

O'REILLY®
Report

Why Distributed SQL Is the Modern Foundation for AI

Simplifying and Accelerating
the Development of Resilient
AI Applications

Steve Suehring

Compliments of



yugabyteDB



Distributed PostgreSQL for the AI Era

Discover YugabyteDB - the AI-native, multi-modal, distributed PostgreSQL database for cloud-native applications.



Global Scalability



Ultra Resilience



Flexible Geo-Distribution

Trusted by industry-leading organizations including **Shopify**, **Kroger**, and **GM**, YugabyteDB accelerates innovation, reduces risk, and delivers enhanced enterprise capabilities.

Book a
YugabyteDB Demo



Sign-up
for Meko



To discuss why YugabyteDB would be a good fit for your cloud native apps, contact a YugabyteDB expert (yugabyte.com/contact) or join our open source community (inviter.co/yugabytedb).

Why Distributed SQL Is the Modern Foundation for AI

*Simplifying and Accelerating
the Development of Resilient
AI Applications*

Steve Suehring

O'REILLY®

Why Distributed SQL Is the Modern Foundation for AI

by Steve Suehring

Copyright © 2026 O'Reilly Media, Inc. All rights reserved.

Published by O'Reilly Media, Inc., 141 Stony Circle, Suite 195, Santa Rosa, CA 95401.

O'Reilly books may be purchased for educational, business, or sales promotional use. Online editions are also available for most titles (<https://oreilly.com>). For more information, contact our corporate/institutional sales department: 800-998-9938 or corporate@oreilly.com.

Acquisitions Editor: Nicole Butterfield

Development Editor: Gary O'Brien

Production Editor: Katherine Tozer

Copyeditor: Piper Content Partners

Proofreader: O'Reilly Media, Inc.

Cover Designer: Susan Brown

Cover Illustrator: Susan Brown

Interior Designer: David Futato

Interior Illustrator: Kate Dullea

July 2026: First Edition

Revision History for the First Edition

2026-07-14: First Release

The O'Reilly logo is a registered trademark of O'Reilly Media, Inc. *Why Distributed SQL Is the Modern Foundation for AI*, the cover image, and related trade dress are trademarks of O'Reilly Media, Inc.

The views expressed in this work are those of the author and do not represent the publisher's views. While the publisher and the author have used good faith efforts to ensure that the information and instructions contained in this work are accurate, the publisher and the author disclaim all responsibility for errors or omissions, including without limitation responsibility for damages resulting from the use of or reliance on this work. Use of the information and instructions contained in this work is at your own risk. If any code samples or other technology this work contains or describes is subject to open source licenses or the intellectual property rights of others, it is your responsibility to ensure that your use thereof complies with such licenses and/or rights.

This work is part of a collaboration between O'Reilly and Yugabyte. See our [statement of editorial independence](#).

979-8-341-67373-1

[LSI]

Table of Contents

1. The AI Application Revolution.	1
Comprehending the Exponential Growth of AI	1
Why Build-to-Last Fails for AI	3
2. Architectural Principles for AI Applications.	7
Outcomes Drive Results	7
Assessing Business Requirements	11
Designing for Human in the Loop	14
Nonfunctional Requirements Become	
System-Level Requirements	17
3. The Four Critical Considerations in Production AI Systems.	21
1. Navigating Rapidly Evolving AI Tech Stacks	21
2. Accommodating Diverse Data Types and Sources	25
3. The Complexity of Handcrafted Data Pipelines	28
4. The Painful Route from Prototype to Production	29
4. Security, Observability, Explainability, and AI-Enhanced Database	
 Operations.	37
Security, Observability, and Explainability Requirements	37
Agentic Database Administration and Performance Tuning	40
Simplified Migration and Operational Tooling	40
5. Conclusions and Making It Work.	43
From Prototype to Production-Ready AI	43
Next Steps for Implementing Distributed SQL for AI	46

The AI Application Revolution

Artificial intelligence (AI) is no longer confined to experimentation at the edges of the enterprise. It has moved into core workflows, decision systems, and customer-facing applications. This integration has reshaped expectations about speed, insight, and automation. As AI capabilities accelerate, the demands placed on underlying systems, particularly databases, increase just as rapidly.

This chapter explores how the exponential growth of AI is influencing enterprise adoption patterns and exposing structural gaps in traditional architectures. The chapter examines how organizations interpret change acceleration and how AI adoption reshapes expectations around build-to-last architectures. The chapter concludes with a look at why modernization efforts must extend beyond the application layer and into the database and infrastructure as well.

Comprehending the Exponential Growth of AI

Organizations struggle with AI adoption because they plan for linear improvement while AI capabilities evolve nonlinearly. Traditional systems improve in predictable increments, which supports steady planning, budgeting, and scaling assumptions. AI does not follow this pattern.

Exponential growth appears gradual at first but accelerates rapidly as each improvement compounds on the last. Early progress may seem incremental, but once capability reaches a threshold,

each subsequent doubling produces significantly greater impact. **Figure 1-1** illustrates this concept.

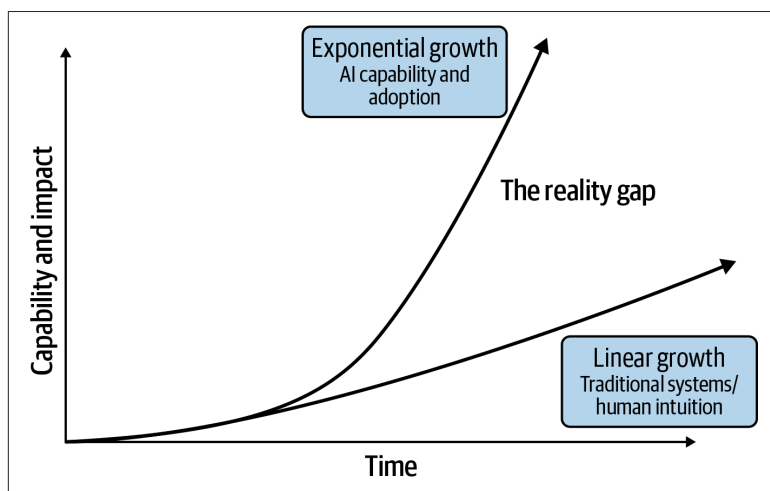


Figure 1-1. Exponential growth versus linear growth; comparing AI systems with traditional systems

The Rapidly Evolving State of Enterprise AI Adoption

AI development reflects this pattern. Advances in data, model size, architecture, and training techniques build on one another, transforming early novelty into practical utility. As a result, enterprise adoption often follows a similar trajectory where limited experimentation gives way to rapid expansion when performance crosses a threshold of real-world value.

For enterprise adoption, the difficulty is not recognizing the potential of AI, but identifying the inflection point where incremental progress becomes exponential change. Organizations that continue to plan for linear growth risk underestimating future demand and designing systems that cannot support the pace of AI-driven transformation.

Managing the Reality Gap

AI-driven progress raises expectations as quickly as it improves capability. Tasks that once took days are now expected to happen in hours, and manual workflows are assumed to be automated. The problem is that processes like planning, governance, and

procurement do not accelerate at the same pace. This creates a gap between what AI enables and what organizations can operationalize.

The gap is not caused by a lack of awareness, but by a mismatch in speed. As new use cases emerge and scope expands, static planning models struggle to keep pace. AI adoption is not a one-time initiative but an ongoing acceleration, and organizations that treat it as a discrete risk are falling behind. Closing the gap requires aligning expectations with systems and structures that can continuously adapt.

Why Build-to-Last Fails for AI

The *build-to-last* (B2L) mindset assumes stable requirements and gradual change. That model no longer works. AI introduces continuous and compounding improvements in capability and demand, making static architectures more of a liability than a strength.

B2L Now Means Designing for Change

Designing for change means treating evolution as a default condition. The goal is not to freeze requirements, but instead to build systems that can absorb new use cases and new patterns of demand without repeated redesign. Iteration becomes a property of the system itself. Modularity is favored over rigidity, automation over manual intervention, and adaptive scaling and resilience as part of the core architecture.

In the context of AI, this also means accounting for model updates, evolving data pipelines, and unpredictable workload patterns driven by agents and real-time inference. Systems must be able to incorporate new capabilities without disrupting existing operations, while adhering to performance criteria and compliance requirements. Designing for change is now about building and sustaining momentum in an environment where advancements push forward rapidly.

Going Beyond Code Changes

AI transformation is frequently treated as an application-layer initiative. Organizations focus on models, APIs, and user-facing functionality. However, intelligent operational applications often depend on real-time analytic processing within the data platform itself, rather than relying on batch-oriented data movement. As

operational analytical workloads converge, the database becomes a central component of AI execution rather than a passive store of record.

This shift exposes the limits of legacy architectures. Systems designed for predictable transaction processing and incremental growth often rely on centralized scaling models and manual operational controls. While migration may relocate these systems to new environments, it does not alter their fundamental characteristics. Structural constraints around scalability, availability, and performance still remain in effect.

Modernization addresses the constraints, but not without introducing new considerations. Distributed architectures provide horizontal scalability by adding nodes rather than vertically expanding a single system. Built-in replication and automated failover support high availability, while geo-distribution brings data closer to users and applications to enable consistent performance across regions. Observability and operational control also become embedded capabilities rather than external add-ons.

However, many organizations encounter a new form of complexity as they modernize. In practice, supporting diverse workloads such as transactional processing, analytics, and AI often requires assembling multiple specialized systems. This results in database sprawl, duplicated data pipelines, and increased operational overhead. Instead of eliminating constraints, modernization efforts can shift them into integration challenges across fragmented data platforms. Reducing this complexity requires a data foundation capable of supporting multiple access patterns and workloads within a unified architecture.

This challenge becomes even more pronounced in the context of AI-driven applications. Modern AI systems depend on heterogeneous data, including structured application data, unstructured documents, embeddings, and vector representations. In many environments, these data types are managed across separate systems, requiring preprocessing pipelines to transform and synchronize data between them. This fragmentation increases latency, complicates governance, and makes it difficult for AI agents to maintain consistent context and memory. [Figure 1-2](#) shows this problem. The reason most multi-agent systems fail isn't due to logic but rather a lack of shared memory and knowledge.

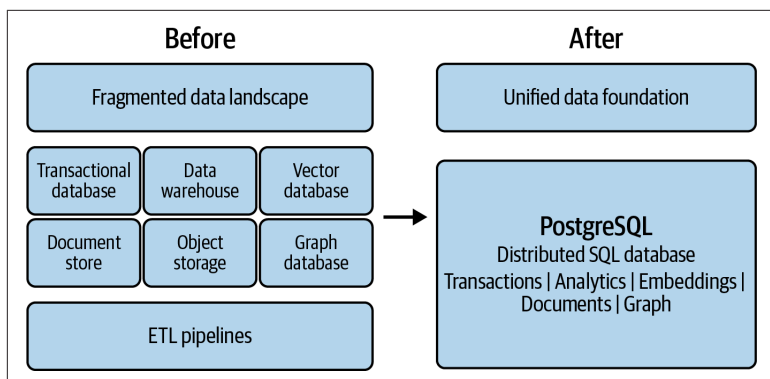


Figure 1-2. Data fragmentation becomes a significant challenge that prevents AI systems from using the latest data

Shared State Inconsistency

A study on multi-agent systems found that 36.9% of failures in agentic AI environments originate with inconsistent shared state between agents. These failures occur when one agent updates the system but that change is not reliably visible to others, leading later agents to act on stale or invalid assumptions. As a result, decisions that were correct at one point in time quickly become misaligned with the current state of the system (Cemri et al., 2025).¹

True modernization is not only about scalability and resiliency but also about unifying how data is stored, accessed, and processed. Platforms that can support relational, analytical, and vector workloads reduce the need for complex data movement and enable more efficient AI execution. With this type of platform and architecture, the database evolves from a system of record into a system of intelligence, where agents can access and update context in real time.

Additional important characteristics such as scalability, replication, availability, and regional distribution are not optional enhancements but are core to the functionality needed to drive the real-time value gains from AI. These capabilities are the operational foundation that enables intelligent applications to scale reliably. Without modernization at the data and infrastructure layer, AI adoption outpaces the systems required to sustain it.

¹ Cemri, M., et al. (2025). "Why Do Multi-Agent LLM Systems Fail?" *arXiv: 2503.13657*, <https://doi.org/10.48550/arxiv.2503.13657>.

Architectural Principles for AI Applications

When architected with intent, AI initiatives bring significant value to the business. Data is the foundation on which successful architectures are created. This chapter looks at the architectural decisions that unlock business value from a well-designed AI-enabled application.

Outcomes Drive Results

To succeed with AI, organizations need to understand their existing systems and data. Architectural decisions made early based on the desired outcome help to ensure success. When incorporating data into early decisions, architects also ensure scalability and enable rapid response, increasing agility for the business.

Business-Driven and Data-Oriented Architecture

Business outcomes determine the real-world constraints of the AI-enabled system, such as response times, fault tolerance, scaling, and other operational considerations of the system. An AI system that needs to support real-time decision making has different architectural requirements than a system designed for asynchronous analysis or offline recommendations. Important decisions like data placement, system boundaries, and deployment characteristics need to be made early.

A data-oriented perspective is essential when architecting AI systems for production. Data provides the true foundational value for AI systems, making decisions around data essential for the overall success of the project. Postponing decisions about data creates risks such as needing to rework the architecture due to performance or reliability problems. Rearchitecting becomes exceedingly difficult for mission-critical systems.

Reducing user burden

Reducing user burden is not simply about automation. From an architectural perspective, reducing user burden means reducing operational friction and cognitive load. An AI system that behaves inconsistently or requires constant verification increases user burden rather than reducing it. This inconsistency is caused not only by slow or unstable system behavior but also by fragmented data that exists across operational, vector, and unstructured systems.

From an architectural perspective, reducing user burden requires explicit design choices that limit the surface area of unpredictability:

Predictable behavior boundaries

AI systems must clearly signal when outputs are reliable and when they are uncertain. Likewise, architectural design choices should help to detect and mitigate partial failures such as latency and capacity pressure. The architecture should be able to surface these conditions and steer around the unhealthy components. Users should never be required to infer system state.

Deterministic system integration

Distributed data architectures create a level of abstraction, hiding things like shard placement, replication, and underlying failover mechanics. This abstraction enables application logic to remain deterministic, even as the underlying system adapts to failure and scaling events. The overall result is a reduction in the operational burden placed on users, developers, and systems staff.

Data pipeline simplification

AI systems often rely on complex preprocessing pipelines to prepare operational, unstructured, and vector data for use in models. When these pipelines are manual, fragmented, or inconsistent, they introduce errors and increase the burden on both developers and operators. Architectural design should

reduce this complexity by minimizing the need for custom preprocessing steps and enabling more integrated, consistent data handling across the system.

Failure containment

Ultra-resilience explicitly prioritizes containment of failures within a given fault domain. This prevents outages from becoming user-facing service degradations. By maintaining service continuity through replication and automated failover, systems avoid forcing users to diagnose or work around infrastructure-level issues.

Enabling rapid response

Rapid response in AI systems is not determined by inference speed alone. End-to-end response time is heavily influenced by data locality, coordination overhead, and consistency guarantees. This is especially important in globally distributed environments. Architecturally, enabling rapid response requires latency budgets, intentional AI logic placement, and time-bounded execution:

Explicit latency budgets

The full request path and all layers, including data access, coordination, and replication, need to have predictable low-latency architectures that minimize coordination delays. Systems that lack explicit constraints around latency often degrade quietly under load, creating a cascade of delays.

Intentional placement of AI logic

Placing applications close to data replicas helps to ensure that time-sensitive functions execute with minimal latency. Distributed SQL systems support locality-aware data placement, which then guards against latency-related issues and cross-region penalties.

Time-bounded execution

An architecture should be built to fail fast, detecting potential issues and then quickly routing around those issues. Degraded regions or nodes should not affect overall application availability, thereby preventing slow components from dominating response times. Additionally, architectures that support persistent agent memory and reuse of prior context can reduce redundant computation, improving response time, increasing accuracy, and lowering token consumption over repeated interactions.

Simplifying tasks

AI systems are often introduced to help simplify complex tasks. But simplification at the interface level can conceal significant architectural complexity. A modern distributed architecture helps to reduce operational complexity by centralizing intelligence in the platform rather than the application.

Three key characteristics are common in architectural simplification efforts:

Clear separation of responsibilities

Application logic and infrastructure are clearly separate. Application developers should not create or add logic around resiliency but rather should be able to rely on the underlying infrastructure to work around any potential issues.

Modular boundaries

AI behavior should evolve separately from core business workflows. Too much integration leads to fragility and difficulty innovating.

Observable decision paths

Built-in metrics and monitoring provide visibility into latency, load distribution, and failure modes. Performance and compliance are then verifiable by operators at the system level and in a repeatable manner.

Ensuring scalability

For AI systems, consider whether the AI component needs to operate synchronously in response to user-facing events or if the AI responds to background workflows. Asynchronous workflows can often be predicted or controlled to work around other demand-related access. However, architects need to define how the system should behave in response to burst traffic and uneven access patterns.

Distributed SQL systems address these issues through automatic sharding, replication, and load balancing. However, architects should be aware that early decisions about these elements and early assumptions about application demand cycles can make it difficult to rearchitect for scaling later.

Reducing costs

Intentional design and architectural decisions reduce cost by avoiding manual intervention and reducing operational overhead. Making these decisions requires visibility into cost drivers, avoidance of duplicated effort and infrastructure, and predictability around scaling behaviors. This is true even if demand cannot always be predicted.

Assessing Business Requirements

A business requirements assessment for AI systems is not a traditional requirements gathering exercise. The purpose is not solely tied to enumerating features but rather to surfacing large scale architectural constraints early, before assumptions create difficult implementations. The goal of the requirements gathering exercise is to determine what kind of system is being built, what guarantees it must provide, and where uncertainty is acceptable. Four dimensions are particularly helpful: the role of AI within the application, the interaction models that the system needs to support, the way in which data is generated and updated, and how users and consumers are distributed.

Is AI the App?

One of the most consequential decisions that needs to be made is whether AI is the app itself or is a capability that is embedded into an existing (or newly built) system. This distinction determines where authority resides and how much variability the system can tolerate.

Conversational systems, AI-native analytics tools, and planning agents are prime examples of AI being the application itself. When created for internally facing systems, probabilistic behavior is often acceptable as long as interactions remain coherent and recoverable within boundaries. Users expect iteration, approximation, and refinement in these systems, and priority is often placed on availability and responsiveness over strict determinism.

By contrast, when AI is added to an existing application or workflow, the AI portion operates within constraints and expectations around correctness and consistency. The AI-produced output cannot become the system of record or redefine correctness for the

system in such cases. Distributed data platforms are explicitly designed to preserve transactional integrity and predictable behavior, even as their dependent systems have failures and service degradations. The distributed data platform provides the foundation for deterministic outcomes around which the AI portion can be built.

Architecturally, this portion of the requirements assessment determines whether AI outputs are:

Advisory

Influencing decisions that are validated elsewhere, such as other portions of the app or external to the app

Authoritative

Providing direct triggers for state changes or actions

Failing to make this architectural distinction often results in systems where AI components silently accumulate authority that they were never designed to have.

Understanding User Interaction Model Options

AI that interacts through a traditional user interface (UI) might show AI-generated recommendations, summaries, and predictions rather than acting autonomously. Chatbots extend the UI interaction into a more conversational format. Like familiar generative AI, chatbots can recall context throughout the conversation. AI chatbots utilize stateful interaction to maintain persistence of context and conversation history. This increases the importance of observability, as chatbots are required to respond even when confidence is low, making it more difficult to identify and diagnose unreliable outputs.

Model Context Protocol (MCP) represents a different type of user interaction. Instead of being human-oriented like a chatbot, MCP is machine-oriented. MCP enables external tools, services, and agents to provide context to AI models. MCP interactions are programmatic rather than conversational.

Agentic behavior is typically the most demanding interaction model. Autonomous agents may invoke AI continuously, chaining actions across multiple systems. From an architectural perspective, autonomous or agentic AI systems require strict scoping of authority and permissions along with bounded execution and rate limiting. Unlike human interaction, agentic AI does not recover gracefully from ambiguous responses and may be unable to infer nuance. As a

result, architectures that utilize agentic AI need to enforce stronger isolation boundaries than those designed as chatbots.

Agentic systems also introduce stateful components, such as memory, knowledge retrieval, and intermediate reasoning steps, that must remain consistent across interactions. When these capabilities are implemented across multiple disconnected systems, coordination becomes more difficult and operational complexity increases. Architectures that unify these stateful elements and maintain persistent agent memory are better able to support reliable, scalable agent behavior.

Data Generation Patterns and Update Frequency

AI systems are structurally dependent on how data is produced, processed, updated, and accessed over time. Unlike traditional applications where schemas and access patterns change slowly, AI systems must contend with high data churn, evolving semantics, and feedback loops. Update frequency and coordination overhead affect latency and scalability decisions; this is particularly important in distributed environments where data needs to be replicated and rebalanced dynamically.

Performance, correctness, and resilience are tightly coupled with data access patterns and update frequency. Systems that are designed for read-heavy analytical access behave differently when compared with systems that support frequent transactional updates requiring low latency and strong consistency.

Early architectural decisions need to address several factors around the data lifecycle:

- Which data sources are authoritative versus derived?
- How fresh does the data need to be for AI-enabled systems and AI-influenced decisions?
- Are data updates continuous, bursty, or batched?
- How will schema changes be managed over time?

Ignoring or deferring these questions means rework later. Ad hoc caching schemes, duplicated pipelines, or application-centric custom synchronization logic will typically need to be added when these questions are not answered. While these add-ons can help to solve problems, doing so adds operational debt and slows innovation.

User Distribution

Modern AI systems are rarely local to the users or the sources of data. Consumers of AI are geographically distributed, including both human and nonhuman agents. With disparate systems involved, partial failures become more common than total outages. Systems need to assume uneven performance and intermittent connectivity or availability of upstream systems and downstream consumers.

Data locality and latency become key architectural decision points. Compliance and data residency are part of the conversation as well. Distributed SQL architectures are designed to present a single logical system while handling things like replication, failover, and locality internally, rather than requiring the application to handle these requirements.

Designing for Human in the Loop

AI systems introduce probabilistic reasoning into environments that are otherwise expected to behave deterministically. In production architectures, especially those backed by strongly consistent distributed SQL platforms, this contrast must be managed deliberately. The data platform guarantees transactional integrity, consistent state, and predictable failure behavior across nodes and regions.

NOTE

Transactional integrity, consistent state, and predictable failure might sound familiar as the core ACID properties for databases: Atomicity, Consistency, Isolation, and Durability.

Designing for human involvement from the beginning ensures that probabilistic outputs are reconciled with deterministic system boundaries. **Figure 2-1** illustrates how this split works.

Human review, escalation paths, and authority constraints should be embedded directly into workflows that are then persisted and enforced by the underlying distributed database. This prevents AI components from silently accumulating authority beyond what the architecture was designed to tolerate.

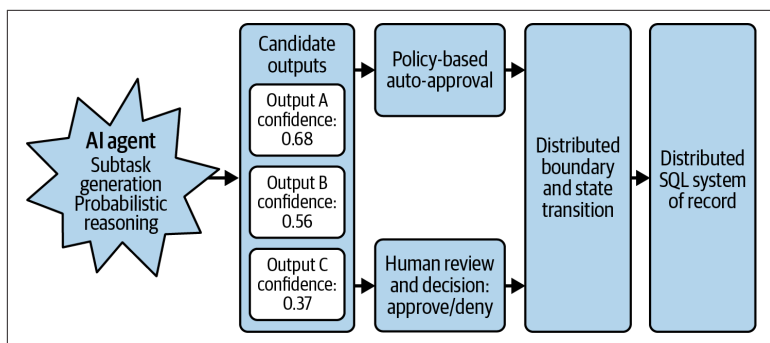


Figure 2-1. A human-in-the-loop workflow with separation at state transition when needed

When human validation is introduced late, it often appears as conditional logic in application code rather than as a durable workflow state. In distributed architectures, this creates ambiguity around where authority resides and how state transitions are committed. Instead, review states, approval gates, and overrides should be modeled as explicit data transitions.

With ACID semantics preserved across regions and replicas, the data transitions remain consistent even during node failures or regional degradation. Human checkpoints then become part of the same deterministic framework governing all other state changes. The database remains the system of record, and AI operates within well-defined transactional boundaries.

Why Transparency and Maintainability Matter

AI systems layered on top of this infrastructure must surface their own signals of quality. Things like confidence levels, input provenance, and decision outcomes should be stored in such a way that they can be audited alongside system state. Audit trails remain intact even during failover events because the data platform guarantees consistency and durability.

Maintainability follows from separation of concerns. The platform manages sharding, replication, locality, and recovery. Application logic does not implement custom resilience mechanisms. By keeping AI logic modular and avoiding tight coupling with infrastructure behavior, model updates and workflow evolution can occur without destabilizing the system of record.

Using AI for Subtasks Instead of Automation

Bounded AI tasks align naturally with architectures that emphasize deterministic enforcement at data boundaries. Classification, ranking, summarization, and anomaly detection produce outputs that can be validated before triggering transactional changes. When AI is constrained to well-defined subtasks, its outputs can be evaluated within strongly consistent transactions. Confidence thresholds, validation rules, and human review steps can be enforced before committing state changes. The distributed SQL layer ensures that once committed, the changes are replicated across regions consistently.

Complete automation removes these validation boundaries. In systems that require correctness guarantees, allowing probabilistic output to directly mutate authoritative data increases the risk of inconsistent or unintended state transitions. Strong consistency at the data layer provides a foundation for correctness, but only if architectural controls preserve the distinction between recommendation and execution.

AI as Augmentation

Augmented systems rely on a stable and authoritative data layer to preserve correctness while AI accelerates analysis and decision support. Distributed SQL platforms are designed to present a single logical system, even while managing important architectural components such as replication, failover, and locality. This stability enables AI to operate as an advisory layer without redefining system truth.

When AI suggestions are stored alongside authoritative records within the same consistent database, users can compare recommendations with the current state, historical context, and related transactions. This helps to preserve both auditability and repeatability. Augmentation becomes architecturally viable when the system of record remains deterministic and resilient. AI enhances workflows but the data platform ensures continuity and correctness under failure conditions.

Nonfunctional Requirements Become System-Level Requirements

Nonfunctional requirements define how the system behaves rather than what it does. What might traditionally be considered nonfunctional requirements around latency, replication timing, and scaling behavior are core structural properties of database systems that support AI-enabled systems. These nonfunctionals must be considered alongside the functional requirements of an AI-enabled system.

Distributed SQL systems manage automatic sharding, replication, and load balancing internally; the operational burden of these systems is reduced. The capabilities are incorporated into architectural planning early. Latency budgets, locality decisions, and replication strategies shape system behavior long before production traffic arrives.

Resilience and Scalability as Primary Choices

Resilience in distributed systems is achieved through properties such as replication, automated failover, and containment of failures within defined fault domains. These properties are built into the data platform itself. Applications that assume a single-node or single-region model cannot retroactively inherit these guarantees without architectural change.

Scalability likewise depends on early decisions about data distribution and locality. Distributed SQL systems rebalance shards and replicas automatically, but application-level assumptions about data placement or coordination patterns can introduce bottlenecks if they are not aligned with the capabilities of the platform.

Attempting to bolt on resilience or scalability after deployment often leads to duplicated services, inconsistent caches, and complex and often fragile synchronization logic. These patterns undermine the deterministic behavior that the distributed platform is designed to provide. Architectural alignment with replication, locality-aware placement, and automatic rebalancing from the outset ensure that growth and failure are handled within the system's intended design.

Cloud-Agnostic and Commodity Hardware Considerations

Distributed SQL architectures are designed to run across regions and infrastructure environments while presenting a unified logical database. This abstraction enables flexibility in deployment topology without requiring that the application manage replication or failover directly. Cloud-agnostic design is not an abstract preference but rather an intentional design choice.

Design choices for cloud architecture affect data placement, compliance posture, and cost predictability. Architectures that rely heavily on proprietary services can constrain mobility across environments and complicate multi-region and multi-cloud deployment strategies. By contrast, systems built on distributed SQL with locality-aware data placement and replication policies can deploy across multiple regions and providers while maintaining strong consistency and a single logical schema.

Data Residency, Compliance, and Cost Factors

Data residency requirements influence where replicas are placed and how traffic is routed. Locality-aware data placement enables specific datasets to reside within defined geographic regions while remaining accessible through a unified logical database interface. Applications do not need complex custom routing logic in order to enforce residency constraints.

Compliance requirements are supported through deterministic transactional guarantees and consistent replication policies. Auditability is preserved across regions due to transactionally committed state transitions with predictable replication patterns. As regulatory frameworks evolve to address GenAI- and LLM-specific concerns, such as data provenance, model training sources, and jurisdictional restrictions on model usage, architectures will need to account for stricter controls around what data can be used, where that data originates, and how it is processed. Systems that maintain clear data lineage and enforce consistent policies across regions will be better positioned to adapt to these emerging compliance requirements.

Architectural decisions and business-level requirements around data placement and replication requirements also directly affect operational expense. Cost considerations are influenced by

replication topology, cross-region traffic, and storage configuration within these boundaries. Designing with awareness of the options available for these parameters prevents cost overruns and expensive reconfiguration later.

Commodity hardware considerations further affect portability and long-term flexibility. Architectures that rely on distributed systems capable of operating across varied infrastructure environments help to preserve options in provider selection, replication, data redundancy, and deployment strategy.

The Four Critical Considerations in Production AI Systems

AI systems behave very differently in production environments than in prototype or pilot environments. Early experiments focus on model quality and feature velocity, but production deployments expose pressures around scale, resilience, governance, and operational control. As AI capabilities become embedded in customer-facing applications and revenue-generating workflows, the database layer carries sustained load from transactional traffic, vector search, real-time analytics, and automated agents operating simultaneously.

This chapter examines four critical considerations that determine whether an AI system can move from experimentation to durable production use. The chapter explores how to navigate rapidly evolving AI technology, accommodate diverse data types, avoid brittle pipelines, and design for scalability.

1. Navigating Rapidly Evolving AI Tech Stacks

The modern AI stack evolves at a pace that exceeds traditional enterprise technology cycles. New foundation models appear often, along with dimensional changes and revisions to inference APIs. Alongside these changes, orchestration frameworks evolve and consolidate rapidly. Organizations with architectures designed for change and iteration absorb these changes without issue.

The deeper challenge is not in selecting the right model or framework but rather in designing a data foundation that remains stable while the layers above it change. Production AI systems must assume change as a design choice without destabilizing transactional integrity, compliance controls, or service levels and resiliency. The assumption of certain evolution shifts the architectural focus downward, toward the database layer. It is within the database layer where long-term durability and standards compliance matter more than short-term velocity.

The Case for PostgreSQL as the Foundation for AI Data Infrastructure

AI systems that move beyond experimentation and into mission-critical workflows place sustained pressure on the data layer. Once models are embedded into customer-facing applications, emphasis on consistency, availability, and compliance across regions becomes more important. In these environments, a distributed SQL database that preserves compatibility with popular database systems like PostgreSQL provides a stable and extensible foundation for long-term production deployment.

PostgreSQL compatibility is important because it protects existing application investments while enabling architectural evolution. Applications continue to use standard PostgreSQL drivers, object-relational mappings (ORMs), migration tooling, and operational monitoring systems without modification. By maintaining SQL compatibility, organizations can introduce AI capabilities into existing systems without needing parallel data stacks or rewriting related service interfaces.

Maintaining separation between the application database and the AI database becomes costly. Embeddings are derived from transactional records and enriched with metadata before being joined back with structured tables at query time. By keeping relational data and vector representations near to each other, for example, in the same PostgreSQL database, the data layer can ensure atomic writes, consistent reads, and simplified resiliency architecture under a single-access control model.

Distributed SQL architecture strengthens this foundation by enabling horizontal scale without sacrificing relational guarantees. As vector volumes grow and retrieval traffic increases, additional nodes can be added incrementally while also preserving standard SQL semantics. This enables the system to scale out rather than needing to scale up in order to meet demand. A distributed SQL architecture helps to avoid operational ceilings that are often encountered with single-node deployments.

Production workflows require a distributed data layer that can tolerate node, zone, and even regional failures. Survivability is equally important once AI features become embedded in those production workflows. Because AI-driven responses often sit directly in user journeys, resilience at the database layer directly affects customer experience.

Geo-distribution further reinforces the need for a unified foundation. Applications increasingly operate across regions with differing latency expectations and data residency requirements. A distributed SQL system that supports geo-partitioning and multi-region replication enables data to be placed deliberately while still appearing as a single logical database to the application layer.

Architectural fragmentation is reduced by consolidating transactional workloads, vector search, and distributed resilience into a single PostgreSQL-compatible system. There is no separate vector data store with its own replication strategy, consistency model, and security surface to reconcile. Instead, AI capabilities extend the system, inheriting its traits such as durability, governance controls, and scaling.

Understanding the Need for Open Standards and Flexibility from Pluggable Vector Stores

Open standards play a stabilizing role for enterprise technology. Considering AI and the data layer specifically, open standards help to ensure that things like drivers, ORMs, and migration tools continue to function even as vector capabilities are introduced or enhanced. By preserving standard interfaces, such as PostgreSQL wire compatibility, organizations avoid creating parallel AI-specific data stacks that would then have their own operational needs and security surface area.

Pluggable vector capabilities further reinforce this insulation by preventing tight coupling between application logic and a single embedding or indexing implementation. As embedding dimensionality shifts or new ANN indexing approaches mature, organizations can introduce new vector indexes alongside existing indexes rather than performing disruptive migrations. This flexibility aligns with the broader architectural principles of assuming change at the AI layer while preserving stability at the data layer.

Flexibility also mitigates the risk of over-optimizing or optimizing too early during experimentation. Teams frequently begin with smaller embedding models and vector volumes but encounter scaling issues when attempting to embed the features into the production-level user journey. A PostgreSQL-compatible distributed SQL foundation that supports vector workloads enables incremental evolution rather than cutover-style migrations to specialized data storage.

Operational simplicity is another advantage of open standards. Separate vector databases almost certainly have their own design paradigms for replication, backup, authentication, and monitoring. Consolidating those properties into a single distributed SQL system reduces overhead. Open interfaces also preserve options around cloud provider deployments. When AI data infrastructure adheres to established PostgreSQL drivers and SQL tooling, it can operate across regions and cloud providers. Both multi-region and multi-cloud capabilities become available and further facilitate requirements around compliance, cost optimization, and latency management.

Finally, pluggable vector stores reinforce long-term sustainability by decoupling retrieval mechanics from business logic. Retrieval-Augmented Generation (RAG) systems evolve as new ranking strategies, metadata filters, and hybrid search techniques emerge. By embedding vector search within a standards-compliant relational system, organizations retain the ability to refine retrieval strategies without destabilizing the surrounding application ecosystem.

2. Accommodating Diverse Data Types and Sources

Production AI systems do not operate on a single, neatly structured dataset. Instead, they draw from transactional records, operational logs, external APIs, document repositories, and unstructured content like PDFs, chat transcripts, and metadata from multimedia sources. Designing for this diversity requires more than simple storage capacity and instead demands a coherent architecture that supports the different data types and access patterns.

The challenge of supporting data types and access patterns increases when AI capabilities are embedded directly into customer-facing workflows. These workflows require high retrieval quality, low latency, and freshness, all of which are challenges requiring effective access to both structured and unstructured data. As a result, the data layer needs to unify access without sacrificing relational guarantees or operational resilience.

Why RAG Requires a Diversity of Data Sources

RAG systems depend on more than static document embeddings. High-quality responses often require combining vector similarity search over unstructured corpora with structured filters drawn from transactional systems. For example, a generated answer may need to reference a policy document while also verifying account status or region-specific constraints that are stored in relational tables. **Figure 3-1** illustrates this concept.

The hybrid nature of many interactions, requiring data from disparate sources, introduces complexity. Embeddings derived from documents need to remain tightly associated with metadata, access controls, and structured attributes to ensure accurate filtering and authorization. By co-locating vector representations with relational data, organizations enable consistent joins and enforce governance policies uniformly.

RAG systems evolve as new external data feeds are incorporated. Market data, telemetry streams, and third-party APIs may need to be embedded alongside internal records. A flexible data foundation enables these sources to be integrated without requiring bespoke ingestion and query layers for each new dataset.

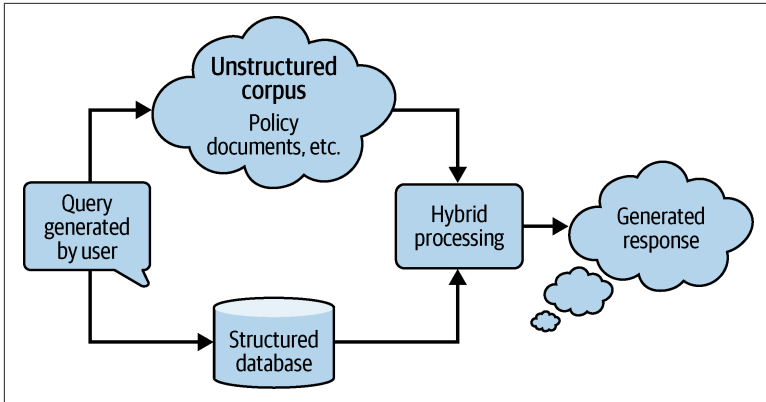


Figure 3-1. Gathering disparate data in order to generate a response based on the latest information

Why You Need Unified Data Access Across Disparate Sources

Unified access reduces both latency and architectural fragility. Response times increase when applications need to query multiple systems, such as one system for vectors, another for transactional data, and a third for operational metadata. Each additional logical retrieval increases the chances of failure that will directly affect user-facing AI features.

A single logical database that supports both relational and vector queries simplifies the coordination. Standard SQL queries combine similarity search with structured predicates, enabling precise and explainable retrieval workflows. This approach eliminates the need for application-level connective logic that attempts to reconcile inconsistent results across multiple data stores.

Governance and compliance considerations further reinforce the value of unified access. Data residency requirements, audit logging, and role-based access control (RBAC) become significantly more complex when AI data is dispersed across specialized services. Consolidation into a PostgreSQL-compatible distributed data layer ensures that established controls are applied uniformly to AI data.

Balancing Transactional, Operational, External, and Unstructured Data

AI systems require reconciliation of competing workload characteristics within a single architecture. Transactional workloads demand low latency writes and strong consistency, while vector search may require compute-intensive similarity calculations over large indexes. External and unstructured data introduce variability around ingestion rates, accuracy, schema flexibility, and update frequency.

Balancing these demands requires a system capable of horizontal scale without sacrificing relational guarantees. Distributed SQL architectures enable nodes to be added incrementally as vector volumes grow, while still preserving semantics around ACID. This prevents the common antipattern of isolating AI workloads in a separate system that eventually drifts away from the source of truth.

The objective is not to collapse all workloads indiscriminately but rather to ensure that the workloads share a consistent and coherent governance and resilience model. By aligning transactional integrity, vector retrieval, and geo-distribution within a unified PostgreSQL-compatible system, organizations create a data layer that evolves alongside AI capabilities. This balance reduces architectural entropy while supporting sustained progress and innovation in production environments.

Consistency becomes more visible when AI outputs directly influence user decisions and automated actions. If embeddings are generated from stale transactional data, retrieval results may contradict the current system state. Ensuring atomic writes between structured records and their associated vector representations reduces the risk and preserves trust in the AI-drive features.

Distributed deployments further complicate consistency expectations. Multi-region replication and geo-partitioning need to be configured to align with latency and data residency requirements while preserving a coherent view of data. A distributed SQL system that maintains strong consistency across nodes enables AI applications to operate confidently even during node or zone failures.

Ultimately, data consistency in production AI is not a theoretical concern but a customer experience requirement. As AI features move from experimentation to core workflows, discrepancies between structured data and retrieved context can erode credibility.

Designing the data layer to enforce relational guarantees across transactional and vector workloads ensures that AI systems remain reliable as they scale.

3. The Complexity of Handcrafted Data Pipelines

As AI initiatives mature, the hidden complexity of heavily customized data pipelines becomes increasingly visible. Early prototypes often rely on scripts that extract data from a handful of sources, transform it informally, and generate embeddings in batch jobs that are loosely monitored, if they're monitored at all. Prototype scripts have a habit of becoming production scripts, but this approach rarely withstands the reliability, scaling, and governance needs of production systems.

The difficulty with customized data pipelines is not simply technical debt but also operational fragility. Each custom data ingestion pipeline becomes a separate failure domain that must be monitored, retried, and audited. When AI-driven features are incorporated into customer-facing applications, even minor pipeline disruptions can surface as incorrect AI answers.

Why Roll-Your-Own Slows Innovation

Roll-your-own pipelines often emerge by necessity due to time constraints rather than through a carefully-architected strategic plan. When time pressure exists, as it often does, teams might stitch together extraction logic, vectorization code, and indexing workflows using ad hoc orchestration. Over time, these bespoke components become tightly coupled to specific models, embedding dimensions, and assumptions. Often, the process is known only to certain individuals, resulting in a lack of visibility when things break.

As models evolve and embedding dimensions change, handcrafted pipelines require coordinated updates across ingestion scripts, storage schemas, and retrieval logic. Without standardized abstractions at the data layer, these updates introduce risk and downtime. Instead of enabling rapid experimentation, the pipeline becomes a constraint.

Negative customer impact follows from this fragility. If ingestion lags behind source-of-truth systems, embeddings become stale and retrieval quality degrades. When failure handling is inconsistent, partial updates may result in vectors that no longer align with their structured metadata, leading to confusion and contradictory or incorrect AI responses.

Why Automation Is Key in Data Preprocessing Pipelines

Automated preprocessing pipelines reduce variability and enforce consistency across ingestion workflows. Rather than relying on manually triggered scripts, production systems require repeatable processes for normalization, metadata enhancement, embedding generation, and indexing. These processes need to integrate with the broader database foundation to ensure atomic updates and consistent reads. Automation also supports observability and auditability by detecting failures, retrying, and logging within established tooling and monitoring.

Most importantly, automated pipelines restore focus on product innovation. Engineers can iterate on models, retrieval strategies, and application logic without needing to continually rework fragile ingestion code. By anchoring preprocessing workflows to a scalable distributed data layer, organizations create a repeatable path from raw data to reliable AI functionality in production environments.

4. The Painful Route from Prototype to Production

The transition from prototype to production is where most AI architectures fall down. The platform functions smoothly in the pilot environment, maybe even under simulated stress testing, but then fails when deployed to production and the real-world stresses inherent in that environment. Underlying assumptions about scale and locality that were true for the pilot are revealed when real users begin working with the application. The friction and failure modes can occur due to model quality but also because of infrastructure and problems that could have been solved earlier.

Production AI systems need to assume growth in both data and user expectations. Vector volumes expand, retrieval traffic becomes unpredictable, and user journeys change and evolve, testing the

limits of latency and accuracy. Moving beyond experimentation and the pilot implementation requires a data foundation that scales horizontally in a natural way, assuming scaling and failure and handling both gracefully.

Distributed Scaling and Resilience

When AI becomes embedded in customer-facing applications, the need for horizontal scaling becomes even more obvious and important. Retrieval workloads increase as more users rely on AI-generated responses, and vector indexes grow as additional documents, events, and embeddings are ingested. A distributed SQL architecture enables scale-out by adding nodes incrementally while preserving relational guarantees and the underlying SQL semantics.

Single-node systems eventually encounter hard ceilings on performance, both storage and compute. When vector search, transactional updates, and analytical queries compete for the same resources, the natural outcome is operational instability and spikes in latency. Distributed deployments avoid these ceilings by sharding data and query processing across nodes without fragmenting the logical database.

Applications often serve users across multiple regions, which imposes its own set of latency constraints along with the potential for data residency requirements. This global distribution introduces another dimension of complexity. A system that supports multi-region replication and geo-partitioning enables deliberate placement of data while maintaining the logical database layer to upper levels such as application logic.

When scaling horizontally, resilience becomes even more important. Assuming and thus planning for node, zone, and regional failures ensures that such an event does not interrupt user-facing AI workflows. Distributed SQL systems designed for survivability handle failures within the infrastructure layer rather than bubbling those failures up to the application.

The architecture needs to expand without requiring a fresh platform deployment while also avoiding fragmentation and downtime. Designing these properties into the database layer early helps to prevent painful rearchitecting and rewriting later.

Resilience maturity is often not achieved in a single architectural leap. Organizations typically move from single-region deployments to multi-zone redundancy and eventually to multi-region and multi-cloud deployments, increasing continuity capabilities and failure tolerance as they move through these deployments. The challenge lies in advancing along this resilience curve without introducing service interruptions, data loss, or architectural fragmentation.

Distributed SQL systems support this progression by enabling incremental expansion of replication factors, geographic distribution, and fault domains without redefining application logic. Because the database presents a single logical abstraction, even as it spans multiple nodes and regions, applications do not require rewrites as the organizational resilience posture improves. This continuity enables teams to strengthen availability guarantees and recovery objectives while maintaining consistent developer interfaces and operational workflows.

Moving up the resilience curve also requires aligning recovery objectives with architectural capability. Recovery time objective and recovery point objective targets need to be supported by appropriate replication strategies that reflect business requirements, whether the replication requirement is synchronous or asynchronous. By embedding these mechanisms directly into the data layer, organizations can increase fault tolerance and geographic redundancy without layering external systems that complicate failover and often introduce new failure surfaces.

Advancing resilience without disruption ultimately depends on designing for change from early in the planning stages. Systems that integrate replication, automated failover, and rolling upgrades natively, within the platform, can then transition from pilot to mission-critical scale without replatforming. Resilience is not retrofitted after outages, but rather is included as part of the core architecture.

Row-Level Geo-Partitioning for Data Residency

As applications expand across regions, data locality becomes a structural requirement rather than a deployment preference. Organizations need to deliver low-latency access to globally distributed users while also satisfying regulatory requirements. The distributed architecture needs to support deliberate data placement without

fragmenting the logical database or introducing region-specific application stacks.

Row-level geo-partitioning provides granular control over the location of data by enabling specific rows and partitions to be pinned to designated regions. Instead of replicating all records uniformly across every geography, the system enforces placement policies at the data level while maintaining a single logical schema and standard SQL interface. **Figure 3-2** shows row-level data placement. This preserves developer simplicity while ensuring that sensitive or regulated data remains within approved jurisdictions.

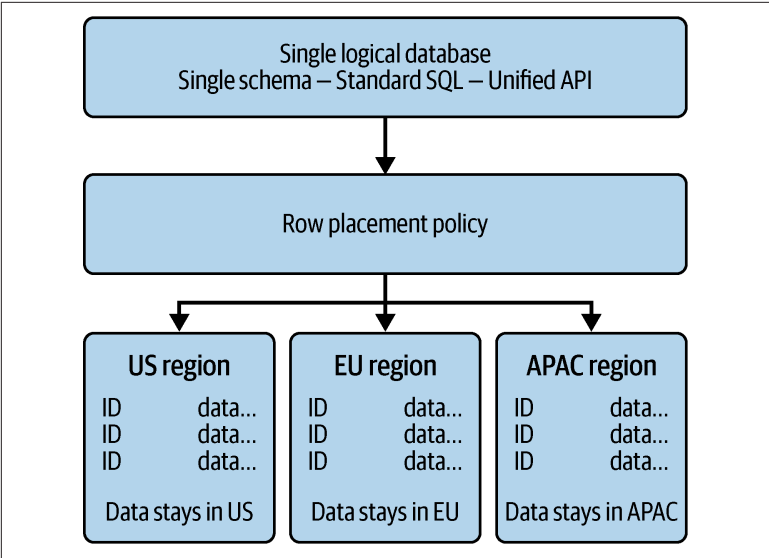


Figure 3-2. Row-level placement of data while preserving a single point of interaction to access the data

Geo-partitioning also reduces architectural fragmentation. By embedding locality and placement rules into the database layer, organizations avoid building custom routing logic or maintaining parallel data stores to satisfy compliance requirements. Data residency, performance optimization, and distributed coordination are handled as native capabilities of the system rather than as external overlays.

Data placement decisions inevitably influence resilience strategy. When leaders and replicas are distributed across multiple regions, the system can tolerate regional outages while continuing to serve

traffic from remaining locations. Geo-partitioning ensures that locality policies remain intact even during failover scenarios, preserving both compliance posture and operational continuity.

Replication strategy further refines the balance between data locality requirements and multi-region fault tolerance. Synchronous replication across designated regions can provide strong consistency and zero data loss for critical workloads, while asynchronous replication supports broader disaster recovery objectives with reduced latency impact. These replication models enable organizations to tune consistency and recovery guarantees according to workload criticality while preserving a unified scheme.

Multi-Cloud, Multi-Region, and Hybrid Deployment Strategies

Modern production systems must assume that infrastructure choices will evolve over time. Organizations rarely operate within a single cloud boundary indefinitely, and many continue to maintain on-premises environments for regulatory, latency, or legacy integration reasons. An AI-ready data architecture needs to support these flexible deployment needs without forcing application rewrites each time the underlying infrastructure changes.

Consistency across environments is essential in such deployments. Application developers should not need to determine whether a query is executing in a public cloud region, a private datacenter, or a secondary cloud provider. Maintaining PostgreSQL compatibility and standard SQL interfaces ensures that applications interact with a consistent data layer even as infrastructure topology evolves beneath it.

Multi-region deployment addresses availability and latency requirements within a single cloud provider, but multi-cloud strategies address a different class of risk. Concentrating critical workloads within one cloud provider introduces exposure to provider-specific outages, pricing changes, and the regional constraints of that provider. A distributed SQL foundation that spans regions and clouds while preserving a single logical database abstraction reduces this dependency and enables infrastructure decisions to remain strategic rather than reactive.

Multi-region topology also plays a central role in latency optimization. Placing data and compute closer to users reduces round-trip times and improves responsiveness for AI-driven features that combine transactional reads, contextual retrieval, and inference. A globally distributed architecture enables leaders and replicas to be strategically positioned across regions while preserving transactional guarantees.

Hybrid deployments often arise from compliance and data residency mandates. Certain data categories may be required to remain within national boundaries or within privately controlled infrastructure. By combining geo-partitioning, replication policies, and distributed coordination, organizations are able to enforce locality requirements while still participating in broader multi-region continuity strategies.

Multi-cloud strategy further strengthens business continuity posture. When clusters can span geographic regions and cloud providers, the system is much better positioned to absorb infrastructure-level disruptions that extend beyond a single availability zone or even an entire region. Distributed consensus and coordinated replication enable workloads to continue operating despite localized failures.

Multi-API Flexibility for Rapidly Changing Environments

Rapidly evolving application stacks demand data platforms that do not constrain interface choices. As organizations introduce new services, AI components, and microservices-based architectures, they frequently need to support different query patterns and developer expectations simultaneously. A data foundation that exposes multiple APIs while preserving transactional integrity enables teams to adapt without redesigning the core system.

Beyond PostgreSQL compatibility, flexibility extends to supporting multiple query paradigms over the same underlying storage engine. A platform should ideally provide both a PostgreSQL-compatible distributed SQL interface and a Cassandra-compatible Cassandra Query Language (CQL) interface, enabling teams to choose relational or document-style access patterns without deploying separate systems. Because both APIs operate on the same distributed storage layer, organizations avoid data duplication while still accommodating diverse workload requirements.

A multi-API approach also accelerates migration and modernization efforts. Support for industry-standard interfaces enables existing PostgreSQL and Cassandra applications to connect without modification, thereby reducing the need for code rewrites. By lowering the barrier to adoption, organizations are able to replace less flexible databases while continuing to fully utilize in-house developer knowledge and existing tooling.

Flexibility is further reinforced through ubiquitous language and driver support across modern development ecosystems. Native drivers and compatibility with established client libraries enable applications written in popular languages to connect using familiar patterns. This ensures that teams can integrate data layer components into their existing stacks, whether building microservices, AI pipelines, or analytics workflows.

Modern AI architectures increasingly require vector search capabilities that coexist with transactional data. Support for the PostgreSQL pgvector extension enables embeddings to be stored and queried directly through SQL. With the pgvector extension, semantic search and RAG are available without deploying a separate vector database. These operations occur within the same PostgreSQL-compatible environment, further enabling vector similarity search to become an extension of the existing API surface rather than a parallel system.

Integration flexibility should also extend to the AI frameworks and natural language interfaces. Compatibility with common frameworks enables external tools to translate natural language into SQL and connect through standard drivers, leveraging metadata exposed by system catalogs. This creates an environment in which AI agents and human developers operate within the same data foundation using established and standardized protocols rather than bespoke integrations.

Operational APIs and infrastructure automation capabilities reinforce this adaptability. Infrastructure-as-code tooling, Kubernetes operators, and declarative provisioning enable database deployments to evolve alongside application releases without introducing drift or manual configuration bottlenecks. As APIs change or services are added, the operational layer can be version-controlled and automated in the same way as application code.

Open source licensing further protects long-term flexibility in rapidly changing environments. With full access to source code and the

ability to self-manage across clouds or on-premises infrastructure, organizations retain the option to move between deployment models without refactoring applications. This architectural portability reduces dependency on any single vendor model while preserving consistent APIs across environments.

Finally, deployment model flexibility ensures that API stability is not constrained by infrastructure choices. The same database engine operates across self-managed clusters, enterprise-managed environments, and fully managed services. This enables organizations to shift operational responsibility as requirements evolve. Applications do not require modification when transitioning between development, testing, production, hybrid, or multi-cloud environments because the underlying APIs remain the same.

Security, Observability, Explainability, and AI-Enhanced Database Operations

AI increases both the value and the potential operational impact of the data platform. As organizations embed models into customer-facing systems and automated decision pipelines, the database becomes the control plane for accuracy, performance, and governance. Reliability and security are no longer background concerns but primary determinants of whether AI initiatives succeed in production.

This chapter examines how AI-enhanced database operations can be integrated with modern distributed SQL platforms. It discusses how database administration can be agentic through the use of telemetry and other existing characteristics of the platform. This chapter also analyzes how access control, encryption, and identity integration provide the secure and observable backbone required for intelligent, policy-aware AI-enhanced operations at enterprise scale.

Security, Observability, and Explainability Requirements

Production AI systems introduce new operational and governance pressures that extend beyond performance and scale. When AI agents, APIs, and human users all access the same operational data, the database becomes both a control point and an accountability

surface. Security, monitoring, and visibility into system behavior are therefore not ancillary features but architectural requirements.

As AI workloads query structured data, generate derived outputs, and automate decision flows, organizations need to ensure that access controls are enforced consistently, activity is traceable, and system behavior remains transparent under load. Governance, monitoring, and operational safeguards are core pillars of a modern distributed data platform and not optional enterprise add-ons.

Traceability is what turns agent behavior from a black box into an auditable system, connecting each decision back to the conversation that prompted it, the plan the agent formed, and the data it used along the way. Without this end-to-end record, teams cannot debug failures, regulators cannot verify compliance, and organizations cannot trust agents to operate in production.

Security

Security begins with precise control over data access. Having the database aligned with standard RBAC models enables enforcement of permissions down to row level. This approach ensures that applications, analysts, and AI service accounts only access the data explicitly granted to them.

Row-level security and column-level controls should be available. These controls further narrow exposure of sensitive data by enabling access restriction based on geography, department, or function while also protecting regulated fields such as those containing personally identifiable information (PII). For AI-driven workloads that query data programmatically, dedicated roles with minimal privileges and full audit logging provide guardrails around automated access.

At the architectural level, strong consistency and distributed replication reinforce security and data integrity. Data is automatically sharded and replicated across nodes, and writes are committed through consensus before acknowledgement, eliminating single points of failure and preventing split-brain inconsistencies. This design reduces risk of data corruption or inconsistency during node, zone, or regional failures.

Another dimension that is sometimes overlooked is the control enabled by open source licensing. Full access to source code reduces vendor lock-in risk and supports compliance reviews. You even retain the ability to modify the source code should the need arise.

Observability

In distributed AI systems, visibility into system behavior is essential. Built-in metrics, slow query logging, and compatibility with standard monitoring systems enable performance tracking, query latency, and operational health. Integration with widely adopted monitoring stacks ensures that database telemetry can be incorporated directly into existing dashboards and alerting pipelines.

Support for OpenTelemetry-based tracing and export to external observability platforms enables end-to-end visibility across distributed components. Metrics can be pushed to observability backends where they are stored as time-series data, thereby supporting dashboards, service-level objectives, and proactive alerting and response. In AI-enhanced applications, where latency or replication lag can directly affect user experience, this operational transparency is critical.

Operational safeguards extend beyond monitoring and into recoverability. Continuous snapshots and transaction log retention enable point-in-time recovery. This protects against logical corruption, accidental deletion, and faulty deployments. Recovery point objectives approach zero for synchronous replicas, and recovery time objectives are measured in seconds for zone-level failovers, ensuring minimal disruption to production systems.

Explainability

Explainability in AI-driven systems depends on consistent definitions of data and metrics. The platform should not impose a built-in semantic layer, but compatibility should enable integration with tools that define metrics, dimensions, and business logic as version-controlled SQL transformations. These semantic models ensure that applications, analysts, and AI agents query against standardized definitions rather than ad hoc calculations.

Materialized views and database-native security controls help enforce consistency and policy alignment. Precomputed aggregations can standardize complex metrics, while row-level security ensures that query results reflect authorized data scopes. This combination supports governed access patterns for both human users and AI agents operating through natural language or programmatic interfaces.

Trustworthy outputs ultimately depend on trustworthy data states. Distributed consensus ensures that every committed write is replicated to a majority of nodes before acknowledgement, providing strong consistency guarantees across regions. When AI applications retrieve data for decision-making or generative tasks, they operate on a foundation designed to prevent data loss, stale reads, and inconsistency under failure conditions.

Agentic Database Administration and Performance Tuning

Performance tuning in a distributed environment also depends on understanding data locality and workload distribution. Having these elements managed in an agentic manner means making sure that the platform tracks how data is placed across regions and how requests are routed to replicas. Intelligent systems can then recommend adjustments to indexing strategies, schema design, and replica placement to improve latency and throughput while preserving transactional guarantees.

Simplified Migration and Operational Tooling

Modernization efforts succeed when migration does not require rewriting the application stack. PostgreSQL compatibility enables existing drivers and operational practices to continue functioning with minimal modification. Teams can then use familiar SQL semantics and tooling while transitioning to a distributed architecture.

This compatibility extends to ecosystem integrations that organizations already rely on for deployment and monitoring. Established processes for backup, schema management, and continuous integration (CI) pipelines can remain largely intact. The result is a migration path that reduces disruption and lowers the operational risk that might otherwise be associated with platform changes.

Operational tooling should be designed to support lifecycle management across distributed clusters. Capabilities such as rolling upgrades and automated replication management enable maintenance activities to occur without extended downtime. Administrators can perform version updates and infrastructure changes while maintaining availability guarantees.

Backup and restore processes should be integrated into the operational model of the database platform. Administrators can create consistent backups and restore data when necessary without complex external orchestration. These capabilities should support both routine operational hygiene and more formal disaster recovery planning.

Conclusions and Making It Work

Creating proof-of-concept AI systems is largely a solved problem. Proof-of-concept AI systems don't require the same architectural decisions as production systems do around scaling because the goals are simply to test the system from end-to-end, with a focus on the main path through the system. When AI reaches production, scaling becomes an issue, and real-world use cases have real-world consequences as the live workload becomes crucial to ensuring accuracy.

Building a solid data layer is vital because that layer becomes the foundation of any AI deployment. The data layer needs flexibility around scaling architectures while also supporting the often unique requirements around accessing data for AI. A data layer that enables modern deployment strategies and the ability to use traditional and established protocols for access makes development seamless. This chapter examines several important factors that help an organization move from prototype to production with an AI deployment.

From Prototype to Production-Ready AI

AI initiatives often begin as narrowly scoped prototypes. Early success is typically measured in relevance metrics or model quality rather than durability, fault tolerance, or operational rigor. The transition from proof of concept to production shifts the center of gravity from model experimentation to data architecture, where consistency and scale matter.

Production AI systems need to support continuous ingestion, transactional updates, concurrent queries, and strict latency requirements while maintaining correctness across distributed nodes. A distributed SQL system should be able to combine horizontal scalability with strong consistency, enabling teams to move from development to production without rewriting application logic. This application-level continuity is what prevents AI projects from stalling when real workloads replace synthetic benchmarks.

A production-ready AI application also requires predictable behavior under failure. Properties and features like synchronous replication, automatic failover, and survivable node outages should be core capabilities rather than optional extensions. For AI workloads used to drive customer-facing experiences, this level of native, built-in resilience is essential to maintaining trust and service continuity.

The shift to production further demands operational transparency. Observability metrics should be built in, enabling teams to understand how vector search, transactional queries, and analytical scans interact in real time. Without this level of insight, scaling AI systems becomes more guesswork than engineering.

Security requirements also evolve at the production stage. Encryption in transit and at rest should have already been built as part of the prototype, though sometimes access controls are not applied the same way for development and production. As AI systems increasingly operate on sensitive enterprise data, these controls become foundational to regulatory compliance and security best practices.

Scaling in Practice: Billion Vector Deployments

Vector search moves from experimental to mission-critical when embeddings are generated at scale and queried under real user traffic. The SQL platform should have native support for vector data types and have indexing integrated directly into the SQL layer. When the SQL platform has these features, it eliminates the need to maintain a separate operational database and vector store. This unified approach simplifies architecture while preserving transactional guarantees.

Scaling to billions of vectors requires more than raw storage capacity. Distributed execution, automatic data sharding, and rebalancing across nodes are core mechanisms that enable horizontal growth without application redesign. As data volumes expand, the

cluster distributes both storage and query processing across nodes to maintain performance.

Indexing strategies are also important at scale. Approximate nearest neighbor indexing can be tuned to balance recall and latency. These configuration controls enable teams to optimize search behavior based on application requirements rather than accepting fixed trade-offs imposed by an external service.

Operational durability remains central even at a massive scale. A consensus-based replication model ensures that vector and relational data remain consistent across replicas, supporting both transactional updates and analytical queries. This consistency model prevents divergence between metadata and embeddings, which is a common risk in loosely coupled architectures.

The result of these scaling efforts is an architecture capable of supporting billion-vector workloads within the same distributed SQL environment that manages application state. Instead of stitching together multiple specialized systems, organizations can scale AI workloads using the same operational and governance controls already in place in their core data stores.

Selecting Flexible Frameworks over Cutting-Edge Components

Early-stage AI development often gravitates toward the newest model or the most specialized vector engine. Production systems, by contrast, benefit from architectural flexibility and adherence to open standards. For example, the use of standard SQL interfaces and compatibility with open platforms like PostgreSQL are strategic choices that protect long-term application portability.

Open standards enable integration with a broad ecosystem of frameworks, drivers, and tools. A platform that presents as a PostgreSQL-compatible endpoint enables developers to use existing ORMs, analytics tools, and data pipelines without custom-built connectors. This compatibility reduces operational friction and shortens time to deployment.

Pluggable vector capabilities further reinforce this flexibility. Rather than forcing a closed ecosystem, the database platform should enable the ability to integrate with external model providers and AI

frameworks. Having this flexibility allows organizations to execute their AI strategy without restructuring due to data layer constraints.

Cloud-agnostic deployments are also important in a modern enterprise infrastructure. By avoiding tight coupling to a single cloud provider or proprietary stack, organizations retain control over cost, residency, and constraints that might otherwise be present with a single cloud deployment. AI workloads can fluctuate significantly, and having flexibility at the data and infrastructure layers helps to scale up and down, in real time.

Next Steps for Implementing Distributed SQL for AI

Moving toward a distributed SQL foundation begins with a clear assessment of workload characteristics and existing state. Evaluation of transactional requirements, expected data growth, and latency targets informs cluster sizing and topology decisions. Early architectural alignment prevents costly redesigns later.

Deployment models should be selected based on operational and compliance requirements. Managed cloud, self-managed clusters, and hybrid configurations should be considered based on needs. Organizations can align infrastructure with governance requirements and can migrate to a flexible platform in phases if needed, so that the application layer is supported by a robust and future-proof, AI-enabled data foundation.

Schema design and data modeling for AI applications should account for both relational metadata and vector embeddings. Ideally, embeddings can be stored alongside structured data within the same schema. This enables hybrid queries that combine filtering and semantic search. Designing for this coexistence from the outset simplifies application logic.

Operational readiness must also include monitoring, backup, and disaster recovery planning. Built-in observability tooling and automated replication features are intended to support continuous operations even during node failures and maintenance windows. Establishing these practices early strengthens reliability as workloads scale.

Finally, organizations should pilot new applications in a production-like manner rather than relying solely on synthetic workloads. Realistic concurrency and data volumes help to validate performance and resilience. By grounding evaluation in real-world conditions, teams can confidently move from prototype to production, and make the data layer work as part of the AI-enabled organization.

About the Author

Steve Suehring is an associate professor of computing at the University of Wisconsin-Stevens Point, where he teaches courses on a variety of topics from development to networking to cybersecurity. Prior to joining the faculty, Steve worked in several roles, including as a technical architect, systems engineer, and data security analyst. Steve was an editor for *LinuxWorld Magazine* and has written several technology books.