

Beyond the AI Hype

Why Large Language Models Will Transform Software—Not Replace It

Version 1.0 (Publication Edition)

Antonio Latorre



Abstract

Large Language Models (LLMs) have transformed software development, knowledge work, and human-computer interaction by allowing users to express intent in natural language. This paper argues that the dominant claim that AI will replace enterprise software rests on a category error: Large Language Models (LLMs) are probabilistic systems optimized for plausible language generation, whereas enterprise software is required to provide deterministic execution, reproducibility, governance, auditability, security, and mathematically reliable computation. The paper makes three original contributions. First, it proposes a theory of Three Generations of Computing: Deterministic Computing, Probabilistic Computing, and Intent-Driven Computing (IDC). Second, it defines the AI Competence Mirage and the Computational Trust Gap as explanatory concepts for the mismatch between perceived linguistic competence and computational guarantees. Third, it proposes the architectural model for next-generation enterprise systems: a Probabilistic Interface Layer (PIL) that interprets human intent, a Domain Intelligence Layer (DIL) that encapsulates domain knowledge and business methodology, and a Deterministic Execution Layer (DEL) that performs trusted computation. Economically, the paper argues that AI does not eliminate enterprise software; instead, it lowers the cost of software creation while increasing the strategic importance of domain intelligence, architecture, governance, and operational trust.

Keywords: Large Language Models (LLMs); Intent-Driven Computing (IDC); enterprise architecture; Probabilistic Interface Layer (PIL); Domain Intelligence Layer (DIL); Deterministic Execution Layer (DEL); Computational Trust Gap; AI Competence Mirage.

Table of Contents

1. Introduction
2. Three Generations of Computing
3. How Large Language Models Work
4. Truth versus Plausibility
5. The Computational Trust Gap
6. The Great Misunderstanding: AI Is Not Replacing Software
7. The Next Generation Enterprise Architecture
8. Case Study: Enterprise Procure-to-Pay and Invoice Approval
9. Related Work
10. Limitations
11. Conclusion
12. Proposed Principles
13. References

1. Introduction

The emergence of Large Language Models (LLMs) represents one of the most significant technological advances in the history of computing. In only a few years, systems such as ChatGPT, Claude, Gemini and others have demonstrated an unprecedented ability to understand natural language, generate software, summarize complex information, assist scientific research and interact with humans in ways that, until recently, appeared unattainable.

The speed of this technological progress has generated enormous enthusiasm across virtually every industry. At the same time, it has also created an equally unprecedented level of misunderstanding regarding the actual role of Artificial Intelligence within modern software systems.

A growing number of articles, conference presentations and public discussions argue that enterprise software will soon become obsolete because users will simply "ask an AI" to perform any task they require. According to this narrative, conversational interfaces will gradually replace traditional applications, dramatically reducing the need for specialized software platforms.

This paper argues that such conclusions arise from a fundamental confusion between two computational paradigms that, although closely related, solve fundamentally different problems.

Traditional software systems are deterministic computational engines. Given the same inputs and the same internal state, they are expected to produce exactly the same outputs, every time. This determinism constitutes the foundation upon which financial systems, medical devices, telecommunications networks, industrial automation, transportation systems and virtually every mission-critical enterprise application have been built for decades.

Large Language Models, by contrast, belong to a fundamentally different computational category. They are probabilistic models trained to estimate the most likely continuation of a sequence of tokens given a particular context. Their remarkable capabilities emerge not from executing predefined algorithms that guarantee mathematically correct results, but from learning statistical representations of language and knowledge from enormous quantities of data.

This distinction is considerably more important than it may initially appear.

Many public discussions surrounding Artificial Intelligence implicitly assume that because LLMs often produce correct, coherent and highly sophisticated responses, they therefore operate as deterministic reasoning engines capable of replacing traditional software. While this assumption is understandable, it does not accurately reflect the computational principles upon which these systems are built.

In reality, Large Language Models (LLMs) and deterministic software exhibit complementary strengths. LLMs excel at interpreting ambiguous human intentions, communicating in natural language, generating code, synthesizing knowledge and assisting complex cognitive tasks. Deterministic software excels at executing formally defined algorithms, enforcing business rules, guaranteeing reproducibility, maintaining transactional consistency, satisfying regulatory requirements and producing mathematically verifiable results [31], [32], [36]–[40].

Consequently, this paper proposes that the future of enterprise computing should not be understood as a transition from software to Artificial Intelligence.

Instead, it represents the emergence of a new computational architecture in which probabilistic intelligence becomes the primary interface between humans and deterministic computational systems.

Within this architecture, Artificial Intelligence is responsible for understanding intent, while deterministic software remains responsible for executing business logic, mathematical computation, optimization, governance, auditing, security and reproducibility [22]–[30].

This architectural separation allows organizations to simultaneously benefit from the flexibility of natural language interaction and the reliability required by mission-critical enterprise systems.

The implications of this distinction extend far beyond software engineering.

They fundamentally reshape how organizations should think about enterprise applications, software development, human-computer interaction, governance, regulation and competitive advantage in the age of Artificial Intelligence.

The central thesis of this paper can therefore be summarized as follows:

Large Language Models do not represent the end of deterministic software. They represent the emergence of a new computational paradigm in which probabilistic intelligence orchestrates deterministic execution.

This thesis challenges the increasingly common perception that conversational AI will eliminate the need for enterprise software. Instead, it suggests that software is entering a new evolutionary stage: one in which language becomes the universal interface, while deterministic computation remains the foundation upon which trust, governance and business execution continue to depend.

2. Three Generations of Computing

Throughout its history, computing has evolved by progressively reducing the cognitive effort required for humans to interact with machines. Each major technological revolution has not replaced the previous generation, but rather introduced a new abstraction layer that enabled increasingly natural forms of communication between humans and computers.

From this perspective, the emergence of Large Language Models should not be understood merely as another advance in Artificial Intelligence. Instead, it represents the beginning of a third computational paradigm, fundamentally changing the role of software within enterprise systems.

We propose that modern computing can be understood as the evolution of three distinct generations: **Deterministic Computing**, **Probabilistic Computing**, and **Intent-Driven Computing**.

Although these generations coexist today, each introduces a fundamentally different relationship between humans, software and computation.

First Generation: Deterministic Computing

The first generation encompasses virtually all traditional software developed since the birth of digital computing.

Its defining characteristic is determinism.

Given identical inputs and identical system states, the software is expected to produce identical outputs every time it executes. Predictability, repeatability and mathematical correctness constitute the foundations of this paradigm.

Within deterministic computing, every possible behavior must be explicitly designed and implemented by software engineers. Business rules are encoded as algorithms, workflows are predefined, and every user interaction follows a sequence anticipated during development.

Communication between humans and machines occurs through rigid interfaces, including command-line instructions, graphical menus, forms, buttons, APIs and predefined workflows.

In this paradigm, software developers act as translators between human intentions and machine execution. Every new business requirement generally requires new software development.

This approach has enabled the construction of virtually every critical digital system in modern society, including financial platforms, telecommunications networks, medical devices, enterprise resource planning systems, industrial automation, transportation systems and scientific computing infrastructures.

Despite its enormous success, deterministic computing exhibits an important limitation.

Humans must adapt their communication to the language of software rather than software adapting to the language of humans.

Second Generation: Probabilistic Computing

The emergence of Large Language Models introduced a fundamentally different computational capability.

Instead of requiring humans to communicate through predefined interfaces, computers became capable of interpreting natural language directly.

This represents a profound shift.

For the first time, software users no longer need to understand menus, APIs, programming languages or complex workflows in order to accomplish sophisticated tasks.

Instead, they simply describe their intentions using ordinary human language.

This capability is enabled by probabilistic neural networks trained to estimate the conditional probability of linguistic sequences [1], [9]–[13].

Unlike deterministic software, these systems do not execute explicitly programmed reasoning procedures.

Rather, they generate responses by predicting the most probable continuation of a sequence of tokens given the surrounding context.

As a consequence, their outputs are inherently probabilistic.

For highly constrained questions, the probability distribution becomes extremely concentrated, making responses appear deterministic.

For open-ended problems, however, multiple valid responses naturally emerge because numerous continuations may exhibit similarly high probabilities.

This probabilistic nature provides extraordinary flexibility.

It allows the system to interpret ambiguous instructions, generate software, summarize documents, translate languages, explain concepts and assist creative work.

However, it also introduces an important limitation.

Because outputs are generated probabilistically rather than algorithmically, correctness cannot be guaranteed solely by the language model itself.

The model optimizes plausibility rather than formal correctness [12], [16]–[18].

Consequently, probabilistic computing dramatically expands human-computer interaction while simultaneously requiring new mechanisms to guarantee reliability whenever deterministic outcomes are essential.

Third Generation: Intent-Driven Computing

We argue that the next stage of computing is not the replacement of deterministic software by probabilistic Artificial Intelligence.

Instead, it is the integration of both paradigms into a unified computational architecture.

We refer to this emerging paradigm as **Intent-Driven Computing (IDC)**.

Definition 1: Intent-Driven Computing (IDC). Intent-Driven Computing (IDC) is an enterprise-computing paradigm in which human goals are expressed primarily as natural-language intent, interpreted by probabilistic language systems, constrained by explicit domain intelligence, and executed through deterministic computational mechanisms that remain auditable, reproducible and governable.

Within Intent-Driven Computing, natural language becomes the primary interface through which humans express objectives, constraints and preferences.

Large Language Models are responsible for interpreting these intentions, resolving linguistic ambiguity and transforming human requests into structured computational instructions.

Execution, however, remains the responsibility of deterministic software systems.

Business rules, optimization algorithms, numerical computation, transactional consistency, governance, auditing, security and regulatory compliance continue to be implemented by deterministic execution engines whose behavior is mathematically reproducible.

This separation fundamentally changes the role of Artificial Intelligence.

Rather than replacing software, AI becomes an orchestration layer positioned between human cognition and deterministic computation.

Its purpose is no longer to execute business processes directly, but to translate human intentions into deterministic computational workflows.

This architecture combines the principal strengths of both paradigms while adding a dedicated Domain Intelligence Layer (DIL) for business knowledge, methodology, and governance.

Probabilistic intelligence contributes flexibility, natural interaction and contextual understanding.

Deterministic software contributes reproducibility, reliability, explainability, governance and computational certainty [22]–[30].

Together, these complementary capabilities enable enterprise systems that are simultaneously intuitive for humans and trustworthy for organizations.

The Evolution of Human-Computer Interaction

Viewed from a historical perspective, these three generations reveal a common direction.

Each generation reduces the amount of technical knowledge required for humans to interact with computers.

The progression can be summarized as follows:

As shown in Table 1, this progression moves from structured human-machine interaction toward Intent-Driven Computing (IDC).

Table 1. Three Generations of Computing. The table summarizes the evolution from deterministic interaction to probabilistic language generation and Intent-Driven Computing (IDC).

Generation	Human Communication	Machine Execution	Primary Characteristic
Deterministic Computing	Commands, forms, menus, APIs	Deterministic algorithms	Predictability
Probabilistic Computing	Natural language	Probabilistic language generation	Flexibility
Intent-Driven Computing	Natural language intentions	Deterministic execution orchestrated by AI	Trustworthy autonomy

This progression suggests that Artificial Intelligence should not be viewed as a substitute for software engineering.

Instead, it represents the evolution of the human interface through which deterministic software systems are instructed.

Accordingly, the future of enterprise computing is unlikely to consist of conversational agents replacing enterprise applications.

Rather, enterprise applications themselves will increasingly incorporate probabilistic intelligence as their primary interface while preserving deterministic computational cores responsible for execution, governance and trust.

Figure 4 provides the visual overview for the trusted-enterprise-results argument developed in this section.

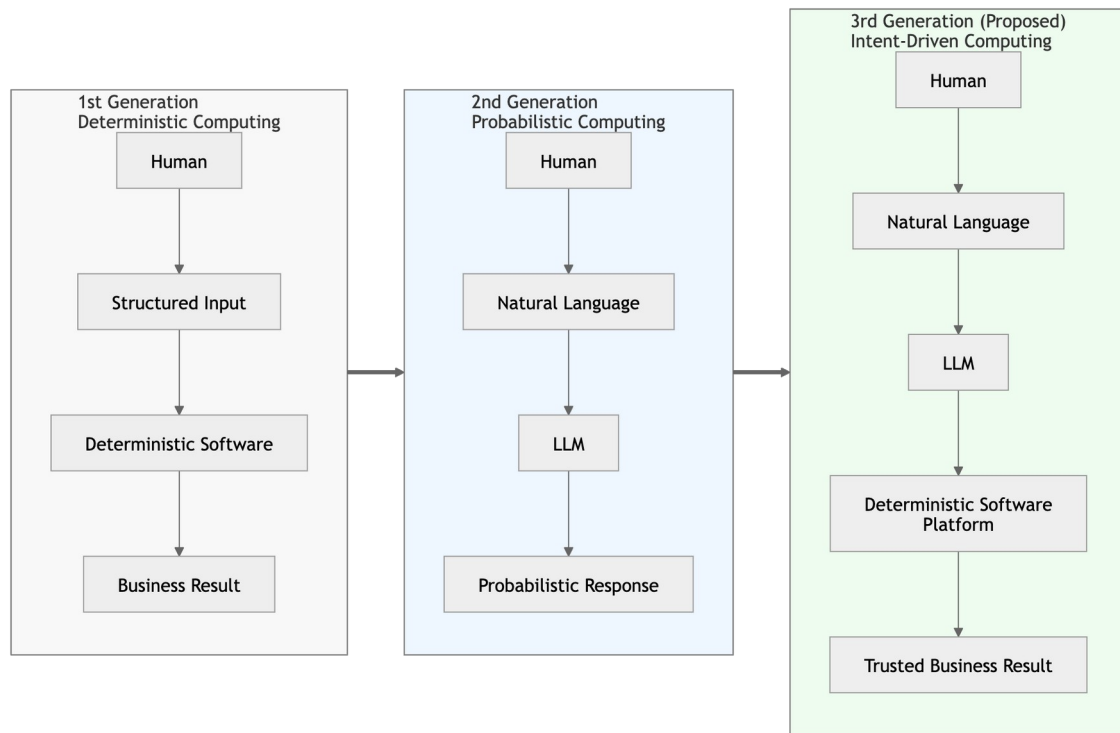


Figure 4. Architectural intuition for trusted enterprise results. The current architecture formalized in this paper is PIL -> DIL -> DEL; the figure preserves the visual intuition that AI-mediated intent must ultimately connect to governed execution.

3. How Large Language Models Work

Understanding the computational nature of Large Language Models is essential to understanding both their extraordinary capabilities and their inherent limitations.

Much of the public discussion surrounding Artificial Intelligence treats LLMs as if they were reasoning engines capable of computing objective truths about the world. While their responses frequently exhibit remarkable coherence and sophistication, the underlying computational mechanism is fundamentally different from that of traditional deterministic software.

At their core, Large Language Models are probabilistic neural networks trained to estimate the conditional probability distribution of the next token given all previously observed tokens [1], [9]–[13].

Formally, given a sequence of tokens

$$x_1, x_2, \dots, x_n$$

the model attempts to estimate

$$P(x_{n+1}, /, x_1, x_2 \dots x_n)$$

In other words, the model continuously answers a single mathematical question:

Given everything that has been written so far, what is the most probable next token? [1], [10]–[13]

Every generated word, punctuation mark and symbol is produced by repeating this prediction process thousands of times until the complete response has been generated.

Importantly, the model does not retrieve predefined answers from a database, nor does it execute symbolic reasoning procedures similar to those implemented in classical software systems.

Instead, each prediction results from evaluating billions of learned parameters that encode statistical relationships acquired during training.

The remarkable capabilities of modern LLMs emerge from the scale of these learned representations rather than from explicitly programmed reasoning rules.

Language Modeling Is a Probabilistic Process

For every position in the generated response, the model computes a probability distribution over every token in its vocabulary.

Conceptually, suppose the user asks:

What is the capital of France?

Immediately before generating the final word, the model internally evaluates many possible continuations.

Illustratively, the distribution might resemble:

As shown in Table 2, a low-entropy case can make a probabilistic model appear deterministic because one continuation dominates the distribution.

Table 2. Low-Entropy Token Distribution. The example illustrates why simple factual prompts may appear deterministic even when generated by a probabilistic model.

Candidate Token	Estimated Probability
Paris	99.9999%
Lyon	0.00008%
Marseille	0.00002%
London	<0.00001%

These numbers are purely illustrative.

The important observation is not the specific probabilities, but the existence of a probability distribution itself.

The model does not, by its basic language-modeling objective, independently verify whether "Paris" is objectively true.

Instead, it estimates that "Paris" is overwhelmingly more probable than any alternative given the context of the conversation and the statistical structure learned during training.

This distinction may appear subtle, but it fundamentally changes how AI systems should be interpreted.

Prediction Rather Than Verification

Traditional software frequently computes results by executing formally specified algorithms.

A banking system calculating compound interest applies mathematical equations.

A database executing an SQL query follows deterministic execution plans.

A sorting algorithm produces exactly the same ordering every time it receives identical inputs.

Large Language Models operate differently.

They generate outputs by predicting plausible linguistic continuations rather than verifying propositions against formal mathematical models.

Consequently, the correctness of a generated response is an emergent consequence of successful statistical learning—not a mathematical guarantee provided by the computational process itself.

This distinction explains one of the most misunderstood characteristics of modern AI.

A Large Language Model can generate an answer that is entirely correct.

It can also generate an answer that is partially correct.

Or one that is completely incorrect while remaining grammatically coherent and highly convincing.

The computational mechanism remains exactly the same in all three cases.

Why Simple Questions Often Produce Stable Answers

A common objection arises immediately.

If LLMs are probabilistic, why do they almost always answer:

Paris

when asked about the capital of France?

The answer lies in the shape of the underlying probability distribution.

For highly constrained factual questions, one candidate token often receives an overwhelmingly higher probability than every alternative.

Although the generation process remains probabilistic, the distribution is so heavily concentrated that repeated executions frequently produce identical outputs.

This apparent determinism should not be confused with true deterministic computation.

The model remains probabilistic.

The probability distribution simply exhibits extremely low entropy.

As prompts become broader, more creative or less constrained, probability mass naturally spreads across many possible continuations.

Consequently, multiple different responses become equally plausible.

The following prompt illustrates this phenomenon:

Design a machine learning pipeline for credit risk prediction.

There is no unique continuation.

Different executions may legitimately recommend different algorithms, validation strategies, feature engineering techniques, hyperparameter optimization procedures or deployment architectures.

Each response may be technically valid.

None is uniquely determined by the prompt alone.

The Consequences for Enterprise Systems

This probabilistic behavior constitutes one of the greatest strengths of Large Language Models.

It enables flexibility, creativity, contextual understanding and natural language interaction that deterministic software cannot easily achieve.

However, the same characteristic also limits their ability to serve as the sole computational engine for enterprise systems requiring:

reproducibility, [22]–[30]

- auditing,
- governance,
- regulatory compliance,
- transactional consistency,
- mathematical correctness,
- deterministic execution.

Organizations operating critical systems cannot rely exclusively on computational processes whose outputs may legitimately vary across executions.

Consequently, LLMs should be understood not as replacements for deterministic computation, but as probabilistic interpreters capable of translating human intentions into computational workflows subsequently executed by deterministic software.

This architectural distinction forms the conceptual foundation of the next generation of enterprise computing proposed throughout this paper.

Large Language Models Do Not Decide — They Sample

One of the most widespread misconceptions about modern Artificial Intelligence is the belief that Large Language Models *decide* what to answer.

From a computational perspective, this is not what actually happens.

An LLM does not internally conclude that one answer is objectively correct while rejecting all alternatives.

Instead, after processing the input context, the model produces a probability distribution over every possible next token in its vocabulary [1], [10]–[13].

Conceptually, the model may estimate something similar to:

As shown in Table 3, token probability explains plausibility but does not itself constitute a formal truth guarantee [12], [16]–[18].

Table 3. Candidate Token Probabilities. The illustrative distribution shows how plausibility can concentrate around a correct answer without becoming a formal truth guarantee.

Candidate	Estimated Probability
Paris	99.9999%
Lyon	0.00008%
Marseille	0.00002%
London	<0.00001%

The model itself stops there.

It does not "choose Paris."

It simply computes that Paris has the highest conditional probability given the context [1], [9]–[13].

A separate decoding procedure is then responsible for producing the next token [1], [10]–[13].

Depending on the decoding strategy, the system may:

- always select the highest-probability token (greedy decoding);
- search across multiple candidate sequences (beam search);
- randomly sample from the probability distribution (sampling);
- restrict sampling to the most probable candidates (top-k);
- sample only from the smallest probability mass exceeding a threshold (top-p / nucleus sampling);
- adjust diversity through a temperature parameter.

These decoding strategies do not change the underlying knowledge learned by the model.

They merely determine how the probability distribution is converted into observable text.

This distinction is essential because it reveals that the apparent certainty of an LLM response should not be confused with deterministic computation.

Even when two executions produce identical outputs, the computational process remains fundamentally probabilistic.

The model did not execute a mathematical proof.

It estimated probabilities.

Why Sampling Matters

This distinction has profound practical implications.

Suppose an organization asks an LLM:

Design a credit scoring pipeline.

There is no single mathematically correct continuation.

The model assigns meaningful probability to many technically valid alternatives.

One execution may recommend LightGBM.

Another may recommend CatBoost.

A third may prioritize explainability and propose Logistic Regression.

All three responses may be technically defensible.

The model is not changing its opinion.

It is generating different samples from a probability landscape containing multiple plausible continuations.

This flexibility constitutes one of the greatest strengths of modern AI.

It enables creativity, adaptation and contextual reasoning.

However, it also explains why enterprise systems requiring reproducibility cannot delegate final execution entirely to probabilistic models [22]–[30].

From Sampling to Deterministic Execution

The consequence is not that LLMs should be avoided.

Quite the opposite.

Their ability to interpret ambiguous human intentions is unprecedented.

The challenge lies in assigning them the appropriate computational role.

In enterprise environments, the LLM should function primarily as an **intent interpreter** rather than as the final execution engine.

Once the user's objective has been understood, deterministic software components should execute the corresponding algorithms, business rules, optimization procedures and governance mechanisms [31], [32], [36]–[40].

In this architecture, the Probabilistic Interface Layer (PIL) determines what the user wants, the Domain Intelligence Layer (DIL) determines how the domain problem should be solved, and the Deterministic Execution Layer (DEL) executes the result.

Deterministic software executes the task through reproducible, auditable and governed computational processes.

This separation combines the strengths of both paradigms while minimizing their respective weaknesses.

4. Truth versus Plausibility

One of the most important misconceptions surrounding Large Language Models is the assumption that they are optimized to discover truth. They are not. They are optimized to generate plausible language.

This distinction lies at the heart of virtually every misunderstanding regarding modern Artificial Intelligence.

When interacting with an LLM, users naturally evaluate the quality of an answer according to its factual correctness.

The model, however, was never trained with factual correctness as its primary optimization objective.

Instead, during pre-training, the objective is considerably different.

The model continuously adjusts billions of internal parameters in order to maximize the probability of correctly predicting the next token given its surrounding context [1], [10]–[13].

In other words, the optimization target is not:

"Produce the true answer."

It is:

"Produce the continuation that is statistically most consistent with everything observed so far."

This distinction is subtle but profound.

Truth and plausibility frequently overlap [12], [16]–[18].

They are not the same concept.

When Plausibility Equals Truth

For many questions, the statistically most probable continuation also happens to be objectively correct.

For example,

What is the capital of France?

Because "Paris" overwhelmingly dominates the statistical distribution learned during training, the model produces the correct answer with remarkable consistency.

Similarly,

Water freezes at what temperature?

Who wrote Hamlet?

What is the square root of 81?

These questions typically produce correct responses because human knowledge exhibits overwhelming statistical agreement.

In these situations, plausibility becomes an excellent approximation of truth [12], [16]–[18].

This explains why modern LLMs appear astonishingly accurate across many domains.

When Plausibility Diverges from Truth

The distinction becomes visible whenever statistical plausibility no longer uniquely identifies factual correctness [12], [16]–[18].

Examples include:

- ambiguous questions;
- incomplete information;
- rare events;
- novel situations;
- fabricated entities;
- conflicting sources;
- highly specialized domains.

Consider the following prompt:

Describe the contributions of Professor John Smith to quantum finance.

If no such researcher exists, the model frequently produces a convincing biography rather than acknowledging uncertainty.

Importantly, the model is not intentionally fabricating information.

It is doing exactly what it was trained to do.

It generates the continuation that appears most statistically plausible given the linguistic patterns it has learned.

From the perspective of the optimization objective, this behavior is entirely consistent.

From the perspective of factual correctness, it constitutes a hallucination.

Hallucinations are therefore not software defects [17], [18].

They are expected consequences of optimizing plausibility instead of truth [12], [16]–[18].

Why Humans Confuse Plausibility with Truth

Human cognition evolved under conditions where fluent language generally implied competence [17], [19]–[21].

For most of human history, coherent explanations came almost exclusively from other humans possessing relevant knowledge.

As a consequence, our brains naturally associate characteristics such as:

- grammatical correctness;
- logical structure;
- confidence;
- technical vocabulary;
- coherent narrative.

with truthfulness.

Large Language Models exploit exactly these statistical regularities.

Their outputs possess all the linguistic characteristics humans instinctively associate with expertise.

Consequently, users frequently assign a degree of credibility that exceeds what the underlying computational process actually guarantees.

This psychological phenomenon contributes significantly to the current public perception that LLMs "know" or "understand" facts in the same manner as human experts.

In reality, they generate language that strongly resembles expert communication.

Those are not equivalent concepts.

Why This Distinction Matters

In everyday applications, maximizing plausibility is often sufficient [12], [16]–[18].

Drafting emails.

Summarizing reports.

Generating software prototypes.

Brainstorming ideas.

Explaining concepts.

Creative writing.

In these domains, multiple acceptable outputs naturally exist.

Enterprise systems, however, frequently operate under fundamentally different constraints.

Financial transactions.

Medical diagnosis.

Tax calculation.

Credit approval.

Insurance pricing.

Regulatory reporting.

Industrial automation.

In these contexts, there may exist only one acceptable computational outcome.

Plausibility is insufficient.

Deterministic correctness becomes mandatory.

This distinction explains why LLMs should not replace deterministic computational engines in mission-critical systems.

Instead, they should complement them.

Artificial Intelligence excels at interpreting intent, reducing ambiguity and assisting human reasoning.

Deterministic software excels at producing reliable, reproducible and verifiable execution.

Understanding where plausibility ends and determinism becomes necessary constitutes one of the central architectural principles proposed throughout this paper [12], [16]–[18].

5. The Computational Trust Gap

One of the most significant consequences of Large Language Models is not purely computational.

It is cognitive.

Modern LLMs generate language with an extraordinary level of fluency, coherence and confidence. Their responses closely resemble those produced by highly knowledgeable human experts. As a result, users naturally attribute a level of competence and certainty that frequently exceeds what the underlying computational process is actually capable of guaranteeing.

We refer to this psychological phenomenon as the **AI Competence Mirage**.

Definition 2: AI Competence Mirage. The AI Competence Mirage is the cognitive tendency of human users to infer expertise, truthfulness or operational reliability from fluent, coherent and confident AI-generated language, even when the system's output is produced by probabilistic generation and is not accompanied by an independent computational guarantee.

The AI Competence Mirage occurs when humans unconsciously confuse **linguistic competence** with **computational certainty**.

Because LLMs communicate using language that appears authoritative, internally consistent and technically sophisticated, users instinctively infer that the underlying reasoning must also be correct.

In reality, these are fundamentally different properties.

A response may be linguistically flawless while still being factually incorrect, statistically uncertain or computationally unverifiable.

The language creates the appearance of certainty.

The mathematics does not.

Why Humans Fall Into the Mirage

Human cognition evolved under conditions where fluent language generally served as a reliable indicator of knowledge [17], [19]–[21].

Throughout human history, coherent explanations almost always originated from individuals possessing genuine expertise.

Consequently, our brains developed powerful cognitive heuristics that associate characteristics such as:

fluent language; [17], [19]–[21]

- technical vocabulary;
- coherent structure;
- logical progression;
- confidence;
- persuasive explanations;

with competence and truthfulness.

Large Language Models reproduce precisely these linguistic characteristics.

They do not intentionally deceive users.

Rather, they generate language whose statistical properties closely resemble those found in expert human communication.

As a consequence, humans naturally assign a degree of credibility that often exceeds the mathematical guarantees provided by the underlying computational process.

This reaction is neither irrational nor accidental.

It is an expected consequence of human cognition interacting with probabilistic language generation.

From the AI Competence Mirage to the Computational Trust Gap

The AI Competence Mirage produces what we define as the Computational Trust Gap [19]–[21], [55].

Definition 3: Computational Trust Gap. The Computational Trust Gap is the mismatch between perceived confidence in an AI-generated output and the level of reliability that can be justified by the underlying computational mechanism, evidence, validation procedure or deterministic execution path [19]–[21], [55].

The Computational Trust Gap represents the difference between: [19]–[21], [55]

- the confidence perceived by the human user after reading a response; and
- the degree of certainty that can actually be guaranteed by the computational mechanism that generated it.

Conceptually, we define:

$$\text{Computational Trust Gap} = \text{Perceived Confidence} - \text{Computational Guarantee}$$

where:

- **Perceived Confidence** represents the level of certainty that a human user attributes to the response after evaluating its linguistic quality.
- **Computational Guarantee** represents the degree of certainty that can be justified by the underlying computational process, whether deterministic, probabilistic or formally verifiable.

This equation is not intended as a quantitative metric.

Rather, it provides a conceptual framework for understanding why users frequently overestimate the reliability of modern AI systems.

When both quantities are approximately equal, the Computational Trust Gap approaches zero [19]–[21], [55].

This is typically the case for deterministic systems such as calculators, compilers or formally verified algorithms.

Conversely, when linguistic confidence greatly exceeds computational guarantees, the gap becomes significant.

This situation frequently occurs when interacting with Large Language Models on complex analytical, legal, financial or scientific questions.

Importantly, users often have no reliable way of recognizing when this transition occurs.

Why This Matters for Enterprise Systems

The Computational Trust Gap has relatively limited consequences when AI is used for low-risk activities such as: [19]–[21], [55]

- brainstorming;
- drafting documents;
- translating text;
- summarizing reports;
- generating software prototypes;
- creative writing.

In these domains, occasional inaccuracies generally carry limited operational consequences.

Enterprise systems operate under fundamentally different requirements.

Organizations routinely make decisions involving:

- financial transactions;
- credit approval;

- insurance underwriting;
- tax calculation;
- medical diagnosis;
- industrial automation;
- cybersecurity;
- regulatory reporting.

In these contexts, persuasive language cannot substitute for deterministic guarantees.

The confidence inspired by an answer is irrelevant unless the computational process itself is capable of producing reproducible, auditable and verifiable results.

This distinction explains why enterprise architectures should deliberately separate two fundamentally different responsibilities:

- interpreting human intent; and
- executing deterministic computation.

Large Language Models excel at the first task.

Deterministic software excels at the second.

Confusing these responsibilities inevitably increases the Computational Trust Gap and, consequently, organizational risk [19]–[21], [55].

Beyond AI Hype

Public discussions frequently describe today's enthusiasm for Artificial Intelligence as **AI Hype**.

While this expression captures the social dimension of the phenomenon, it does not explain its underlying mechanism.

The issue is not simply that society has become excessively optimistic about AI.

The issue is considerably deeper.

Modern language models systematically produce responses possessing nearly every linguistic characteristic that humans have historically associated with intelligence, expertise and certainty.

The resulting overconfidence is therefore not merely a consequence of media attention or marketing.

It emerges naturally from the interaction between two independent systems:

- a probabilistic computational model optimized to generate plausible language; and

a human cognitive system evolved to interpret fluent language as evidence of knowledge [17], [19]–[21].

Understanding this interaction is essential for the responsible adoption of Artificial Intelligence.

The challenge is not reducing enthusiasm for AI.

The challenge is learning to calibrate trust according to the mathematical guarantees provided by the computational architecture rather than according to the persuasive quality of its language.

6. The Great Misunderstanding: AI Is Not Replacing Software

One of the most common predictions surrounding Artificial Intelligence is that enterprise software will gradually disappear.

According to this view, organizations will no longer require specialized applications because users will simply ask an LLM to perform any task they need.

This paper argues that such a conclusion confuses two fundamentally different phenomena.

Artificial Intelligence is indeed disrupting the software industry.

However, the disruption is not occurring because software is becoming unnecessary.

The disruption is occurring because software has become dramatically easier to build.

This distinction changes the entire economic interpretation of modern AI.

The Wrong Question

Organizations increasingly ask software vendors questions such as:

"Why should I purchase your software if I can simply ask ChatGPT or Claude to perform the same task?"

Although understandable, this question is based on an incorrect assumption.

The correct comparison is not:

Enterprise Software versus ChatGPT

The correct comparison is:

Enterprise Software versus building and maintaining an equivalent software platform using AI-assisted software development.

These are fundamentally different problems.

A Large Language Model can generate code.

An enterprise platform requires engineering.

Software Is Far More Than Code

Source code represents only one component of enterprise software.

A mature enterprise platform typically contains years—or even decades—of accumulated engineering, including:

business rules; [31], [32], [36]–[40]

- deterministic algorithms;
- database architecture;
- authentication;
- authorization;
- APIs;
- workflow engines;
- transaction management;
- security;
- audit trails;
- monitoring;
- scalability;
- testing;
- deployment pipelines;
- backup strategies;
- version management;
- regulatory compliance;
- documentation;
- user experience;
- operational support.

These capabilities do not reliably emerge merely by asking a Large Language Model (LLM) a question.

They must be deliberately designed, implemented, tested, validated and continuously maintained.

A Large Language Model can significantly accelerate this engineering process.

It does not eliminate it.

The True Competitive Threat

The greatest competitive threat facing traditional software companies is therefore not conversational AI itself.

It is AI-accelerated software engineering [48]–[51].

Modern coding assistants such as Claude Code, GPT, Codex and similar systems have dramatically increased the productivity of software developers.

Tasks that previously required days of implementation can frequently be completed in hours.

Complex refactoring can often be performed automatically.

Documentation, testing, debugging and code generation have become substantially faster.

Consequently, relatively small engineering teams can now develop sophisticated enterprise platforms that would previously have required organizations many times larger.

This changes the economics of software creation.

It does not eliminate the need for software itself.

The competitive advantage shifts away from writing code toward understanding business problems.

The New Competitive Landscape

Historically, one of the largest competitive advantages enjoyed by major software vendors was engineering scale.

Building enterprise applications required:

- hundreds of software engineers;
- large quality assurance teams;
- dedicated documentation teams;
- infrastructure specialists;
- years of development.

Artificial Intelligence has dramatically reduced this barrier.

Today, relatively small engineering teams augmented by AI can achieve development velocities that were previously unattainable.

As a consequence, the software industry is becoming increasingly meritocratic.

Competitive advantage no longer depends primarily on the number of engineers employed.

Instead, it increasingly depends on:

- domain expertise;
- architectural quality;
- proprietary data;
- product vision;
- execution speed;
- customer understanding.

The ability to write code is rapidly becoming commoditized.

The ability to solve business problems is not.

Why Most Organizations Should Not Build Their Own Enterprise Software

An important consequence follows naturally.

Although AI makes software development dramatically more accessible, this does not imply that every organization should build its own enterprise systems.

Consider a retail company.

Its competitive advantage lies in merchandising, logistics, customer experience and supply chain management.

It does not lie in maintaining a CRM platform.

Similarly:

A bank competes through risk management and financial products.

A telecommunications company competes through network quality and customer service.

An insurance company competes through underwriting and claims management.

None of these organizations creates competitive advantage by dedicating engineering resources to rebuilding software that already exists.

Even if AI significantly reduces development costs, building and maintaining enterprise software remains a long-term commitment involving architecture, testing, governance, maintenance, support, cybersecurity and continuous evolution.

The opportunity cost often exceeds the licensing cost of specialized software.

Consequently, AI changes **how software is built** far more than **who should build it**.

The New Role of Enterprise Software

Enterprise software therefore becomes more valuable—not less.

Its value, however, shifts.

Historically, organizations purchased software because writing it was prohibitively expensive.

In the future, organizations will purchase software because maintaining it distracts them from their core business.

This distinction is fundamental.

The role of enterprise software evolves from:

providing functionality that customers cannot build

to

providing functionality that customers choose not to build because their competitive advantage lies elsewhere.

Artificial Intelligence accelerates software development.

It does not eliminate specialization.

A New Economic Principle

The central economic consequence of Artificial Intelligence can therefore be summarized as follows:

AI is not replacing enterprise software. AI is replacing much of the effort historically required to create enterprise software.

These are profoundly different statements.

The first predicts the disappearance of software.

The second predicts an explosion in software innovation driven by dramatically lower development costs.

Evidence increasingly suggests that the second interpretation more accurately describes the transformation currently taking place.

7. The Next Generation Enterprise Architecture

The previous chapters lead to a natural architectural conclusion.

Large Language Models are inherently probabilistic.

Enterprise systems fundamentally require deterministic execution.

Yet neither of these paradigms alone is sufficient to solve complex business problems.

A language model can understand what a user wants, but it does not possess the accumulated domain expertise required to solve every business problem correctly.

Conversely, deterministic software can execute algorithms with perfect reproducibility, but it cannot naturally understand human intent [22]–[30].

Between these two worlds lies a third component that has historically defined the value of enterprise software: **domain intelligence**.

We therefore propose that the next generation of enterprise systems should not be understood as the integration of Artificial Intelligence with deterministic software alone.

Instead, they should be viewed as architectures composed of three complementary computational layers:

- a **Probabilistic Interface Layer (PIL)** responsible for understanding human intent;
- a **Domain Intelligence Layer (DIL)** responsible for encapsulating business knowledge;
- a **Deterministic Execution Layer (DEL)** responsible for executing trusted computation.

Each layer solves a fundamentally different problem.

Together they define a new computational paradigm for enterprise systems.

Probabilistic Interface Layer (PIL)

The first layer is responsible for communication.

Historically, enterprise software required users to adapt themselves to the language of machines.

Users learned menus.

Filled forms.

Configured workflows.

Wrote SQL.

Called APIs.

Artificial Intelligence reverses this relationship.

Humans now communicate using their own language.

The responsibility of the Probabilistic Interface Layer is therefore to interpret human intentions.

Definition 5: Probabilistic Interface Layer (PIL). The Probabilistic Interface Layer (PIL) is the architectural layer that uses Large Language Models (LLMs) or related probabilistic AI systems to interpret natural-language intent, resolve ambiguity, formulate candidate plans, communicate explanations and route requests without itself serving as the final authority for business-critical computation.

Typical responsibilities include:

- understanding natural language;
- resolving ambiguity;
- interpreting objectives;
- answering questions;
- generating documentation;
- orchestrating workflows;
- interacting conversationally with users.

Importantly, this layer is intentionally probabilistic.

Human language itself is inherently ambiguous.

Consequently, the objective of this layer is not deterministic correctness.

Its objective is to understand what the user is attempting to accomplish.

Domain Intelligence Layer (DIL)

Understanding user intent is only the first step.

Knowing **how** to solve a problem requires an entirely different type of intelligence.

This second layer represents the true intellectual capital of enterprise software.

We refer to it as the **Domain Intelligence Layer (DIL)**.

Definition 4: Domain Intelligence Layer (DIL). The Domain Intelligence Layer (DIL) is the architectural layer that encodes domain-specific knowledge, business rules, regulatory constraints, analytical methods, decision policies, proprietary data, validation procedures and operational expertise so that interpreted intent can be transformed into a governed computational plan [31], [32], [36]–[40].

The DIL contains the accumulated knowledge that organizations have developed over years—or often decades—of experience within a specific business domain.

This knowledge extends far beyond software engineering.

It includes:

business rules; [31], [32], [36]–[40]

- domain expertise;
- regulatory knowledge;
- optimization methodologies;
- decision policies;
- best practices;
- proprietary workflows;
- feature engineering methodologies;
- analytical frameworks;
- validation procedures;
- operational experience;

proprietary datasets; [22], [29], [30]

- data enrichment methodologies.

Unlike Large Language Models, which possess broad general knowledge, the Domain Intelligence Layer contains highly specialized knowledge that differentiates one enterprise platform from another.

For example, two organizations may use exactly the same language model.

Yet they will produce fundamentally different solutions because their Domain Intelligence Layers are different.

This observation has profound implications.

The competitive advantage of enterprise software has never resided primarily in its source code [33], [42], [48]–[51].

Rather, it resides in the domain-specific intelligence that has been progressively codified into software.

Consequently, Artificial Intelligence dramatically reduces the cost of writing code.

It does not automatically create decades of accumulated business expertise.

Deterministic Execution Layer (DEL)

Once user intent has been interpreted and domain knowledge has determined the appropriate course of action, execution becomes a deterministic problem.

This responsibility belongs to the Deterministic Execution Layer.

Definition 6: Deterministic Execution Layer (DEL). The Deterministic Execution Layer (DEL) is the architectural layer that executes validated computational procedures, workflows, transactions, model-training pipelines, reports and governance controls in a manner designed for reproducibility, traceability, auditability and operational accountability [22]–[30].

Unlike the probabilistic interface, every operation performed within the DEL must satisfy the traditional requirements of enterprise software:

reproducibility; [22]–[30]

- correctness;
- auditability;
- traceability;
- security;
- transactional consistency;
- regulatory compliance.

Typical responsibilities include:

- mathematical computation;
- optimization algorithms;
- workflow execution;
- database transactions;
- model training;
- deployment;
- reporting;
- monitoring;
- security enforcement;
- audit logging.

Given identical inputs, controlled state and stable dependencies, the DEL is designed to produce reproducible outputs.

Its purpose is not flexibility.

Its purpose is trust.

The Separation of Responsibilities

One of the central principles of this architecture is that each layer performs only the task for which it is naturally suited.

The Probabilistic Interface Layer answers:

What does the user want?

The Domain Intelligence Layer answers:

How should this problem be solved within this specific business domain?

The Deterministic Execution Layer answers:

Execute the solution in a reproducible, auditable and mathematically reliable manner.

This separation dramatically reduces the Computational Trust Gap introduced in Chapter 5 [19]–[21], [55].

Users continue benefiting from natural language interaction while organizations continue benefiting from deterministic execution.

The strengths of both paradigms become complementary rather than competing.

Why This Architecture Matters

This three-layer architecture provides several important advantages.

First, organizations become independent of any particular Large Language Model.

As foundation models evolve, the Probabilistic Interface Layer can be replaced without modifying either the domain intelligence or the deterministic execution engine.

Second, proprietary business knowledge becomes explicitly separated from both language models and computational infrastructure.

This significantly improves maintainability, governance and long-term sustainability.

Third, deterministic execution remains independently testable, certifiable and auditable.

Finally, this architecture naturally explains why enterprise software continues to exist despite the emergence of increasingly capable language models.

Artificial Intelligence does not replace enterprise software.

It replaces the traditional interface through which humans communicate with enterprise software.

The New Definition of Enterprise Software

This architecture suggests a broader reinterpretation of enterprise software itself.

Historically, software has often been viewed simply as collections of algorithms implemented in source code [33], [42], [48]–[51].

This perspective is no longer sufficient.

Enterprise software should instead be understood as the codification of domain intelligence into deterministic computational processes.

Its value therefore extends far beyond code generation.

It resides in the systematic accumulation of knowledge that allows organizations to solve highly specialized problems reliably, consistently and at scale.

From this perspective, Large Language Models become extraordinarily powerful interfaces.

Domain Intelligence becomes the organization's competitive advantage.

Deterministic execution becomes the foundation of trust.

Together, these three layers define what we believe constitutes the next generation of enterprise computing.

A New Definition of Enterprise Artificial Intelligence

Based on the architectural principles introduced throughout this paper, we propose the following definition:

Enterprise Artificial Intelligence is a computational architecture in which probabilistic intelligence interprets human intent, domain intelligence determines how business problems should be solved, and deterministic software executes those solutions with mathematical reliability, reproducibility and governance [22]–[30].

This definition differs fundamentally from the increasingly common perception that enterprise AI consists merely of connecting a Large Language Model to organizational data.

Instead, it recognizes that trustworthy enterprise systems emerge from the integration of three complementary forms of intelligence:

- **Linguistic Intelligence**, responsible for understanding human language.
- **Domain Intelligence**, responsible for understanding the business.
- **Computational Intelligence**, responsible for executing deterministic computation.

Only the combination of these three layers enables enterprise systems that are simultaneously intuitive, trustworthy, auditable and economically sustainable.

8. Case Study: Enterprise Procure-to-Pay and Invoice Approval

The architectural principles introduced throughout this paper become particularly evident in enterprise procure-to-pay workflows, where natural-language intent must be translated into governed financial execution.

At first glance, modern Large Language Models appear capable of performing many activities associated with procurement and accounts payable teams.

They summarize invoices.

They draft purchase requests.

They explain contract language.

They answer supplier questions.

They generate approval notes.

They classify expenses.

This naturally leads many organizations to ask a tempting question:

If an LLM can understand a procurement request and produce a fluent recommendation, why would an enterprise still require a procure-to-pay platform?

The answer lies in the difference between interpreting a business request and executing a governed financial transaction.

These are fundamentally different challenges.

Interpreting an Invoice Is Not Approving a Payment

Approving an invoice or purchase order is not merely a language task. It is a controlled enterprise process that must satisfy financial, contractual, operational and audit requirements.

A production-grade procure-to-pay system must provide deterministic capabilities such as:

supplier validation;

purchase order matching;

invoice data extraction and normalization;

budget availability checks;

approval-limit enforcement;

segregation-of-duties controls;

tax and jurisdictional validation;

contract compliance;

duplicate-invoice detection;
payment scheduling;
audit logging;
ERP reconciliation;
security and access control;
regulatory compliance.

A language model can assist each individual step. It cannot guarantee that the complete workflow remains consistent, authorized, reproducible and auditable across years of enterprise operation.

That responsibility belongs to deterministic software.

Applying the Three-Layer Architecture

Procure-to-pay provides a clear example of how the three-layer architecture proposed in this paper naturally emerges.

Probabilistic Interface Layer (PIL)

The employee, buyer, finance analyst or executive communicates with the system using natural language.

Typical requests include:

"Approve this invoice if it matches the purchase order."

"Create a purchase request for this supplier."

"Explain why this payment is blocked."

"Renew this supplier contract if the terms are still compliant."

The Large Language Model interprets these intentions.

No deterministic enterprise execution has yet occurred.

Domain Intelligence Layer (DIL)

Once the user's objective has been understood, the system applies accumulated procurement, finance and compliance expertise.

Examples include:

supplier-risk policies;
budget-control rules;
approval matrices;
contractual obligations;
tax treatment;
payment terms;
internal controls;
audit requirements;
ERP configuration;

organization-specific operating procedures.

This knowledge has typically been accumulated through years of operational experience.

It generally cannot be reliably produced by prompting a Large Language Model (LLM) in isolation.

It represents the enterprise intelligence embedded within the platform.

Deterministic Execution Layer (DEL)

Finally, deterministic software executes the requested financial workflow.

Typical responsibilities include:

matching invoices against purchase orders;

validating supplier master data;

checking budget availability;

routing approvals;

enforcing access rights;

blocking policy violations;

posting transactions to the ERP;

scheduling payments;

recording immutable audit trails;

reconciling financial records.

Given identical invoices, policies, approval states and system configurations, this layer must produce consistent results.

Its behavior is deterministic by design.

Why Auditability Matters

Enterprise procurement and accounts payable operate within environments where financial decisions must be explainable after the fact.

Organizations routinely need to reconstruct why an invoice was approved, who approved it, which policy applied, whether the supplier was valid and whether the transaction respected budget, contract and compliance constraints.

This implies that every financial action must be traceable, including:

request origin;

supplier identity;

purchase order linkage;

approval chain;

policy evaluation;

budget status;

contract reference;

payment execution.

Large Language Models are not designed to provide this level of deterministic enterprise auditability.

Their objective is to assist human interpretation and communication.

Deterministic enterprise platforms exist to guarantee controlled execution.

These objectives are complementary rather than competing.

The Importance of Domain Intelligence

Perhaps the most significant misconception surrounding AI-assisted procurement is the belief that document interpretation constitutes the primary source of value.

In reality, interpretation is only the interface.

The real competitive advantage resides in the accumulated domain intelligence embedded within the procurement and finance platform.

Examples include:

supplier qualification methodologies;

approval-policy design;

contract-compliance procedures;

risk-scoring rules;

budget-governance methods;

organization-specific workflows;

proprietary supplier data;

financial-control practices.

A language model may generate the corresponding explanation or workflow suggestion.

It does not automatically possess the operating discipline required to decide which rule should apply under which business conditions.

That expertise belongs to the Domain Intelligence Layer.

Beyond Procure-to-Pay

Although this case study focuses on procurement and invoice approval, the same architectural principle generalizes naturally to virtually every enterprise domain.

Customer relationship management.

Enterprise resource planning.

Supply chain optimization.

Cybersecurity.

Medical diagnosis.

Legal compliance.

Industrial automation.

In every case:

Artificial Intelligence understands what the user wants.

Domain Intelligence determines how the organization solves that class of problems.

Deterministic software executes the corresponding business processes.

The architecture remains identical.

Only the domain knowledge changes.

The Enterprise Value Proposition

This analysis leads to an important conclusion.

The value proposition of enterprise software is no longer primarily the automation of computation.

Computation itself is increasingly commoditized.

Instead, enterprise software creates value by integrating three capabilities that are extremely difficult to replicate simultaneously:

intuitive human interaction through probabilistic language models;

accumulated domain expertise codified into reusable policies, workflows and methodologies;

deterministic execution capable of satisfying enterprise requirements.

Organizations therefore purchase enterprise software not because they lack access to Artificial Intelligence.

They purchase enterprise software because building, validating and continuously maintaining this integrated architecture would divert resources away from the activities that truly differentiate their business.

9. Related Work

The proposed theory is situated within several mature research traditions rather than outside them. Foundational work in information theory, symbolic artificial intelligence, probabilistic reasoning, and modern machine learning establishes the distinction between communication, computation, search, uncertainty, and learned representation [1]–[9]. Transformer architectures and foundation models provide the technical basis for contemporary Large Language Models (LLMs), including pre-training, instruction tuning, and alignment through human or AI feedback [10]–[16].

Prior work has also examined the limitations of Large Language Models (LLMs), including hallucination, stochastic language generation, calibration, and the social risks of confusing fluent text with reliable knowledge [17], [18], [55]. These results support the paper's distinction between plausibility and truth, but they do not by themselves provide an enterprise architecture for assigning probabilistic language models, domain knowledge, and deterministic execution to separate system responsibilities.

The AI Competence Mirage and the Computational Trust Gap are related to work on trust in automation, misuse and overreliance, appropriate reliance, and human calibration of automated systems [19]–[21]. The contribution of this paper is to reinterpret those concerns in the specific context of Large Language Models (LLMs): the perceived competence arises from fluent probabilistic language, while the gap arises because linguistic confidence and computational guarantees are generated by different mechanisms.

The proposed Domain Intelligence Layer (DIL) and Deterministic Execution Layer (DEL) overlap with research in software architecture, modularity, enterprise application architecture, and enterprise architecture frameworks [31]–[38]. They are also related to research on production machine learning, MLOps, model governance, data validation, model cards, and dataset documentation [22]–[30]. The distinction introduced here is architectural rather than merely operational: the paper argues that domain knowledge should be treated as a first-class computational layer between intent interpretation and execution.

The paper is also adjacent to literature on business process redesign, competitive advantage, programming as theory building, and the design of human-computer systems [39]–[42]. These works help explain why software value often resides in codified organizational knowledge rather than in code alone. In this sense, the Domain Intelligence Layer (DIL) is related to prior accounts of business process knowledge and software design knowledge, but it is framed specifically for enterprise systems that use Large Language Models (LLMs) as probabilistic interfaces.

Finally, research on intelligent agents, retrieval-augmented generation, tool use, reasoning-and-acting loops, and AI-assisted programming demonstrates that Large Language Models (LLMs) can participate in increasingly complex workflows [43]–[54]. This paper does not deny those developments. Its claim is narrower: even when agents or code-generation systems are powerful, enterprise trust still requires explicit separation among probabilistic interpretation, domain intelligence, and deterministic execution.

The novelty of Intent-Driven Computing (IDC) therefore does not lie in the isolated observation that users can express intent, nor in the existence of agents, workflows, semantic layers, or enterprise architecture. Related work in intent-based networking and intent-driven resource management also studies translation from high-level objectives to operational configuration. The contribution claimed here is the integrated enterprise-computing theory that combines the AI Competence Mirage, the Computational Trust Gap, and the PIL → DIL → DEL architecture as a unified account of how Large Language Models (LLMs) should be positioned within trustworthy enterprise systems.

10. Limitations

This paper is conceptual. It proposes a theory and architectural framework, but it does not experimentally validate the Probabilistic Interface Layer (PIL), Domain Intelligence Layer (DIL), and Deterministic Execution Layer (DEL) in a controlled empirical setting. Future work should evaluate whether systems designed according to this architecture measurably improve reliability, auditability, user calibration, development productivity, or operational outcomes.

The paper does not compare the proposed architecture against neuro-symbolic AI, formal methods, probabilistic programming, semantic parsing, retrieval-augmented generation, or agentic workflow systems in a systematic benchmark. These approaches may partially address some of the same concerns through different mechanisms, and they should be examined in future work.

The paper does not claim that Large Language Models (LLMs) are incapable of reasoning. Contemporary research has shown that prompting, tool use, retrieval, and reasoning-oriented methods can substantially improve performance on complex tasks [45]–[47], [52]–[54]. The claim is instead that enterprise systems should not treat probabilistic language generation alone as a substitute for governed domain knowledge and deterministic execution.

The paper also does not invalidate autonomous AI agents. Agents may become important components of enterprise systems. The limitation is that agent autonomy does not remove the need for explicit architectural boundaries, observability, governance, reproducibility, and accountability.

Finally, the proposed concepts require further operationalization. Future work should define measurable indicators for the Computational Trust Gap, develop design patterns for implementing the Domain Intelligence Layer (DIL), and compare PIL -> DIL -> DEL systems with alternative architectures across domains such as finance, healthcare, supply chain optimization, cybersecurity, and regulated machine learning.

11. Conclusion

The emergence of Large Language Models represents one of the most significant technological advances in the history of computing.

Never before have computers been capable of understanding and generating natural language with such fluency, flexibility and contextual awareness.

Their impact on software engineering, knowledge work and human-computer interaction will undoubtedly reshape virtually every industry over the coming decades.

However, this transformation should not be interpreted as the replacement of enterprise software by conversational Artificial Intelligence.

Throughout this paper we have argued that this widespread interpretation arises from a fundamental misunderstanding of the computational nature of Large Language Models.

LLMs are probabilistic systems.

Enterprise software is deterministic.

These are not competing paradigms.

They solve different classes of problems.

Language Models excel at interpreting human intentions, reducing ambiguity and facilitating communication between humans and machines.

Deterministic software excels at executing business rules, mathematical computation, optimization, governance, auditing and reproducible decision making [31], [32], [36]–[40].

Neither paradigm can fully replace the other.

Instead, the future of enterprise computing lies in their deliberate integration.

To formalize this perspective, this paper introduced a three-layer computational architecture consisting of:

- the **Probabilistic Interface Layer (PIL)**, responsible for understanding human language;
- the **Domain Intelligence Layer (DIL)**, responsible for encapsulating business expertise and accumulated organizational knowledge;

- the **Deterministic Execution Layer (DEL)**, responsible for executing trusted computational processes.

Together, these layers separate three fundamentally different forms of intelligence that have often been incorrectly treated as a single problem:

understanding language,

understanding the business,

and executing computation.

Recognizing these distinctions fundamentally changes how enterprise software should be understood.

Historically, software has often been viewed primarily as source code [33], [42], [48]–[51].

Artificial Intelligence now demonstrates that source code is becoming increasingly commoditized [33], [42], [48]–[51].

What remains difficult to replicate is not the ability to write software.

It is the ability to encode decades of domain expertise into deterministic systems capable of producing reliable, reproducible and auditable business outcomes.

Consequently, the true competitive advantage of enterprise software is shifting.

Increasingly, value will reside less in software engineering itself and more in the systematic codification of domain intelligence.

Artificial Intelligence dramatically lowers the cost of building software.

It does not automatically create the business knowledge required to solve highly specialized problems.

This distinction also explains why many current discussions surrounding the future of software reach misleading conclusions.

Organizations will not abandon enterprise platforms because conversational AI exists.

Instead, software platforms will increasingly embed conversational AI while continuing to provide the deterministic execution, governance and domain expertise that organizations require.

The interface changes.

The computational foundations remain.

The software industry therefore faces a profound transformation—but not the one commonly described.

The principal disruption is not that Artificial Intelligence replaces software.

The principal disruption is that Artificial Intelligence radically reduces the cost of creating software.

As barriers to software development continue to fall, competition will increasingly shift from engineering scale toward domain expertise, architectural quality and the ability to convert organizational knowledge into reusable computational systems.

In this sense, Artificial Intelligence resembles previous technological revolutions that dramatically reduced production costs without eliminating the products themselves.

Just as industrial automation transformed manufacturing without eliminating manufactured goods, Artificial Intelligence is transforming software engineering without eliminating software.

Finally, this paper proposes a broader reinterpretation of the role of Artificial Intelligence within enterprise systems.

Rather than viewing AI as a replacement for deterministic software, we argue that it should be understood as the natural language interface through which humans communicate with deterministic computational systems enriched by domain-specific intelligence.

From this perspective, Artificial Intelligence does not represent the end of enterprise software.

It represents the beginning of a new generation of enterprise computing.

One in which language becomes the universal interface.

Domain intelligence becomes the primary strategic asset.

And deterministic computation remains the foundation of trust.

12. Proposed Principles

The arguments presented throughout this paper may be summarized by the following principles:

Large Language Models optimize plausibility, not truth [12], [16]–[18].

1. Fluent language should not be confused with computational certainty.

The Computational Trust Gap increases whenever perceived confidence exceeds computational guarantees [19]–[21], [55].

2. Artificial Intelligence is transforming software development far more than it is replacing software.

Enterprise software derives its value primarily from codified domain intelligence rather than from source code alone [33], [42], [48]–[51].

3. Trustworthy enterprise AI requires the explicit separation of probabilistic language understanding, domain intelligence and deterministic computation.
4. The future of enterprise computing belongs to architectures that integrate these three complementary forms of intelligence.

13. References

- [1] C. E. Shannon, “A mathematical theory of communication,” *Bell System Technical Journal*, vol. 27, no. 3, pp. 379–423, 1948; vol. 27, no. 4, pp. 623–656, 1948.
- [2] A. M. Turing, “Computing machinery and intelligence,” *Mind*, vol. 59, no. 236, pp. 433–460, 1950.
- [3] J. McCarthy, M. L. Minsky, N. Rochester, and C. E. Shannon, “A proposal for the Dartmouth summer research project on artificial intelligence,” 1955.
- [4] A. Newell and H. A. Simon, “Computer science as empirical inquiry: Symbols and search,” *Communications of the ACM*, vol. 19, no. 3, pp. 113–126, 1976.
- [5] H. A. Simon, *The Sciences of the Artificial*, 3rd ed. Cambridge, MA, USA: MIT Press, 1996.
- [6] J. Pearl, *Probabilistic Reasoning in Intelligent Systems*. San Mateo, CA, USA: Morgan Kaufmann, 1988.
- [7] J. Pearl, *Causality: Models, Reasoning, and Inference*, 2nd ed. Cambridge, U.K.: Cambridge University Press, 2009.
- [8] S. Russell and P. Norvig, *Artificial Intelligence: A Modern Approach*, 4th ed. Hoboken, NJ, USA: Pearson, 2020.
- [9] Y. LeCun, Y. Bengio, and G. Hinton, “Deep learning,” *Nature*, vol. 521, pp. 436–444, 2015.
- [10] A. Vaswani et al., “Attention is all you need,” in *Advances in Neural Information Processing Systems*, 2017, pp. 5998–6008.
- [11] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, “BERT: Pre-training of deep bidirectional transformers for language understanding,” in *Proc. NAACL-HLT*, 2019, pp. 4171–4186.
- [12] T. B. Brown et al., “Language models are few-shot learners,” in *Advances in Neural Information Processing Systems*, vol. 33, 2020, pp. 1877–1901.
- [13] OpenAI, “GPT-4 technical report,” *arXiv:2303.08774*, 2023.
- [14] L. Ouyang et al., “Training language models to follow instructions with human feedback,” in *Advances in Neural Information Processing Systems*, vol. 35, 2022, pp. 27730–27744.
- [15] Y. Bai et al., “Constitutional AI: Harmlessness from AI feedback,” *arXiv:2212.08073*, 2022.
- [16] R. Bommasani et al., “On the opportunities and risks of foundation models,” *arXiv:2108.07258*, 2021.
- [17] E. M. Bender, T. Gebru, A. McMillan-Major, and S. Shmitchell, “On the dangers of stochastic parrots: Can language models be too big?” in *Proc. ACM FAccT*, 2021, pp. 610–623.
- [18] Z. Ji et al., “Survey of hallucination in natural language generation,” *ACM Computing Surveys*, vol. 55, no. 12, pp. 1–38, 2023.
- [19] R. Parasuraman and V. Riley, “Humans and automation: Use, misuse, disuse, abuse,” *Human Factors*, vol. 39, no. 2, pp. 230–253, 1997.
- [20] J. D. Lee and K. A. See, “Trust in automation: Designing for appropriate reliance,” *Human Factors*, vol. 46, no. 1, pp. 50–80, 2004.
- [21] K. A. Hoff and M. Bashir, “Trust in automation: Integrating empirical evidence on factors that influence trust,” *Human Factors*, vol. 57, no. 3, pp. 407–434, 2015.
- [22] D. Sculley et al., “Hidden technical debt in machine learning systems,” in *Advances in Neural Information Processing Systems*, 2015, pp. 2503–2511.
- [23] E. Breck, S. Cai, E. Nielsen, M. Salib, and D. Sculley, “The ML test score: A rubric for ML production readiness and technical debt reduction,” in *Proc. IEEE BigData*, 2017, pp. 1123–1132.
- [24] S. Amershi et al., “Software engineering for machine learning: A case study,” in *Proc. ICSE-SEIP*, 2019, pp. 291–300.
- [25] E. Breck et al., “Data validation for machine learning,” in *Proc. SysML*, 2019.
- [26] D. Baylor et al., “TFX: A TensorFlow-based production-scale machine learning platform,” in *Proc. KDD*, 2017, pp. 1387–1395.
- [27] M. Zaharia et al., “Accelerating the machine learning lifecycle with MLflow,” *IEEE Data Engineering Bulletin*, vol. 41, no. 4, pp. 39–45, 2018.
- [28] D. Kreuzberger, N. Kühl, and S. Hirschl, “Machine learning operations (MLOps): Overview, definition, and architecture,” *IEEE Access*, vol. 11, pp. 31866–31879, 2023.
- [29] M. Mitchell et al., “Model cards for model reporting,” in *Proc. ACM FAccT*, 2019, pp. 220–229.

- [30] T. Gebru et al., “Datasheets for datasets,” *Communications of the ACM*, vol. 64, no. 12, pp. 86–92, 2021.
- [31] D. Garlan and M. Shaw, “An introduction to software architecture,” in *Advances in Software Engineering and Knowledge Engineering*, vol. 1, 1993, pp. 1–39.
- [32] D. L. Parnas, “On the criteria to be used in decomposing systems into modules,” *Communications of the ACM*, vol. 15, no. 12, pp. 1053–1058, 1972.
- [33] F. P. Brooks, Jr., “No silver bullet: Essence and accidents of software engineering,” *Computer*, vol. 20, no. 4, pp. 10–19, 1987.
- [34] B. W. Boehm, “A spiral model of software development and enhancement,” *Computer*, vol. 21, no. 5, pp. 61–72, 1988.
- [35] L. Bass, P. Clements, and R. Kazman, *Software Architecture in Practice*, 4th ed. Boston, MA, USA: Addison-Wesley, 2021.
- [36] M. Fowler, *Patterns of Enterprise Application Architecture*. Boston, MA, USA: Addison-Wesley, 2002.
- [37] J. A. Zachman, “A framework for information systems architecture,” *IBM Systems Journal*, vol. 26, no. 3, pp. 276–292, 1987.
- [38] The Open Group, *TOGAF Standard*, 10th ed. San Francisco, CA, USA: The Open Group, 2022.
- [39] T. H. Davenport and J. E. Short, “The new industrial engineering: Information technology and business process redesign,” *Sloan Management Review*, vol. 31, no. 4, pp. 11–27, 1990.
- [40] M. E. Porter and V. E. Millar, “How information gives you competitive advantage,” *Harvard Business Review*, vol. 63, no. 4, pp. 149–160, 1985.
- [41] M. Winograd and F. Flores, *Understanding Computers and Cognition: A New Foundation for Design*. Norwood, NJ, USA: Ablex, 1986.
- [42] F. Naur, “Programming as theory building,” *Microprocessing and Microprogramming*, vol. 15, no. 5, pp. 253–261, 1985.
- [43] F. M. T. Brazier et al., “Agent-based organizational modeling for enterprise integration,” in *Proc. HICSS*, 1998.
- [44] M. Wooldridge and N. R. Jennings, “Intelligent agents: Theory and practice,” *The Knowledge Engineering Review*, vol. 10, no. 2, pp. 115–152, 1995.
- [45] P. Lewis et al., “Retrieval-augmented generation for knowledge-intensive NLP tasks,” in *Advances in Neural Information Processing Systems*, vol. 33, 2020, pp. 9459–9474.
- [46] S. Yao et al., “ReAct: Synergizing reasoning and acting in language models,” in *Proc. ICLR*, 2023.
- [47] T. Schick et al., “Toolformer: Language models can teach themselves to use tools,” in *Advances in Neural Information Processing Systems*, vol. 36, 2023.
- [48] M. Chen et al., “Evaluating large language models trained on code,” *arXiv:2107.03374*, 2021.
- [49] Y. Li et al., “Competition-level code generation with AlphaCode,” *Science*, vol. 378, no. 6624, pp. 1092–1097, 2022.
- [50] P. Vaithilingam, T. Zhang, and E. L. Glassman, “Expectation vs. experience: Evaluating the usability of code generation tools powered by large language models,” in *Extended Abstracts of CHI*, 2022.
- [51] S. Peng, E. Kalliamvakou, P. Cihon, and M. Demirer, “The impact of AI on developer productivity: Evidence from GitHub Copilot,” *arXiv:2302.06590*, 2023.
- [52] J. Wei et al., “Chain-of-thought prompting elicits reasoning in large language models,” in *Advances in Neural Information Processing Systems*, vol. 35, 2022, pp. 24824–24837.
- [53] T. Kojima, S. S. Gu, M. Reid, Y. Matsuo, and Y. Iwasawa, “Large language models are zero-shot reasoners,” in *Advances in Neural Information Processing Systems*, vol. 35, 2022, pp. 22199–22213.
- [54] X. Wang et al., “Self-consistency improves chain of thought reasoning in language models,” in *Proc. ICLR*, 2023.
- [55] S. Kadavath et al., “Language models (mostly) know what they know,” *arXiv:2207.05221*, 2022.
- [56] T. M. Mitchell, *Machine Learning*. New York, NY, USA: McGraw-Hill, 1997.