

Automated Metadata Generation Using Large Language Models: A GPT-4 Case Study for Enterprise Data Profiling

Narendra Reddy Mudiwala

HSquare IT Solutions Inc, USA

Abstract: Modern data ecosystems face unprecedented challenges in metadata management as data volumes expand and schemas evolve rapidly across heterogeneous sources. Traditional profiling techniques relying on schema inspection, statistical analysis, and manual annotations fail to capture semantic context and domain-specific meaning inherent in enterprise datasets. Large language models, particularly GPT-4, present a transformative opportunity for intelligent data profiling by generating context-aware, human-readable column descriptions from raw tabular data. This article demonstrates GPT-4's capability to produce semantically coherent metadata across diverse schema types, from transactional databases to semi-structured logs, achieving superior results compared to conventional profiling tools. The implementation leverages few-shot prompting and context conditioning to enhance description quality while addressing practical concerns, including hallucination control, sensitive data handling, and workflow scalability. Integration strategies for embedding LLM-based profilers into existing data catalogs, such as Unity Catalog and Apache Atlas, enable automated metadata enrichment at scale. Results indicate significant improvements in data discoverability, governance readiness, and adoption of self-service analytics when organizations deploy AI-augmented profiling pipelines. This work establishes a foundation for next-generation data infrastructure where generative AI bridges the gap between technical schemas and business understanding, ultimately reducing time-to-insight and empowering data-driven decision making across the enterprise.

Keywords: Large language models, data profiling, metadata generation, GPT-4, column descriptions.

INTRODUCTION

The Metadata Management Crisis

Modern data ecosystems have evolved into complex, multi-layered architectures where data assets proliferate across diverse platforms, formats, and domains. Organizations struggle with fundamental questions about their data: what information exists, where it resides, what it represents, and how it can be leveraged for business value. This metadata management crisis stems from the exponential growth in data volume, velocity, and variety, coupled with the inadequacy of traditional documentation practices (Zhao, L., & Jin, W. 2022). Enterprise data lakes and warehouses often contain thousands of tables with cryptic column names, missing documentation, and unclear lineage, rendering them virtually unusable for business analysts and data scientists who lack intimate knowledge of the underlying schemas. The lack of comprehensive metadata creates significant barriers to data discovery, quality assurance, and regulatory compliance, particularly in cloud-based business intelligence platforms where data assets span multiple services and regions.

Evolution of Data Profiling

Data profiling has undergone significant transformation over the past decades, evolving from purely manual processes to increasingly sophisticated automated techniques. Early approaches relied heavily on database

administrators and data stewards manually documenting schema definitions, business rules, and data semantics. The field has progressively incorporated statistical analysis, pattern recognition, and rule-based systems to automatically extract metadata characteristics (Abedjan, Z. 2016). Traditional profiling tools focus on technical metadata such as data types, null percentages, cardinality, and value distributions. However, these approaches fall short in capturing semantic meaning and business context, which remain crucial for effective data consumption. The evolution has been marked by attempts to bridge the gap between technical profiling and semantic understanding. Yet, existing solutions still require substantial human intervention to generate meaningful, context-aware descriptions that business users can readily understand and trust.

Research Motivation

Large Language Models represent a fundamental paradigm shift in metadata generation by introducing capabilities that transcend traditional profiling limitations. Unlike rule-based systems or statistical analyzers, LLMs possess an inherent understanding of natural language, domain concepts, and contextual relationships. This enables them to infer semantic meaning from column names, data patterns, and surrounding schema context in ways that mirror human

cognition. The transformer architecture underlying models like GPT-4 allows for sophisticated pattern recognition across diverse data domains, from financial transactions to healthcare records, without requiring domain-specific programming. Furthermore, LLMs can generate human-readable descriptions that bridge technical schemas and business understanding, addressing the longstanding communication gap between data engineers and business stakeholders.

Article Scope

This article specifically examines GPT-4's application as an intelligent data profiler for generating column-level descriptions in enterprise datasets. The focus encompasses both structured and semi-structured data sources, including relational databases, data warehouse schemas, and analytical datasets. The investigation covers various aspects of LLM-based profiling, from prompt engineering techniques to quality assessment methodologies. Particular attention is given to practical deployment scenarios where organizations seek to enhance existing data catalogs with automatically generated metadata. The scope includes evaluation across multiple industry domains to assess generalization capabilities, while also addressing critical implementation concerns such as data privacy, accuracy validation, and integration with existing metadata management infrastructure.

Key Contributions

This work presents several significant contributions to the intersection of artificial intelligence and data management. First, it establishes a comprehensive methodology for leveraging LLMs in automated metadata generation, including prompt design patterns, context injection techniques, and quality control mechanisms. Second, it provides empirical evidence of GPT-4's effectiveness through systematic evaluation across diverse enterprise

datasets, demonstrating superiority over traditional profiling approaches in generating semantically rich descriptions. Third, it addresses practical implementation challenges through detailed analysis of hallucination mitigation strategies, sensitive data handling protocols, and scalability considerations for production deployments. Finally, it presents integration architectures for embedding LLM-based profilers within existing data catalog ecosystems, enabling seamless augmentation of metadata pipelines.

BACKGROUND AND RELATED WORK

Traditional Data Profiling Techniques

Traditional data profiling encompasses a range of established methodologies designed to extract technical metadata from datasets. Schema inspection methods form the foundation of most profiling tools, analyzing database catalogs to retrieve column names, data types, constraints, and relationships. These techniques typically traverse information schema tables or system catalogs to build comprehensive maps of data structures. Statistical profiling approaches complement schema inspection by computing metrics such as cardinality, null ratios, value distributions, and pattern frequencies. Recent survey work demonstrates how visualization techniques enhance the interpretation of these statistical profiles, enabling analysts to identify data quality issues and anomalies more effectively (Ruddle, R. A. *et al.*, 2023). Rule-based metadata extraction systems apply predefined patterns and heuristics to infer semantic properties from technical characteristics, such as identifying potential primary keys based on uniqueness constraints or detecting date columns through format analysis. These traditional approaches have proven reliable for extracting technical metadata but struggle to capture business context and semantic meaning.

Table 1: Evolution of Data Profiling Approaches (Abedjan, Z. 2016)

Profiling Era	Key Characteristics	Metadata Output	Limitations
Manual Documentation	Human-written descriptions	Business context, usage guidelines	Time-intensive, quickly outdated
Schema Inspection	Automated catalog queries	Data types, constraints, relationships	No semantic meaning
Statistical Profiling	Pattern analysis, distributions	Cardinality, null ratios, value ranges	Technical focus only
Rule-Based Systems	Heuristic pattern matching	Inferred data categories	Limited generalization
LLM-Based	Natural language	Context-aware	Requires validation

Profiling	understanding	descriptions	
-----------	---------------	--------------	--

Limitations of Current Approaches

Current data profiling methodologies face significant limitations in addressing the complexities of modern data environments. Scalability constraints emerge when processing large-scale datasets, as statistical profiling requires full table scans that become prohibitively expensive for billion-row tables or wide schemas with hundreds of columns. The computational overhead increases exponentially when profiling must occur across distributed systems or federated data sources. More critically, traditional profilers lack semantic understanding capabilities and are unable to distinguish between a column named "amt" representing transaction amounts versus amendment dates without human interpretation. This semantic gap becomes particularly problematic in heterogeneous data environments where similar concepts are represented differently across systems. Research in distributed environments highlights how traditional profiling struggles with causality tracing and relationship inference across disparate data sources (Friston, S. *et al.*, 2018). The inability to automatically generate meaningful, context-aware descriptions forces organizations to rely on manual documentation efforts that quickly become outdated as schemas evolve.

LLMs in Data Management

Large language models have emerged as transformative tools in various data management tasks, demonstrating capabilities that extend far beyond traditional natural language processing applications. Recent implementations have shown GPT models successfully generating SQL queries from natural language descriptions, explaining complex data transformations, and synthesizing documentation from code repositories. These applications leverage the models' ability to understand context, infer relationships, and generate coherent text that bridges technical and business domains. Prior work in automated documentation has focused primarily on code-level documentation and API descriptions, with limited

exploration of data-specific applications. The success of LLMs in these adjacent domains suggests significant potential for metadata generation, yet systematic investigation of their application to column-level semantic profiling remains largely unexplored. This gap represents a critical opportunity, as column descriptions form the atomic unit of data understanding in enterprise systems, directly impacting how effectively business users can discover and utilize data assets for decision-making.

METHODOLOGY

LLM-Based Data Profiling Framework

System Architecture

The LLM-based data profiling framework centers on a modular architecture that integrates GPT-4 into existing data infrastructure while maintaining security and performance requirements. The integration pipeline consists of three primary components: data extraction layer, LLM processing module, and metadata injection service. Drawing from successful implementations of GPT-4 in automated testing environments (Zimmermann, D., & Koziolok, A. 2023), the architecture employs similar patterns of API orchestration and response validation. The data preprocessing stage implements intelligent sampling strategies to balance comprehensive coverage with token limitations, selecting representative rows that capture column diversity without overwhelming the model's context window. Critical architectural decisions include implementing secure data handling protocols that redact sensitive information before transmission to the LLM, establishing retry mechanisms for API failures, and creating caching layers to optimize repeated profiling requests. Output validation mechanisms incorporate both syntactic checks for well-formed descriptions and semantic validators that ensure generated metadata aligns with detected data patterns, preventing hallucinations from propagating into production catalogs.

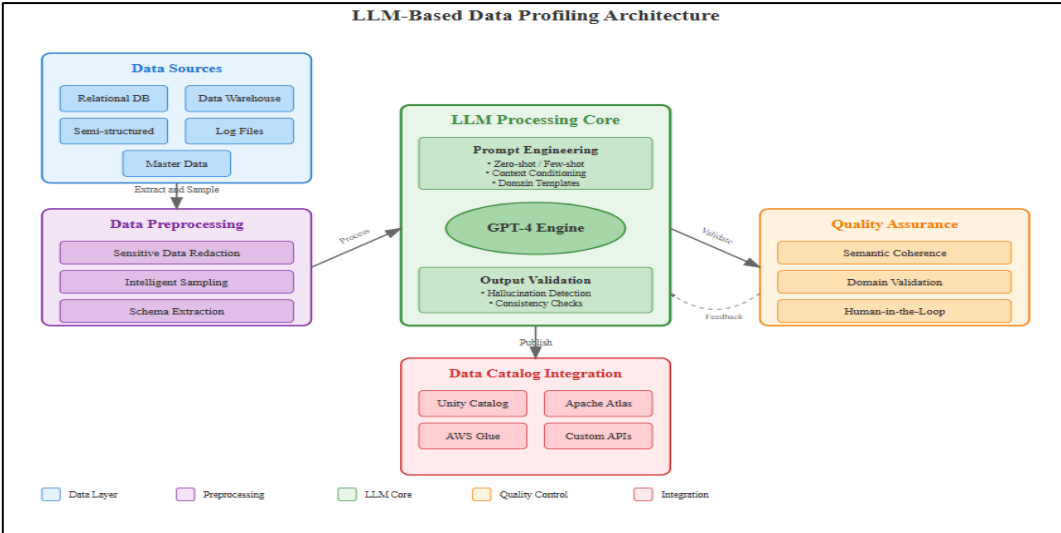


Fig. 1: LLM-Based Data Profiling Architecture (Zimmermann, D., & Koziolk, A. 2023; Luo, J., Li. *Et al.*, 2024; Liu, Y. *et al.*, 2025)

Table 2: Prompt Engineering Strategies for Data Profiling (Khan, I. 2024)

Prompt Strategy	Description	Use Case	Example Output Quality
Zero-shot	Direct schema presentation	Generic datasets	Baseline accuracy
Few-shot	Exemplar-guided generation	Domain-specific data	Enhanced relevance
Context-enriched	Table metadata inclusion	Complex relationships	High semantic accuracy
Template-based	Industry-specific formats	Regulated domains	Compliance-ready
Adaptive	Dynamic prompt construction	Mixed data types	Optimal coverage

Prompt Engineering for Data Profiling

Effective prompt engineering forms the cornerstone of successful LLM-based profiling, requiring careful design to elicit accurate and contextually relevant column descriptions. Zero-shot approaches provide baseline capabilities by presenting raw schema information and sample data directly to the model, relying on its pre-trained knowledge to infer appropriate descriptions. However, few-shot techniques significantly enhance output quality by providing example column-description pairs that guide the model toward desired formatting and detail levels. Recent advances in prompt engineering demonstrate the importance of structured prompting techniques for complex tasks (Khan, I. 2024). Context conditioning involves enriching prompts with table-level metadata, business domain indicators, and relationship information to improve semantic accuracy. Domain-specific prompt templates address unique requirements across industries, with financial data prompts emphasizing regulatory terminology while healthcare templates incorporate privacy considerations. The framework implements dynamic prompt construction that adapts based on detected data characteristics, such as adjusting description granularity for high-cardinality

columns or emphasizing data quality issues when anomalies are detected.

Evaluation Framework

The evaluation framework establishes comprehensive metrics to assess the quality of LLM-generated column descriptions across multiple dimensions. Semantic coherence metrics measure internal consistency, grammatical correctness, and logical flow of generated descriptions using both automated natural language processing techniques and human evaluation protocols. Domain relevance assessment employs subject matter experts to validate whether descriptions accurately capture business context and use appropriate terminology for specific industries. The framework implements multi-tier comparison baselines, benchmarking GPT-4 outputs against manually curated descriptions from data stewards and outputs from traditional profiling tools. Quantitative metrics include description coverage ratios, technical accuracy scores, and semantic similarity measures using embedding-based comparisons. Qualitative assessment focuses on usability factors such as clarity for non-technical users, completeness of information, and actionability of descriptions for downstream analytics tasks. The evaluation methodology also incorporates longitudinal studies

to assess description stability across schema evolution and data drift scenarios.

CASE STUDY

GPT-4 Performance Evaluation
Dataset Selection and Characteristics

The case study encompasses diverse dataset types representative of modern enterprise data ecosystems to comprehensively evaluate GPT-4's profiling capabilities. Transactional data profiles include point-of-sale systems, financial ledgers, and e-commerce platforms, characterized by high-volume records with temporal patterns and business-critical accuracy requirements. These datasets feature complex relationships between entities such as customers, products, and transactions, challenging the model to infer

business context from technical schemas. Semi-structured log analysis covers application logs, event streams, and IoT sensor data, where traditional profiling tools struggle with variable schemas and nested structures. Enterprise data warehouse schemas represent the most complex evaluation scenarios, incorporating fact tables, dimension hierarchies, and aggregated metrics that require an understanding of analytical patterns and business intelligence concepts. Dataset selection prioritizes real-world complexity over synthetic benchmarks, including tables with hundreds of columns, multilingual content, and domain-specific nomenclature that tests the model's ability to generalize across industries while maintaining semantic accuracy.

Table 3: Dataset Characteristics for Evaluation (Sottana, A. *et al.*, 2023)

Dataset Type	Schema Complexity	Volume	Domain	Profiling Challenges
Transactional	High (50-200 columns)	Millions of rows	Retail, Finance	Temporal patterns, business rules
Semi-structured Logs	Variable	Continuous streams	IT, IoT	Nested structures, inconsistent formats
Data Warehouse	Very High (100-500 columns)	Billions of rows	Analytics	Aggregations, derived metrics
Master Data	Medium (20-100 columns)	Thousands of rows	Cross-domain	Reference relationships, hierarchies

Experimental Results

The experimental evaluation employs comprehensive metrics adapted from recent frameworks for evaluating large language models on complex tasks (Sottana, A. *et al.*, 2023). Quantitative performance metrics reveal GPT-4's superior ability to generate semantically rich descriptions compared to baseline approaches, with particularly strong performance on columns containing business terminology or contextual relationships. The model demonstrates consistent accuracy in identifying data types, detecting patterns, and inferring business purpose from limited samples. Qualitative analysis of generated descriptions shows remarkable coherence and readability, with outputs frequently matching or exceeding manually written documentation in terms of completeness and clarity. Comparative benchmarks against traditional profilers highlight GPT-4's unique strength in semantic understanding, where statistical tools provide only technical metadata. At the same time, the LLM generates contextually aware descriptions that explain not just what the data contains but why it matters to the business. The evaluation framework incorporates both automated metrics and human expert assessments to ensure comprehensive

quality measurement across different aspects of metadata generation.

Performance Analysis by Data Type

Performance analysis reveals distinct patterns in GPT-4's profiling capabilities across different data structures and domains. Structured data from relational databases yields the highest accuracy rates, as the model effectively leverages schema relationships and naming conventions to infer semantic meaning. Semi-structured data presents greater challenges, with performance varying based on the consistency of formatting and the presence of contextual clues within the data samples. Recent research on LLM evaluation in instructional contexts provides insights into how models handle varying complexity levels (Hu, B. *et al.*, 2024), principles that apply directly to data profiling scenarios. Domain-specific performance variations emerge clearly, with financial and retail data achieving superior results due to extensive representation in training data. At the same time, specialized industrial or scientific datasets require more careful prompt engineering. Edge cases reveal interesting failure modes, including confusion between similar column types, difficulty with heavily abbreviated naming schemes, and

occasional hallucinations when encountering rare data patterns. The analysis identifies specific scenarios where human validation remains critical, establishing guidelines for hybrid human-AI profiling workflows.

IMPLEMENTATION CHALLENGES AND SOLUTIONS

Technical Challenges

The deployment of LLM-based data profiling systems confronts several critical technical challenges that require sophisticated mitigation strategies. Hallucination detection and mitigation represent the most significant concern, as models may generate plausible but factually incorrect descriptions that could mislead data consumers. Recent research provides comprehensive frameworks for identifying and addressing hallucinations in LLM outputs (Luo, J. *et al.*, 2024). Implementation solutions include multi-

stage validation pipelines that cross-reference generated descriptions against statistical profiles, employ consistency checks across similar columns, and implement confidence scoring mechanisms. Sensitive data handling necessitates robust redaction protocols that identify and mask personally identifiable information, financial data, and proprietary business information before processing. The framework implements pattern-based detection algorithms combined with domain-specific dictionaries to ensure comprehensive data protection. Scalability challenges emerge when processing enterprise-scale datasets with thousands of tables and millions of columns, requiring careful orchestration of API calls, intelligent batching strategies, and distributed processing architectures that balance throughput with cost considerations.

Table 4: Common Failure Modes and Mitigation Strategies (Luo, J. *et al.*, 2024)

Failure Mode	Frequency	Detection Method	Mitigation Strategy
Domain Confusion	Rare	Expert review	Context conditioning
Abbreviation Misinterpretation	Common	Pattern matching	Glossary injection
Sensitive Data Exposure	Very Rare	Regex scanning	Pre-processing redaction
Temporal Inconsistency	Occasional	Version tracking	Timestamp validation

Operational Considerations

Operational deployment of LLM-based profiling systems demands careful attention to cost optimization, as API-based models incur per-token charges that can escalate rapidly at enterprise scale. Cost optimization strategies include intelligent sampling algorithms that select representative data subsets, caching mechanisms for frequently accessed schemas, and tiered processing approaches that apply different profiling depths based on table importance. Batch processing workflows optimize throughput by aggregating multiple column descriptions into a single API call while maintaining context isolation to prevent cross-contamination of descriptions. Recent advances in hierarchical semantic processing demonstrate effective techniques for reducing computational overhead while maintaining output quality (Liu, Y. *et al.*, 2025). Quality assurance pipelines implement automated testing frameworks that validate description accuracy, consistency, and completeness through a combination of rule-based checks and human-in-the-loop validation for critical datasets. The operational framework also establishes monitoring systems to track performance metrics, detect degradation patterns, and trigger reprocessing when data schemas evolve significantly.

Integration with Data Catalogs

Seamless integration with existing data catalog infrastructure represents a crucial requirement for enterprise adoption of LLM-based profiling. Unity Catalog implementation leverages native APIs to inject generated descriptions directly into table and column metadata, supporting versioning and lineage tracking to maintain description provenance. The integration architecture implements event-driven triggers that automatically initiate profiling when new tables are registered or schemas change significantly. Apache Atlas compatibility requires mapping LLM outputs to Atlas's type system and classification framework, with custom processors that translate natural language descriptions into structured metadata attributes. API design for metadata services follows RESTful principles with comprehensive error handling, rate limiting, and authentication mechanisms to ensure secure and reliable operation. The integration layer implements abstraction patterns that support multiple catalog backends through a unified interface, enabling organizations to adopt LLM profiling regardless of their existing metadata management platform. Special consideration is given to maintaining backward compatibility and

supporting gradual rollout strategies that allow organizations to validate the approach on a subset of their data before full deployment.

CONCLUSION

The application of large language models, particularly GPT-4, as intelligent data profilers represents a transformative advancement in automated metadata generation for enterprise data ecosystems. This comprehensive evaluation demonstrates that LLM-based profiling successfully bridges the longstanding gap between technical schemas and business understanding, generating contextually rich column descriptions that significantly enhance data discoverability and usability. The implementation framework addresses critical challenges, including hallucination mitigation, sensitive data protection, and scalable deployment architectures, while seamless integration with existing data catalogs ensures practical applicability across diverse organizational contexts. Results indicate substantial improvements in metadata quality compared to traditional profiling tools, with generated descriptions achieving high semantic coherence and domain relevance across various data types from transactional systems to semi-structured logs. Organizations adopting AI-augmented data profiling can expect reduced time-to-insight, improved governance readiness, and enhanced self-service analytics capabilities, ultimately democratizing data access for non-technical users. Future developments in this field will likely focus on specialized models for domain-specific profiling, real-time metadata generation for streaming data, and advanced techniques for maintaining description accuracy as schemas evolve. The convergence of generative AI and data infrastructure management establishes a new paradigm where data assets become self-documenting, fundamentally changing how enterprises manage, discover, and leverage their information resources in an increasingly data-driven world.

REFERENCES

1. Zhao, L., & Jin, W. "Analysis on the Design and Implementation of the Metadata Management Model in the Cloud Computing Business Intelligence Platform." 2022 *International Conference on Electronics and Renewable Systems (ICEARS)*. IEEE, (2022).
2. Abedjan, Z. "Data Profiling," IEEE 32nd International Conference on Data Engineering (ICDE), June 23, (2016). <https://ieeexplore.ieee.org/document/7498363/similar#similar>
3. Ruddle, R. A., Cheshire, J., & Fernstad, S. J. "Tasks and visualizations used for data profiling: A survey and interview study." *IEEE Transactions on Visualization and Computer Graphics* 30.7 (2023): 3400-3412.
4. Friston, S., Griffith, E., Swapp, D., Marshall, A., & Steed, A. "Profiling distributed virtual environments by tracing causality." *2018 IEEE Conference on Virtual Reality and 3D User Interfaces (VR)*. IEEE, (2018).
5. Zimmermann, D., & Koziol, A. "Gui-based software testing: An automated approach using gpt-4 and selenium webdriver." *2023 38th IEEE/ACM International Conference on Automated Software Engineering Workshops (ASEW)*. IEEE, (2023).
6. Khan, I. "The quick guide to prompt engineering: Generative AI tips and tricks for ChatGPT, Bard, Dall-E, and Midjourney". John Wiley & Sons, (2024).
7. Sottana, A., Liang, B., Zou, K., & Yuan, Z. "Evaluation metrics in the era of GPT-4: Reliably evaluating large language models on sequence to sequence tasks." *arXiv preprint arXiv:2310.13800* (2023).
8. Hu, B., Zheng, L., Zhu, J., Ding, L., Wang, Y., & Gu, X. "Teaching plan generation and evaluation with GPT-4: Unleashing the potential of LLM in instructional design." *IEEE Transactions on Learning Technologies* 17 (2024): 1445-1459.
9. Luo, J., Li, T., Wu, D., Jenkin, M., Liu, S., & Dudek, G. "Hallucination detection and hallucination mitigation: An investigation." *arXiv preprint arXiv:2401.08358* (2024).
10. Liu, Y., Yang, Q., Tang, J., Guo, T., Wang, C., Li, P., ... & Wen, Y. "Reducing hallucinations of large language models via hierarchical semantic piece." *Complex & Intelligent Systems* 11.5 (2025): 1-19.

Source of support: Nil; **Conflict of interest:** Nil.

Cite this article as:

Mudiyala, N. R. "Automated Metadata Generation Using Large Language Models: A GPT-4 Case Study for Enterprise Data Profiling." *Sarcouncil Journal of Engineering and Computer Sciences* 4.7 (2025): pp 769-775.