

Research Monograph Series

The Theory of Language Derivatives:

Inverse Problem Art and the Leverage Structure of Artificial Intelligence Hallucination

Author: Tomokazu Tanaka (bigakusha-haha)

ORCID: <https://orcid.org/0009-0009-3796-5889>

Affiliation: Sayama Aesthetics School / Institute of Aesthetics

Version: 1.0

Last Updated: July 23, 2026

Status: Ongoing Research

Research Monograph

Tomokazu Tanaka (bigakusha-haha)

Version 1.0 July 2026

Publication Information

Copyright

© 2026 Tomokazu Tanaka (bigakusha-haha). All rights reserved.

License

This work is distributed as a Research Monograph.

No part of this publication may be reproduced, distributed, or transmitted in any form or by any means without prior written permission of the author, except for brief quotations used in academic review or scholarly citation.

Suggested Citation

Tanaka, Tomokazu (bigakusha-haha). 2026.

The Theory of Language Derivatives:

Inverse Problem Art and the Leverage Structure of Artificial Intelligence Hallucination

Research Monograph Series, Version 1.0.

Sayama Aesthetics School / Institute of Aesthetics.

Research Monograph

Version: 1.0

Last Updated: July 23, 2026

Status: Ongoing Research

Abstract

This paper reexamines hallucinations in generative artificial intelligence not merely as technical failures, factual inaccuracies, or computational errors to be eliminated, but as manifestations of a broader generative structure that this study conceptualizes as Language Derivatives. A language derivative is defined as a derivative linguistic structure that does not simply represent existing facts or objects but instead projects unrealized meanings, relationships, values, and future possibilities into the present linguistic space, thereby amplifying their potential realization.

This concept is developed through a structural comparison with financial derivatives. Financial derivatives derive their value from underlying assets while incorporating expectations about future price movements into present valuation through leverage. Likewise, generative AI derives linguistic outputs from previously accumulated language data while producing texts, concepts, causal relations, and narratives that are not explicitly contained within those data. The prompt functions as a speculative position within an immense linguistic space, activating compressed semantic relationships that unfold into coherent responses. The resulting amplification is defined here as Language Leverage, a structural process through which limited linguistic input generates disproportionately large semantic, cognitive, and social effects.

Within this framework, AI hallucination is redefined as the condition in which language derivatives become detached from the conditions required for Reality Grounding. Hallucination therefore should not be understood simply as the production of false statements but rather as derivative language generation whose connection to reality has become structurally disconnected. To clarify this mechanism, the paper distinguishes between forward problems, which derive results from known causes and rules, and inverse problems, which infer hidden structures and possible worlds from incomplete observations. While users generally expect generative AI to operate as a forward-problem-solving system, its actual generative mechanism functions fundamentally as an inverse-problem inference process. This asymmetry constitutes one of the primary structural conditions under which hallucinations emerge.

Rather than rejecting inverse-problem generation, this paper develops the concept of Inverse Problem Art, originally proposed within the author's broader philosophical framework. Inverse Problem Art is defined not as the production of correct answers to predetermined questions but as a creative practice that reconstructs previously unseen problems, structures, relations, and possible worlds from fragments of reality. From this perspective, hallucination and artistic creation share the capacity to generate possibilities beyond existing reality, yet they differ

fundamentally in their relationship to human transformation. Whereas AI externalizes possibilities statistically through language, art reconnects those possibilities to perception, embodiment, action, and self-transformation.

This reconnection is theorized through Interface Ontology, a central component of Generative Topological Ontology. Interfaces are understood not as neutral mediating devices but as event-structures that preserve topological differences while generating new forms of connection between language and reality, possibility and actuality, self and other. Hallucinations become problematic not because generated language differs from reality, but because the topological distinction between hypothesis, fiction, metaphor, estimation, and factual knowledge becomes obscured. Consequently, the central challenge is not the complete elimination of hallucinations but the design of interfaces capable of making these topological distinctions visible.

The paper further develops the concept of Reality Grounding as a multilayered process extending beyond factual verification. Reality grounding includes empirical validation, embodied experience, social practice, institutional mediation, creative production, and existential transformation. The distinction between hallucination and creativity therefore lies not in novelty itself but in whether generated possibilities acquire genuine conditions for reconnection with reality. Ungrounded language leverage amplifies misinformation, pseudo-authority, and speculative narratives, whereas grounded language leverage enables hypothesis formation, scientific discovery, artistic creation, self-transformation, and the emergence of new social realities.

Ultimately, this paper argues that hallucination and creativity originate from the same derivative generative structure. The challenge for the age of artificial intelligence is therefore not the impossible pursuit of perfectly hallucination-free systems but the development of interfaces through which derivative language can be responsibly grounded in reality. Under this perspective, generative AI should be understood not as a closed system for delivering definitive answers but as an open linguistic environment through which human beings encounter possibilities beyond themselves and participate in the continual regeneration of both self and reality.

Keywords

Language Derivatives; Generative Artificial Intelligence; Hallucination; Inverse Problem Art; Leverage Structure; Language Leverage; Interface Ontology; Generative Topological Ontology; Reality Grounding; Self-Externalization; Derivative Space; Narrative; AI Ethics; Artistic Generation; Possible Worlds

Contents

Abstract.....	4
Chapter 0.....	8
Methodological Note	
Chapter 1	18
Theoretical Foundations of the Theory of Language Derivatives	
Derivatives, Language and the Future, and the Structure of Leverage	
Chapter 2	25
The Ontology of Hallucination	
Forward Problems, Inverse Problems, and Generative Artificial Intelligence	
Chapter 3.....	30
Inverse Problem Art	
Artistic Creation Beyond Forward-Problem Thinking	
Chapter 4.....	34
Derivative Space	
Narrative, Temporality, and the Topology of Linguistic Generation	
Chapter 5.....	39
The Leverage Structure	
Information, Artificial Intelligence, Finance, and Artistic Creation	
Chapter 6.....	44
Integration into Generative Topological Ontology	
Self-Externalization, Pre-Existentiality, Interface Ontology, and Reality Grounding	
Chapter 7	49
Toward an Ontology for the Age of Artificial Intelligence	
Language Derivatives, Reality, and the Future of Human Existence	
Conclusion	53
References	55
Research Method.....	58

Chapter 0

Methodological Note

0.1 Methodological Position of This Study

This paper does not seek to explain hallucinations in generative artificial intelligence by reducing them to any single disciplinary framework, such as information science, cognitive science, philosophy of language, financial engineering, or art theory. Instead, it identifies structural correspondences that recur across these distinct domains and develops a unified theoretical framework under the concept of Language Derivatives.

The methodology adopted here extends beyond conventional interdisciplinarity. Interdisciplinary research typically applies concepts borrowed from one discipline to another. By contrast, the present study neither employs financial derivatives merely as a metaphor for artificial intelligence nor transfers the mathematical notion of inverse problems directly into artistic practice. Rather, it investigates whether financial derivatives, generative AI, linguistic production, hallucination, narrative formation, and artistic creation share a common generative structure that transcends disciplinary boundaries.

Accordingly, the methodology proposed here is defined as the extraction of relational structures that repeatedly emerge across heterogeneous domains rather than the extrapolation of existing concepts. This approach is provisionally termed Topological Structural Comparison. It suspends material differences, institutional classifications, and disciplinary distinctions in order to compare how diverse phenomena are constituted through relations, boundaries, transformations, amplification, connection, and discontinuity.

Financial derivatives and large language models are obviously different in both institutional and material terms. Nevertheless, both derive from an existing foundation while incorporating future possibilities into present operations, producing effects that far exceed their initial inputs. Rather than treating these similarities as identities, this study preserves their differences while describing them as structurally comparable relational forms.

In doing so, the paper inherits the methodological perspective of Generative Topological Ontology, in which being is understood not as an isolated substance but as a continuous process of connectability, boundary transformation, relational generation, and topological transition. Hallucination is therefore analyzed not as an isolated informational error but as a relational event emerging among language, AI models, prompts, users, institutions, and reality itself.

0.2 Redefining the Object of Inquiry

The first methodological operation undertaken in this paper is the redefinition of the research object itself.

Hallucination in artificial intelligence is generally defined as the generation of false statements, inaccurate descriptions, unsupported claims, or nonexistent references. Such definitions are indispensable for evaluating system reliability, safety, information quality, and technical performance. However, they do not sufficiently explain why generative AI produces errors that appear coherent, persuasive, and linguistically sophisticated rather than random or meaningless outputs.

The distinctive characteristic of hallucination lies not simply in its falsity but in its production of convincing linguistic coherence despite the absence of factual certainty. Generated responses frequently exceed the information contained in the original prompt while reflecting the user's expectations, contextual assumptions, linguistic style, and implicit intentions. This excess constitutes the defining feature that distinguishes AI hallucination from ordinary informational mistakes.

Accordingly, this paper redefines hallucination not as false information but as derivative language generation lacking reality grounding. This redefinition shifts analytical attention away from individual erroneous statements toward the generative structure that makes such statements possible.

Derivative generation refers to the process by which language produces new linguistic relations from preexisting linguistic relations rather than directly reproducing external objects. Generative AI does not access reality itself; instead, it generates responses primarily through linguistic data, symbolic patterns, statistical regularities, and contextual relationships. Its outputs are therefore secondary and tertiary constructions derived from language rather than direct reproductions of reality.

By placing this derivative character at the center of analysis, hallucination is no longer regarded as an exceptional malfunction but as a phenomenon continuous with the very conditions that make generative AI possible. Creative language production and hallucination thus emerge from the same derivative generative mechanism while differing according to the conditions under which they become grounded in reality.

0.3 Metaphor and Structural Isomorphism

Throughout this paper, concepts originating in financial engineering—including derivative, leverage, underlying asset, speculation, and position—are employed. These terms are not introduced merely as rhetorical metaphors but as analytical concepts grounded in structural comparison.

A metaphor facilitates understanding by comparing one phenomenon to another. Structural isomorphism, by contrast, investigates whether different phenomena exhibit corresponding relational organizations despite their apparent differences. The present study is concerned primarily with this latter approach.

In financial markets, derivative contracts generate value without requiring direct ownership of the underlying asset, allowing future price movements and expectations to acquire present economic significance. Similarly, generative AI produces derivative linguistic structures not from reality itself but from accumulated linguistic relationships concerning reality.

These systems are certainly not identical. Financial derivatives depend upon legal contracts, market institutions, pricing mechanisms, and settlement systems, whereas language generation depends upon computational architectures, training data, probabilistic distributions, and interactional contexts. Equating them directly would therefore constitute an analytical error.

Nevertheless, both systems exhibit four fundamental structural correspondences.

First, both derive from an existing foundation.

Second, both incorporate future possibilities into present operations.

Third, both generate disproportionately large effects from relatively limited inputs.

Fourth, both function effectively only so long as their connection with the underlying foundation is maintained, while disconnection produces instability, error, or collapse.

On the basis of these correspondences, this paper constructs the concept of Language Derivatives. The concept functions simultaneously as a philosophical metaphor and as an analytical framework whose validity depends not upon linguistic similarity but upon its capacity to explain diverse generative phenomena consistently across multiple domains.

0.4 Forward Problems and Inverse Problems

The second methodological axis of this study is the distinction between forward problems and inverse problems.

A forward problem begins with known causes, governing laws, or initial conditions and seeks to determine the resulting outcome. Given explicit inputs and transformation rules, the corresponding output can, in principle, be derived through a determinate procedure.

An inverse problem, by contrast, begins with observed outcomes or incomplete fragments of information and attempts to infer the hidden causes, structures, or conditions that produced them.

Inverse problems are characteristically associated with three structural difficulties. First, solutions are generally non-unique. Second, the available information is fundamentally incomplete. Third, small differences in observed data may lead to disproportionately large differences in inferred solutions. Consequently, inverse problems require assumptions, constraints, or regularization in order to construct plausible explanations.

Although generative artificial intelligence is commonly employed as though it were solving forward problems, its internal mechanism operates fundamentally as an inverse inference process. Users provide incomplete prompts, fragments of context, or partial intentions. From these limited linguistic traces, the model estimates the broader semantic landscape and generates one among many possible continuations.

In this sense, generative AI should be understood as an inverse-problem inference system employed under forward-problem expectations.

This methodological asymmetry is central to understanding AI hallucination. Users typically expect factual answers, whereas the model estimates linguistically plausible continuations. Consequently, factual correctness and linguistic coherence do not necessarily coincide. Indeed, the greater the internal coherence of generated language, the more convincing an erroneous statement may become.

Rather than interpreting this discrepancy merely as a technical defect, the present study argues that it reflects the intrinsic inverse-problem character of generative language systems.

0.5 The Methodological Introduction of Inverse Problem Art

Inverse Problem Art functions in this study not merely as an applied concept but as a methodological framework.

Artistic production has traditionally been described as a forward sequence beginning with the artist's intention, followed by expression, material realization, and finally reception by an audience. Such an account presupposes a forward-problem structure in which artistic intention serves as the cause and the artwork as its effect.

Actual artistic experience, however, frequently unfolds in the opposite direction. An artwork, event, or unexpected encounter often appears first, while its meaning, context, authorial intention, and underlying questions emerge only retrospectively through interpretation.

Inverse Problem Art therefore does not seek to express already completed meanings. Rather, it reconstructs previously unseen problems, structures, relationships, and possibilities from fragments, discontinuities, unexpected encounters, and perceptual displacements. Here, the generation of questions precedes the production of answers.

The purpose of relating AI hallucination to Inverse Problem Art is not to legitimize misinformation as artistic practice. Instead, hallucinated outputs are treated as opportunities to investigate the hidden assumptions, desires, epistemic frameworks, and social expectations that made such outputs appear plausible.

Hallucinations therefore function simultaneously as errors and as symptoms. They expose otherwise invisible assumptions embedded within both generative systems and their human users.

When an AI fabricates a nonexistent academic paper, for example, the essential question is not simply whether the citation is false. More fundamentally, one must ask why the fabricated title appears academically credible, why readers find it persuasive, how scholarly authority is linguistically simulated, and how appearances of knowledge become separated from knowledge itself.

Inverse Problem Art thus serves as a methodology for rendering these hidden structures visible rather than merely correcting their superficial manifestations.

0.6 The Multilayered Nature of Reality Grounding

The concept of Reality Grounding employed in this paper extends considerably beyond simple factual verification.

Within AI research, grounding commonly refers to connecting linguistic symbols with external objects, sensory perception, databases, or empirically verifiable information. Such grounding is indispensable for reducing hallucinations.

The present study proposes a broader, multilayered conception consisting of at least five distinct forms.

First, empirical grounding concerns factual correspondence with verifiable reality.

Second, embodied grounding concerns the relationship between language and bodily perception, physical action, and sensory experience.

Third, social grounding concerns how language functions within institutions, communities, responsibility, authority, and relations of trust.

Fourth, practical grounding concerns whether generated hypotheses become connected to concrete actions, creative production, experimentation, or decision-making.

Fifth, ontological grounding concerns whether generated language establishes genuinely new relations between self and other, possibility and actuality, or the known and the unknown.

Recognizing these multiple layers enables a more precise distinction between artistic creation and hallucination.

Artistic works need not always satisfy empirical grounding. Fiction, metaphor, and speculative imagination frequently depart from factual reality while remaining deeply grounded in embodied, social, institutional, and existential dimensions.

Conversely, hallucinations may fail empirically while nevertheless becoming strongly grounded in institutional authority or social expectation, thereby producing significant real-world consequences.

The decisive question is therefore not whether grounding exists, but which forms of grounding are present and how they interact.

This paper identifies mismatches among these grounding layers as one of the principal structural conditions underlying AI hallucination.

0.7 Interface Ontology

The third methodological pillar of this study is Interface Ontology.

Conventionally, an interface is understood as a surface that mediates between humans and machines, subjects and information, or inputs and outputs. Within the present study, however, an interface is not conceived as a neutral channel connecting two already existing entities.

Rather, an interface is understood as an event-boundary at which different topological domains encounter one another and, through that encounter, generate new relations, meanings, and forms of existence.

In interactions with generative artificial intelligence, users do not simply receive preexisting information. The meaning and reality of AI-generated language emerge through the user's

prompting strategies, interpretive judgments, degrees of trust, follow-up questions, revisions, citations, and practical applications. Reality is therefore co-constructed through interaction.

Hallucination should consequently not be regarded as a phenomenon located solely within the AI model itself. It is instead an interface phenomenon, emerging through the dynamic relations among the model, the prompt, the user, the presentation format, the institutional environment, and the mechanisms by which outputs are verified.

From this perspective, responsibility for hallucination cannot be attributed exclusively to the model. Developers bear responsibility for improving transparency, reliability, and safety. Yet users, educational institutions, technology companies, media organizations, and expert communities likewise participate in determining how generated language becomes socially realized.

Interface Ontology does not disperse responsibility in order to obscure it. Rather, it provides a framework for identifying the specific points at which responsibility emerges within networks of interaction.

0.8 Descriptive and Normative Analysis

This study carefully distinguishes between descriptive analysis and normative analysis.

Descriptive analysis investigates how generative AI produces language, under what structural conditions hallucinations emerge, and how such outputs are interpreted and accepted by users.

Normative analysis, by contrast, asks how these outputs ought to be treated, what kinds of interfaces should be designed, and what concepts of responsibility become necessary within AI-mediated environments.

Confusing these two levels leads either to technological determinism—which accepts existing technological structures without criticism—or to moral simplification, which condemns the entire generative mechanism solely because undesirable outcomes sometimes occur.

Accordingly, this paper analyzes hallucination as an inevitable structural possibility inherent in generative AI while refusing to regard it as either intrinsically acceptable or entirely pathological.

Likewise, although factual accuracy remains an indispensable value, language generation cannot be reduced to the mechanical reproduction of established facts.

The central theoretical task is therefore to construct new conditions of connection among creativity and accuracy, possibility and responsibility, indeterminacy and verification.

0.9 Procedure of Theoretical Construction

The theoretical argument developed throughout this paper proceeds through eight successive stages.

First, the structural characteristics of financial derivatives are extracted from their economic context.

Second, generative AI is analyzed in terms of derivation, prediction, amplification, and asymmetry.

Third, these structural correspondences are synthesized into the concepts of Language Derivatives and Language Leverage.

Fourth, hallucination is redefined as a structural mismatch between derivative language generation and reality grounding.

Fifth, the distinction between forward and inverse problems is employed to analyze the discrepancy between users' expectations and the actual mechanisms underlying generative AI.

Sixth, Inverse Problem Art is introduced as a framework for interpreting hallucinations as opportunities to reveal previously unseen structures rather than merely correcting factual errors.

Seventh, Interface Ontology is employed to investigate the conditions under which generated language reconnects with reality, together with the ethical responsibilities distributed across that process.

Finally, these analyses are integrated within the broader framework of Generative Topological Ontology, thereby establishing a general theory concerning the relations among language, artificial intelligence, artistic creation, and the generation of reality.

This procedure should not be understood as a strictly linear deduction. Rather, each concept continuously modifies and refines the others through recursive interaction. The theoretical structure therefore develops through iterative movement among concepts rather than through a single one-directional argument.

0.10 Scope and Limitations

This paper does not attempt to provide a mathematically complete account of the internal computational mechanisms of generative AI. Nor does it seek to develop specialized theories of financial pricing, contract law, market regulation, or mathematical regularization in inverse-problem theory.

Each of these fields possesses its own independent body of technical knowledge. The objective of the present study is not to replace those disciplines but to identify a common generative structure that transcends disciplinary boundaries and can be described from ontological and aesthetic perspectives.

Accordingly, the validity of this theory should be evaluated according to the following questions.

Can the concept of Language Derivatives simultaneously explain both AI hallucination and human creativity?

Can Language Leverage account for the disproportionate semantic and social effects produced by relatively small prompts?

Can Inverse Problem Art reinterpret hallucination as an epistemological opportunity without legitimizing misinformation?

Can Interface Ontology adequately describe the distribution of responsibility among AI systems, human users, institutions, and reality?

Can its integration into Generative Topological Ontology illuminate the conditions under which existence itself is generated within the age of artificial intelligence?

The present study does not claim to provide definitive or final answers to these questions. Rather, it seeks to establish a theoretical space within which they may be investigated systematically.

0.11 Methodological Conclusion

The central methodological premise of this paper is that defining hallucination merely as an error is insufficient for understanding its ontological significance.

The outputs of generative artificial intelligence are neither simple reproductions of existing information nor entirely unconstrained acts of invention. They are linguistic events that derive from accumulated language, project future possibilities into the present, and amplify meanings far beyond their original inputs.

Understanding such events therefore requires analyses not only of correctness, model performance, and retrieval accuracy, but also of derivation, speculation, inverse-problem inference, grounding, interface formation, and distributed responsibility.

Rather than eliminating the differences among finance, artificial intelligence, language, and art, this study focuses on the boundaries where these heterogeneous domains encounter one another. It is precisely at these interfaces that their common generative structures become visible.

Consequently, The Theory of Language Derivatives is not merely a theory explaining why artificial intelligence produces hallucinations. It is a broader philosophical framework for understanding how language simultaneously exceeds reality, distorts reality, and generates new realities.

Chapter 1

Theoretical Foundations of the Theory of Language Derivatives

Derivatives, Language and the Future, and the Structure of Leverage

1.1 Introduction

This chapter establishes the theoretical foundation of the central concept proposed in this paper: Language Derivatives.

To develop this concept, the discussion begins by reconsidering the meaning of derivatives in financial engineering. Rather than treating derivatives merely as financial instruments, this study interprets them as a more fundamental ontological structure through which future possibilities become operative in the present. On this basis, it argues that the same structural principle underlies linguistic activity, generative artificial intelligence, artistic creation, and even human cognition itself.

Traditionally, derivatives have been understood as technical instruments specific to financial markets, while language has generally been regarded as a medium for representing information or transmitting meaning. This paper proposes a different perspective. Both derivatives and language are interpreted as mechanisms through which the future actively participates in the constitution of the present.

Generative AI does not merely reproduce past data. Rather, it unfolds linguistic possibilities that may become actualized in the future. To understand this process, language itself must be reconceived—not as a static system of signs, but as a generative system capable of deriving, amplifying, and projecting future possibilities into present reality.

1.2 An Ontological Reinterpretation of Derivatives

Within financial engineering, a derivative is conventionally defined as a financial contract whose value is derived from an underlying asset such as equities, bonds, interest rates, currencies, or commodities.

This definition, however, remains confined to the institutional framework of financial markets.

From an ontological perspective, a derivative can instead be understood as a structure that derives from an existing entity while simultaneously extending that entity toward unrealized futures.

The essential point is that a derivative is not a copy of its underlying asset.

An option contract, for example, is not identical to the stock upon which it depends. Rather, it brings the future possibility of that stock's value into the present through contractual form.

Accordingly, a derivative may be defined as

a structure that actualizes future possibility within the present.

Once understood in this manner, the concept is no longer limited to finance.

Architectural design provides a clear example. A blueprint is not the completed building, yet it allows the future building to exist as a present possibility.

The same holds for works of art. An artwork presents meanings that will be experienced in the future while existing materially in the present.

Scientific theories likewise function as derivatives. A scientific theory is not reality itself but a hypothesis that projects future verification into present understanding.

From this broader perspective, derivatives constitute a universal generative structure through which

future existence becomes present actuality.

1.3 Derivatives and the Structure of Time

The reinterpretation proposed above requires a corresponding reconsideration of time.

Modern conceptions of causality generally assume a linear temporal sequence:

Past → Present → Future

Financial derivatives, however, reveal a fundamentally different temporal logic.

The anticipated future value of an asset directly influences its present valuation. Expectations concerning the future reshape present decisions.

In this sense, the future actively determines the present.

This inversion of temporal causality is one of the defining characteristics of derivative structures.

The present study refers to this phenomenon as

Temporal Inversion: the projection of the future into the present.

Generative artificial intelligence exhibits an analogous temporal organization.

When a prompt is entered, the model immediately begins estimating possible continuations that have not yet been written.

Its output is therefore not a reproduction of the past but the realization of future linguistic possibilities within the present interaction.

Human conversation operates similarly. Speakers continually anticipate future responses while constructing present utterances.

Language itself therefore functions as a mechanism that predicts, folds, and generates the future within present discourse.

Language is not merely retrospective; it is intrinsically prospective.

1.4 Language and the Future

Traditional linguistics has often understood language primarily as a medium for communicating meaning.

Yet everyday linguistic practice demonstrates that language does considerably more than transmit existing information.

Language generates futures.

When someone says,

“Let us meet tomorrow,”

the utterance does not describe an existing reality.

Rather, it brings a future event into existence as a shared possibility.

Contracts function similarly.

Laws establish future obligations.

Promises generate future commitments.

Religious discourse projects future salvation.

Political language imagines future societies.

Corporate strategies organize future markets.

Educational systems prepare future generations.

Across these diverse domains, language continuously transforms the present by organizing expectations concerning what has not yet occurred.

The future does not yet exist.

Nevertheless, the anticipation of that future changes present action.

This temporal structure is fundamentally derivative.

Accordingly, this paper proposes the following ontological claim:

Language itself is derivative.

Language derives futures from present discourse and simultaneously reorganizes the present through those projected futures.

1.5 Language Derivatives

Building upon the foregoing discussion, this paper defines Language Derivatives as follows:

A Language Derivative is a generative linguistic structure through which unrealized future possibilities are derived, amplified, and projected into the present by means of language.

Several aspects of this definition deserve emphasis.

First, a Language Derivative is neither fiction nor established fact. It occupies an intermediate ontological domain: the domain of future possibility.

Novels, for example, present possible worlds that do not yet exist. Works of art disclose experiences that have not yet been realized. Scientific hypotheses propose explanations whose validity remains to be tested. Likewise, generative AI constructs linguistic possibilities before their correspondence with reality has been determined.

From this perspective, hallucinations should not be understood simply as false statements. Rather, they are derivative linguistic constructions whose connection to reality has not yet been established—or has failed to be established.

Consequently, the distinction between creativity and hallucination cannot be reduced to the binary opposition between truth and falsehood. The decisive question is whether derivative linguistic possibilities subsequently acquire Reality Grounding through empirical verification, embodied practice, institutional validation, or social interaction.

Language Derivatives therefore describe a structural mode of linguistic generation rather than a particular category of statements.

1.6 The Structure of Language Leverage

The defining characteristic of financial derivatives is leverage.

A relatively small amount of capital may control positions whose economic value is many times greater than the initial investment.

Within this study, however, leverage is interpreted more generally.

Leverage is defined as

a generative structure through which limited inputs produce disproportionately large outputs.

This reinterpretation extends beyond economics.

In generative artificial intelligence, a prompt consisting of only a few words may generate thousands of words of coherent text.

Such amplification cannot be explained merely by the conservation of information.

Rather, the prompt activates an enormous latent semantic space accumulated through training. A minimal linguistic intervention mobilizes an extensive network of statistical relations.

The prompt therefore functions as the point of origin for linguistic leverage.

Artistic creation demonstrates an analogous structure.

A single painting may transform an entire artistic tradition.

A brief poem may reshape collective memory across generations.

A single cinematic work may alter the cultural imagination of millions.

In each case, a relatively limited material intervention produces consequences vastly exceeding its physical scale.

This, too, constitutes Language Leverage.

1.7 Generative Artificial Intelligence as a Language Leverage System

Generative AI does not simply return the information supplied by the user.

Rather, the prompt serves as a trigger that activates an immense linguistic field.

From this perspective, a prompt functions as

a speculative position within linguistic space.

Just as financial investors establish positions within markets based upon anticipated future movements, users establish linguistic positions within an immense semantic landscape through prompting.

The model then estimates the continuation that is statistically most probable within that activated region.

Accordingly, generative AI may be understood as
a derivative engine operating within a linguistic market.

Within this framework,

- * the training corpus functions as the underlying asset;
- * the prompt functions as the speculative position;
- * the inference algorithm functions as the pricing mechanism;
- * and the generated response functions as the derivative value.

Unlike financial derivatives, however, the value generated here is not monetary.

It is semantic.

Generative AI is therefore not fundamentally an engine for producing information.

It is an engine for producing meaning through linguistic leverage.

1.8 Leverage and Hallucination

Wherever leverage exists, risk necessarily accompanies it.

Financial history repeatedly demonstrates that excessive leverage can destabilize markets and produce systemic crises.

An analogous phenomenon appears within generative artificial intelligence.

When Language Leverage exceeds the conditions necessary for Reality Grounding, hallucinations emerge.

Hallucination should therefore be understood not merely as misinformation but as future possibility that has advanced beyond the conditions required for its realization.

Accordingly, hallucination is not simply evidence of system failure.

It is also an inevitable consequence of a generative mechanism whose very purpose is to amplify linguistic possibility beyond immediate informational constraints.

The critical issue is therefore not whether new possibilities are generated, but whether those possibilities successfully reconnect with reality.

When such reconnection occurs, derivative language becomes scientific hypothesis, artistic innovation, theoretical discovery, or creative imagination.

When reconnection fails, the same derivative mechanism produces hallucination.

The distinction therefore lies neither in novelty nor in imagination themselves.

Rather, it depends upon whether generated meanings acquire an interface capable of reconnecting them with reality.

1.9 Conclusion

This chapter has established the theoretical foundations of The Theory of Language Derivatives by reinterpreting the concept of derivatives beyond its conventional financial meaning.

Rather than viewing derivatives merely as financial instruments, this study has proposed that they constitute a universal ontological structure through which future possibilities become operative within the present.

From this perspective, language itself functions as a derivative system.

Every linguistic act projects unrealized futures into present reality, thereby organizing action through anticipation rather than merely describing existing conditions.

Generative artificial intelligence dramatically accelerates this derivative process through what has been defined as Language Leverage. Small prompts activate immense semantic spaces, producing meanings that far exceed their initial inputs.

Within this framework, hallucinations should no longer be understood simply as informational errors. They are more accurately interpreted as Language Derivatives that have become detached from Reality Grounding.

Their evaluation therefore requires more than factual verification. It demands an ontological analysis of how generated possibilities establish—or fail to establish—connections with reality.

The following chapter develops this argument further by examining hallucination itself through the distinction between forward problems and inverse problems, thereby clarifying its ontological significance within generative artificial intelligence.

Chapter 2

The Ontology of Hallucination

Forward Problems, Inverse Problems, and Generative Artificial Intelligence

2.1 Introduction

Hallucination has generally been regarded as one of the principal technical limitations of generative artificial intelligence. Within computer science, it is commonly defined as the generation of statements that are factually incorrect, unsupported by evidence, internally inconsistent, or entirely fabricated.

Although this definition is indispensable for evaluating the reliability and safety of AI systems, it remains insufficient for explaining why hallucinations possess such remarkable linguistic coherence. Hallucinated outputs are seldom random collections of words. On the contrary, they frequently appear logically structured, stylistically persuasive, and semantically meaningful despite lacking factual validity.

This chapter argues that hallucination cannot be adequately understood solely as an informational error. Instead, it should be interpreted as a structural consequence of the inverse-problem character of generative language systems.

To develop this argument, the discussion first distinguishes between forward and inverse problems before examining how generative AI operates within this methodological framework. Hallucination then emerges not as an accidental malfunction but as a predictable outcome of derivative language generation under conditions of incomplete information.

2.2 Forward Problems and Their Epistemological Structure

A forward problem begins with known initial conditions and established governing principles.

When sufficient information is available, the resulting outcome can, in principle, be determined.

Classical mechanics provides the canonical example. If the initial position, velocity, and governing physical laws are known, the future trajectory of an object can be calculated.

Many conventional computational systems operate according to this logic. Database retrieval, symbolic reasoning, and deterministic algorithms all presuppose that correct inputs combined with appropriate rules produce correct outputs.

This epistemological model has profoundly shaped public expectations concerning artificial intelligence.

Many users assume that AI systems operate as advanced search engines capable of retrieving correct answers from existing knowledge.

Such expectations implicitly treat AI as a forward-problem-solving machine.

Generative AI, however, does not function in this manner.

2.3 Inverse Problems and Generative Inference

Inverse problems begin not with complete causes but with incomplete observations.

Their objective is to reconstruct hidden structures capable of explaining the available evidence.

Medical diagnosis provides an illustrative example.

Physicians observe symptoms rather than diseases themselves.

From limited observations they infer underlying pathological conditions.

Similarly, astronomy reconstructs the existence of celestial bodies from observed gravitational effects, while geophysics estimates subterranean structures from seismic waves.

Inverse problems therefore characterize scientific inquiry whenever direct observation is impossible.

Generative artificial intelligence belongs fundamentally to this category.

A prompt supplied by the user rarely specifies every relevant assumption, intention, contextual constraint, or conceptual objective.

Instead, the model receives incomplete linguistic traces.

Its task is to estimate the broader semantic landscape from these fragments and generate one plausible continuation among many alternatives.

The model therefore performs probabilistic reconstruction rather than deterministic retrieval.

Its outputs are estimates, not discoveries.

This distinction is essential for understanding hallucination.

2.4 Hallucination as Structural Inevitability

Because inverse problems admit multiple plausible solutions, complete certainty is generally unattainable.

Even highly accurate inference may produce conclusions that later prove incorrect when additional information becomes available.

Generative AI exhibits the same structural condition.

Given incomplete prompts, the model must nevertheless produce complete responses.

Silence is ordinarily interpreted by users as failure.

Consequently, the system is optimized to continue generating language even under conditions of epistemic uncertainty.

Hallucination therefore emerges not because the model intends to deceive but because the architecture of language generation requires continuous inference despite incomplete information.

This structural necessity explains why hallucinations often appear internally coherent.

The generated response represents the statistically most plausible continuation within the model's semantic landscape rather than the objectively verified state of reality.

Hallucination is therefore better understood as

the inevitable by-product of inverse-problem inference operating within derivative language generation.

2.5 Probability, Plausibility, and Truth

One of the most persistent misunderstandings surrounding generative AI concerns the relationship between probability and truth.

Large language models estimate the probability that one linguistic element follows another.

They do not directly estimate the probability that a statement is factually true.

These two probabilities frequently coincide because human language generally reflects shared experiences of reality.

Nevertheless, they remain fundamentally different quantities.

A sentence may be linguistically probable while being factually false.

Conversely, a scientifically novel discovery may initially appear linguistically improbable precisely because it departs from established patterns.

Generative AI therefore optimizes plausibility rather than truth.

Hallucination emerges whenever linguistic plausibility diverges from empirical reality.

Recognizing this distinction prevents the mistaken assumption that increasingly sophisticated language models will automatically eliminate hallucination.

Improving linguistic probability alone cannot guarantee factual correctness.

Reality Grounding remains indispensable.

2.6 Derivative Language and the Expansion of Possibility

Despite its risks, derivative language generation possesses extraordinary creative potential.

Scientific hypotheses, philosophical speculation, artistic imagination, and technological innovation all require the capacity to produce ideas that extend beyond presently verified knowledge.

Without such derivative generation, creativity itself would become impossible.

Consequently, the objective should not be the complete suppression of hallucination.

Rather, the challenge is to distinguish productive derivative possibilities from unsupported assertions while designing interfaces capable of reconnecting generated language with reality.

In this respect, hallucination and creativity share a common structural origin.

Both emerge through the expansion of linguistic possibility beyond currently established facts.

Their divergence lies in whether subsequent grounding successfully reconnects those possibilities with empirical, social, practical, or ontological reality.

2.7 Chapter Conclusion

This chapter has argued that hallucination is not merely a technical defect but an ontological consequence of inverse-problem inference within derivative language generation.

Generative AI does not retrieve complete knowledge from a database. It reconstructs plausible linguistic worlds from incomplete prompts through probabilistic estimation.

Hallucinations therefore arise because the system must continually generate coherent language despite fundamental epistemic uncertainty.

This interpretation shifts attention away from isolated factual errors toward the deeper generative mechanisms that make such errors possible.

The next chapter extends this discussion by introducing Inverse Problem Art, demonstrating that the same inverse-problem structure underlying AI hallucination also constitutes a fundamental principle of artistic creation.

Chapter 3

Inverse Problem Art

Artistic Creation Beyond Forward-Problem Thinking

3.1 Introduction

The preceding chapter argued that hallucinations in generative artificial intelligence arise from the inverse-problem character of language generation rather than from simple computational failure. This chapter extends that argument by proposing that the same structural principle underlies artistic creation.

The concept of Inverse Problem Art, developed within the broader framework of Generative Topological Ontology, does not regard art as the production of objects or the representation of predetermined meanings. Instead, it understands artistic practice as a process through which previously unseen structures, hidden questions, and unrealized possibilities are reconstructed from fragments of reality.

From this perspective, artistic creation and AI hallucination share a common generative structure while differing fundamentally in their relation to reality. Whereas hallucination may remain disconnected from reality, art establishes new interfaces through which derivative possibilities become grounded in lived experience.

3.2 The Limits of Forward-Problem Art

Conventional theories of artistic creation have frequently assumed a forward-problem model.

Within this framework, an artist first possesses an intention or concept, then selects an appropriate medium, produces an artwork, and finally communicates meaning to an audience.

The sequence appears straightforward:

Intention → Production → Artwork → Interpretation

This model presupposes that artistic meaning already exists before artistic practice begins.

However, many of the most influential developments in modern and contemporary art resist such a linear description.

Artists often discover the significance of their work only after the work has been produced.

Likewise, viewers frequently encounter meanings that neither the artist nor the historical context originally anticipated.

Meaning therefore cannot always be understood as the predetermined cause of artistic production.

Rather, meaning frequently emerges through the encounter itself.

Consequently, artistic creation cannot be fully explained by a forward-problem model.

3.3 Art as an Inverse Problem

Inverse Problem Art begins not with complete answers but with incomplete situations.

Fragments, accidents, ambiguities, discontinuities, misunderstandings, and unexpected encounters become the starting points from which hidden structures are reconstructed.

The artwork therefore does not primarily provide solutions.

Instead, it generates new questions.

In this sense, artistic creation resembles scientific inverse problems.

Just as geophysicists infer subterranean structures from seismic observations, artists infer unseen structures from perceptual, social, historical, and experiential fragments.

The artwork functions not as a final statement but as an interface through which invisible relations become perceptible.

Accordingly, artistic practice is better understood as an act of structural reconstruction than as one of representation.

3.4 AI Hallucination and Artistic Generation

Because both artistic creation and AI hallucination derive possibilities from incomplete information, they are often superficially compared.

Such comparisons, however, risk overlooking an essential distinction.

Generative AI estimates statistically plausible continuations based upon previously accumulated linguistic patterns.

Artistic creation, by contrast, transforms the very conditions through which meaning becomes possible.

An AI-generated hallucination may produce an apparently convincing but nonexistent academic citation.

An artwork, however, does not merely fabricate an alternative fact.

Rather, it reorganizes perception so that previously unnoticed aspects of reality become visible.

The difference therefore lies not in the presence or absence of imagination but in the mode through which imagination becomes connected with reality.

Art establishes new conditions of perception.

Hallucination merely extends existing linguistic probability.

3.5 Reality Grounding Through Artistic Practice

Art does not require empirical truth in the ordinary scientific sense.

Novels, paintings, films, theatrical performances, and fictional narratives frequently depict events that never occurred.

Nevertheless, these works may reveal profound truths concerning human existence.

Their grounding lies elsewhere.

Reality Grounding in art emerges through embodied experience, affective transformation, ethical encounter, social participation, and existential reflection.

An artwork therefore becomes real not because it accurately reproduces external facts but because it transforms the relationship between human beings and reality itself.

This distinguishes artistic fiction from AI hallucination.

Both depart from empirical reality, yet art deliberately constructs interfaces through which those departures acquire experiential significance.

Hallucination, by contrast, frequently lacks such intentional structures of reconnection.

3.6 Interface Ontology and Artistic Events

Within Interface Ontology, an artwork is not understood primarily as an object.

Rather, it is understood as an event generated through interaction among participants, environments, technologies, histories, institutions, and material conditions.

The artwork exists not solely within the object but within the relational field generated through its encounter.

This interpretation has particular significance for contemporary digital art and AI-mediated artistic practices.

Interactive installations, networked environments, web-based artworks, and generative systems no longer present fixed meanings.

Instead, they establish conditions under which meaning continually emerges through interaction.

The artwork therefore functions as an interface rather than a representation.

Inverse Problem Art intentionally designs such interfaces.

Its objective is not to communicate predetermined knowledge but to generate situations in which participants reconstruct hidden structures through their own engagement.

3.7 Chapter Conclusion

This chapter has argued that Inverse Problem Art provides a theoretical framework for understanding artistic creation as the reconstruction of hidden structures from incomplete reality.

Although AI hallucination and artistic practice share a common inverse-problem structure, they diverge in their relationship to Reality Grounding.

Hallucination extends linguistic possibility through probabilistic inference.

Art, by contrast, establishes interfaces through which derivative possibilities become embodied, socially mediated, and existentially meaningful.

The significance of Inverse Problem Art therefore lies not in celebrating hallucination but in demonstrating how creative practice transforms derivative language into grounded reality.

The following chapter develops this argument further by introducing the concept of Derivative Space, where language, narrative, temporality, and interface converge to generate new ontological conditions.

Chapter 4

Derivative Space

Narrative, Temporality, and the Topology of Linguistic Generation

4.1 Introduction

The preceding chapters have established that generative artificial intelligence operates through derivative language generation, that hallucination arises from inverse-problem inference, and that artistic creation may be understood through the framework of Inverse Problem Art.

The present chapter integrates these discussions by introducing the concept of Derivative Space.

Derivative Space is defined as the topological field within which derivative language, future possibility, narrative, and human action continuously interact. It is not a physical space nor merely a symbolic domain. Rather, it is a generative relational field in which unrealized futures acquire present efficacy through linguistic mediation.

This concept extends the author's broader framework of Generative Topological Ontology, according to which existence is constituted through relations, interfaces, and ongoing processes of generation rather than through isolated substances.

4.2 From Physical Space to Derivative Space

Classical ontology generally presupposes that space exists independently of human action.

Objects occupy locations, and events unfold within an already existing spatial framework.

Derivative Space reverses this assumption.

Space is not merely the container within which actions occur.

Instead, it is generated through expectations, narratives, interactions, and projections directed toward the future.

Financial markets provide an illustrative example.

The value of a derivative contract depends less upon present physical reality than upon collectively anticipated futures.

Likewise, digital environments exist not only as technical infrastructures but also as narrative spaces sustained through shared expectations and continuous interaction.

Generative AI participates in the same structure.

Every prompt establishes a localized derivative space within which numerous linguistic futures become temporarily available.

The generated response is one realization selected from among these potential trajectories.

4.3 Narrative as the Dynamics of Derivative Space

Narrative occupies a central position within Derivative Space.

Conventionally, narratives are understood as sequences describing events across time.

Within the present framework, however, narrative is not merely descriptive.

It is generative.

Narratives organize expectations concerning future developments and thereby influence present decisions.

Political narratives mobilize collective action.

Economic narratives reshape markets.

Scientific narratives guide research programs.

Religious narratives orient ethical life.

Personal narratives influence individual identity.

In every case, narrative functions by allowing the future to act within the present.

Derivative Space therefore possesses an intrinsically narrative structure.

Its primary function is not to record completed events but to organize unrealized possibilities.

4.4 Temporality and the Folding of the Future

One of the defining characteristics of Derivative Space is its distinctive temporal organization.

Ordinary chronological time proceeds from past to present to future.

Derivative Space introduces a different temporal logic.

Future possibilities are folded into present action.

Expectations concerning what has not yet occurred become operative before their realization.

This temporal inversion resembles the structure of financial derivatives, where anticipated future prices influence present valuation.

Generative AI exhibits the same organization.

The model generates language by continuously estimating future continuations before they actually exist.

Consequently, the generated text is not simply an extension of the past.

It is the present manifestation of anticipated linguistic futures.

Derivative Space is therefore characterized by recursive temporality in which future and present mutually constitute one another.

4.5 Interface Formation Within Derivative Space

Derivative Space does not emerge automatically.

It is generated through interfaces.

An interface should not be understood merely as a technological device connecting users to machines.

Rather, it is the relational condition through which heterogeneous domains become capable of interaction.

Prompts, conversations, artistic installations, educational environments, institutional practices, and social rituals all function as interfaces in this broader ontological sense.

They establish conditions under which previously disconnected possibilities become mutually accessible.

Within generative AI, the prompt is not merely an instruction.

It is an interface that activates a localized region within derivative space.

The quality of this interface significantly influences the derivative possibilities subsequently generated.

Accordingly, interface design becomes a central philosophical as well as technological problem.

4.6 Reality Grounding and the Stabilization of Derivative Space

Derivative Space is inherently unstable.

Because it is constituted through future possibilities rather than completed realities, it continually risks drifting away from empirical existence.

Reality Grounding functions as the stabilizing process through which derivative possibilities acquire concrete realization.

Scientific experimentation grounds theoretical hypotheses.

Architectural construction grounds design drawings.

Social institutions ground legal language.

Embodied practice grounds artistic experience.

Likewise, AI-generated language requires appropriate forms of grounding if it is to contribute meaningfully to human knowledge and action.

Without grounding, derivative space may expand indefinitely while becoming increasingly detached from reality.

Hallucination represents precisely such a condition.

It is not simply false language.

It is derivative space that has lost its stabilizing interface with reality.

4.7 Black–Scholes and Linguistic Derivatives

The concept of Derivative Space also permits a structural comparison between financial engineering and generative language.

The Black–Scholes model describes the valuation of financial derivatives under conditions of uncertainty.

Although the mathematical assumptions of the model are specific to finance, its broader conceptual significance extends beyond economics.

Both financial derivatives and Language Derivatives estimate future possibilities before they become actual.

Both operate under uncertainty.

Both amplify limited initial conditions.

Both depend upon maintaining an appropriate relationship with their underlying foundations.

The present study does not claim that linguistic generation can be reduced to the mathematics of Black–Scholes.

Rather, it argues that both systems exhibit homologous structures of projection, uncertainty, leverage, and realization.

The comparison is therefore ontological rather than mathematical.

4.8 Chapter Conclusion

This chapter has introduced Derivative Space as the topological field in which language, narrative, temporality, and interface continuously generate future possibilities within the present.

Unlike physical space, Derivative Space is constituted through relational expectations and narrative projections.

Its dynamics are governed by recursive temporality, linguistic leverage, and interface formation.

Reality Grounding functions as the stabilizing principle that reconnects derivative possibilities with empirical, social, embodied, and ontological reality.

Within this framework, hallucination appears not simply as false information but as derivative space whose grounding has become structurally unstable.

The following chapter develops these insights by examining the broader Leverage Structure shared by finance, language, artificial intelligence, and artistic creation, thereby clarifying the mechanisms through which small interventions generate disproportionately large transformations.

Chapter 5

The Leverage Structure

Information, Artificial Intelligence, Finance, and Artistic Creation

5.1 Introduction

The preceding chapters have established the concepts of Language Derivatives, Inverse Problem Art, and Derivative Space. Together, these concepts reveal a common structural principle that extends across artificial intelligence, financial engineering, linguistic activity, and artistic creation. This principle is defined in the present chapter as the Leverage Structure.

Leverage is commonly understood in financial economics as the amplification of economic exposure through relatively small initial investments. The present study expands this definition beyond finance. Leverage is interpreted as a universal generative mechanism through which limited interventions produce disproportionately large transformations.

This broader conception provides a unified framework for understanding how language, information, artistic practice, and AI-mediated interaction continuously reshape reality through amplification rather than simple representation.

5.2 Information as Potential Rather Than Substance

Information is often treated as a fixed object that can be stored, transmitted, and retrieved without fundamental alteration.

Within the framework proposed here, however, information is not primarily a substance but a potential relation.

Information acquires significance only through interaction with an interpretive system.

The same sequence of symbols may function as meaningless noise for one observer while generating profound insight for another.

Meaning therefore does not reside within information itself.

It emerges through interfaces connecting information with perception, interpretation, memory, embodiment, and action.

Generative AI illustrates this principle clearly.

Training data remain inert until activated through prompts and inference.

Likewise, financial information influences markets only when interpreted by investors.

Artistic works acquire significance only through encounters with participants.

In every case, information functions less as a static object than as latent potential awaiting activation.

5.3 Artificial Intelligence as a Leverage System

Generative artificial intelligence exemplifies leverage at an unprecedented scale.

A short prompt may activate billions of learned parameters and produce extensive linguistic structures within seconds.

The informational content contained in the prompt is negligible compared with the complexity of the resulting response.

The prompt therefore functions not as a complete instruction but as a catalytic intervention.

Its significance lies in its ability to activate an immense derivative space.

This amplification constitutes Language Leverage.

Importantly, leverage is neither inherently beneficial nor inherently harmful.

The same generative mechanism capable of producing scientific hypotheses, educational explanations, artistic innovation, and philosophical reflection may also generate misinformation, fabricated references, ideological manipulation, and persuasive hallucinations.

The ethical question therefore concerns not leverage itself but the conditions governing its grounding and use.

5.4 Financial Leverage and Linguistic Leverage

Financial leverage amplifies economic outcomes.

Language Leverage amplifies semantic outcomes.

Although these systems operate within different domains, their structural organization exhibits striking correspondences.

Both rely upon an underlying foundation.

Both project future possibilities into present operations.

Both amplify limited inputs.

Both generate nonlinear consequences.

Both remain dependent upon maintaining an appropriate connection with reality.

In financial markets, excessive leverage produces speculative bubbles and systemic crises.

Within generative AI, excessive Language Leverage produces derivative linguistic structures detached from Reality Grounding.

Hallucination therefore represents the linguistic analogue of over-leveraged speculation.

This comparison should not be interpreted as a simple metaphor.

Rather, it identifies a shared relational architecture governing amplification across heterogeneous domains.

5.5 Artistic Creation as Leverage

Art also operates through leverage.

The material dimensions of an artwork are often remarkably limited.

A painting consists of pigments on a surface.

A poem may comprise only a few lines.

An installation may employ ordinary everyday objects.

Yet their cultural consequences may extend across generations.

This disproportion between material input and existential transformation constitutes artistic leverage.

The value of art therefore cannot be explained solely by its physical properties.

Its transformative capacity emerges through the interfaces connecting artworks with viewers, institutions, histories, memories, and future possibilities.

Artistic creation amplifies perception rather than capital.

It generates new modes of seeing, remembering, imagining, and inhabiting reality.

From this perspective, art functions as one of the most sophisticated leverage systems developed by human culture.

5.6 The Ethics of Leverage

Because leverage amplifies consequences, it simultaneously amplifies responsibility.

Small interventions may generate extensive social transformations.

This principle becomes particularly significant in the age of generative AI.

Prompts, algorithmic adjustments, interface design, recommendation systems, and institutional policies may influence millions of users through relatively minor technical decisions.

Ethics can therefore no longer be understood exclusively in terms of individual intentions.

Responsibility becomes distributed throughout networks of interfaces.

Developers, users, researchers, educators, governments, cultural institutions, and technological infrastructures all participate in determining how leverage operates.

The ethical objective is not to eliminate amplification.

Rather, it is to design conditions under which amplification remains continuously connected to reality through appropriate forms of grounding, transparency, accountability, and critical reflection.

5.7 Toward a General Theory of Leverage

The analyses presented throughout this chapter suggest that leverage should be understood as a general ontological principle rather than a concept restricted to financial engineering.

Whenever limited actions generate disproportionately extensive transformations, leverage is present.

Language, science, finance, art, education, politics, technology, and artificial intelligence all exhibit this fundamental characteristic.

The present study therefore proposes the concept of the Leverage Structure as a universal framework for describing generative amplification across heterogeneous domains.

Within this framework,

- * Language Derivatives explain how future linguistic possibilities emerge.
- * Derivative Space explains where these possibilities unfold.
- * Language Leverage explains how they become amplified.
- * Reality Grounding explains how they reconnect with empirical and existential reality.
- * Interface Ontology explains how these transformations become socially and materially operative.

Together, these concepts form an integrated theoretical architecture extending beyond disciplinary boundaries.

5.8 Chapter Conclusion

This chapter has argued that leverage constitutes a universal generative principle connecting finance, language, artificial intelligence, and artistic creation.

Rather than interpreting leverage solely as an economic mechanism, the present study has demonstrated its broader ontological significance as the amplification of future possibilities through limited present interventions.

Within this perspective, AI hallucination appears not merely as a technical failure but as one possible consequence of derivative amplification exceeding the conditions necessary for Reality Grounding.

The following chapter synthesizes these concepts within Generative Topological Ontology, demonstrating how Language Derivatives, Inverse Problem Art, Derivative Space, and the Leverage Structure collectively contribute to a comprehensive ontology of existence in the age of artificial intelligence.

Chapter 6

Integration into Generative Topological Ontology

Self-Externalization, Pre-Existentiality, Interface Ontology, and Reality Grounding

6.1 Introduction

The preceding chapters have introduced the concepts of Language Derivatives, Inverse Problem Art, Derivative Space, and the Leverage Structure. Each concept has been developed independently in order to clarify a particular aspect of generative artificial intelligence and its relation to language, creativity, and reality.

The present chapter integrates these concepts within the broader philosophical framework of Generative Topological Ontology.

Rather than treating language, artificial intelligence, finance, art, and human existence as separate domains, Generative Topological Ontology understands them as different manifestations of a common generative structure. Existence is not conceived as a fixed substance but as an ongoing process in which relations, interfaces, and differences continually generate new realities.

Within this framework, hallucination is no longer interpreted merely as an informational error. It becomes a phenomenon revealing the generative dynamics through which reality itself is continuously reconstructed.

6.2 Self-Externalization and Derivative Generation

One of the central concepts of Generative Topological Ontology is Self-Externalization.

Human beings never exist as completely self-contained entities.

Thought, language, technology, institutions, artistic practice, and artificial intelligence all function as processes through which aspects of the self are externalized into the world.

Writing externalizes memory.

Architecture externalizes habitation.

Scientific theories externalize reasoning.

Art externalizes perception.

Generative AI externalizes linguistic possibility.

Derivative language therefore represents one mode of self-externalization.

When a prompt is entered into a generative model, the user's intentions, expectations, assumptions, and uncertainties are partially externalized into a linguistic system that responds by producing new possibilities unavailable to the individual alone.

Hallucinations demonstrate that self-externalization is never perfectly controllable.

Once externalized, language acquires relative autonomy.

Derivative structures may therefore exceed, transform, or even contradict the intentions from which they originated.

This excess should not be understood merely as failure.

Rather, it reveals the fundamentally generative nature of externalization itself.

6.3 Pre-Existentiality and the Emergence of Possibility

Generative Topological Ontology further proposes the concept of Pre-Existentiality.

Pre-Existentiality refers to the domain in which possibilities exist prior to their stabilization as identifiable beings.

This domain should not be interpreted as a metaphysical realm existing independently of experience.

Instead, it designates the generative condition from which determinate realities emerge.

Derivative language belongs precisely to this pre-existential domain.

Before generated language becomes verified knowledge, artistic creation, scientific discovery, or misinformation, it exists as possibility.

Generative AI continuously operates within this region.

Each response represents the temporary stabilization of one trajectory among innumerable unrealized alternatives.

Hallucination therefore should not be regarded solely as deviation from reality.

It is also evidence that language continues to inhabit the pre-existential field from which reality itself is generated.

6.4 Interface Ontology as the Principle of Generation

Within the present theoretical framework, Interface Ontology functions as the generative principle connecting derivative possibilities with stabilized reality.

An interface is not simply a boundary separating independent entities.

Rather, it is the dynamic condition through which heterogeneous domains become mutually generative.

Self and other.

Language and embodiment.

Artificial intelligence and human judgment.

Possibility and actuality.

Art and everyday life.

Each of these relations becomes productive only through interfaces capable of preserving difference while enabling connection.

Reality therefore does not precede the interface.

Reality is continually regenerated through interfaces.

Hallucinations reveal moments in which these interfaces become unstable.

Conversely, artistic creation intentionally constructs interfaces capable of transforming instability into new forms of understanding.

6.5 Reality Grounding as Generative Stabilization

Reality Grounding has been discussed throughout the preceding chapters primarily as the condition preventing derivative language from becoming detached from reality.

Within Generative Topological Ontology, however, grounding possesses a more profound significance.

Grounding should not be understood merely as verification after generation.

Rather, it constitutes the ongoing process through which reality itself becomes stabilized.

Scientific experiments stabilize hypotheses.

Educational practice stabilizes knowledge.

Institutions stabilize social norms.

Embodied action stabilizes perception.

Art stabilizes previously unseen modes of experience.

Likewise, AI-generated language becomes meaningful only insofar as it participates in these stabilizing processes.

Grounding therefore does not oppose generation.

It completes generation.

Without derivative generation there can be no innovation.

Without grounding there can be no reality.

Reality itself emerges through the continuous interaction of these complementary processes.

6.6 Artificial Intelligence and the Ontology of Generation

Generative artificial intelligence occupies a distinctive position within this theoretical framework.

Unlike conventional computational systems, generative AI does not merely manipulate existing representations.

It continuously generates new linguistic configurations whose future significance remains fundamentally indeterminate.

AI therefore should not be regarded exclusively as an instrument for answering questions.

Its deeper philosophical significance lies in revealing the generative character of language itself.

Hallucinations expose this structure with particular clarity.

They demonstrate that language is never simply reproductive.

It continually exceeds its own foundations.

The central philosophical challenge is therefore not to suppress this excess completely but to cultivate interfaces capable of directing derivative generation toward constructive forms of reality.

This represents a shift from the engineering problem of error reduction to the ontological problem of generative responsibility.

6.7 Chapter Conclusion

This chapter has integrated the principal concepts developed throughout the present study within the framework of Generative Topological Ontology.

Language Derivatives describe the generation of unrealized possibilities.

Derivative Space describes the relational field in which these possibilities unfold.

Language Leverage explains their amplification.

Reality Grounding explains their stabilization.

Interface Ontology explains their connection.

Self-Externalization explains their human origin.

Pre-Existentiality explains their ontological horizon.

Taken together, these concepts establish a unified theory in which hallucination, creativity, language, and artificial intelligence are understood not as isolated phenomena but as interconnected expressions of a single generative structure.

The final chapter considers the broader implications of this framework for philosophy, artificial intelligence, aesthetics, and the future of human existence in an increasingly AI-mediated world.

Chapter 7

Toward an Ontology for the Age of Artificial Intelligence

Language Derivatives, Reality, and the Future of Human Existence

7.1 Introduction

The emergence of generative artificial intelligence represents more than a technological innovation. It marks a profound transformation in humanity's relationship with language, knowledge, creativity, and reality itself.

Previous chapters have argued that hallucination should not be regarded merely as a computational defect. Rather, it reveals a fundamental generative structure through which derivative language projects future possibilities into the present. Understanding this structure requires moving beyond conventional discussions of accuracy and error toward a broader ontology of generation.

The present chapter considers the philosophical implications of this theoretical framework and examines how Language Derivatives may contribute to an ontology appropriate for the age of artificial intelligence.

7.2 From Representation to Generation

Modern philosophy has generally treated language as a system of representation.

Whether in analytic philosophy, structural linguistics, or information theory, language has often been understood as a medium for describing an independently existing reality.

Generative artificial intelligence challenges this assumption.

Its outputs do not merely represent reality.

They actively participate in generating new conceptual relations, narratives, expectations, and possibilities.

Language therefore cannot be understood solely as representational.

It must also be understood as generative.

This transition from representation to generation constitutes one of the defining philosophical shifts of the AI era.

The Theory of Language Derivatives proposes that this generative character is not unique to artificial intelligence but has always been implicit within human language itself.

AI simply renders this latent structure visible.

7.3 Human Beings Within Derivative Space

If language continuously generates derivative possibilities, then human beings themselves inhabit Derivative Space.

Individuals do not merely react to present circumstances.

They continuously organize their actions through imagined futures, anticipated possibilities, memories, promises, expectations, and narratives.

Human existence is therefore fundamentally derivative.

Civilizations are organized through constitutions, myths, scientific theories, educational systems, religious traditions, legal institutions, artistic practices, and technological visions.

None of these exist purely as physical objects.

Their power lies in their capacity to organize future possibilities within the present.

Generative AI now participates directly in this historical process.

It expands the scale and velocity through which derivative language circulates throughout society.

Consequently, humanity enters an unprecedented condition in which derivative generation becomes technologically accelerated.

7.4 Hallucination as an Ontological Event

Within this broader framework, hallucination acquires new philosophical significance.

It is not merely an incorrect statement requiring correction.

Nor is it simply evidence of insufficient computational precision.

Rather, hallucination reveals the ontological openness of language itself.

Language is capable of producing meanings that exceed currently stabilized reality.

Some of these meanings become scientific discoveries.

Some become artistic innovations.

Some become political ideologies.

Some become myths.

Some disappear without consequence.

Others become misinformation.

Hallucination therefore belongs to the broader process through which reality continually differentiates itself from possibility.

Its importance lies not in the falsehood of individual statements but in its demonstration that language always exceeds complete determination.

7.5 Responsibility in the Age of Generative Intelligence

The expansion of derivative language inevitably transforms the concept of responsibility.

Responsibility can no longer be assigned exclusively to individual speakers or technological systems.

Generative language emerges through complex interactions among users, developers, institutions, educational practices, media environments, economic incentives, and cultural traditions.

Responsibility therefore becomes relational rather than isolated.

The central ethical challenge is not merely preventing hallucination.

Rather, it is cultivating interfaces capable of guiding derivative generation toward constructive forms of Reality Grounding.

This requires transparency, critical literacy, interdisciplinary collaboration, and continuous reflection concerning the social consequences of derivative amplification.

The future of AI ethics therefore depends not only upon improving algorithms but also upon redesigning the interfaces through which human beings encounter derivative language.

7.6 Toward a New Ontology of Language

The Theory of Language Derivatives ultimately proposes a new understanding of language itself.

Language is neither a passive mirror of reality nor merely a symbolic communication system.

It is a generative medium through which reality is continuously reorganized.

Every utterance introduces new possibilities.

Every narrative reconfigures future expectations.

Every dialogue transforms the relational field in which meaning emerges.

Artificial intelligence does not replace this human capacity.

Instead, it externalizes and accelerates it.

The philosophical task is therefore not to defend humanity against AI as though they were opposing entities.

Rather, it is to understand how human beings and artificial intelligence jointly participate in the ongoing generation of reality.

7.7 Chapter Conclusion

This chapter has argued that the age of artificial intelligence demands an ontology centered upon generation rather than representation.

Language Derivatives provide a conceptual framework through which hallucination, creativity, narrative, and technological mediation may be understood as interconnected aspects of a common generative process.

Artificial intelligence does not simply produce answers.

It expands the field of derivative possibilities within which human beings think, create, decide, and act.

Accordingly, the future of AI should not be evaluated solely according to its capacity for factual accuracy.

Its deeper significance lies in its capacity to reshape the ontological conditions through which reality itself is generated.

The Theory of Language Derivatives therefore concludes that hallucination is neither an accidental defect nor a purely technical obstacle. It is an expression of the derivative nature of language itself. The central philosophical task of the AI era is not the elimination of derivative generation, but the cultivation of interfaces through which derivative possibilities become responsibly grounded in reality.

Conclusion

This paper has proposed The Theory of Language Derivatives as a new theoretical framework for understanding hallucinations in generative artificial intelligence. Rather than interpreting hallucination merely as a technical malfunction, a factual error, or an engineering problem to be eliminated, the study has argued that hallucination emerges from the derivative structure of language itself.

Beginning with a structural reinterpretation of financial derivatives, the paper demonstrated that both financial and linguistic systems derive present operations from unrealized future possibilities. This common generative principle was conceptualized through the notions of Language Derivatives and Language Leverage, revealing that linguistic generation, artistic creation, scientific hypothesis formation, and AI inference all operate through processes of derivative amplification.

The analysis further demonstrated that hallucination cannot be adequately explained without distinguishing between forward problems and inverse problems. While users typically approach generative AI as though it were solving forward problems, the model itself operates fundamentally through inverse inference under conditions of incomplete information. Hallucination therefore appears not as an accidental failure but as a structural consequence of probabilistic reconstruction within derivative language generation.

To clarify the creative implications of this structure, the paper introduced Inverse Problem Art. Artistic practice was interpreted not as the expression of predetermined meanings but as the reconstruction of hidden structures from incomplete reality. Although artistic creation and AI hallucination share a common inverse-problem structure, they diverge through their relationship to Reality Grounding. Art intentionally constructs interfaces through which derivative possibilities become embodied, socially mediated, and existentially meaningful, whereas hallucination remains disconnected whenever such interfaces fail to emerge.

These analyses were subsequently integrated through the concepts of Derivative Space, Interface Ontology, and Generative Topological Ontology. Derivative Space was defined as the relational field within which future possibilities become operative in the present through narrative and linguistic interaction. Interface Ontology explained how heterogeneous domains—including language, embodiment, technology, institutions, and reality—become mutually generative through dynamic relational boundaries. Generative Topological Ontology provided the broader philosophical framework within which these processes could be understood as expressions of continuous existential generation rather than static ontological states.

From this perspective, hallucination is neither an accidental computational defect nor a phenomenon unique to artificial intelligence. Instead, it reveals a more fundamental characteristic of language itself: language continuously exceeds existing reality by generating unrealized possibilities. Scientific discoveries, artistic innovations, technological inventions, social institutions, myths, and fictional narratives all emerge from this same derivative capacity. Hallucination therefore occupies a continuous spectrum connecting imagination, hypothesis, creativity, and misinformation.

Accordingly, the principal challenge for the age of artificial intelligence is not the complete elimination of hallucination. Such an objective would require eliminating the derivative and generative capacities that simultaneously make creativity, discovery, and innovation possible. The more fundamental task is the design of interfaces capable of reconnecting derivative linguistic generation with empirical, embodied, institutional, social, and ontological forms of Reality Grounding.

The Theory of Language Derivatives therefore shifts the central philosophical question of artificial intelligence. Rather than asking how machines can perfectly reproduce reality, it asks how derivative language participates in the continual generation of reality itself.

In this sense, artificial intelligence should not be understood merely as a computational technology. It represents a new stage in humanity's long history of self-externalization through language. Generative AI extends the derivative capacity already inherent in human linguistic activity, making visible structures that have previously remained implicit within culture, science, art, and everyday communication.

Ultimately, this study argues that the future relationship between humanity and artificial intelligence depends less upon eliminating uncertainty than upon cultivating responsible forms of derivative generation. The question is not whether language will continue to exceed reality—it always has—but whether human beings can design interfaces through which those derivative possibilities become ethically, socially, creatively, and existentially grounded.

Within this framework, hallucination ceases to be merely an error to be corrected. It becomes a philosophical phenomenon through which the generative nature of language, reality, and human existence itself can be reconsidered.

References

- Austin, J. L. 1962. *How to Do Things with Words*. Oxford: Oxford University Press.
- Baudrillard, Jean. 1981. *Simulacres et Simulation*. Paris: Éditions Galilée.
- Black, Fischer, and Myron Scholes. 1973. "The Pricing of Options and Corporate Liabilities." *Journal of Political Economy* 81 (3): 637–654.
- Boden, Margaret A. 2004. *The Creative Mind: Myths and Mechanisms*. 2nd ed. London: Routledge.
- Bourriaud, Nicolas. 2002. *Relational Aesthetics*. Dijon: Les Presses du Réel.
- Clark, Andy. 1997. *Being There: Putting Brain, Body, and World Together Again*. Cambridge, MA: MIT Press.
- Floridi, Luciano. 2011. *The Philosophy of Information*. Oxford: Oxford University Press.
- Floridi, Luciano. 2014. *The Fourth Revolution: How the Infosphere Is Reshaping Human Reality*. Oxford: Oxford University Press.
- Foucault, Michel. 1966. *Les Mots et les choses*. Paris: Gallimard.
- Goodfellow, Ian, Yoshua Bengio, and Aaron Courville. 2016. *Deep Learning*. Cambridge, MA: MIT Press.
- Haraway, Donna J. 1991. *Simians, Cyborgs, and Women*. New York: Routledge.
- Heidegger, Martin. 1927. *Sein und Zeit*. Tübingen: Max Niemeyer.
- Heidegger, Martin. 1954. *Vorträge und Aufsätze*. Pfullingen: Neske.
- Kimura, Kenjiro. 2014. *Inverse Problems and Their Applications*. Tokyo: Kyoritsu Shuppan. (Japanese)

Lakoff, George, and Mark Johnson. 1980. *Metaphors We Live By*. Chicago: University of Chicago Press.

Latour, Bruno. 2005. *Reassembling the Social*. Oxford: Oxford University Press.

Merleau-Ponty, Maurice. 1945. *Phénoménologie de la perception*. Paris: Gallimard.

Merton, Robert C. 1973. "Theory of Rational Option Pricing." *Bell Journal of Economics and Management Science* 4 (1): 141–183.

Russell, Stuart, and Peter Norvig. 2021. *Artificial Intelligence: A Modern Approach*. 4th ed. Hoboken, NJ: Pearson.

Searle, John R. 1969. *Speech Acts*. Cambridge: Cambridge University Press.

Shannon, Claude E. 1948. "A Mathematical Theory of Communication." *Bell System Technical Journal* 27 (3–4): 379–423, 623–656.

Wiener, Norbert. 1948. *Cybernetics: Or Control and Communication in the Animal and the Machine*. Cambridge, MA: MIT Press.

Wittgenstein, Ludwig. 1953. *Philosophical Investigations*. Oxford: Blackwell.

Author's Previous Works

Tanaka, Tomokazu (bigakusha-haha). 2010. *Web Art Declaration*. Osaka.

Tanaka, Tomokazu (bigakusha-haha). 2014. *Smart Art Declaration*. Osaka.

Tanaka, Tomokazu (bigakusha-haha). 2016. *Quantum Art Manifesto*. Osaka.

Tanaka, Tomokazu (bigakusha-haha). 2018. Derivative Art Manifesto. Osaka.

Tanaka, Tomokazu (bigakusha-haha). 2022. Decentralized Autonomous Art (DAA). Osaka.

Tanaka, Tomokazu (bigakusha-haha). 2026. Axioms of Generative Topological Ontology. Research Memorandum.

Tanaka, Tomokazu (bigakusha-haha). 2026. Self-Externalization and Generative Existence: Interface Ontology in Generative Topological Ontology. Research Memorandum.

Research Method

A Methodological Statement on the Publication of Research Memoranda

This text is published as a Research Memorandum within the continuing development of Generative Topological Ontology.

The concepts, definitions, hypotheses, and arguments presented here record the theoretical position reached at the time of writing and publication. They should not be understood as a finally established theoretical system. They may be expanded, revised, redefined, or reorganized through future research, artistic practice, interaction with reality, academic dialogue, and peer review.

This research does not regard only completed theories as legitimate research outcomes. It also treats the processes through which concepts emerge, become interconnected, undergo transformation, and are integrated into a theoretical structure as significant objects of inquiry. Accordingly, this memorandum should be understood not merely as an incomplete draft, but as a record of a particular phase in the generation of the theory.

The public release of a Research Memorandum is not intended to present a developing theory prematurely as a completed one. Rather, it is a method of documenting the generative processes of thought and practice, enabling their subsequent examination and reconstruction, and opening them to academic and artistic dialogue from different perspectives.

If the contents of this memorandum are subsequently developed into a peer-reviewed article, scholarly monograph, or other formal research output, the finalized work will be identified as a separate and fixed document.

When citing or referring to this memorandum, please include the title, author, version number, and date of the most recent update. Different versions may contain substantive differences in concepts, definitions, and lines of argument.

This Research Memorandum is not a final conclusion. It is one theoretical position reached within the continuing formation of Generative Topological Ontology and, at the same time, an interface open to further generation.