RAG is particularly useful for tasks that require incorporating external information, such as question answering or content creation.

Integration with GxP Environment

To ensure compliance, all AI-generated findings are recorded in the audit trail, alongside model version, prompt template, and input document references. The reviewer can trace each statement back to the original SAP paragraph, mock shell page, and TLF file. This preserves explainability and ensures that AI support never replaces regulated human judgment.



USE CASE SCENARIO

The goal was to determine whether a language-model–driven comparison could correctly interpret the analytical intent expressed in the SAP and mock shells and identify meaningful deviations in the final outputs.

In a conventional process, a reviewer must visually compare the output against both documents to identify this inconsistency. In the AI-assisted workflow, the system automatically detects the deviation, generates a narrative explanation, and categorizes the issue as a “statistical method mismatch”.

The reviewer receives an annotated report with direct links to the SAP section and the output file. With this triaged summary, human QC can focus on meaningful discrepancies rather than routine verification. Over multiple studies, the accumulated data from these reports can inform continuous improvement, helping teams identify deviations or unclear SAP language that trigger false positives.

Across the following examples, the AI-assisted validation pipeline successfully demonstrated alignment with SAP language patterns, quantitative fields, and population references. It also identified meaningful discrepancies without relying on rigid templates.

Figure 1 – entimICE FastTrack® generating traceable Subject Disposition summary

The screenshot displays the entimICE FastTrack application interface. The main window shows a table titled "Table 10.1-1 Subject Disposition" with three columns: "Category", "TRS003, 3 mg/kg IV Infusion Dose (A)", and "China-Approved Bevacizumab, 3 mg/kg IV Infusion Dose (B)". The table lists various categories such as Screened, Screening Failures, Not Enrolled, Enrolled, Dosed, Not Dosed, Completed, Number of Subjects Discontinued, and Primary Reason for Discontinuation, along with their respective counts and percentages for both groups.

On the right side, the "FastTrack AI" panel is visible, showing a "Summarize key findings from Table 10.1-1 Subject Disposition using only the values provided in the table." prompt. Below the prompt, there are "Rules" and "Key findings from the table:" sections. The "Key findings" section lists four points: 1. 93 subjects were screened in total, with 6 (6.5%) Screening Failures. 2. 3 subjects were not enrolled in the study, 1 in the TRS003 group and 2 in the China-Approved Bevacizumab group. 3. 84 subjects were enrolled, with 27 subjects (90.0%) in the TRS003 group, 29 (90.6%) in the China-Approved Bevacizumab group, and 28 (90.3%) in the US Licensed Avastin® group. 4. 84 subjects were dosed, with 27 subjects (90.0%) in the TRS003 group, 29 (90.6%) in the China-Approved Bevacizumab group, and 28 (90.3%) in the US Licensed Avastin® group.

Prompt Details

Relevant documents:

Path	
InfectiousDiseases/enticyp/CYP_1234/Cyp_Data/work/work_xmja/ai_test/Table_10.1-1 Subject Disposition.pdf	

Subject Disposition narratives are frequently reused across CSR sections, protocol amendments, and data-cleaning discussions. Automating this step ensures consistency while allowing statisticians to focus on interpretation rather than transcription. In multi-study programs, the standardized output also enables cross-study comparison and reuse of summary templates. Because the summarization is fully grounded in the TLF and controlled by deterministic prompts, it meets traceability expectations while reducing the cognitive load on reviewers.

Figure 2 – entimICE FastTrack® ingests both the SAP and the corresponding TLF (e.g., Table 14.3.1-1) and extracts numerical incidence values for each SOC and PT. These values are compared against the incidence-based ordering rules defined in SAP_004 using a hybrid RAG and structured-extraction pipeline. The AI Assistant produces a natural-language compliance assessment, citing the source documents used for each inference step, and delivers a final classification (Fully Compliant, Partially Compliant, or Non-Compliant).

The screenshot displays the entimICE FastTrack AI interface. On the left, a table titled "Table 14.3.1-1 Frequency of Subjects Experiencing Treatment-Emergent Adverse Events and Number of Events" is shown. The table compares data for TRS003, 3 mg/kg IV Infusion (N=27) and China-Approved Bevacizumab, 3 mg/kg IV Infusion (N=29). The table lists various System Organ Classes (SOCs) and Preferred Terms (PTs) with their respective incidence rates.

On the right, the "FastTrack AI" section shows a prompt and the AI's response. The prompt asks to validate Table 14.3.1-1 against SAP_004 ordering rules. The AI's response provides a detailed compliance assessment, including the following information:

- Overall Incidence Rates:**
 - Infections and infestations: 9 (10.7%)
 - Nervous system disorders: 8 (9.5%)
 - Gastrointestinal disorders: 7 (8.3%)
 - General disorders and administration site conditions: 1 (1.2%)
- Preferred Terms (PTs) within SOC:**
 - Infections and infestations: 5 (6.0%)
 - Nasopharyngitis: 5 (6.0%)
 - Upper respiratory tract infection: 4 (4.8%)
- Conclusion:** The table aligns with the SAP_004 ordering rules, with SOC and PTs correctly ordered by incidence rate. There are no deviations from the expected order.
- Compliance Status:** Fully Compliant

A "Prompt Details" window is also visible, showing the relevant documents used for the assessment:

- Path: InfectiousDiseases/enticyc/CYP_1234/Cyp_Data/work_xmja/ai-test/TEAE_Table_14.3.1_1.pdf
- Path: InfectiousDiseases/enticyc/CYP_1234/Cyp_Data/work_xmja/ai-test/SAP_004.pdf

The purpose of including GPT-4o is to show that the validation workflow and RAG architecture are model-agnostic, while the governed and supported SCE implementation remains Mistral-based. The current AI Assistant provides document-level provenance through the Prompt Details window, ensuring that all retrieved context is limited to validated study artefacts stored within the SCE. To further strengthen auditability and align with ongoing AI governance expectations, additional metadata capture capabilities are planned. These include:

- **Page-level provenance**, allowing the system to record exactly which pages of each PDF were retrieved and used during an inference.
- **Prompt versioning**, ensuring that any changes to system prompts or template structures are governed and traceable.
- **Model version tagging**, where every inference is automatically linked to the specific model build and configuration active at the time.

Figure 3 – AI validation of Table 14.3.1-1 Frequency of Subjects Experiencing Treatment-Emergent Adverse Events

The LLM model (Mistral) verifies internal consistency between the total “Subjects with TEAEs” counts and the summed values across all System Organ Classes (SOCs). For each treatment group (TRS003, China-Approved Bevacizumab, and US-Licensed Avastin), the overall counts are less than or equal to the SOC totals, and rounding deviations are within 0.1%. The table is assessed as Compliant with SAP_004 specifications.

The screenshot shows the entimICE FastTrack interface. The main window displays Table 14.3.1-1, titled "Frequency of Subjects Experiencing Treatment-Emergent Adverse Events and Number of Events Summarized per Treatment Group". The table has four columns: MedDRA® System Organ Class, MedDRA® Preferred Term, Statistic, and three treatment groups: TRS003, 3 mg/kg IV Infusion (N=27), China-Approved Bevacizumab, 3 mg/kg IV Infusion (N=29), and US Licensed Avastin® (Bevacizumab) (N=26). The rows list various adverse events, including Infections and infestations, Nervous system disorders, Gastrointestinal disorders, and General disorders and administration site conditions. To the right of the table, the "FastTrack AI" panel shows a prompt and the model's response. The prompt asks to verify the overall "Subjects with TEAEs" count for each treatment group against the sum of subjects across all SOC classes. The model's response confirms that the overall counts are less than or equal to the SOC totals and that rounding deviations are within 0.1%, concluding that the table is compliant.

MedDRA® System Organ Class	MedDRA® Preferred Term	Statistic	TRS003, 3 mg/kg IV Infusion Dose (N=27)	China-Approved Bevacizumab, 3 mg/kg IV Infusion Dose (N=29)	US Licensed Avastin® (Bevacizumab) Dose (N=26)
	Number of TEAEs	E	8	7	9
	Subjects with TEAEs	n (%)	7 (25.9%)	6 (20.7%)	9
Infections and infestations		n (%) E	3 (11.1%) 3	2 (6.9%) 2	4
Nasopharyngitis		n (%) E	2 (7.4%) 2	1 (3.4%) 1	2
Upper respiratory tract infection		n (%) E	1 (3.7%) 1	1 (3.4%) 1	2
Nervous system disorders		n (%) E	3 (11.1%) 3	2 (6.9%) 2	3
Headache		n (%) E	2 (7.4%) 2	1 (3.4%) 1	2
Dizziness		n (%) E	1 (3.7%) 1	1 (3.4%) 1	1
Gastrointestinal disorders		n (%) E	2 (7.4%) 2	2 (6.9%) 2	3
Nausea		n (%) E	1 (3.7%) 1	1 (3.4%) 1	2
Diarrhea		n (%) E	1 (3.7%) 1	1 (3.4%) 1	1
General disorders and administration site conditions		n (%) E	0 (0.0%) 0	1 (3.4%) 1	0
Injection site pain		n (%) E	0 (0.0%) 0	1 (3.4%) 1	0
Fatigue		n (%) E	0 (0.0%) 0	0 (0.0%) 0	0

Programming Notes:
1) SOC will be presented in descending order of overall incidence rate in terms of frequency of subjects and then in frequency of events (alphabetical order will be used in case of equal rates). For each SOC, PT will be
2) Refer to footnotes for additional instructions.

AI Validation Prompt:
Verify that the overall "Subjects with TEAEs" count for each treatment group in Table 14.3.1-1 is less than or equal to the sum of subjects across all SOC classes (since subjects may appear in multiple SOC classes). Flag only if the overall count exceeds the sum or if percentage rounding deviates by > 0.1 %.

AI Validation Response:
[mistral] Verify that the overall "Subjects with TEAEs" count for each treatment group in Table 14.3.1-1 is less than or equal to the sum of subjects across all SOC classes (since subjects may appear in multiple SOC classes). Flag only if the overall count exceeds the sum or if percentage rounding deviates by > 0.1 %.

AI Validation Conclusion:
Answer: The overall count for "Subjects with TEAEs" in each treatment group is: TRS003 (22), China-Approved Bevacizumab (20), US Licensed Avastin® (26). The sum of subjects across all SOC classes is 84. The overall counts are all less than or equal to the sum of subjects across all SOC classes. No rounding deviations exceed 0.1 %. The table is compliant.

Prompt Details

Relevant documents:

Path
...ata/work/work_xmja/ai_test/Table 14.3.1-1 Frequency of Subjects Experiencing Treatment-Emergent Adverse Events.pdf

For this example, the AI Assistant operates under a constrained prompting architecture designed specifically for validated statistical workflows. Unlike open-ended conversational models, FastTrack's implementation restricts LLM behaviour through:

- Schema-locked system prompts that prevent speculative reasoning.
- Structured output templates enforcing precise compliance categories.
- Document-bound context retrieval, ensuring the model cannot rely on external knowledge.
- Numeric grounding, where extracted values must be echoed back rather than inferred.
- Error-sensitive phrasing rules, prohibiting high-risk language such as definitive guarantees without supporting evidence.

These behavioural constraints are essential in a GxP-regulated setting, where the role of the LLM is to support human reviewers, not to produce autonomous statistical conclusions. The example in Figure 3 illustrates how constrained reasoning reinforces traceability and audit reliability.

Figure 4: AI cross-validation of “Subjects with TEAEs” between Table 14.3.1-1 and Table 14.1-1

The LLM model (Mistral) compares the “Subjects with TEAEs” (n %) from Table 14.3.1-1 against the corresponding Safety Population N from Table 14.1-1 to confirm internal consistency. Calculated percentages of total subjects experiencing TEAEs per treatment group (TRS003, China-Approved Bevacizumab, and US-Licensed Avastin) align precisely with the Safety Population denominators, showing no rounding deviations. The tables are assessed as **Fully Compliant**.

The screenshot displays the entimICE FastTrack AI interface. The main window shows a table titled "Table 14.3.1-1 Frequency of Subjects Experiencing Treatment-Emergent Adverse Events and Number of Events Summarized per Treatment Group". The table has columns for MedDRA System Organ Class, MedDRA Preferred Term, Statistic, and three treatment groups: TRS003, 3 mg/kg IV Infusion Dose; China-Approved Bevacizumab, 3 mg/kg IV Infusion Dose; and US Licensed Avastin® (Bevacizumab), 3 mg/kg IV Infusion Dose. The table lists various adverse events such as Infections and infestations, Nasopharyngitis, Upper respiratory tract infection, Nervous system disorders, Headache, Dizziness, Gastrointestinal disorders, Nausea, Diarrhea, General disorders and administration site conditions, Injection site pain, and Fatigue. A prompt details window is open, showing the relevant documents used for the AI analysis. The window lists two documents: "...ata/work/work_xmja/ai_test/Table 14.3.1-1 Frequency of Subjects Experiencing Treatment-Emergent Adverse Events.pdf" and "...234/Cyp_Data/work/work_xmja/ai_test/Table 14.1-1 Summary of Demographics of Subjects in the Safety Population.pdf". The window also shows the calculated percentages of total subjects experiencing TEAEs per treatment group, which are consistent with the corresponding Safety Population N in Table 14.1-1, resulting in a "Fully Compliant" status.

MedDRA® System Organ Class MedDRA® Preferred Term	Statistic	TRS003, 3 mg/kg IV Infusion Dose	China-Approved Bevacizumab, 3 mg/kg IV Infusion Dose	US Licensed Avastin® (Bevacizumab), 3 mg/kg IV Infusion Dose
Number of TEAEs	E			
Subjects with TEAEs	n (%)			
Infections and infestations	n (%) E			
Nasopharyngitis	n (%) E			
Upper respiratory tract infection	n (%) E			
Nervous system disorders	n (%) E			
Headache	n (%) E			
Dizziness	n (%) E			
Gastrointestinal disorders	n (%) E			
Nausea	n (%) E			
Diarrhea	n (%) E			
General disorders and administration site conditions	n (%) E	0 (0.0%) 0	1 (3.4%) 1	
Injection site pain	n (%) E	0 (0.0%) 0	1 (3.4%) 1	
Fatigue	n (%) E	0 (0.0%) 0	0 (0.0%) 0	

Programming Notes:
1) SOC will be presented in descending order of overall incidence rate in terms of frequency of subjects and then in frequency of events (alphabetical order will be used in case of equal rates). For each SOC, PT will be presented in descending order of overall incidence rate.
2) Refer to footnotes for additional instructions.

E: Number of TEAEs; N: Number of subjects dosed; n (%): Number and percent of subjects with TEAE; MedDRA®: Medical Dictionary for Regulatory Activities, version 21.0; TEAEs: Treatment-Emergent Adverse Events

All AI-generated observations included traceable citations to the originating SAP and table sections, enabling transparent reviewer audit and alignment with GxP expectations for explainability and documented evidence trails. The system did not operate as a generative black box; instead, it grounded each finding in the source study documents, preserving validation integrity.

RESULTS AND DISCUSSION

Preliminary internal simulations using anonymized SAP–mock–output pairs suggest that an AI-assisted system could automatically pre-screen approximately **75% of validation checks** (in internal simulation benchmarks) typically performed manually. The greatest value lies in early-stage triage: AI can confirm structural conformity (titles, variable order, denominators) and detect inconsistencies in descriptive statistics, footnote text, and population labelling.

Equally important is explainability. Each deviation is accompanied by a rationale generated in plain language, enabling reviewers to quickly determine whether the issue is genuine or due to formatting differences.

From a quality standpoint, this approach introduces consistency across reviewers. Rather than depending on individual interpretation of SAP sections, the AI enforces uniform logic derived from structured templates. Over time, a library of validated comparison patterns could further improve reliability.

Nevertheless, several challenges remain. SAP language can be ambiguous or narrative in areas where statistical specifications are expected. Mock shells vary widely across sponsors, and automated parsing of tables with highly customized formatting remains technically demanding. Additionally, LLM behaviour in statistical contexts requires continuous monitoring, domain-specific prompt engineering, and controlled execution environments to ensure reproducibility. Therefore, the approach is best as assistive technology, pre-screening and contextualizing findings, rather than full automation.

COMPLIANCE AND VALIDATION CONSIDERATIONS

Integrating AI within a GxP-compliant SCE introduces governance requirements. Each model and prompt template must undergo formal qualification with defined acceptance criteria. Outputs must remain deterministic within version control, ensuring that re-running the same input yields identical results.

The validation approach aligns with principles outlined in EMA's Reflection Paper on AI in GxP (2024 draft) and FDA's Good Machine Learning Practice (GMLP). These frameworks emphasize that AI systems in regulated environments must be transparent, traceable, and under human oversight.

Within *entimICE FastTrack*®, model artifacts are treated like code modules: versioned, reviewed, and locked under change control. Reviewer sign-off remains mandatory for any AI-suggested discrepancy.

CONCLUSION

This paper presented a conceptual design for AI-assisted validation of TLFs against their SAP and mock shell specifications. By embedding explainable AI reasoning within a compliant SCE, the proposed workflow reduces manual effort, enhances traceability, and introduces a new level of consistency in statistical quality control.

While still theoretical, the concept represents a natural evolution of metadata-driven automation in clinical programming. The next phase will focus on building a limited prototype using de-identified SAPs and mock shells to benchmark accuracy and usability. Long-term, this capability could extend to cross-validation of ADaM datasets and integrated submission package review.

Ultimately, AI's role is not to replace the human reviewer but to amplify their precision and focus, ensuring that statistical outputs truly reflect their intended design, every time.

FUTURE WORK

Ongoing efforts at Entimo will evaluate real-world pilot implementations of this concept within *entimICE FastTrack*®. Future enhancements will include:

- Expansion to ADaM-to-TLF lineage validation.
- Integration of visual analysis for figure comparison.
- UI/UX enhancements for discrepancy visuals and reviewer feedback.
- Continuous model improvement based on validated findings.

CONTACT INFORMATION

Contact the authors at: spa@entimo.de, mja@entimo.de

REFERENCE

- [1] IBM. What is a RAG explained.
<https://www.ibm.com/architectures/hybrid/genai-rag>
- [2] Video explaining how to improve your RAG.
<https://www.youtube.com/watch?v=TRjq7t2Ms5I>
- [3] What is a vector database?
<https://www.ibm.com/think/topics/vector-database>
- [4] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric P. Xing, Hao Zhang, Joseph E. Gonzalez, Ion Stoica (2023). Judging LLM-as-a-Judge with MT-Bench and Chatbot Arena
<https://arxiv.org/abs/2306.05685> [cs.CL]
- [5] ClinicalTrials.gov. Study NCT03882424: SAP_004 — Statistical Analysis Plan. Available at:
https://cdn.clinicaltrials.gov/large-docs/24/NCT03882424/SAP_004.pdf