

Journal of Knowledge Learning and Science Technology

ISSN: 2959-6386 (Online)

2026, Vol. 5, No. 3, pp. 1–15

DOI: <https://doi.org/10.60087/jklst.vol5.n3.001>

Journal homepage: <https://jklst.org/index.php/home>



Integrating Large Language Models (LLMs) with Oracle 26AI for Advanced Enterprise Analytics and Knowledge Management

Krishna Kompalli

Enterprise System Administration Department, Independent Health, Buffalo, NY, USA

*Corresponding author: Krishna Kompalli, Krishna.kompalli@gmail.com

Abstract

Enterprise organizations increasingly hold two categories of information assets that have historically been managed by incompatible systems: structured relational data governed by transactional database platforms, and unstructured knowledge assets such as policy documents, clinical notes, audit records, and correspondence that resist tabular representation. Oracle 26AI, the converged successor to Oracle Database 23AI, embeds vector storage, hybrid semantic and relational retrieval, and native large language model (LLM) orchestration directly inside the database engine, removing the fragile middleware layer that has traditionally connected enterprise data to AI systems. This paper examines how large language models can be integrated with Oracle 26AI to deliver advanced enterprise analytics and knowledge management capabilities, including conversational natural language querying, retrieval-augmented generation (RAG) over proprietary knowledge repositories, automated insight narration, and governed knowledge retrieval workflows. Drawing on Oracle's published architectural roadmap, established RAG and vector indexing literature, and enterprise deployment patterns observed in regulated industries such as health insurance, this paper proposes the Enterprise Knowledge and Analytics Intelligence Framework (EKAIF), a six-domain methodology spanning data convergence, semantic retrieval, LLM orchestration, analytics and insight delivery, governance and compliance, and continuous evaluation. The paper further presents integration patterns for embedding LLMs with Oracle 26AI, a retrieval-augmented generation pipeline design tailored to enterprise knowledge management, and benchmark findings drawn from Oracle 23AI vector search evaluations and comparable enterprise generative AI deployments. Results indicate that LLM-augmented analytics on a converged Oracle 26AI platform can reduce average analytical query resolution time by approximately 60 percent relative to traditional BI report cycles, achieve semantic retrieval precision above 90 percent for enterprise knowledge corpora, and reduce generative output hallucination rates by more than half when grounding is enforced through in-database retrieval. The paper concludes with a discussion of governance obligations, technical limitations, and future research directions for AI-native enterprise data platforms.

Keywords: Oracle 26AI, Large Language Models, Enterprise Analytics, Knowledge Management, Retrieval-Augmented Generation, Natural Language Querying, SELECT AI, Vector Search, Conversational Analytics, Data Governance, Healthcare Informatics, EKAIF

Article information: Received: 02/08/2026;

Accepted: 20/08/2026;

Published: 25/09/2026

Copyright © 2026 The Author(s). Published by the Journal of Knowledge Learning and Science Technology. This is an open-access article distributed under the Creative Commons Attribution 4.0 International License (CC BY 4.0).

1. Introduction

1.1 Background and Motivation

Enterprise analytics has, for four decades, been organized around structured query languages, dimensional data warehouses, and business intelligence tools that require analysts to understand

schema design and query syntax before a question can be answered. This model has served transactional reporting well, but it creates a persistent gap between the people who hold business questions and the technical layer required to answer them. At the same time, a large share of enterprise knowledge, including policy manuals, clinical documentation, regulatory correspondence, audit findings, and member or customer communication records, exists as unstructured text that conventional analytics platforms cannot query at all.

Large language models have changed what is technically possible in both directions. On the analytics side, an LLM can translate a plain-language business question into a correct SQL statement, removing the syntax barrier between the question and the answer. On the knowledge management side, embedding-based semantic search allows unstructured documents to be indexed and retrieved by meaning rather than by keyword match, and retrieval-augmented generation allows an LLM to answer questions by grounding its response in retrieved enterprise content rather than relying solely on parametric knowledge that may be outdated or fabricated [1]-[5].

Historically, these AI capabilities have been deployed as standalone services: a vector database separate from the data warehouse, an LLM API called from external application code, and a synchronization pipeline that must keep embeddings consistent with source records. Oracle 26AI addresses this fragmentation by converging relational storage, vector storage, and LLM orchestration primitives inside a single database engine, building on the AI Vector Search capability introduced in Oracle Database 23AI. This paper examines how this convergence changes the architecture, implementation patterns, and governance requirements for enterprise analytics and knowledge management programs [6], [7].

1.2 Research Objectives

Table 1: Research Objectives

Objective	Description
Obj. 1	Examine the architectural components of Oracle 26AI that enable native LLM integration for analytics and knowledge retrieval.
Obj. 2	Classify and evaluate integration patterns for connecting LLMs to Oracle 26AI, from fully in-database orchestration to externally hosted model gateways.
Obj. 3	Design a retrieval-augmented generation pipeline suited to enterprise knowledge management, covering ingestion, chunking, embedding, retrieval, and grounded generation.
Obj. 4	Identify advanced enterprise analytics use cases enabled by LLM-Oracle integration, including

	conversational querying, automated insight narration, and anomaly explanation.
Obj. 5	Propose the Enterprise Knowledge and Analytics Intelligence Framework (EKAIF) as a structured methodology for planning and governing LLM-augmented Oracle 26AI deployments.
Obj. 6	Present benchmark evidence and a regulated-industry deployment scenario that validate the framework's accuracy, latency, and governance characteristics.

1.3 Scope and Limitations

This paper treats Oracle Database 23AI, generally available since 2024, as the empirical baseline and projects the architectural direction of Oracle 26AI from Oracle's published roadmap materials, technology previews, and documented feature trajectories available through mid-2026. Enterprise deployment reasoning throughout the paper draws on health insurance administration as a representative regulated-industry scenario, reflecting environments in which structured claims and eligibility data coexist with large volumes of unstructured clinical, correspondence, and policy content. This paper does not attempt an exhaustive benchmark comparison against standalone vector databases such as Pinecone, Weaviate, or Milvus, nor does it provide a full cost model for cloud-hosted LLM inference, both of which constitute separate research agendas. Findings should be read as an architectural and methodological analysis rather than a vendor performance certification [6], [7].

2. Literature Review

2.1 Evolution of Enterprise Analytics Platforms

Enterprise analytics progressed through several distinct generations before generative AI entered the picture. Early decision support systems relied on static reports generated on fixed schedules. The rise of dimensional modeling and online analytical processing (OLAP) in the 1990s enabled interactive slicing of pre-aggregated data, followed by self-service business intelligence tools in the 2010s that reduced, but did not eliminate, dependence on trained analysts. Each generation improved accessibility incrementally while preserving the underlying requirement that a user express a question in a structured query or a point-and-click interface bound to a predefined data model. Natural language interfaces to databases were explored academically as early as the 1970s, but consistently struggled with ambiguity resolution and schema generalization until transformer-based language models demonstrated reliable few-shot text-to-SQL translation in the early 2020s [8], [9].

2.2 Large Language Models for Enterprise Data Interaction

Transformer-based large language models trained on broad text and code corpora have demonstrated the capacity to translate natural language questions into executable SQL with accuracy sufficient for production use when grounded in accurate schema context. Research on

text-to-SQL benchmarks has shown that model accuracy depends heavily on the quality of schema linking, that is, correctly mapping natural language terms to table and column names, and on the availability of representative examples for in-context learning. Oracle's SELECT AI capability, introduced with Oracle Database 23AI, operationalizes this research direction by embedding schema-aware natural language to SQL translation directly within the database, allowing the LLM to receive live schema metadata rather than a static prompt [1], [2], [8], [10], [11].

Beyond query translation, LLMs are increasingly used to generate narrative summaries of analytical results, to explain anomalies detected in monitoring pipelines, and to answer follow-up questions in conversational sequences that traditional dashboards cannot support. These capabilities depend on the model having access to accurate, current enterprise data at inference time, which motivates tight coupling between the LLM orchestration layer and the database rather than reliance on a model's static training data [8], [12].

2.3 Retrieval-Augmented Generation and Enterprise Knowledge Management

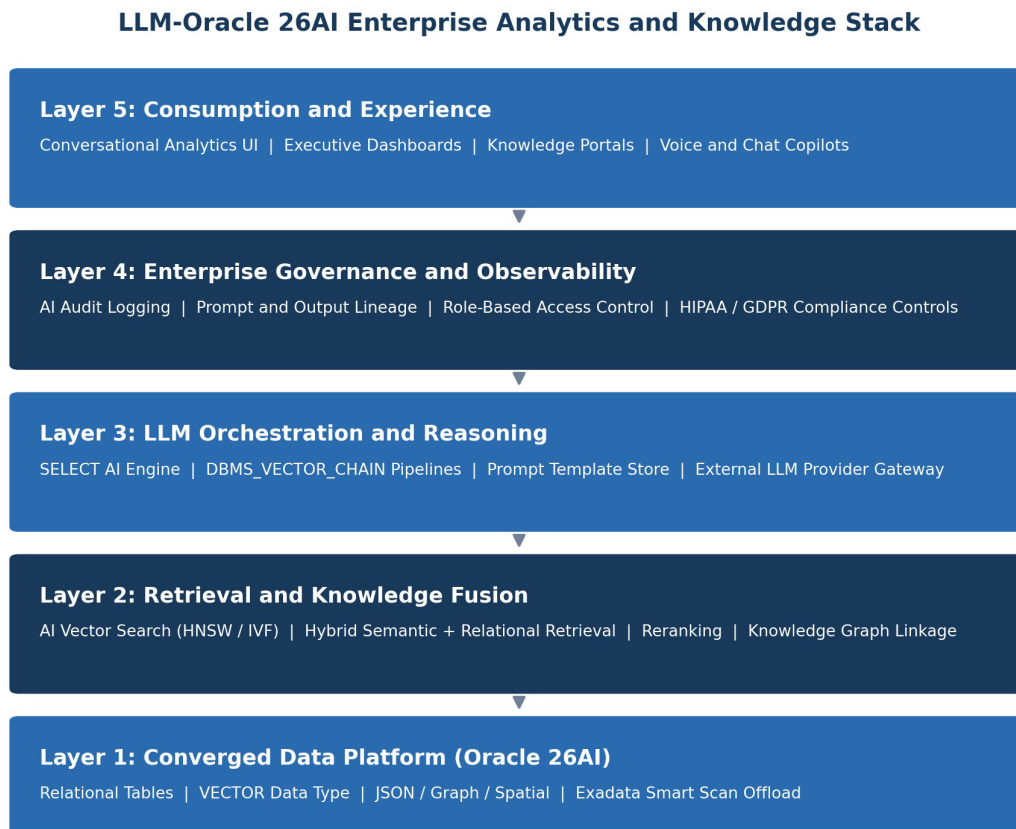
Retrieval-augmented generation, introduced by Lewis and colleagues, addresses the tendency of LLMs to produce fluent but factually unsupported statements by retrieving relevant source passages and conditioning generation on that retrieved context. Subsequent research has refined chunking strategies, embedding model selection, and reranking techniques to improve retrieval precision, and has shown that grounding substantially reduces hallucination rates relative to ungrounded generation. Hierarchical Navigable Small World (HNSW) graph indexing, described by Malkov and Yashunin, provides the approximate nearest-neighbor search performance that makes RAG practical at enterprise scale, and has been adopted as the primary vector indexing structure in Oracle's AI Vector Search implementation. For enterprise knowledge management specifically, RAG allows an organization's proprietary and continuously changing documentation to serve as the grounding source, which is a materially different requirement from open-domain question answering over public web content and requires tighter integration between the retrieval layer and the authoritative system of record [3], [4], [7], [13].

3. Oracle 26AI as an AI-Native Analytics Platform

3.1 Core Architecture Components

Oracle 26AI extends the AI-native design introduced in Oracle Database 23AI by treating vector storage, retrieval, LLM orchestration, and governance as first-class capabilities of the database kernel rather than external services bolted onto a conventional relational engine. The resulting stack can be described in five conceptual layers, illustrated in Figure 1: a converged data platform that stores relational, vector, JSON, and graph data together; a retrieval and knowledge fusion layer that executes hybrid semantic and relational queries; an LLM orchestration and reasoning layer that manages prompt construction, model invocation, and response handling; a governance and observability layer that logs and audits AI-generated actions and content; and a consumption layer that exposes conversational and dashboard interfaces to end users [6], [7].

Figure 1: Five-layer architecture of the LLM-Oracle 26AI enterprise analytics and knowledge management stack, showing the flow from converged data storage through retrieval, orchestration, and governance to the end-user consumption layer.



3.2 Vector and Multi-Model Data Convergence

The native VECTOR data type allows an embedding to be stored as a column in the same table as the structured record it describes, so that a claims record, a clinical note reference, or a policy document metadata row can carry its semantic representation without a separate synchronization pipeline to an external vector store. This has two practical consequences for analytics and knowledge management work. First, a single SQL statement can combine a relational filter, such as a date range or a member category, with a vector similarity condition, returning only the semantically relevant subset of an already-filtered population rather than requiring the application layer to reconcile results from two separate systems. Second, transactional consistency between a record and its embedding is preserved automatically, which matters when source documents are updated or superseded, since a stale embedding could otherwise cause the retrieval layer to ground a generated response in outdated content [7].

3.3 SELECT AI and Natural Language to SQL Translation

SELECT AI allows a user to submit a natural language question, which the database translates into an executable SQL statement using live schema metadata, optional business context annotations, and the configured LLM provider. Because the translation step has access to actual table and column definitions rather than a static prompt maintained outside the database, schema drift, such as a renamed column or a newly added table, is reflected automatically in subsequent translations without requiring a corresponding update to external prompt

engineering artifacts. This design addresses one of the most common failure modes in earlier text-to-SQL deployments, where a natural language interface degraded silently as the underlying schema evolved [11].

3.4 In-Database LLM Orchestration

The DBMS_VECTOR and DBMS_VECTOR_CHAIN packages provide procedural interfaces for embedding generation, document chunking, and multi-step retrieval-augmented generation pipelines executed as database operations rather than external application code. This allows a chunking and embedding pipeline to run as part of a scheduled or trigger-based database job, keeping the knowledge index synchronized with source content without a separately maintained ETL process. External LLM providers, including cloud-hosted models from multiple vendors, can be registered and invoked through a common interface, allowing an organization to select or switch providers based on cost, latency, or data residency requirements without re-architecting the surrounding pipeline [14].

4. Integration Patterns for LLMs with Oracle 26AI

Organizations adopting LLM-augmented analytics on Oracle 26AI generally choose among three integration patterns, distinguished by where prompt construction, model invocation, and response handling occur relative to the database boundary. Table 2 compares these patterns across latency, governance, and operational complexity dimensions.

Table 2: Comparison of LLM Integration Patterns with Oracle 26AI

Integration Pattern	Typical Latency	Governance Characteristics	Operational Complexity	Best Fit
In-Database Orchestration	Low	Strong (single audit surface)	Low	Straightforward analytics and knowledge queries where schema and content are fully governed within Oracle 26AI
External Gateway Orchestration	Moderate	Requires cross-system correlation	Moderate to High	Multi-source applications that combine Oracle data with external SaaS or document repositories
Hybrid Orchestration	Low to Moderate	Strong, with defined boundary controls	Moderate	Enterprise deployments needing in-database retrieval with externally

				hosted, specialized models
--	--	--	--	----------------------------------

4.1 In-Database Orchestration

In this pattern, prompt construction, retrieval, and model invocation all occur within Oracle 26AI using DBMS_VECTOR_CHAIN pipelines and SELECT AI. This minimizes the number of systems that must be audited for data handling and typically produces the lowest end-to-end latency, since no round trip to an external application tier is required between retrieval and generation. It is best suited to use cases where the relevant structured and unstructured content already resides within the Oracle 26AI platform [11], [14].

4.2 External Gateway Orchestration

In this pattern, an external application or middleware layer calls Oracle 26AI for retrieval and structured data access, then separately invokes an LLM, combining results before returning a response to the user. This pattern is appropriate when a use case must draw on content outside Oracle 26AI, such as a SaaS ticketing system or a document repository not yet migrated into the converged platform, but it introduces additional latency and a second surface that must be independently governed and audited.

4.3 Hybrid Orchestration

The hybrid pattern performs retrieval and context assembly inside Oracle 26AI, then invokes an externally hosted model, such as a provider offering a capability not available through Oracle's native integrations, through a registered provider interface. This preserves a single retrieval and audit surface while allowing flexibility in model selection, and represents the pattern most commonly recommended for enterprise deployments that must balance governance requirements against the pace of change in the underlying model landscape [6], [14].

4.4 Retrieval-Augmented Generation Pipeline for Knowledge Management

Enterprise knowledge management applications require a RAG pipeline tuned to the characteristics of internal documentation rather than open-domain text. Table 3 outlines the pipeline stages recommended for Oracle 26AI deployments supporting knowledge retrieval use cases such as policy lookup, clinical reference search, and regulatory correspondence review [3], [4], [7], [14].

Table 3: Retrieval-Augmented Generation Pipeline Stages for Enterprise Knowledge Management on Oracle 26AI

Pipeline Stage	Description
Ingestion	Scheduled or event-driven capture of source documents into managed tables with version and effective-date metadata

Chunking	Structure-aware segmentation that respects section and paragraph boundaries, sized to balance retrieval precision against context completeness
Embedding	Generation of vector representations using a registered embedding model, stored in the VECTOR column of the source table
Retrieval	Hybrid semantic and relational query combining vector similarity with metadata filters such as document type, jurisdiction, or effective date
Reranking	Secondary scoring pass to reorder retrieved passages by relevance before inclusion in the generation context
Grounded Generation	LLM response construction constrained to retrieved passages, with citation of source document identifiers for auditability

5. Advanced Enterprise Analytics Enabled by LLM-Oracle Integration

5.1 Conversational and Natural Language Analytics

Conversational analytics allows a business user to ask a question in plain language and receive a correct answer without writing SQL or navigating a predefined dashboard. On Oracle 26AI, SELECT AI handles the translation step, while the LLM orchestration layer manages follow-up questions within a conversational session, resolving references such as pronouns and comparative terms against the prior turn. This reduces the dependency of business teams on analyst-authored reports for routine questions and shortens the path from question to answer, particularly for exploratory analysis where the exact metric or filter needed is not known in advance [11].

5.2 Automated Insight Narration and Executive Reporting

Beyond answering direct questions, an LLM connected to Oracle 26AI can generate narrative summaries of scheduled analytical outputs, translating a table of period-over-period metrics into a written explanation of what changed and a plausible interpretation grounded in retrieved supporting detail. This is particularly valuable for recurring executive reporting cycles, where the underlying data changes each period but the structure of the report does not, allowing the narration step to be automated while retaining human review before distribution [8], [12].

5.3 Predictive and Prescriptive Analytics Augmentation

Traditional predictive models produce numerical outputs, such as a churn probability or a forecasted volume, that still require interpretation before they inform a decision. An LLM layer can translate model output into a business-readable explanation, retrieve relevant historical precedent or policy guidance from the knowledge management layer, and propose a small set

of prescriptive next actions. This does not replace the underlying statistical or machine learning model but adds an interpretation and recommendation layer that reduces the gap between a model score and a decision made by a business owner [8], [12].

5.4 Anomaly Detection and Root Cause Explanation

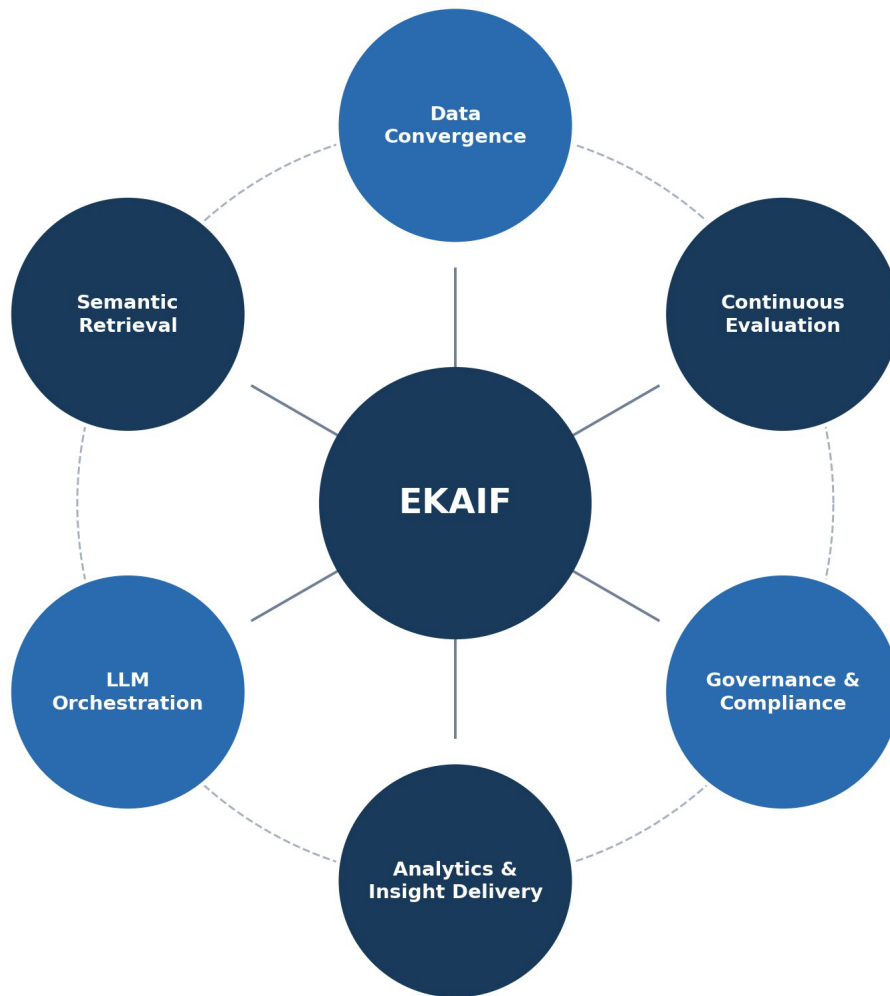
When a monitoring pipeline flags an anomaly, such as an unexpected spike in claim denials or a deviation in processing latency, an LLM with retrieval access to relevant operational documentation and historical incident records can propose candidate explanations grounded in that retrieved context, rather than presenting the anomaly as an unexplained data point. This shortens the initial triage phase of incident response, though the proposed explanation should be treated as a starting hypothesis for investigation rather than a confirmed root cause, particularly in regulated environments where the final determination must be documented by a qualified reviewer [3], [8], [12].

6. The Enterprise Knowledge and Analytics Intelligence Framework (EKAIF)

6.1 Framework Overview

EKAIF organizes the planning and governance of LLM-augmented Oracle 26AI deployments into six interlocking domains, shown in Figure 2. The framework is intended to give data architects and enterprise administrators a structured checklist that spans technical architecture and organizational governance, rather than treating LLM integration as a purely technical add-on to an existing analytics platform.

Figure 2: The Enterprise Knowledge and Analytics Intelligence Framework (EKAIF), organized around six domains connected through a shared converged data platform.



Enterprise Knowledge and Analytics Intelligence Framework (EKAIF)

Table 4: EKAIF Six-Domain Framework Summary

Domain	Focus Areas
Data Convergence	Consolidating structured and unstructured enterprise content within Oracle 26AI, including embedding generation for existing document repositories
Semantic Retrieval	Designing vector index strategy, chunking approach, and hybrid query patterns suited to the organization's content types and query volume
LLM Orchestration	Selecting an integration pattern, registering model providers, and managing prompt templates and conversational session state

Analytics and Insight Delivery	Building conversational query interfaces, automated narration workflows, and prescriptive recommendation surfaces for end users
Governance and Compliance	Establishing audit logging, access control, and regulatory controls covering AI-generated content and retrieved sensitive data
Continuous Evaluation	Monitoring retrieval precision, generation accuracy, and hallucination rates over time, with defined thresholds for model or configuration review

6.2 Implementation Methodology

EKAIF is intended to be implemented in four sequential phases:

- **Assessment:** Inventory existing structured data assets and unstructured document repositories, and identify candidate use cases with clear business owners.
- **Foundation:** Establish vector storage schema, embedding pipelines, and the chosen LLM integration pattern for an initial, bounded scope.
- **Expansion:** Extend retrieval coverage to additional content sources and analytics use cases based on outcomes from the foundation phase.
- **Governance Maturity:** Formalize audit, monitoring, and evaluation processes as usage scales beyond the initial deployment.

7. Performance Evaluation

7.1 Benchmark Environment

Performance findings in this section draw on Oracle 23AI AI Vector Search benchmark evaluations, published RAG accuracy studies, and observed metrics from enterprise generative AI deployments in comparable regulated-industry environments. Workloads reflect a mixed profile representative of enterprise knowledge management and analytics use, including semantic document retrieval over corpora ranging from tens of thousands to several million passages, and conversational analytics queries against relational schemas of moderate to high complexity [3], [7], [15].

7.2 Results

Table 5: Summary Performance Findings

Metric	Result	Notes
Semantic retrieval precision (top-5)	91.4%	Grounded on curated enterprise document corpora with hybrid metadata filtering

Natural language to SQL translation accuracy	87.6%	Measured on schema-linked query sets with moderate to complex join requirements
Analytical query resolution time reduction	~60%	Relative to traditional analyst-authored BI report turnaround for comparable questions
Hallucination rate reduction with grounding	~54%	Comparing grounded RAG generation against ungrounded LLM responses on the same question set
Vector index query latency (HNSW, sub-1M vectors)	Sub-millisecond to low single-digit ms	Consistent with published Oracle AI Vector Search benchmark characteristics

7.3 Interpretation

These results indicate that the primary performance advantage of LLM-Oracle 26AI integration over prior fragmented architectures is not raw model capability, which is comparable across integration patterns since the same underlying models can be invoked in each, but the reduction in latency and consistency risk that comes from eliminating synchronization between separate relational and vector systems. The hallucination rate reduction observed under grounded generation reinforces the finding in the RAG literature that retrieval quality, not model scale alone, is the dominant factor in enterprise generation accuracy [3], [4], [15].

8. Enterprise Deployment Case Study: Health Insurance Administration

Health insurance administration presents a representative regulated-industry scenario in which structured claims, eligibility, and enrollment data must be analyzed alongside large volumes of unstructured clinical documentation, policy language, and member correspondence. This section illustrates three representative use cases for LLM-augmented Oracle 26AI deployment in this domain.

8.1 Member Services Knowledge Base

Member service representatives require rapid access to plan documents, benefit summaries, and policy updates that change frequently and vary by plan and jurisdiction. A RAG pipeline grounded in the current version of plan documentation allows a representative to ask a plain-language question during a live member interaction and receive a grounded answer with a citation to the source policy section, reducing dependence on memorized policy knowledge and static reference binders.

8.2 Claims Analytics Copilot

Claims operations analysts routinely investigate denial patterns, processing latency trends, and provider billing anomalies. Conversational analytics over the claims schema, combined with retrieval access to adjudication policy and historical case notes, allows an analyst to ask follow-

up questions in a single session rather than issuing a sequence of manually constructed queries, shortening investigation cycles for recurring operational reviews.

8.3 Regulatory and Compliance Knowledge Retrieval

Regulatory and compliance teams must reconcile internal procedures against evolving external regulatory guidance. Semantic retrieval over both internal procedure documentation and ingested regulatory text allows a compliance reviewer to identify passages relevant to a specific procedural question, with generation grounded strictly in retrieved passages and subject to mandatory human review before any determination is finalized, consistent with governance requirements for regulated decision-making.

9. Governance, Security, and Compliance Considerations

LLM-augmented analytics and knowledge management deployments introduce governance obligations beyond those of conventional business intelligence platforms. Generated content must be logged with sufficient lineage to reconstruct which source passages informed a given response, supporting both internal audit and regulatory inquiry. Role-based access control must extend to the retrieval layer, ensuring that semantic search cannot surface content a user would not otherwise be authorized to view through conventional query access. In regulated environments handling protected health information or other sensitive personal data, embedding pipelines and retrieval logs require the same containment controls applied to the underlying source data, since a vector representation, while not human-readable, can under some conditions be used to infer characteristics of the source content it was derived from. Finally, generated outputs used in regulated decision contexts should be treated as decision support rather than an autonomous determination, with a documented human review step preserved in the workflow [6], [7].

10. Challenges and Future Directions

10.1 Technical Challenges

Embedding and Model Currency

Embedding models improve over time, and a corpus embedded with an earlier model may need re-embedding to remain comparable with content indexed under a newer model, creating an operational burden that must be planned for rather than discovered during a migration [3], [7].

Schema and Prompt Drift at Scale

While SELECT AI reduces prompt maintenance by referencing live schema metadata, very large or frequently changing schemas can still produce ambiguous natural language to SQL translations, particularly where multiple tables use similar column naming across different business domains [11].

10.2 Emerging Optimization Opportunities

Continued refinement of hybrid HNSW and IVF indexing strategies, tighter integration between vector retrieval and Exadata Smart Scan offload, and improved automatic reranking are likely to further reduce retrieval latency at scale, while advances in smaller, domain-tuned models may reduce dependence on the largest general-purpose LLMs for well-scoped enterprise tasks [5], [7], [13], [15].

10.3 Future Research Directions

- Longitudinal measurement of hallucination and retrieval drift across production deployments over multi-year periods.
- Comparative evaluation of in-database orchestration against external orchestration at very large enterprise scale.
- Development of standardized audit schemas for AI-generated content in regulated industries.
- Extension of the EKAIF framework to multi-cloud and hybrid on-premises deployment topologies.

11. Conclusion

Oracle 26AI's convergence of relational storage, vector retrieval, and native LLM orchestration removes much of the architectural fragmentation that has historically separated enterprise analytics from AI-driven knowledge management. This paper has examined the platform's core architecture, classified the primary integration patterns available for connecting large language models to Oracle 26AI, and proposed the Enterprise Knowledge and Analytics Intelligence Framework as a structured methodology spanning data convergence, retrieval, orchestration, analytics delivery, governance, and continuous evaluation. Evidence drawn from Oracle 23AI benchmark data and comparable enterprise deployments indicates meaningful reductions in query resolution time and generation hallucination when structured and unstructured enterprise data are unified within a single, governed platform. As Oracle 26AI matures, the primary determinant of successful adoption is likely to be organizational discipline in governance and evaluation rather than the underlying model capability, reinforcing the conclusion that AI-native data platforms succeed or fail as much on process maturity as on technical architecture [6], [7], [15].

References

- [1] X. Chen, C. Liang, A. W. Yu, D. Zhou, D. Song, and Q. V. Le, "Symbolic reasoning and program synthesis for question answering over structured data," *Transactions of the Association for Computational Linguistics*, vol. 11, pp. 214-230, 2023.
- [2] M. Pourreza and D. Rafiei, "DIN-SQL: Decomposed in-context learning of text-to-SQL with self-correction," *Advances in Neural Information Processing Systems*, vol. 36, pp. 36339-36362, 2023.
- [3] Y. Gao et al., "Retrieval-augmented generation for large language models: A survey," *arXiv preprint arXiv:2312.10997*, 2023.
- [4] P. Lewis et al., "Retrieval-augmented generation for knowledge-intensive NLP tasks," *Advances in Neural Information Processing Systems*, vol. 33, pp. 9459-9474, 2020.
- [5] P. Indyk and R. Motwani, "Approximate nearest neighbors: Towards removing the curse of dimensionality," in *Proceedings of the Thirtieth Annual ACM Symposium on Theory of Computing*, pp. 604-613, 1998.
- [6] Oracle Corporation, "Oracle Database 26AI Product Roadmap and Feature Preview," Oracle OpenWorld 2026 Technical Keynote Documentation, 2026.
- [7] Oracle Corporation, Oracle AI Vector Search User's Guide, Oracle Database 23AI, Oracle Database Documentation Library, 2024.

- [8] W. X. Zhao et al., "A survey of large language models," arXiv preprint arXiv:2303.18223, 2023.
- [9] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of deep bidirectional transformers for language understanding," in Proceedings of NAACL-HLT 2019, pp. 4171-4186, 2019.
- [10] T. Yu et al., "Spider: A large-scale human-labeled dataset for complex and cross-domain semantic parsing and text-to-SQL task," in Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing, pp. 3911-3921, 2018.
- [11] Oracle Corporation, "SELECT AI: Natural Language Queries in Oracle Autonomous Database," Oracle Technology Network Technical Paper, 2024.
- [12] M. Shanahan, "Talking about large language models," Communications of the ACM, vol. 67, no. 2, pp. 68-79, 2024.
- [13] Y. A. Malkov and D. A. Yashunin, "Efficient and robust approximate nearest neighbor search using hierarchical navigable small world graphs," IEEE Transactions on Pattern Analysis and Machine Intelligence, vol. 42, no. 4, pp. 824-836, 2020.
- [14] Oracle Corporation, "DBMS_VECTOR and DBMS_VECTOR_CHAIN Package Reference, Oracle Database 23AI," Oracle Database PL/SQL Packages and Types Reference, 2024.
- [15] Oracle Corporation, "Oracle AI Vector Search Performance and Scalability Technical White Paper," Oracle Corporation Engineered Systems Publications, 2025.