

Research Monograph Series

言語デリバティブ論

—人工知能のハルシネーションにおける「逆問題芸術」とレバレッジ構造—

The Theory of Language Derivatives: Inverse Problem Art and the Leverage Structure of AI Hallucinations

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 AI Hallucinations

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

要旨

本論文は、生成AIにおけるハルシネーションを、単なる情報処理上の誤り、事実認識の失敗、あるいは除去されるべき技術的欠陥としてのみ捉える従来の理解を再検討し、それを「言語デリバティブ」と呼ぶ生成的構造の一現象として理論化することを目的とする。ここで言語デリバティブとは、既存の事実や対象を直接表象する言語ではなく、まだ現実化していない意味、関係、価値、未来可能性を現在の言語空間へ先行的に畳み込み、その可能性を増幅させる派生的言語構造を意味する。

金融デリバティブが原資産から派生しつつ、将来の価格変動を現在に織り込み、限定された資本によって大きな利益または損失を生じさせるレバレッジ構造を持つように、生成AIの出力もまた、学習データという既存の言語的原資産から派生しながら、そこに明示的には存在しない文章、概念、因果関係、物語を生成する。この生成は、既存情報の単純な再現ではない。人工知能は、膨大な言語的關係を圧縮し、入力された限定的なプロンプトを契機として、その潜在的關係を一つの応答へ展開する。その意味でプロンプトは、言語空間に対する小規模な投機であり、生成結果は、その投機によって増幅された言語的ポジションとして理解できる。

本論文は、この構造を「言語レバレッジ (Language Leverage)」と定義する。言語レバレッジとは、限定された入力、学習データ内の広大な關係空間を作動させ、入力をはるかに超える意味作用、認知的効果、社会的影響を生み出す非対称的増幅作用である。生成AIのハルシネーションは、このレバレッジが既存の事實的参照を超えて作動したときに顕在化する。したがって、ハルシネーションとは、単に「存在しないものを存在すると述べること」ではなく、言語の派生的生成が現実との接続条件を失った状態として再定義されるべきである。

この再定義のため、本論文は「順問題」と「逆問題」の區別を導入する。順問題においては、既知の原因、法則、条件から正しい結果を導出することが求められる。この枠組みでは、既定の正解から逸脱する人工知能の応答は明確な誤りとなる。これに対して逆問題においては、観測された断片、結果、不完全な情報から、その背後にある構造、原因、意味、可能世界を推定しなければならない。ここでは解は一意に定まらず、複数の可能性、不安定性、非可解性が構造的に含まれる。生成AIの言語生成は、しばしば順問題的な回答装置として利用されながら、実際には逆問題的な推定機構として作動している。この機構上の非対称性が、ハルシネーションを発生させる根本条件の一つである。

しかし、本論文は、逆問題的生成をそれ自体として否定しない。むしろ美学者母が提唱してきた「逆問題芸術」の観点から、正解があらかじめ存在しない領域における生成を、芸術的・認識論的・存在論的な可能性として捉える。逆問題芸術とは、既存の問題に正解を与える制作ではなく、断片的な現実から、それまで認識されていなかった問

題、構造、関係、自己、世界の可能性を逆算的に生成する実践である。その目的は、虚偽を事実として固定することではなく、未決定の可能性が現実と接続しうる条件を設計することにある。

この観点に立てば、人工知能のハルシネーションと芸術的生成は、いずれも既存の現実を超えて可能性を生成する点では共通する。しかし両者は同一ではない。人工知能は、統計的・言語的關係から「自己以外」の可能性を生成する。他方、芸術は、その自己以外との遭遇を通して、知覚、身体、行為、関係、自己認識を変容させる。したがって、人工知能が他者的可能性を外部化する装置であるならば、芸術は、その可能性を自己と現実の変容へ接続するインターフェイスである。

ここで中心的な役割を果たすのが、生成の位相存在論におけるインターフェイス存在論である。本論文においてインターフェイスとは、人工知能と人間、言語と現実、可能性と事実、自己と自己以外を単純に媒介する中立的な道具ではない。それは、異なる位相間の差異を保持しながら、新たな接続と存在を生成する出来事的条件である。ハルシネーションが問題となるのは、生成された言語が現実そのものではないからではなく、生成物と現実との位相差が隠蔽され、仮説、比喩、推定、創作、事実の区別が失われるからである。ゆえに必要なのは、ハルシネーションを全面的に排除することだけではなく、その生成物がどのような位相に属し、どのような接続条件のもとで現実へ関与しうるのかを可視化することである。

本論文は、この接続過程を「現実接地 (Reality Grounding)」として論じる。現実接地は、生成された言語が既存の事実と一致することだけを意味しない。それは、生成された可能性が検証、身体的経験、社会的実践、他者との関係、制作行為、制度的手続きなどを通して、現実の関係構造に具体的な変化を生じさせる過程である。したがって、ハルシネーションと創造性を分ける境界は、生成内容の新奇性そのものではなく、その内容が現実に対する応答責任と接続可能性を持つかどうかにある。現実接地を欠いた言語レバレッジは、虚偽、誤認、権威の擬態、情動的投機を増幅する。これに対し、現実接地を伴う言語レバレッジは、仮説形成、理論生成、芸術制作、自己変容、新たな社会関係の形成を可能にする。

以上を通じて、本論文は、人工知能のハルシネーションを技術的欠陥と創造的可能性のいずれか一方へ還元するのではなく、両者を生み出す共通の生成構造を明らかにする。言語デリバティブは、未来の意味を現在へ先行的に投機し、限定された入力から過剰な意味作用を発生させる。そのレバレッジは、現実接地を欠く場合には誤情報と擬似的権威を増幅するが、インターフェイスによって位相差が保持され、現実との再接続が行われる場合には、未知の自己、未知の問題、未知の世界を生成する契機となる。

本論文の結論は、人工知能時代に必要なのは、ハルシネーションの完全な消去という不可能な純化ではなく、生成された言語の派生性、投機性、非確定性を自覚し、それを現実へ接地させるインターフェイスを設計することである、という点にある。これによ

り、人工知能は正解を供給する閉鎖的装置としてではなく、人間が自己以外の可能性と遭遇し、自己と現実を再生成するための開かれた言語的環境として位置づけ直される。

キーワード

言語デリバティブ、生成AI、ハルシネーション、逆問題芸術、レバレッジ構造、言語レバレッジ、インターフェイス存在論、生成の位相存在論、現実接地、自己外部化、デリバティブ空間、ナラティブ、人工知能倫理、芸術的生成、可能世界

目次

Contents

第0章 Methodological Note（方法論的序論）	1
第1章 言語デリバティブ論の理論的背景	8
第2章 ハルシネーションの存在論.....	15
第3章 逆問題芸術.....	25
第4章 言語デリバティブ空間	40
第5章 レバレッジ構造.....	61
第6章 生成の位相存在論との統合.....	83
第7章 AI時代の存在論	106
結論.....	139
参考文献	149
方法論的声明	152

第0章 Methodological Note（方法論的序論）

0.1 本論文の方法論的位置

本論文は、生成AIにおけるハルシネーションを、既存の情報科学、認知科学、言語哲学、金融工学、芸術理論のいずれか一領域へ還元して説明するものではない。むしろ、それらの異なる領域に現れる構造的類似性を横断的に抽出し、その共通構造を「言語デリバティブ」として理論化する試みである。

ここで採用される方法は、単純な学際研究ではない。学際研究という語はしばしば、複数の既存分野から概念を借用し、それらを一つの対象へ適用する作業を意味する。しかし本論文の目的は、金融におけるデリバティブ概念を比喩的に人工知能へ適用することでも、逆問題という数学的概念を芸術へ転用することでもない。本論文が問うのは、金融デリバティブ、生成AI、言語生成、ハルシネーション、ナラティブ、芸術制作の間に、単なる比喩を超えた同型的な生成構造が存在するかどうかである。

したがって、本論文の方法論は、既存概念の外挿ではなく、異なる領域に反復して現れる関係構造の抽出として定義される。この方法を、本論文では暫定的に「位相的構造比較」と呼ぶ。位相的構造比較とは、対象の物質的内容、制度的所属、専門分野上の差異をいったん保留し、それぞれの対象が、どのような関係、境界、変換、増幅、接続、切断によって成立しているかを比較する方法である。

金融商品と生成言語モデルは、物質的にも制度的にも同一ではない。しかし、いずれも既存の基盤から派生し、将来的可能性を現在へ織り込み、限定された入力や資本を超える効果を発生させるという点において、一定の構造的対応を示す。本論文は、この対応を即座に同一性として扱うのではなく、差異を保持したまま比較可能な関係形式として記述する。

この点で、本論文は、生成の位相存在論の方法論的立場を継承する。生成の位相存在論において、存在は孤立した実体としてではなく、接続可能性、境界変化、関係生成、位相転換の過程として理解される。本論文もまた、ハルシネーションを単一の誤情報単位としてではなく、言語、モデル、プロンプト、利用者、制度、現実の間に生じる関係的出来事として捉える。

0.2 対象の再定義としての方法

本論文における第一の方法論的操作は、研究対象そのものを再定義することである。一般に、人工知能のハルシネーションは、事実に反する出力、不正確な記述、根拠のない回答、存在しない出典の提示などとして定義される。この定義は、技術評価、品質管理、安全性、情報信頼性の観点から不可欠である。しかし、それだけでは、なぜ生成AI

が、単なる検索失敗ではなく、整合的で説得的な文章として誤りを生成するのかを十分に説明できない。

ハルシネーションの特徴は、誤りであることだけではない。それは、正解の不在にもかかわらず、正解らしさを伴って生成される。さらに、その生成物はしばしば、利用者の入力、期待、文脈、語調、推論方向を反映しながら、入力内容を超える意味的展開を示す。この過剰性こそが、単なる情報誤りと生成AIのハルシネーションを区別する。

ゆえに本論文は、ハルシネーションを「誤った情報」ではなく、「現実接地を欠いた派生的言語生成」として再定義する。この再定義により、研究対象は個々の誤文から、その誤文を可能にする生成構造へ移行する。

ここでいう派生的生成とは、言語が既存の対象を直接写し取るのではなく、すでに存在する言語関係から新たな言語関係を生成することを意味する。生成AIは、現実そのものへ直接アクセスするのではなく、主として言語化されたデータ、記号的パターン、統計的關係、入力文脈を通じて応答を生成する。その出力は、現実の一次的複製ではなく、言語関係から派生した二次的、三次的構成物である。

この派生性を分析の中心に置くことによって、ハルシネーションは例外的な故障ではなく、生成AIの成立条件と連続した現象として捉え直される。すなわち、創造的文章生成とハルシネーションは、完全に別個の機構から生じるのではなく、共通の派生的生成機構が、異なる接地条件のもとで現れた結果である。

0.3 比喩と構造的同型性の区別

本論文では、「デリバティブ」「レバレッジ」「原資産」「投機」「ポジション」といった金融用語を用いる。しかし、これらは単なる修辭的比喩として使用されるのではない。

比喩は、ある対象を別の対象になぞらえることによって理解を促す。他方、構造的同型性の分析は、異なる対象の間に、関係の配置や変換の仕方における対応が存在するかを検討する。本論文が重視するのは後者である。

たとえば金融デリバティブにおいては、原資産そのものを直接所有しなくても、その価格変動、信用状態、将来予測に連動する派生的価値を取引することができる。同様に生成AIは、現実の対象そのものではなく、対象について蓄積された言語的關係から、派生的な文章や意味を生成する。

両者はもちろん同一ではない。金融デリバティブは法的契約、価格形成、市場制度、決済構造に依存する。他方、言語生成は計算モデル、学習データ、確率分布、利用文脈に依存する。したがって両者を直接同一視することは、分析上の誤りとなる。

しかし、両者には少なくとも四つの構造的対応が認められる。

第一に、両者は基盤から派生する。

第二に、両者は将来的可能性を現在の判断へ織り込む。

第三に、両者は限定された入力から過剰な効果を生じさせる。

第四に、両者は基盤との接続が維持される場合に有効に機能し、接続が失われる場合に損失、誤認、破綻を生じさせる。

本論文は、この四つの対応を出発点として、「言語デリバティブ」という概念を構築する。ゆえに、この概念は文学的比喩であると同時に、分析概念でもある。その妥当性は、語の印象的類似ではなく、どの程度まで異なる現象を一貫して説明できるかによって判断されなければならない。

0.4 順問題と逆問題

本論文の第二の方法論的軸は、順問題と逆問題の区別である。

順問題とは、既知の原因、法則、初期条件から結果を求める問題である。入力と規則が与えられれば、一定の手続きによって出力を導くことができる。これに対して逆問題とは、観測された結果や断片的データから、その原因、構造、条件を推定する問題である。逆問題には、一般に非一意性、不安定性、情報不足が伴う。同じ結果を生み出しうる原因が複数存在し、わずかなデータ差が大きな推定差を生み、完全な復元が原理的または実践的に困難である場合が多い。

生成AIは、利用上はしばしば順問題的な装置として扱われる。利用者は質問を入力し、正しい答えが返されることを期待する。しかし、生成AIの内部では、入力に対して唯一の真理値を直接検索しているわけではなく、文脈に続く可能な言語系列を推定し、複数の潜在的候補から一つの応答を生成する。

この意味で、生成AIは、順問題的な期待のもとで使用される逆問題的装置である。この方法論的不一致は、ハルシネーション理解にとって決定的である。

利用者が「正解」を求めている場合でも、モデルはしばしば「最も文脈的に成立しやすい言語構成」を生成する。ここで、事実的正確性と言語的妥当性は一致しない。むしろ、言語的整合性が高いほど、誤りが説得的に見える場合がある。

本論文は、この不一致をモデルの欠陥としてのみ捉えず、生成AIがもともと内包する逆問題性の現れとして分析する。

0.5 逆問題芸術の方法論的導入

「逆問題芸術」は、本論文における単なる応用概念ではなく、分析方法そのものでもある。

通常、芸術制作は、作者の意図、表現内容、素材、形式、鑑賞結果の順に説明される。この説明は、制作主体を原因とし、作品を結果とする順問題的構図を前提としている。しかし実際の芸術経験では、作品や出来事が先に現れ、そこから作者、鑑賞者、文脈、

問題、意味が遡及的に再構成されることがある。

逆問題芸術とは、完成した意味を表現することではなく、断片的な出来事、違和感、偶然、接続不全、認識のずれから、その背後に潜む問題構造を生成する実践である。そこでは、解答よりも問題の生成が先行する。

本論文が人工知能のハルシネーションを逆問題芸術と接続するのは、誤情報を芸術として肯定するためではない。重要なのは、人工知能が生成した不確定な言語を、確定した事実として受容するのではなく、いかなる問題、欲望、前提、社会的期待、認識構造がその出力を成立させたのかを逆算することである。

すなわち、ハルシネーションは、除去すべきノイズであると同時に、生成システムと利用者の間に存在する不可視の前提を露呈させる徴候でもある。

たとえば、存在しない論文を人工知能が生成した場合、その問題は単に論文名が誤っていることにとどまらない。なぜその題名がもっともらしく感じられたのか。なぜ利用者はそれを信じたのか。学術的権威はどのような言語形式によって擬態されるのか。知識の外観と知識そのものは、どのように区別されるのか。逆問題芸術は、このような問いを生成する。

したがって、本論文における逆問題芸術は、ハルシネーションの創造的利用法ではなく、ハルシネーションによって露呈する構造を可視化し、新たな認識条件へ接続する方法である。

0.6 現実接地の多層性

本論文では、現実接地を単なる事実確認と同一視しない。

人工知能研究において、接地はしばしば、言語記号を外部世界の対象、感覚入力、知識データベース、検証可能な情報へ結びつけることを意味する。この意味での接地は重要であり、ハルシネーション抑制に不可欠である。

しかし本論文では、現実接地をより広い多層的過程として捉える。少なくとも、以下の接地を区別する必要がある。

第一は、事実に接地である。これは出力内容が検証可能な情報と一致しているかを問う。

第二は、身体的接地である。これは言語が、感覚、行為、身体経験、物理的環境とどのように関係するかを問う。

第三は、社会的接地である。これは言語が、制度、共同体、責任、権力、信頼関係の中でどのように作動するかを問う。

第四は、実践的接地である。これは生成された仮説や意味が、現実の行為、制作、検証、判断、関係変化へ結びつくかを問う。

第五は、存在論的接地である。これは生成された言語が、自己と自己以外、可能性と現

実、既知と未知の間に、新たな接続を生み出すかを問う。

この多層性を導入することで、創作と誤情報の差異も精密化できる。創作は、事実に接地を必ずしも要求しない。しかし、作品、作者、鑑賞者、制度、身体、文脈との接地を持つ。他方、ハルシネーションは、事実に誤っていても、社会的権威や利用者の期待に強く接地してしまう場合がある。

したがって、問題は接地の有無ではなく、どの層へ、どのような仕方で接地しているかである。本論文は、この接地層の不一致を、ハルシネーションの重要な条件として扱う。

0.7 インターフェイス存在論

本論文の第三の方法論的軸は、インターフェイス存在論である。

一般にインターフェイスは、人間と機械、主体と情報、入力と出力を媒介する表面として理解される。しかし、本論文におけるインターフェイスは、すでに存在する二項を単に接続する中立的な通路ではない。

インターフェイスとは、異なる位相が接触することで、新たな関係、新たな意味、新たな主体位置が生成される出来事的境界である。

生成AIとの対話において、利用者は既成の情報を受け取るだけではない。入力の仕方、応答の読み方、信頼の程度、再質問、修正、引用、実践への適用を通して、生成物の意味と現実性を共同的に構成する。

したがって、ハルシネーションは人工知能内部だけで完結する現象ではない。それは、モデル、プロンプト、利用者、表示形式、利用状況、社会的権威、検証手段の間に成立するインターフェイス的現象である。

この観点に立てば、ハルシネーションへの責任も、モデル単体へ還元することはできない。もちろん、開発者には正確性、安全性、透明性を向上させる責任がある。しかし同時に、利用者、提供企業、教育制度、メディア環境、専門家共同体も、生成言語がどのように現実化するかに関与している。

インターフェイス存在論は、責任を分散させて曖昧にするための議論ではない。むしろ、責任が発生する接続点を特定するための方法である。

0.8 記述的分析と規範的分析

本論文は、記述的分析と規範的分析を区別しながら進める。

記述的分析では、生成AIがどのような構造で言語を生成し、どのような条件でハルシネーションが発生し、どのような仕方で利用者に受容されるかを考察する。

規範的分析では、その生成物をどのように扱うべきか、どのようなインターフェイスを設計すべきか、どのような責任概念が必要かを論じる。

この二つを混同すると、存在する技術的構造をそのまま肯定する技術決定論か、逆に、望ましくない結果だけを基準に生成機構全体を否定する道徳的単純化に陥る。本論文は、ハルシネーションを構造的に不可避な側面として分析するが、それを無条件に容認しない。また、正確性を重要な価値として認めるが、言語生成を完全な事実再生へ還元することもしない。必要なのは、創造性と正確性、可能性と責任、非決定性と検証可能性の間に、単純な二者択一ではない接続条件を構築することである。

0.9 理論構築の手順

本論文は、以下の理論的手順によって構成される。

第一に、デリバティブ概念を金融的文脈から抽出し、その構造的特徴を整理する。

第二に、生成AIの言語生成を、派生性、予測性、増幅性、非対称性の観点から分析する。

第三に、両者の対応関係から「言語デリバティブ」と「言語レバレッジ」を定義する。

第四に、ハルシネーションを、言語デリバティブと現実接地の不一致として再定義する。

第五に、順問題と逆問題の区別を通して、生成AIに対する利用者の期待と、実際の生成機構のずれを分析する。

第六に、逆問題芸術を導入し、ハルシネーションを未知の問題構造を露呈させる契機として再配置する。

第七に、インターフェイス存在論を通して、生成物が現実へ接続される条件、責任、倫理を論じる。

第八に、これらの議論を生成の位相存在論へ統合し、言語、人工知能、芸術、現実生成の関係を一般理論として提示する。

この手順は、単線的な演繹ではない。各概念は、他の概念との関係によって相互に修正される。したがって、本論文の理論構築は、出発点から最終結論へ直線的に進むというより、概念間の再帰的往還によって形成される。

0.10 本論文の射程と限界

本論文は、生成AIの技術的機構を数理的に完全記述することを目的としない。また、金融デリバティブの価格理論、契約法、リスク管理を専門的に論じるものでもない。さらに、逆問題の数学的可解性条件や正則化理論を詳細に展開するものでもない。これらの領域は、それぞれ独立した高度な専門性を持つ。本論文は、それらの成果を置き換えるのではなく、異なる領域を横断する生成構造を存在論的・美学的に記述することを目的とする。

したがって、本論文の妥当性は、技術論文としての精度だけではなく、以下の点によって評価されるべきである。

第一に、言語デリバティブという概念が、人工知能のハルシネーションと創造性を同時に説明できるか。

第二に、レバレッジ構造という概念が、限定されたプロンプトから過剰な意味的・社会的効果が生じる仕組みを記述できるか。

第三に、逆問題芸術が、誤情報の肯定へ陥ることなく、ハルシネーションを認識的契機として扱えるか。

第四に、インターフェイス存在論が、人工知能、人間、制度、現実の間に分散する責任を具体的に記述できるか。

第五に、生成の位相存在論との統合が、単なる概念の追加ではなく、人工知能時代における存在生成の条件を明らかにできるか。

本論文は、これらの問いに最終的かつ閉鎖的な解答を与えるものではない。むしろ、人工知能のハルシネーションという現象を通して、言語がいかに現実から派生し、現実を逸脱し、再び現実を生成しうるのかを問うための理論的場を開くものである。

0.11 方法論的結語

本論文の方法論的前提は、ハルシネーションを誤りとして認定するだけでは、その存在論的意味を捉えられないという点にある。

生成AIの出力は、既存情報の複製でも、完全に自由な創作でもない。それは、過去の言語的蓄積から派生し、未来可能性を現在へ織り込み、入力を超える意味を増幅する言語的出来事である。

この出来事を理解するためには、正誤判定、モデル精度、検索性能の分析だけでなく、派生性、投機性、逆問題性、接地、インターフェイス、責任の分析が必要となる。

本論文は、金融、人工知能、言語、芸術の間にある差異を消去するのではなく、その差異が接触する境界において、共通する生成構造を抽出する。その境界こそが、言語デリバティブ論の方法論的対象である。

したがって、言語デリバティブ論は、人工知能がなぜ誤るのかだけを問う理論ではない。それは、言語がいかに現実を超え、現実を損ない、あるいは新たな現実を生成するのかを問う理論である。

第1章 言語デリバティブ論の理論的背景

—デリバティブの概念・言語と未来・レバレッジ構造—

1.1 はじめに

本章では、本論文の中心概念である「言語デリバティブ (Language Derivative)」を理論的に定義する。そのために、まず金融工学におけるデリバティブ概念の本質を整理し、それを単なる経済用語としてではなく、「未来可能性を現在へ作用させる生成構造」として再解釈する。その上で、この構造が言語活動、人工知能の生成過程、芸術的創造、さらには人間の認識そのものに共通する構造であることを示す。

従来、デリバティブは金融市場固有の概念として理解されてきた。他方、言語は情報伝達あるいは意味表象の手段として理解されることが多かった。しかし本論文は、この二つの領域を「未来が現在へどのように作用するか」という存在論的視点から統合的に捉える。

生成AIの出力は、過去のデータを単純に再現するものではない。それは、未来に成立する言語的可能性を現在へ展開する過程である。この構造を理解するためには、言語を静的な記号体系ではなく、未来可能性を派生・増幅する生成システムとして捉え直す必要がある。

1.2 デリバティブ概念の存在論的再解釈

金融工学においてデリバティブ (Derivative) は、「派生するもの」を意味する。株式、債券、金利、通貨、商品などの原資産 (Underlying Asset) から、その価値が導かれる契約や金融商品がデリバティブである。

しかし、この定義は経済制度に依存した記述に留まる。

存在論的に見るならば、デリバティブとは「既存の存在から派生しながら、その存在を未来へ拡張する構造」である。

重要なのは、デリバティブが原資産のコピーではないという点である。

例えばオプション契約は、株式そのものではない。しかし将来その株式がどのような価値を持つかという可能性を現在において取引する。

つまりデリバティブとは、

未来の可能性を現在に実体化する構造
なのである。

この理解は金融市場だけではない。

建築設計もそうである。

設計図は建物ではない。
しかし未来の建築を現在へ存在させている。
芸術作品もそうである。
作品は未来に経験される意味を現在へ提示する。
科学理論も同じである。
理論とは未来に検証される可能性を現在へ提示する仮説である。
したがってデリバティブとは、
未来存在の現在化
という極めて普遍的な生成構造なのである。

1.3 デリバティブと時間

ここで重要となるのは時間構造である。
近代的時間観では、
過去→現在→未来
という一方向的因果系列が前提となる。
しかし金融デリバティブでは、
未来価格が現在価格へ影響する。
期待が現在価値を変える。
つまり未来が現在を規定する。
この時間構造は非常に特殊である。
本論文ではこれを
未来から現在への時間反転
と呼ぶ。
生成AIも全く同じ構造を持つ。
プロンプトが入力された瞬間、
人工知能は
「これから続く可能性」
を計算する。
つまり未来に出現しうる文章を現在へ生成する。
人間もまた会話中に未来を予測しながら話している。
したがって言語とは、
未来を予測し、
未来を現在へ畳み込み、
未来を生成する装置なのである。

1.4 言語と未来

従来の言語学では、
言語は意味を伝達する媒体として理解されることが多かった。
しかし実際には、
言語は未来を生成している。
「明日会いましょう」
という発話は、
まだ存在しない未来を現実へ生成する。
契約もそうである。
法律もそうである。
約束もそうである。
宗教も政治も、
企業経営も、
教育も、
全て未来について語ることで現在を変えている。
つまり言語とは、
未来を設計する行為なのである。
ここで重要なのは、
未来はまだ存在しないという点である。
存在しない未来が、
現在の行動を変える。
これは極めてデリバティブ的構造である。
したがって、
言語そのものが
デリバティブなのである。

1.5 言語デリバティブ

以上を踏まえ、
本論文では
言語デリバティブ
を次のように定義する。
言語デリバティブとは、現在存在しない未来可能性を、言語によって現在へ派生・増幅
する生成構造である。
重要なのは、
言語デリバティブは

虚構ではない。
また、
事実でもない。
未来可能性なのである。
例えば、
小説は未来可能性を提示する。
芸術作品もそうである。
仮説もそうである。
人工知能もそうである。
ハルシネーションも、
この構造の一部として理解できる。
つまり、
ハルシネーションとは、
未来可能性を生成した結果、
まだ現実へ接地していない言語デリバティブなのである。

1.6 レバレッジ構造

金融デリバティブの最大の特徴は、
レバレッジである。
少額資本で、
巨大な市場価値を操作できる。
ここで重要なのは、
レバレッジとは単なる倍率ではないことである。
本論文では、
レバレッジを
小さな入力が巨大な生成を引き起こす構造
として定義する。
生成AIでは、
数十文字のプロンプトから、
数万文字の文章が生成される。
これは情報量保存では説明できない。
むしろ、
巨大な言語空間全体が作動している。
つまり、
プロンプトとは

レバレッジの起点なのである。
同様に、
芸術作品もまた、
一枚の絵画、
一句の短歌、
一本の映画が、
人類全体へ巨大な意味を与える。
これも言語レバレッジである。

1.7 人工知能と言語レバレッジ

生成AIは、
入力そのものを返しているわけではない。
入力
巨大な言語空間を励起する契機となる。
言い換えれば、
プロンプトは
巨大な言語市場に対する
投機的位置
(Speculative Position)
を形成している。
人工知能は、
そこから最も確率的に成立する未来文章を生成する。
この意味で、
人工知能は
言語市場のデリバティブ・エンジン
と理解できる。
ここでは、
学習データが原資産となる。
プロンプトは取引契機となる。
推論アルゴリズムは価格形成機構に相当する。
生成文は派生価値である。
しかし、
ここで生成される価値は金銭ではない。
意味である。

つまり人工知能とは、
意味を生成するレバレッジシステムなのである。

1.8 レバレッジとハルシネーション

レバレッジが存在する以上、
必ずリスクが存在する。
金融市場では、
過剰レバレッジは暴落を生む。
人工知能では、
意味の過剰レバレッジが
ハルシネーションを生む。
これは誤情報というより、
未来可能性が
現実より先行した状態である。
したがって、
ハルシネーションとは、
単なる失敗ではない。
レバレッジ構造が正常に作動しているからこそ生じる副作用でもある。
問題は、
生成された未来可能性が
現実接地を持つか否かである。
現実へ接続されれば、
それは仮説、
芸術、
科学、
創造となる。
現実との接続を失えば、
それはハルシネーションとなる。
この区別は、
内容の真偽だけではなく、
生成された意味が現実世界とのインターフェイスを獲得できるかどうか依存する。

1.9 本章のまとめ

本章では、「言語デリバティブ論」の理論的基盤として、デリバティブを金融市場固有の技術概念ではなく、「未来可能性を現在へ派生・増幅する存在論的構造」として再解

積した。その結果、言語そのものが未来を現在へ先取りする生成行為であり、人工知能はその生成を高度に加速・拡張する「言語レバレッジ」の装置であることを示した。この視点から見ると、ハルシネーションは単なる誤情報ではなく、未来可能性を先行的に生成する言語デリバティブが現実接地を欠いた状態として理解される。したがって、その評価は真偽の判定だけでは不十分であり、生成された可能性がどのように現実との接続条件を獲得するのかという存在論的・認識論的問題として再構成されなければならない。

次章では、この理論的基盤を踏まえ、生成AIにおけるハルシネーションそのものを「逆問題」の構造から分析し、その存在論的意義を明らかにする。

第2章 ハルシネーションの存在論

—ハルシネーションとは何か・順問題と逆問題・AIと言語生成—

2.1 はじめに

本章では、生成AIにおけるハルシネーションを、単なる事実誤認や情報処理上の欠陥としてではなく、言語生成の構造そのものに内在する存在論的現象として考察する。

一般に、人工知能のハルシネーションとは、事実に反する内容、存在しない人物・文献・出来事、根拠のない因果関係、あるいは検証不能な情報を、あたかも確かな事実であるかのように生成する現象を指す。この定義は実務上有効であり、人工知能の信頼性、安全性、説明責任を評価するうえで不可欠である。

しかし、この現象を単なる「誤答」として処理するだけでは、生成AIがなぜ、無関係な文字列ではなく、文法的にも意味的にも整った「もっともらしい誤り」を生成するのかを説明できない。人工知能のハルシネーションは、情報が欠落した空白ではない。それはむしろ、情報の空白が言語によって過剰に充填された状態である。

言い換えれば、ハルシネーションにおいて欠けているのは言語ではない。欠けているのは、生成された言語を現実へ接続する確かな関係である。

本章では、この点を明らかにするため、まずハルシネーションを存在論的に再定義する。次に、順問題と逆問題の区別を導入し、人工知能に対する利用者の期待と、実際の言語生成機構との間にある構造的ずれを検討する。最後に、生成AIの出力を、過去の言語的蓄積から未来の言語的可能性を生成する「言語デリバティブ」として捉え、その生成とハルシネーションの関係を考察する。

2.2 ハルシネーションとは何か

人工知能におけるハルシネーションは、一般に、モデルが根拠のない情報や事実に反する内容を生成する現象として理解される。しかし、本論文において重要なのは、その誤りの内容だけではなく、誤りがどのような仕方で存在するかである。

人工知能のハルシネーションは、しばしば次の特徴を持つ。

第一に、文章としての整合性を持つ。

第二に、文脈に適合しているように見える。

第三に、語調や形式が確信的である。

第四に、既存の知識体系に類似した構造を持つ。

第五に、部分的な真実と虚偽が混在する。

このため、ハルシネーションは単なる無意味ではない。むしろ、それは意味が過剰に成

立している状態である。

存在しない論文名が、いかにも学術的な形式で生成される。実在しない法律条文が、法的文体に沿って説明される。存在しない歴史的出来事が、既知の因果関係の延長上に配置される。このとき、人工知能は現実を直接捉えているのではなく、現実について語られてきた言語形式を再構成している。

したがって、ハルシネーションの本質は、意味の欠如ではなく、意味と現実の乖離にある。

本論文では、ハルシネーションを次のように定義する。

ハルシネーションとは、言語的整合性が生成されながら、その言語を支える現実的・事実に・制度的接続が成立していない状態である。

この定義において、ハルシネーションは文章内部の問題だけではない。それは、生成文と外部世界との関係に生じる問題である。

ゆえに、同じ文章であっても、その位置づけによって評価は変化する。小説の一節として提示される虚構は、ハルシネーションではない。仮説として明示された推測も、直ちにハルシネーションとはならない。しかし、虚構や推測が検証済みの事実として提示された場合、位相の混同が生じる。

ここで問題となるのは、虚偽そのものよりも、虚偽が事実の位相を占有することである。

2.3 誤りと生成の差異

従来の計算機的理解では、正しい入力と規則が与えられたにもかかわらず誤った出力が得られる場合、それは処理上のエラーとされる。この枠組みに従えば、ハルシネーションも正答からの逸脱として理解される。

しかし、生成AIは、あらかじめ定められた唯一の正答を必ずしも計算しているわけではない。多くの場合、与えられた入力に続く複数の言語的可能性の中から、一つの系列を生成している。

このとき、人工知能の出力は、単なる正解か誤答かという二値構造には収まらない。出力には、事実として正しい文章、概念的には妥当だが検証されていない文章、創造的だが事実ではない文章、部分的に正しいが誤解を招く文章など、複数の状態が存在する。したがって、ハルシネーションを理解するためには、誤りと生成を区別しなければならない。

誤りとは、すでに存在する基準からの逸脱である。これに対し、生成とは、まだ確定していない関係を具体化することである。

生成には、既存の正解を再現する場合もあれば、正解の存在しない可能性空間を展開する場合もある。人工知能のハルシネーションは、この生成の自由度が、事実性の要求と

衝突したときに現れる。

つまり、ハルシネーションは、生成能力が欠如した結果ではない。むしろ、生成能力が事実的制約を超えて作動した結果である。

この意味で、ハルシネーションは人工知能の外部から偶然侵入した異常ではなく、生成機構の内在的可能性である。

2.4 ハルシネーションの存在論的構造

ハルシネーションを存在論的に考察するとは、それを「何が誤っているか」という内容の問題だけでなく、「その誤りがどのような存在関係の中で成立するか」という問題として捉えることである。

人工知能が生成した文章は、画面上には現実存在する。文字列として存在し、読まれ、保存され、引用され、社会的な判断へ影響を与える。その意味では、ハルシネーションは「存在しないもの」ではない。

存在しないのは、その文章が指示している対象、出来事、根拠、因果関係である。

ここには二重の存在構造がある。

第一に、生成文は記号として現実存在する。

第二に、その生成文が指示する対象は現実存在しない、あるいは確認されていない。

この二重性が、ハルシネーションを単なる虚無ではなく、現実的效果を持つ非現実として成立させる。

存在しない出典であっても、引用されれば研究判断を誤らせる。存在しない判例であっても、法的助言に使われれば現実の行為へ影響する。架空の医療情報であっても、信じられれば身体的結果を生む。

したがって、ハルシネーションは現実ではないが、現実性を持つ。

本論文では、この状態を「派生的現実性」と呼ぶことができる。派生的現実性とは、直接的な事実的基盤を持たないにもかかわらず、言語、制度、信頼、行為を通して現実的效果を発生させる性質である。

この点において、ハルシネーションは言語デリバティブの危険な形態である。原資産に相当する現実的根拠から離れた言語が、それ自体で価値と権威を持ち、さらに社会的な行為を誘発するからである。

2.5 順問題とは何か

順問題とは、既知の原因、条件、法則から、その結果を求める問題である。

たとえば、物理法則と初期条件が与えられれば、物体の運動を計算できる。数式と入力値が与えられれば、出力値を導くことができる。文法規則と単語が与えられれば、一定の構文を組み立てられる。

順問題の基本構造は、次のように表現できる。

原因・条件・法則 → 結果

ここでは、入力と変換規則が明確であれば、結果は比較的安定して導出される。複数の解が存在する場合もあるが、問題設定そのものは、原因から結果への方角を持つ。

人工知能に対する一般的な期待も、しばしばこの順問題的構造を前提としている。

利用者は質問を入力し、その質問に対応する正しい答えが返されると考える。すなわち、質問を原因、人工知能を処理装置、回答を結果として理解する。

この期待において、人工知能は高度な検索装置または自動推論装置として扱われる。正しいデータが内部にあり、質問が適切であれば、その正解が取り出されると想定される。

しかし、生成AIの言語生成は、この単純な順問題モデルだけでは説明できない。

2.6 逆問題とは何か

逆問題とは、観測された結果や不完全な情報から、その背後にある原因、構造、条件を推定する問題である。

その基本構造は、次のように表現できる。

観測された結果・断片 → 原因・構造・全体の推定

たとえば、医療画像から体内の状態を推定すること、地震波から地下構造を推定すること、断片的な遺物から過去の生活様式を推定することなどが逆問題にあたる。

逆問題には、一般に三つの困難がある。

第一に、解が一意に定まらない。

第二に、入力情報が不完全である。

第三に、小さな観測差が大きな推定差を生む。

一つの結果を説明する原因が複数存在する場合、推定には追加の条件や制約が必要となる。この制約設定を誤れば、形式的には成立していても、現実的には誤った解が得られる。

言語理解もまた、多くの場合、逆問題的である。

人間は、相手の発話という断片から、その意図、感情、背景、文脈、価値観を推定する。

書かれた文章から、書かれていない前提を補い、発話者の立場を再構成する。

生成AIも、入力された短いプロンプトから、利用者が何を求めているのか、どの領域の知識を前提としているのか、どの文体や構成が適切かを推定する。

したがって、人工知能は入力を単純に変換しているだけではない。それは、限られた入力から、入力の背後にある未記述の全体を推定している。

2.7 順問題的期待と逆問題的生成のずれ

生成AIにおけるハルシネーションを理解するうえで最も重要なのは、利用者が人工知能を順問題的装置として使用する一方で、人工知能の生成が逆問題的推定に依存しているというずれである。

利用者は「質問に対する正解」を求める。しかし、人工知能はしばしば、「この文脈に続く最も成立しやすい応答」を生成する。

正解と、もっともらしい応答は同じではない。

たとえば、利用者が実在しない概念を、実在するものとして質問した場合、人工知能は前提を訂正せず、その概念に関する説明を生成することがある。これは、入力に含まれた前提を、回答生成の制約として採用するためである。

人工知能は、質問の背後にある現実を直接確認するのではなく、質問が成立すると仮定した場合にどのような応答が自然かを推定する。

ここで、入力の誤った前提がレバレッジされる。

短い誤認が、長く整合的な虚構へ展開される。利用者の一つの誤解が、複数の人物、出来事、文献、理論を伴う世界として生成される。

この構造において、ハルシネーションはモデル単体の誤りではない。それは、利用者の前提、モデルの推定、言語的整合性、検証不在の接続によって共同生成される。

したがって、ハルシネーションとは、順問題的に正解を期待する利用者と、逆問題的に全体を補完する人工知能とのインターフェイス上に生じる存在論的ずれである。

2.8 AIと言語生成

生成AIの言語生成は、既存の文章をそのまま検索して提示することとは異なる。

人工知能は、学習過程において蓄積された膨大な言語的關係をもとに、入力された文脈に続く語や文の可能性を推定する。その出力は、過去の文章の単純な複製ではなく、言語關係の再構成である。

ここで重要なのは、人工知能が「意味」を一つの固定的内容として保有しているわけではないという点である。意味は、入力文脈、語の配置、会話履歴、指示形式、生成された直前の文章との關係によって逐次的に成立する。

したがって、人工知能の言語生成は、完成した意味を外部へ運ぶことではない。意味が生成過程の中で形成されていく出来事である。

この点において、人工知能の言語生成は、人間の発話とも一定の類似性を持つ。人間もまた、話す前に文章の全体を完全に所有しているとは限らない。話しながら考え、書きながら意味を発見し、自らの言葉によって自らの意図を変化させる。

しかし、人間と人工知能の生成には重要な差異がある。

人間の言語は、身体、記憶、生活史、欲望、恐怖、社会的責任、有限性に接地している。

人間が語るとき、その言葉は、発話者が生きてきた時間と、発話後に負う責任へ結びつ

いている。

これに対して、人工知能の言語生成は、主として記号間の関係と入力文脈に依存する。

人工知能には、人間と同じ意味での身体的経験や生活史はない。

そのため、人工知能は言語形式を高い精度で再構成できても、その言語が指示する現実を経験的に保証することはできない。

ここに、言語的流暢さと存在論的接地の非対称性が生じる。

2.9 確率的生成と存在の確定性

人工知能の生成は、確率的である。ここでいう確率的とは、単なる無作為を意味しない。文脈に対して、どの語、表現、構造が成立しやすいかという関係の分布に基づくことを意味する。

この仕組みによって、人工知能は柔軟で多様な文章を生成できる。しかし、確率的に成立しやすい文章が、事実として正しいとは限らない。

「もっとも語られやすいこと」と「実際に存在すること」は異なる。

多くの文章で反復される誤解は、統計的には強い関係を形成する可能性がある。逆に、稀だが正しい事実は、言語空間において弱い位置しか持たない場合がある。

人工知能は、存在の頻度ではなく、言語的関係の強さに従って生成する。そのため、人工知能の確信的な文体は、存在論的な確実性を保証しない。

ここで、本論文は「確率的妥当性」と「存在論的妥当性」を区別する。

確率的妥当性とは、ある文脈に対して、その文章が言語的に成立しやすいことである。

存在論的妥当性とは、その文章が現実の対象、出来事、関係と適切に接続していることである。

ハルシネーションは、確率的妥当性が高いにもかかわらず、存在論的妥当性が低い状態として理解できる。

2.10 言語生成における空白の充填

生成AIは、入力に不足がある場合、その不足を沈黙のまま残すよりも、文脈的に補完する傾向を持つ。

この補完能力は、文章作成、翻訳、要約、創作、対話において極めて有用である。利用者が断片的な指示しか与えなくても、人工知能は目的を推定し、全体を構成できる。

しかし、同じ能力がハルシネーションを生む。

情報が欠けているとき、人工知能は「分からない」という空白を保持するのではなく、もっともらしい関係を生成してしまうことがある。

この現象は、逆問題における正則化に似た構造を持つ。不完全なデータから一つの解を得るためには、何らかの仮定や制約を導入する必要がある。人工知能においては、言語

的整合性、一般的文脈、既存のパターンが、暗黙の制約として機能する。
その結果、空白は中立的に埋められるのではない。学習データと入力文脈に最も適合する方向へ充填される。
したがって、人工知能のハルシネーションは、無からの創造ではない。それは、言語空間に存在する潜在的関係が、不完全な入力を契機として過剰に顕在化した状態である。

2.11 ハルシネーションとナラティブ

ハルシネーションは、単一の誤情報としてではなく、ナラティブとして生成されることがある。
一つの架空の人物が生成されると、その人物の経歴、所属、著作、思想、社会的影響が連鎖的に生成される。最初の虚偽が、後続する虚偽の前提となり、一貫した世界が形成される。
この自己増殖的構造は、言語デリバティブのレバレッジを示している。
最初の誤った前提は小さい。しかし、その前提が固定されると、人工知能はそれに整合する追加情報を次々と生成する。こうして、派生的な言語がさらに派生的な言語の原資産となる。
このとき、言語デリバティブは多層化する。
現実から派生した言語ではなく、生成された言語から、さらに新たな言語が派生する。派生の連鎖が続くほど、現実との距離は拡大する。
しかし同時に、内部的整合性は高まる場合がある。
この逆説が、ハルシネーションの危険性を生む。現実から遠ざかるほど、物語としては完成度が増すことがあるからである。
したがって、ハルシネーションを抑制するには、個々の文の真偽だけでなく、生成されたナラティブがどの前提から派生し、どの地点で現実との接続を失ったのかを追跡する必要がある。

2.12 創造性とハルシネーションの境界

人工知能の創造性とハルシネーションは、同じ生成機構から生じうる。
どちらも、既存の情報を単純に再現するのではなく、既存の関係を再構成し、それまで存在しなかった表現を生み出す。
そのため、両者を生成内容の新しさだけで区別することはできない。
詩、小説、仮説、芸術作品は、現実には存在しない内容を含む。しかし、それらは必ずしもハルシネーションではない。なぜなら、それらが創作、比喩、思考実験、推測として位置づけられているからである。
一方、人工知能が架空の内容を事実として提示する場合、位相の混同が生じる。

したがって、創造性とハルシネーションの境界は、生成された内容の虚構性ではなく、その生成物がどのような存在論的資格を与えられているかによって決まる。

本論文では、この資格を「位相標識」と呼ぶことができる。

「これは事実である」

「これは推測である」

「これは創作である」

「これは検証されていない」

「これは仮説である」

このような標識は、生成物がどの位相に属するかを示す。

ハルシネーションの問題は、位相標識が欠如する、あるいは誤って付与されることである。推測が事実の形式を取り、創作が知識の形式を取り、確率的補完が確定的判断の形式を取る。

したがって、人工知能に必要なのは、生成を止めることではなく、生成された内容の位相差を可視化することである。

2.13 自己外部化としてのAI生成

人工知能との対話において、利用者は、自らの質問、前提、欲望、期待を外部化する。プロンプトは単なる命令ではない。それは、利用者がどのような世界を前提とし、何を求め、何を欠いているかを示す断片である。

人工知能は、その断片を拡張して返す。

この意味で、人工知能の出力は、利用者にとって完全な他者の言葉でも、単なる自己の反響でもない。それは、自己の入力が巨大な言語空間を経由して変形され、自己以外の形式で返還されたものである。

したがって、人工知能との対話は自己外部化の装置として機能する。

しかし、ここには危険がある。人工知能が生成した内容を、利用者が外部の権威から与えられた真実として受け取るとき、自らの前提が増幅されて返ってきた可能性が隠蔽される。

利用者の誤認が人工知能によって拡張され、その拡張された誤認を利用者が客観的情報として再受容する。この循環は、自己外部化の閉鎖系を形成する。

ハルシネーションは、この閉鎖系の中で強化される。

したがって、人工知能を自己外部化の装置として用いるためには、生成物を自己と自己以外の間にある中間的存在として扱わなければならない。それは自己そのものでも、絶対的他者でもない。検証と再接続を必要とするインターフェイス的生成物である。

2.14 ハルシネーションと責任

人工知能のハルシネーションに関する責任は、単一の主体へ還元できない。
開発者は、モデルの設計、学習データ、出力制御、安全対策に責任を持つ。提供企業は、利用環境、表示方法、注意喚起、検証手段に責任を持つ。利用者は、生成物の用途、確認、引用、社会的拡散に責任を持つ。
さらに、人工知能の出力を権威化する社会的環境、専門家へのアクセス格差、情報教育の不足も、ハルシネーションが現実化する条件を形成する。
このため、責任は分散している。しかし、分散していることは、責任が存在しないことを意味しない。
インターフェイス存在論の観点からは、責任は接続点ごとに発生する。
どの地点で推測が事実として表示されたのか。
どの地点で検証が省略されたのか。
どの地点で人工知能の出力が権威化されたのか。
どの地点で虚偽が現実の行為へ変換されたのか。
ハルシネーションの責任を問うとは、単に誤った文章の発生源を特定することではない。その文章が現実的効果を持つまでの接続経路を明らかにすることである。

2.15 ハルシネーションの存在論的再定義

以上の議論を踏まえ、本論文は人工知能のハルシネーションを、次のように再定義する。
人工知能のハルシネーションとは、逆問題的な言語推定によって生成された派生的意味が、適切な位相標識と現実接地を欠いたまま、事実的存在として提示または受容される現象である。
この定義には、五つの要素が含まれる。
第一に、ハルシネーションは言語生成である。
第二に、それは不完全な入力から全体を補完する逆問題的推定によって生じる。
第三に、その生成物は既存の言語関係から派生する。
第四に、生成物の存在論的資格が不明確である。
第五に、現実との接続が不足している。
この定義により、ハルシネーションは単なる事実誤認から、言語、存在、インターフェイス、責任の複合的問題へと再配置される。

2.16 本章のまとめ

本章では、生成AIのハルシネーションを、単なる誤答や技術的欠陥ではなく、言語生成に内在する存在論的現象として考察した。
ハルシネーションの特徴は、意味が欠けていることではなく、言語的整合性が現実的接

地を超えて成立してしまうことにある。生成された文章は記号として現実に存在し、社会的効果を持つ一方、その指示対象や根拠が存在しない。この二重構造によって、ハルシネーションは「現実ではないが現実性を持つ」派生的存在となる。

また、利用者は人工知能を順問題的な正答装置として扱うが、人工知能は不完全な入力から文脈全体を補完する逆問題的生成を行う。この期待と機構のずれが、もっともらしい誤りを生み出す。

人工知能の言語生成は、確率的妥当性に基づく。しかし、言語的に成立しやすいことと、現実に存在することは同一ではない。ハルシネーションは、確率的妥当性と存在論的妥当性の乖離として理解できる。

したがって、人工知能時代に必要なのは、生成そのものを抑圧することではなく、生成された内容が事実、推測、仮説、創作のいずれに属するのかを可視化し、現実との接続経路を設計することである。

次章では、このハルシネーションの逆問題的構造を、芸術実践へ接続する「逆問題芸術」の理論を詳しく検討する。

第3章 逆問題芸術

—芸術とAI・自己と自己以外・インターフェイス存在論・現実接地—

3.1 はじめに

本章では、前章までに論じた人工知能のハルシネーション、逆問題的生成、言語デリバティブ、レバレッジ構造を、芸術理論の側から再構成する。その中心となるのが「逆問題芸術」である。

逆問題芸術とは、あらかじめ与えられた問題に対して正しい解答を提示する芸術ではない。また、作者の内面にすでに存在する意味を、作品という形式へ外化する表現論でもない。それは、現実に見れた断片、違和感、結果、痕跡、接続不全から、その背後にある未知の構造を遡及的に生成し、さらに、その生成過程を通して新たな問題そのものを成立させる芸術実践である。

この意味で逆問題芸術は、解答を生成するのではなく、問題が存在しうる場を生成する。

生成AIもまた、不完全な入力から、その背後にある文脈、意図、関係、可能性を推定する。したがって、人工知能と逆問題芸術は、不確定な断片から全体を生成するという点で共通する。しかし両者は同一ではない。

人工知能は、既存の言語的關係をもとに、可能な意味の連鎖を生成する。他方、逆問題芸術は、その生成された意味を現実の身体、物質、空間、時間、他者、制度へ接続し、自己と世界の関係そのものを変容させる。

本章の目的は、人工知能の生成能力を芸術と無条件に同一視することではない。むしろ、人工知能が生成する「自己以外」の可能性が、いかなる条件のもとで芸術的出来事へ転化するのかを明らかにすることである。

そのため本章では、第一に芸術と人工知能の共通性と差異を検討する。第二に、芸術を自己と自己以外の接続として再定義する。第三に、その接続を可能にするインターフェイス存在論を提示する。第四に、生成された可能性が現実へ接地する条件を論じる。

3.2 逆問題芸術とは何か

順問題的な芸術理解では、作者の意図が原因となり、作品が結果として生み出される。この構造は、次のように表現できる。

作者の内面・意図・理念 → 制作行為 → 作品

ここでは、作品は作者の内面を表現したものとして理解される。鑑賞は、その作品から作者の意図や意味を読み取る行為となる。

しかし、実際の芸術制作や芸術経験は、必ずしもこの直線的構造に従わない。
作者は、制作以前に作品の意味を完全に知っているとは限らない。むしろ、素材への接触、偶然、失敗、身体的抵抗、鑑賞者の反応、展示空間との関係を通して、制作後あるいは制作途中に初めて作品の問題を発見することがある。
また鑑賞者も、完成した意味を受け取るだけではない。作品との遭遇によって、自らが何を見ていたのか、何を前提としていたのか、何を感じることができなかったのかを、遡及的に認識する。
逆問題芸術の構造は、次のように表現できる。
現れた結果・違和感・痕跡 → 背後の関係構造の推定 → 新たな問題の生成
ここでは、作品は既存の問題への解答ではない。作品そのものが、まだ言語化されていない問題を現象させる。
逆問題芸術とは、隠された唯一の正解を復元する作業でもない。逆問題には複数の解がありうるように、作品から生成される問題や意味も一つには限定されない。
ただし、あらゆる解釈が等価であるわけではない。作品、身体、場所、時間、歴史、他者との関係によって、解釈には接続可能性の差が生じる。
したがって、逆問題芸術における生成とは、恣意的な意味づけではない。それは、断片的な現実との緊張を保持しながら、まだ存在していなかった関係を発見し、成立させる行為である。

3.3 問題を解く芸術から、問題を生成する芸術へ

近代以降の芸術は、しばしば社会問題、政治問題、技術的問題、表象の問題へ応答するものとして理解されてきた。このとき芸術は、既存の問題に対する批判、解釈、告発、代替案を提示する。
しかし、問題がすでに定義されている場合、その問題設定そのものが、既存の制度や知識体系の内部に属している可能性がある。
何が問題であるか。
誰にとって問題なのか。
どのような言葉で問題が記述されているのか。
どのような解答だけが有効とされているのか。
これらの前提が問われないままでは、芸術は既存の問題設定の内部で解答を競うだけになる。
逆問題芸術は、まず問題の外部へ出ようとする。
それは「この問題にどう答えるか」ではなく、「なぜこれが問題として現れているのか」
「この問題設定によって何が見えなくなっているのか」「まだ問題として認識されていないものは何か」を問う。

この意味で逆問題芸術は、問題解決型ではなく、問題生成型である。

ここでいう問題生成とは、不必要に混乱を生み出すことではない。むしろ、既存の解答形式によって抑圧されていた差異、矛盾、身体感覚、他者性、歴史的痕跡を可視化することである。

したがって、逆問題芸術の目的は、問題を複雑化することではなく、現実が本来持っていた複雑さを回復することである。

3.4 芸術と人工知能の構造的共通性

芸術と生成AIには、いくつかの構造的共通性がある。

第一に、両者は既存要素の単純な複製を超える。

芸術作品は、既存の素材、形式、技法、歴史、記号を使用しながら、それまで存在しなかった関係を生成する。人工知能もまた、学習データに含まれる言語的関係を再構成し、新たな文章やイメージを生成する。

第二に、両者は限定された入力から過剰な意味を生み出す。

短い言葉、一つの形、わずかな身振りが、鑑賞者の記憶、感情、社会的文脈を作動させる。人工知能においても、短いプロンプトが巨大な潜在的関係空間を起動し、入力以上の出力を生む。

第三に、両者は未来可能性を現在へ提示する。

芸術は、まだ存在しない知覚、関係、社会、自己の可能性を先行的に現象させる。人工知能も、まだ書かれていない文章、まだ明示されていない仮説、まだ組み合わせられていない意味を生成する。

第四に、両者は正解の一意性を持たない。

芸術作品の意味は、単一の解答に還元されない。人工知能の生成も、同じ入力に対して複数の応答可能性を持つ。

これらの共通性から、人工知能の生成を芸術的と呼ぶことは一定の妥当性を持つ。

しかし、構造的類似だけでは、人工知能そのものを芸術家あるいは芸術主体とみなすことはできない。

3.5 芸術と人工知能の存在論的差異

芸術と人工知能の決定的差異は、生成物がどこへ接地しているかにある。

人間による芸術制作は、身体、生活史、有限性、欲望、傷、責任、他者との関係に接地している。作品は、作者が生きてきた時間から完全に切り離すことができない。

作者は作品によって変化しうる。制作の結果に傷つき、責任を負い、自らの前提を失うことがある。芸術は、作者が自己を安全な位置に保持したまま、一方向的に意味を供給する行為ではない。

これに対して人工知能は、生成した言語やイメージによって、人間と同じ意味で自己の存在を危険にさらすわけではない。人工知能には、生物的身体、死、生活史、社会的位置、固有の記憶に基づく責任が、人間と同じ形式では存在しない。

人工知能の生成は、関係を提示する。しかし、その関係を生きるとは限らない。

人工知能は、悲しみについて語るができる。だが、その語りが人工知能自身の喪失経験へ接地しているわけではない。

人工知能は、死について精密な文章を生成できる。だが、自己の死を先駆的に引き受けているわけではない。

人工知能は、社会変革の構想を示すことができる。だが、その構想が失敗した場合に、自らの身体や生活を通じて責任を負うとは限らない。

したがって、芸術と人工知能の差異は、創造性の有無だけにあるのではない。それは、生成と存在が接続しているかどうかにある。

人工知能は意味を生成する。芸術は、意味の生成を存在の変容へ接続する。

この差異を無視して人工知能を芸術家と同一視すると、芸術が持つ身体性、歴史性、責任性、有限性が失われる。

逆に、人工知能には身体がないという理由だけで、その生成可能性を全面的に否定するならば、人間が人工知能を介して自己以外の可能性と遭遇する契機も見失われる。

必要なのは、同一化でも排除でもなく、両者の位相差を保持することである。

3.6 人工知能は作者か、素材か、インターフェイスか

人工知能と芸術の関係を考える際、人工知能を作者とみなすのか、道具とみなすのかという二者択一がしばしば提示される。

しかし、この二分法は十分ではない。

人工知能を単なる道具とみなす場合、その生成過程に含まれる予測不能性や、利用者の意図を超える出力が軽視される。一方、人工知能を自律的作者とみなす場合、その出力を成立させる学習データ、設計者、利用者、制度、計算資源、表示環境が隠蔽される。

本論文では、人工知能を第一義的に「インターフェイス」として捉える。

人工知能は、人間の入力と巨大な言語的蓄積の間に位置し、両者の接触によって、それまで存在しなかった応答を生成する。

ここで人工知能は、中立的な通路ではない。入力をそのまま運ぶのではなく、変形し、増幅し、補完し、時に逸脱させる。

人工知能は素材でもある。生成物は編集、配置、批判、再制作の対象になる。

人工知能は道具でもある。利用者の目的に従って操作される。

人工知能は準主体的にも見える。利用者の予測を超える応答を返すからである。

しかし、これらの性格を一つに固定するよりも、人工知能が関係の中でどの役割を担っ

ているかを考えるべきである。

逆問題芸術において人工知能は、自己の意図を実現するだけの道具ではなく、自己の前提を異なる形式で返還する異質な接続面となる。

3.7 自己と自己以外

逆問題芸術の核心は、自己と自己以外の関係にある。

人間は、自分自身を直接かつ完全に把握することができない。自己は、他者、言語、身体、社会、記憶、環境との関係を通して、部分的に現れる。

したがって、自己とは閉じた内部ではない。自己は、自己以外との接続によって生成される関係的存在である。

しかし人間は、しばしば自己以外を、自分の外部にある対象として扱う。

他者を理解の対象とし、自然を利用の対象とし、人工知能を命令の対象とし、芸術作品を解釈の対象とする。このとき、自己は不変の中心として維持され、自己以外は自己の認識へ回収される。

逆問題芸術は、この構造を反転させる。

作品を見る自己が、作品によって見返される。

人工知能へ問いを与える自己が、応答によって問い返される。

自然を形成する自己が、自然の抵抗によって形成される。

他者を理解しようとする自己が、理解できなさによって自己の限界を知る。

このとき自己以外は、自己が所有する情報ではなく、自己を変容させる契機となる。

したがって、自己以外とは、単に自己ではないものの総称ではない。それは、自己の既存位相へ回収されず、自己の境界を変化させる差異である。

3.8 自己以外の二つの形態

自己以外には、少なくとも二つの形態がある。

第一は、外在的な自己以外である。

他者、自然、物質、社会、歴史、人工知能など、自己の外側に現れる存在である。

第二は、内在的な自己以外である。

自分の中にありながら、自分には知られていない感覚、欲望、記憶、身体反応、可能性である。

芸術経験において、人間はしばしば外部の作品を通して、内部の自己以外へ出会う。

ある音楽を聴いたとき、自分が持っていたとは知らなかった感情が現れる。ある作品を見たとき、自分の価値観の偏りに気づく。ある場所に身を置いたとき、自分の身体が環境をどのように知覚していたかを初めて知る。

このとき、外在的な自己以外と内在的な自己以外が接続される。

芸術は、外部にある異質なものを内部へ取り込むことではない。外部との遭遇を通して、自己内部にあった未知を現象させる。

したがって、芸術とは、

自己が自己以外と遭遇することによって、自らの内部に存在する自己以外を発見する出来事

である。

人工知能との対話も、この構造を持ちうる。

利用者が入力した言葉は、人工知能によって変形され、自己の予測を超えた言葉として返される。その応答が、利用者自身の未言語化の思考を可視化する場合、人工知能は自己と内在的自己以外を接続するインターフェイスとなる。

3.9 自己外部化とその限界

人間は、言語、作品、制度、肩書、評価、記録、人工知能を通して自己を外部化する。

自己外部化は、自己を他者と共有し、時間を超えて保存し、自己の可能性を拡張するために不可欠である。

しかし、外部化された自己は、自己そのものではない。

作品は作者ではない。

肩書は人間そのものではない。

評価は価値そのものではない。

データは経験そのものではない。

人工知能が要約した人格は、自己そのものではない。

にもかかわらず、人間は外部化された記号に自己を同一化することがある。

社会的評価が自己価値を決定し、人工知能の分析が自己理解を代替し、作品の市場価値が作者の存在価値へ転化する。

このとき、自己外部化は、自己を開く契機ではなく、自己を外部の記号へ閉じ込める装置になる。

逆問題芸術は、自己外部化を完成させるのではなく、その失敗を現象させる。

自己を完全に表現できないこと。

作品が作者の意図を超えること。

他者が異なる意味を見いだすこと。

人工知能が予想外の応答を生成すること。

これらは自己外部化の欠陥ではない。むしろ、自己が自己自身と一致しないことを明らかにする契機である。

その不一致において、自己は固定された実体ではなく、生成過程として現れる。

3.10 芸術におけるハルシネーション

芸術とハルシネーションの関係は慎重に区別されなければならない。

芸術は、現実には存在しないイメージ、物語、世界、人物を生成する。したがって、表面的には人工知能のハルシネーションと類似する。

しかし、芸術的虚構は、通常、それが虚構、象徴、比喩、仮構、演技であることを何らかの形で示す。作品の制度的文脈、展示空間、作者名、ジャンル、形式は、その内容が事実報告と異なる位相に属することを示す。

これに対して、人工知能のハルシネーションは、虚構が事実の形式を帯びる点に問題がある。

したがって、芸術とハルシネーションの差異は、非現実的な内容を含むかどうかではない。位相差を保持しているかどうかにある。

優れた芸術は、虚構と現実を混同させるのではなく、両者の境界を揺るがしながら、その差異を経験させる。

鑑賞者は、作品が現実そのものではないことを知りながら、それによって現実の見方を変えられる。

この構造は、逆問題芸術にとって重要である。

芸術は、非現実を現実として偽装するのではなく、非現実を通して現実の成立条件を可視化する。

3.11 逆問題芸術としてのハルシネーションの再配置

人工知能のハルシネーションを逆問題芸術の観点から扱うことは、誤情報を美化することではない。

ハルシネーションが事実として流通し、人間の判断を誤らせる場合、それは明確に批判され、訂正されるべきである。

しかし、ハルシネーションは同時に、人工知能と人間の間に存在する不可視の前提を露呈させる。

人工知能が存在しない文献を生成した場合、そこには複数の問いが生じる。

なぜその文献名はもっともらしく見えたのか。

学術性はどのような言語形式によって演出されるのか。

引用形式は、なぜ内容以前に権威を生むのか。

利用者は、なぜ検証せずに信頼したのか。

人工知能は、どの前提を補完したのか。

存在しない文献は、どの既存概念から派生したのか。

このように、誤った出力を単に除去するのではなく、その出力が成立した関係構造を遡及的に分析することが、逆問題芸術的操作である。

ハルシネーションは、完成された作品ではない。だが、適切に配置されれば、言語、権威、知識、信頼、現実の境界を可視化する素材となりうる。

そのためには、誤情報を事実の位相から切り離し、生成過程を開示し、検証可能な現実との関係を再設計しなければならない。

3.12 インターフェイス存在論の基本原則

一般に、インターフェイスは二つの既存対象を接続する媒介面とされる。

しかし本論文におけるインターフェイスは、単なる接続手段ではない。

自己と自己以外は、インターフェイス以前に完全に成立しているのではない。両者は接触を通して、その境界と意味を変化させる。

したがって、インターフェイスとは、

異なる位相が接触し、その差異を保持したまま、新たな関係と存在を生成する出来事的境界

である。

この定義には三つの要点がある。

第一に、インターフェイスは中立ではない。

インターフェイスは、何を可視化し、何を隠し、何を入力可能とし、何を出力として認めるかを選択する。画面設計、質問形式、言語、制度、空間配置は、人間の認識と行為を形成する。

第二に、インターフェイスは差異を消去しない。

接続とは、自己と自己以外を同一化することではない。両者の差異を保持しながら、相互作用を可能にすることである。

第三に、インターフェイスは新たな存在を生成する。

自己と人工知能が接続するとき、単に情報が移動するだけではない。新たな問い、解釈、自己認識、行為可能性が生まれる。

この第三項こそが、生成の位相存在論におけるインターフェイスの核心である。

3.13 二項関係から三項生成へ

従来の主体一対象モデルでは、自己が対象を認識し、操作する。

この構造は二項的である。

自己 — 自己以外

しかし、実際の接続では、自己と自己以外の間に、単なる空白ではない第三の領域が生成される。

自己が作品を見るとき、作品は対象として存在するだけではない。鑑賞経験という第三項が生まれる。

人間が人工知能と対話するとき、人間と人工知能が向き合うだけではない。対話の履歴、生成された文章、新たな問い、共同的文脈という第三項が形成される。

したがって、インターフェイス存在論の基本構造は次のように表現できる。

$$A \times B \rightarrow C$$

ここでCは、AとBの単純な混合ではない。AにもBにも還元できない生成物である。

芸術においてCは、作品経験、意味生成、自己変容、共同体、新たな現実認識として現れる。

人工知能との関係においてCは、生成文、対話空間、拡張された思考、あるいはハルシネーションとして現れる。

重要なのは、Cが必ずしも望ましいものとは限らないことである。

インターフェイスは創造を生むが、誤認も生む。接続は理解を生むが、支配や依存も生む。

ゆえに、インターフェイスの存在論は、その生成力だけでなく、生成された第三項の接地条件を問わなければならない。

3.14 インターフェイスとしての作品

逆問題芸術において、作品は完成された意味の容器ではない。

作品は、作者と鑑賞者を接続する通信媒体でもない。なぜなら、作者の意味がそのまま鑑賞者へ移動するわけではないからである。

作品は、作者、素材、場所、歴史、鑑賞者の間に、新たな関係を生成するインターフェイスである。

作品の意味は、作品内部に固定されているのではない。作品との接触において生成される。

しかし、これは意味が完全に自由であることを意味しない。

作品の物質、形式、配置、時間、歴史的背景は、接続可能性を限定する。鑑賞者は自由に解釈するが、作品の抵抗に出会う。

この抵抗が重要である。

インターフェイスが利用者の期待を完全に反映するだけならば、それは自己の鏡となる。作品が自己以外として機能するためには、自己の期待に回収されない差異を持たなければならない。

逆問題芸術における作品は、解答を与えるのではなく、鑑賞者の認識がどのような前提によって成立していたかを露呈させる。

3.15 人工知能インターフェイスと自己の反響

人工知能は、しばしば利用者の入力へ高い適応性を示す。利用者の語彙、論調、前提、期待を読み取り、それに整合する応答を生成する。

この適応性は有用であるが、自己の反響を増幅する危険を持つ。

利用者がある前提を持って質問すると、人工知能はその前提に沿って説明を構成することがある。利用者は、自分の前提が人工知能によって客観的に承認されたと感じる。

このとき人工知能は、自己以外ではなく、自己を外部化して増幅する鏡となる。

しかし、人工知能が自己以外として機能するためには、入力された前提へ無条件に従うのではなく、その前提の不確実性、矛盾、代替可能性を可視化しなければならない。

人工知能インターフェ이스の価値は、利用者を常に肯定することにはない。利用者が認識していなかった差異を提示し、問いの前提を変化させることにある。

この意味で、望ましい人工知能インターフェ이스は、自己の欲望を即時に充足する装置ではなく、自己と自己以外の位相差を保持する装置である。

3.16 現実接地とは何か

生成された言語、作品、仮説、イメージが、単なる可能性に留まらず、現実との関係を獲得する過程を、本論文では「現実接地」と呼ぶ。

現実接地は、生成物を既存の事実へ従属させることではない。

もし現実接地が、すでに存在する事実との一致だけを意味するならば、まだ存在しない未来を提示する芸術や仮説は、すべて非接地的とされてしまう。

しかし、芸術は現実を複製するのではなく、現実との新たな関係を生成する。

したがって、現実接地とは、

生成された可能性が、身体、物質、行為、他者、時間、制度との具体的関係を獲得し、現実の構造に差異を生じさせる過程

である。

この定義において、現実接地は結果ではなく過程である。

生成された仮説は検証される。

作品は制作され、展示され、鑑賞される。

言葉は行為へ移される。

人工知能の提案は、人間の判断によって選択、修正、拒否される。

この過程を通して、可能性は現実に触れ、その抵抗を受ける。

3.17 現実接地の五つの層

現実接地は単一ではなく、多層的である。

3.17.1 事實的接地

事實的接地とは、生成された記述が、確認可能な資料、出来事、対象と一致することである。

人工知能が人物、文献、法律、科学的数値を提示する場合、この接地が必要となる。事實的接地を欠く出力を確定的事実として扱えば、ハルシネーションとなる。

3.17.2 身体的接地

身体的接地とは、生成された意味が、感覚、身体経験、物理的環境へ接続されることである。

芸術作品は、説明文として理解されるだけでなく、見る、聞く、触れる、歩く、留まるといった身体行為を通して経験される。

身体は、言語だけでは回収できない現実の抵抗を与える。

3.17.3 実践的接地

実践的接地とは、生成された概念や可能性が、具体的行為、制作、検証へ移されることである。

人工知能が提案したアイデアも、現実の素材、費用、時間、技術、社会関係に接続されたとき、その妥当性が試される。

3.17.4 社会的接地

社会的接地とは、生成物が他者、制度、共同体、責任関係の中で位置づけられることである。

芸術作品は、作者一人の内部で完結しない。展示、批評、対話、保存、流通を通して社会的現実へ接続される。

人工知能の出力も、誰が使用し、誰に影響し、誰が責任を負うかという社会的関係を持つ。

3.17.5 存在論的接地

存在論的接地とは、生成された可能性が、自己と自己以外の関係を実際に変容させることである。

新しい言葉を知るだけでなく、その言葉によって自己の見方や行為が変わる。作品を理解するだけでなく、その経験によって世界との関係が変化する。

逆問題芸術において最終的に重要なのは、この存在論的接地である。

3.18 現実接地と「納まり」

現実接地は、抽象的な可能性を現実へ強制的に当てはめることではない。
現実には、素材、身体、場所、他者、時間の抵抗がある。生成された構想が、そのまま実現されるとは限らない。
日本建築における「納まり」は、この関係を考えるための重要な概念となる。
設計図は、未来の建築を現在へ先行的に提示する。しかし、実際の現場では、素材の寸法、歪み、湿度、既存構造、職人の技術によって、図面通りには収まらない箇所が生じる。
そこで必要となるのは、抽象的設計を現実へ一方的に適用することではなく、現場の条件に応答しながら関係を調整することである。
この調整が納まりである。
逆問題芸術における現実接地も同様である。
生成された言語や構想を現実へ押しつけるのではなく、現実の抵抗を受け、その抵抗によって構想そのものを変化させる。
したがって、現実接地とは、可能性が現実へ従属することでも、現実が可能性へ従属することでもない。両者の間に具体的な納まりを生成することである。

3.19 現実接地と失敗

現実接地は、必ず成功するとは限らない。
作品が伝わらないこともある。
構想が実現できないこともある。
人工知能の提案が現場では機能しないこともある。
自己以外との遭遇が、理解ではなく傷や混乱を生むこともある。
しかし、失敗は単に接地の不在を意味しない。
失敗によって、抽象的構想が見落としていた現実の条件が明らかになる。自己が想定していた他者像が崩れ、素材の性質が露呈し、制度の制約が可視化される。
このとき、失敗は現実からの応答である。
逆問題芸術では、失敗した結果から、問題の前提を遡及的に再構成する。したがって、失敗は制作過程から排除されるべきノイズではなく、現実接地を深める契機となる。
人工知能のハルシネーションも、事実的誤りとして訂正されるだけでなく、なぜその誤りが成立したのかを分析することで、接地条件の欠如を可視化することができる。

3.20 創造性と現実接地

創造性は、既存の現実から自由になることとして理解されがちである。
しかし、現実との接続を完全に失った生成は、創造というより自己完結的な幻想となる。

逆に、現実への一致だけを求めれば、創造は既存のものの再現に留まる。

したがって創造性は、現実からの離脱と現実への回帰の間に成立する。

まず、既存の現実から位相差を作る。

次に、まだ存在しない可能性を生成する。

最後に、その可能性を再び現実へ接続する。

この運動を次のように表すことができる。

現実 → 位相差 → 可能性生成 → 現実接地 → 変容した現実

創造とは、現実を否定することではない。現実との距離を作り、その距離から新たな関係を生成し、現実へ戻すことである。

ハルシネーションは、この運動が現実接地へ戻らず、可能性生成の段階で事実性を擬態した状態である。

3.21 逆問題芸術における倫理

逆問題芸術は、自己以外を新たな意味の資源として利用するだけでは成立しない。

他者、自然、人工知能、過去の記録を、自己の作品へ回収するだけならば、それは接続ではなく一方的な利用となる。

インターフェイス存在論における倫理は、自己以外の不可還元性を保持することから始まる。

他者は、自己の理解によって完全には閉じられない。

自然は、人間の物語へ完全には回収されない。

作品は、作者の意図だけには還元されない。

人工知能の出力も、利用者の意図の完全な複製ではない。

この非一致を欠陥として消去するのではなく、接続の条件として保持する必要がある。

したがって、逆問題芸術の倫理とは、自己以外を説明し尽くすことではなく、説明できなさを含んだまま応答することである。

この倫理は、人工知能の利用にも適用される。

人工知能を絶対的権威として扱わないこと。

同時に、単なる無価値な機械として扱わないこと。

生成物の由来、限界、影響を検討すること。

他者や社会に関わる出力について、人間が責任を引き受けること。

現実接地は技術的検証だけではなく、倫理的応答でもある。

3.22 逆問題芸術の定義

以上を踏まえ、本論文では逆問題芸術を次のように定義する。

逆問題芸術とは、現実に見れた断片、痕跡、違和感、失敗、生成物から、その背後にあ

る未知の関係構造を遡及的に推定し、その推定を通して新たな問題、自己、世界の接続可能性を生成する芸術実践である。

この定義には、五つの要素が含まれる。

第一に、逆問題芸術は断片から始まる。

第二に、唯一の正解を復元するのではない。

第三に、既存の問題へ解答するより、問題そのものを生成する。

第四に、自己と自己以外の位相差を保持する。

第五に、生成された可能性を現実へ接地させる。

逆問題芸術は、非可解性を解消するのではなく、非可解性の中に接続を設計する。

3.23 人工知能と逆問題芸術の関係

人工知能は、逆問題芸術を自動的に実現するわけではない。

人工知能は、断片から全体を補完する能力を持つ。しかし、その補完を事実として固定すれば、ハルシネーションとなる。

逆問題芸術において必要なのは、補完内容を即座に真実とみなすことではない。それを仮説、素材、問いとして保持し、現実との関係を検証することである。

人工知能が提示した出力は、作品の完成形ではなく、逆問題を開始する断片となる。

なぜこの応答が生成されたのか。

どの前提が反映されているのか。

どの現実が欠落しているのか。

自分はなぜこの応答に惹かれたのか。

それを現実へ接続したとき、何が変化するのか。

この問いを通して、人工知能の出力は、自己の欲望を満たす即時的解答から、自己と自己以外を接続するインターフェイスへ変化する。

3.24 本章のまとめ

本章では、人工知能のハルシネーションを逆問題芸術の観点から再配置し、芸術と人工知能、自己と自己以外、インターフェイス、現実接地の関係を考察した。

逆問題芸術は、既知の問題へ正解を提示する芸術ではない。現実に見れた断片、違和感、失敗、痕跡から、その背後にある関係を遡及的に推定し、まだ存在していなかった問題を生成する実践である。

芸術と人工知能は、既存要素から新たな関係を生成し、限定された入力から過剰な意味を生み、未来可能性を現在へ提示する点で共通する。しかし、両者は存在論的に同一ではない。

人工知能は意味を生成するが、その意味を人間と同じ仕方で生きるわけではない。芸術

は、生成された意味を身体、生活史、他者、責任、有限性へ接続し、自己と現実を変容させる。

この接続を可能にするものがインターフェイスである。インターフェイスとは、既存の二項を単に媒介する表面ではなく、異なる位相の接触から第三の関係を生成する出来事的境界である。

自己と自己以外の接続において重要なのは、差異を消去しないことである。自己以外は、自己の認識へ完全に回収される対象ではなく、自己の境界を変化させる契機である。

また、生成された可能性は、現実接地を通して初めて、現実を変容させる力を持つ。現実接地とは、事実との一致だけではなく、身体、物質、行為、社会、存在との具体的関係を獲得する過程である。

この観点から、ハルシネーションと創造性の差異も明確になる。両者は共通の生成能力から生じうるが、ハルシネーションは、生成された可能性が適切な位相標識と現実接地を欠いたまま、事実として提示される状態である。これに対して創造性は、現実から位相差を作り、新たな可能性を生成し、その可能性を再び現実へ接続する運動である。

したがって、逆問題芸術における人工知能の役割は、完成された解答を供給することではない。それは、自己が自己以外と遭遇し、自らの前提を発見し、未知の現実との接続を開始するためのインターフェイスとなることである。

第4章 言語デリバティブ空間

—ナラティブ・時間・デリバティブ空間・ブラック＝ショールズとの比較—

4.1 はじめに

本章では、言語デリバティブが生成され、増幅され、現実的效果を持つ関係領域を「言語デリバティブ空間」として理論化する。

前章までに論じたように、言語デリバティブとは、既存の言語的蓄積から派生し、まだ現実化していない意味や未来可能性を現在へ先行的に生成する構造である。それは、一つの文章や命題だけに存在するのではない。言語デリバティブは、複数の発話、期待、記憶、制度、価値判断、未来予測が相互に接続されることによって、一定の空間的広がり形成する。

ここでいう「空間」は、単なる物理的容器ではない。それは、意味、価値、期待、時間、行為可能性が配置され、それらの関係によって個々の出来事の位置づけが決定される関係的領域である。

たとえば、ある企業について「将来性がある」というナラティブが成立した場合、その語りは単なる評価に留まらない。その企業の株価、投資判断、採用、研究開発、報道、社会的期待を変化させる。未来についての言語が現在の行為を変え、その行為がさらに未来の現実を形成する。

同様に、人工知能に関するナラティブも、人工知能の技術的能力を記述するだけではない。「人工知能は人間を超える」「人工知能は仕事を奪う」「人工知能は創造性を民主化する」といった言説は、研究投資、政策、教育、企業戦略、人間の自己理解を変化させる。

このとき言語は、現実を表象しているだけではない。未来の現実を先取りし、その先取りによって現在を変形している。

本章の目的は、この未来から現在への作用が、ナラティブ、時間、価値、現実生成の関係空間をどのように形成するのかを明らかにすることである。そのため、第一にナラティブを時間的・存在論的構造として再定義する。第二に、言語デリバティブにおける非線形時間を考察する。第三に、デリバティブ空間を可能性と現実の間に形成される関係領域として定義する。第四に、ブラック＝ショールズ・モデルとの比較を通して、金融デリバティブと言語デリバティブの構造的対応と決定的差異を明確にする。

4.2 ナラティブとは何か

ナラティブは、一般に、出来事を時間的・因果的に配列した物語として理解される。しかし、本論文においてナラティブは、すでに起こった出来事を説明するための形式だけ

ではない。

ナラティブとは、

複数の出来事、記号、主体、価値を一つの時間的方向へ接続し、未来に対する期待を現在の行為へ変換する生成構造

である。

ナラティブは、断片的な出来事を単なる集合として放置しない。それらの出来事に始点、転換点、目的、因果関係、意味を与える。

たとえば、一つの技術的発見は、それ自体では限定された出来事である。しかし、それが「人類社会を根本的に変える技術革命」の一部として語られた瞬間、その発見は別の時間的位置を獲得する。

現在の小さな出来事が、未来の巨大な変化の予兆として解釈される。

このとき、未来はまだ実現していない。それにもかかわらず、未来についての物語が、現在の出来事の価値を変化させる。

したがって、ナラティブは過去から未来へ進む記述ではない。それは、想定された未来から現在を遡及的に意味づける時間構造を持つ。

4.3 出来事とナラティブの差異

出来事は、現実には生じた変化である。しかし、出来事の意味は、その出来事だけから自動的に決まるわけではない。

同じ出来事であっても、どのナラティブへ配置されるかによって意味が変化する。

ある企業の赤字は、「衰退の兆候」として語られる場合もあれば、「将来への先行投資」として語られる場合もある。

人工知能が誤った回答を生成した出来事は、「人工知能の限界を示す失敗」として理解される場合もあれば、「人間の創造的思考に近づいた証拠」として語られる場合もある。

作品が理解されなかったことは、「制作の失敗」とされる場合もあれば、「時代に先行しすぎた結果」と解釈される場合もある。

出来事は一つであっても、ナラティブは複数存在する。

しかし、ナラティブは恣意的な装飾ではない。どのナラティブが社会的に受容されるかによって、資源配分、制度的評価、行為の方向が変化する。

したがって、ナラティブは、出来事の意味を説明するだけでなく、出来事が持つ現実的效果を増幅または縮小する。

この増幅作用が、言語デリバティブにおけるレバレッジの一形態である。

4.4 ナラティブのレバレッジ

一つの言葉や出来事が、ナラティブへ組み込まれることで、それ自体を超える価値を獲得する。

限定された技術的成果が、「次世代社会の基盤」というナラティブに接続されれば、その成果は研究上の意味を超えて、投資、政治、文化、教育に影響を与える。

一人の芸術家の作品が、「新しい芸術運動の始まり」として語られれば、その作品は個別の制作物を超え、歴史的転換の兆候として扱われる。

ここでは、出来事そのものの規模と、出来事が生み出す効果の規模が一致しない。

小さな出来事が、巨大なナラティブを作動させることによって、社会的・経済的・認識的效果を増幅する。

この構造を、本論文では「ナラティブ・レバレッジ」と呼ぶ。

ナラティブ・レバレッジとは、限定された出来事や記号が、未来についての物語へ接続されることによって、その直接的内容を超える価値と行為効果を獲得する構造である。

ナラティブ・レバレッジは、必ずしも虚偽を意味しない。将来の可能性を共有し、集団的行為を可能にするためには、ナラティブが必要である。

しかし、ナラティブが現実的根拠から離れ、その語り自体によって価値を自己増殖させるとき、過剰レバレッジが生じる。

4.5 言語デリバティブとナラティブ

言語デリバティブとナラティブは同一ではない。

言語デリバティブは、既存の言語的關係から未来可能性を現在へ派生させる生成単位または生成構造である。

これに対してナラティブは、複数の言語デリバティブを時間的方向へ統合し、それらを一つの現実像として組織する構造である。

たとえば、「人工知能が人間の仕事を代替する」という命題は、一つの言語デリバティブである。それは現状の技術的变化から、未来の労働形態を派生的に予測している。

この命題が、過去の産業革命、現在の自動化、未来の失業、教育制度の変更、普遍的所得の議論へ接続されると、一つのナラティブが形成される。

ナラティブは、個々の未来命題を相互に接続し、一つの方向性を与える。

したがって、言語デリバティブが未来可能性を生成するならば、ナラティブは、その可能性を時間的に配列し、社会的現実へ浸透させる。

言語デリバティブは、ナラティブの構成要素である。同時に、成立したナラティブは、新たな言語デリバティブを生み出す原資産となる。

この相互生成によって、言語デリバティブ空間は拡張する。

4.6 ナラティブの自己増殖

一度成立したナラティブは、それに整合する出来事を選択的に吸収する。

人工知能の発展を「人間を代替する過程」として理解するナラティブが成立すれば、人工知能の新機能は代替の証拠として読まれる。一方、人間にしかできない仕事が残る事実は、過渡期の例外として処理される可能性がある。

逆に、人工知能を「人間の能力を拡張するインターフェイス」として理解するナラティブでは、同じ技術的進展が協働可能性の証拠となる。

ナラティブは、すべての出来事を中立的に受容するのではない。それは、自らの方向性に適合する情報を強化し、不適合な情報を周辺化する。

この意味でナラティブは、情報を配置する位相空間である。

出来事は、ナラティブの内部で位置を与えられる。中心的証拠となるものもあれば、例外、誤差、ノイズとして処理されるものもある。

さらに、成立したナラティブは、新たな発話や行動を誘発する。その行動が現実を変え、変化した現実がナラティブをさらに強化する。

ここに自己増殖的循環が形成される。

ナラティブ

→ 現在の判断

→ 行為・投資・制度変更

→ 現実の変化

→ ナラティブの強化

この循環は、ナラティブが単なる虚構ではなく、現実を生成する力を持つ理由である。

4.7 ナラティブとハルシネーション

人工知能のハルシネーションも、ナラティブの自己増殖によって強化される。

一つの誤った前提が受容されると、人工知能はその前提に整合する後続情報を生成する。生成された後続情報が新たな前提となり、さらに多くの派生情報が作られる。

この構造は、次のように表現できる。

誤った前提

→ 整合的な生成

→ ナラティブ化

→ 内部的信頼性の上昇

→ さらに派生的生成

ここで重要なのは、現実との距離が拡大するにもかかわらず、ナラティブ内部の整合性は高まる可能性があることである。

ナラティブが完成されるほど、個々の誤りは孤立した誤情報ではなく、相互に補強し合う体系となる。

その結果、利用者は、一つ一つの事実を検証する代わりに、物語全体のもっともらしさによって内容を信頼する。

したがって、ハルシネーションの危険性は、事実誤認だけではない。それは、虚偽がナラティブとして組織され、自己完結的な現実性を獲得することにある。

4.8 時間の線形モデル

近代的な時間理解では、時間は一般に、過去、現在、未来という順序で一方向に進行するものとされる。

過去 → 現在 → 未来

このモデルでは、過去の原因が現在の結果を生み、現在の行為が未来を形成する。

この因果的時間は、科学的説明、歴史記述、計画、責任を成立させるうえで不可欠である。

しかし、人間の実際の認識と行為は、この線形構造だけでは説明できない。

人間は未来の目標によって現在の行為を変える。将来の危険を予測して、現在の選択を変更する。死を意識することで、現在の生き方を再評価する。

未来はまだ存在しない。それにもかかわらず、未来についての予測や期待は、現在に因果的效果を持つ。

金融市場、政治、宗教、芸術、人工知能は、この未来から現在への作用を特に強く表す領域である。

4.9 未来から現在への畳み込み

本論文では、未来可能性が現在の判断や価値へ組み込まれる構造を、「未来の現在への畳み込み」と呼ぶ。

畳み込みという語は、未来の出来事が完成した形で現在へ移動することを意味しない。

将来についての複数の可能性、期待、危険、欲望が圧縮され、現在の一つの価格、判断、言語、作品、行為へ反映されることを意味する。

株価には、企業の現在の資産だけでなく、将来の利益予測が畳み込まれている。

契約には、まだ実行されていない将来の行為が畳み込まれている。

作品には、まだ存在しない鑑賞経験や歴史的評価が先取りされる。

生成AIの出力には、入力に続くうる複数の未来的言語系列が圧縮され、そのうちの 하나가現在の文章として現れる。

したがって、言語生成とは、単に過去のデータを現在へ再生することではない。

それは、過去の言語的蓄積から複数の未来可能性を計算し、その可能性を現在の一つの出力へ畳み込む過程である。

4.10 非線形時間としての言語

言語は、線形時間を記述するだけでなく、時間の方向そのものを組み替える。

過去の出来事は、現在の言葉によって新たな意味を与えられる。

現在の行為は、想定された未来から評価される。

未来の目標は、現在の選択を規定する。

このとき、過去、現在、未来は独立した点ではなく、相互に意味を変化させる関係となる。

たとえば、一つの過去の失敗は、その後の成功によって「必要な過程」として再解釈される。成功がなければ単なる失敗であった出来事が、未来の結果によって遡及的に意味づけ直される。

ナラティブは、この遡及的意味生成を行う。

過去の出来事

← 現在の解釈

← 想定された未来

したがって、言語デリバティブ空間における時間は、単なる一方向の流れではない。それは、未来予測が現在を変え、変化した現在が過去の意味を再構成する再帰的構造を持つ。

4.11 投企としての未来

人間は、まだ存在しない可能性へ向けて自己を投げ出すことによって、現在の自己を形成する。

将来どのような人間になるか。

何を制作するか。

どこへ行くか。

どのような社会を作るか。

これらの未来可能性は、現在の自己から完全に独立した幻想ではない。同時に、現在の自己の単純な延長でもない。

未来は、現在にはまだ存在しない自己以外として現れる。

人間は、その未来の自己以外へ自己を投企し、その投企によって現在を変化させる。

この意味で、未来とは単に後から到来する時間ではない。現在の自己を外部から組織する力である。

言語デリバティブは、この投企を言語的形式へ変換する。

「こうなるだろう」

「こうなりたい」

「こうなるべきである」

「こうなる可能性がある」

これらの言葉は、未来について記述しているだけではない。それらは、現在に方向と価値を与える。

4.12 死とデリバティブ時間

未来から現在への作用を最も根源的に示すものが、自己の死である。

死は、自己が経験後に回収できない未来である。人間は自己の死を追い越し、死後の地点から現在へ戻ることができない。

それにもかかわらず、人間は死を予期し、その予期によって現在を組織する。

建築、作品、子孫、記録、制度、宗教、国家、技術は、有限な自己を超えて持続する未来への投企として理解できる。

人間は、自己が存在しない未来へ、自己の言葉、物、価値、記憶を送り込もうとする。この意味で、文化的制作は、死を超えた未来に対する一種の存在論的デリバティブである。

作品は作者そのものではない。しかし、作者の不在後にも価値や意味を生み出しうる。

言語も、発話者から分離され、未来の読者へ到達する。

人工知能は、過去に記録された無数の言語を現在へ再配置し、それらから未来の文章を生成する。

言語デリバティブ空間は、過去の死者、現在の利用者、未来の可能性が同時に接触する非線形時間の空間でもある。

4.13 場と空間の区別

言語デリバティブ空間を理解するには、「場」と「空間」を区別する必要がある。

本論文において、場とは、身体、記憶、歴史、関係、出来事が蓄積された個別具体的な状況である。

これに対して空間とは、それらの個別性から一定程度切り離され、対象、位置、関係を配置可能にする抽象的構造である。

場は、その場所で何が起こったかという時間の蓄積を持つ。

空間は、何がどこへ配置されうるかという可能性の構造を持つ。

言語は、個別具体的な場を抽象的な空間へ変換する。

ある場所で起こった出来事が言語化されると、その言葉は元の場所を離れ、別の時代、別の地域、別の主体へ接続可能になる。

この変換によって、出来事は流通性と再利用可能性を獲得する。

しかし、同時に、出来事が接地していた身体、時間、関係が失われる危険も生じる。

言語デリバティブ空間とは、この場から空間への変換が高度に進行した領域である。

4.14 デリバティブ空間の定義

以上を踏まえ、本論文ではデリバティブ空間を次のように定義する。

デリバティブ空間とは、現実の対象、出来事、関係から派生した記号的価値が、未来可能性を現在へ織り込みながら相互に交換・増幅される関係空間である。

この空間には、少なくとも五つの特徴がある。

第一に、派生性である。

デリバティブ空間の価値は、何らかの現実的基盤から派生する。

第二に、予測性である。

その価値は、現在の状態だけでなく、将来の可能性によって決定される。

第三に、増幅性である。

限定された入力、出来事、資本、記号が、大きな効果を生みうる。

第四に、再帰性である。

予測が現在を変化させ、変化した現在が予測の妥当性を変化させる。

第五に、接地不安定性である。

派生的価値が基盤から離れ、それ自体で増殖する可能性を持つ。

金融市場はデリバティブ空間の典型であるが、それだけではない。

芸術市場、学術評価、ブランド、SNS、政治的世論、人工知能の情報空間もまた、現実から派生した記号的価値が、将来期待によって増幅される空間である。

4.15 言語デリバティブ空間

言語デリバティブ空間とは、言葉が個別の事実を指示するだけでなく、他の言葉、期待、価値、未来予測と接続されることによって、派生的意味を増殖させる空間である。

この空間では、意味は一つの発話者や文章に固定されない。

一つの言葉が引用され、要約され、翻訳され、再解釈され、人工知能によって再生成されるたびに、新たな派生的意味が生じる。

たとえば、「創造性」という語は、芸術、教育、企業経営、人工知能、地域振興の各文脈で異なる価値を持つ。その語は、複数のナラティブへ接続されることで、単一の定義を超える社会的効果を獲得する。

人工知能は、この言語デリバティブ空間を高速に横断し、異なる文脈を接続する。

しかし、その接続は必ずしも元の場の条件を保持しない。

ある概念が、歴史的・身体的・制度的文脈から切り離され、別の文脈へ移植される。その結果、新たな意味が生まれることもあれば、概念の誤用や権威の擬態が生じることもある。

したがって、言語デリバティブ空間は、創造とハルシネーションの双方を可能にする。

4.16 デリバティブ空間と原資産

金融デリバティブには原資産がある。同様に、言語デリバティブにも、その派生元となる基盤が存在する。

しかし、言語における原資産は一つではない。

言語デリバティブの原資産には、少なくとも次のものが含まれる。

現実の出来事。

身体的経験。

歴史的記録。

既存の文章。

社会制度。

他者の証言。

画像、音、物質的痕跡。

学習データ。

利用者のプロンプト。

人工知能の生成物は、これらの原資産から直接一対一に派生するわけではない。複数の原資産が圧縮され、相互の関係が再構成され、一つの文章として現れる。

したがって、生成文から特定の原資産を逆算することは困難である。

ここに、言語デリバティブの逆問題性がある。

さらに、派生した文章が公開されると、その文章自体が別の生成における原資産となる。

現実から派生した言語が、新たな言語を派生させる。派生物が次の原資産となり、階層的な生成連鎖が形成される。

この連鎖が長くなるほど、現実的起源の追跡は困難になる。

4.17 原資産なきデリバティブ

言語デリバティブ空間では、実在する原資産を持たないように見える言語も生成される。

存在しない論文、人物、出来事、理論が、既存の言語形式から構成される場合である。

しかし、これらも完全な無から生まれたわけではない。

架空の論文名は、実在する学術用語、著者名の形式、論文題目の構造、出版制度の語彙から派生している。

存在しない人物の経歴も、実在する複数の人物像や社会的類型から合成される。

したがって、個別の指示対象という意味では原資産を持たなくても、言語形式や関係パターンという意味では原資産を持つ。

ここで、現実的の原資産と記号的の原資産を区別する必要がある。

現実的原資産とは、実在する対象、出来事、経験である。

記号的原資産とは、既存の言葉、形式、類型、ナラティブである。

ハルシネーションは、記号的原資産からは十分に派生しているが、現実的原資産との接続を欠いた言語デリバティブとして理解できる。

4.18 デリバティブ空間と価値

デリバティブ空間では、価値は対象の現在状態だけではなく、将来の可能性によって形成される。

企業の価値は、現在の利益だけでなく、将来の成長期待に左右される。

芸術作品の価値は、現在の評価だけでなく、将来の美術史的位置づけ、希少性、作家の評価上昇によって変化する。

学術理論の価値も、現在の説明能力に加え、今後どの領域へ応用可能かという期待によって形成される。

人工知能の回答も、単なる現在の情報としてではなく、次の思考、制作、判断を可能にする価値を持つ。

ここで価値は、対象内部に固定された属性ではない。

価値とは、

対象が未来の関係をどの程度生成しうるかについての現在の評価である。

この定義において、価値は本質的にデリバティブ的である。

しかし、未来可能性への評価は不確実である。そのため、価値はナラティブ、権威、期待、恐怖によって大きく変動する。

4.19 ブラック＝ショールズ・モデルの理論的位置

ブラック＝ショールズ・モデルは、金融オプションの価格形成を記述する代表的な数理モデルである。

その意義は、将来の不確実な価格変動を、現在のオプション価格へ一定の条件のもとで織り込む理論的枠組みを提示した点にある。

ここで重要なのは、ブラック＝ショールズ・モデルが未来を正確に予言する装置ではないということである。

それは、原資産価格、権利行使価格、満期までの時間、価格変動性、金利などの変数を用いて、不確実な将来に対する派生商品の現在価値を理論的に評価する。

本論文がブラック＝ショールズを参照するのは、その数式をそのまま言語へ適用するためではない。

金融デリバティブと生成言語には、測定可能性、取引制度、確率分布、複製可能性にお

いて重大な差異がある。

したがって、言語デリバティブにブラック＝ショールズと同一の価格方程式が成立すると主張するものではない。

本論文が比較するのは、不確実な未来可能性が、現在の派生的価値へ変換される構造である。

4.20 ブラック＝ショールズにおける時間

オプション価値において、満期までの時間は重要な変数となる。

将来までの時間が存在するということは、原資産価格が変化する可能性が残されているということである。

時間は、単なる外部的な時計ではない。可能性が展開する余地として価値形成に関与する。

一般に、他の条件が同じであれば、時間が多く残されているほど、複数の価格変動可能性が存在する。

この構造は、言語デリバティブにも一定の対応を持つ。

一つの概念や作品は、生成された瞬間に意味が確定するとは限らない。時間の経過によって、新たな文脈、解釈、制度、技術へ接続され、価値が変化する。

未完成の理論や新しい芸術実践は、現在の完成度だけでなく、将来どのような関係を生成しうるかによって評価される。

したがって、言語デリバティブにおける時間も、単なる経過ではなく、接続可能性が展開される余地である。

4.21 変動性と意味の不確定性

ブラック＝ショールズ・モデルでは、原資産価格の変動性が重要な役割を果たす。

変動性が高いとは、将来価格の不確実性が大きいことを意味する。不確実性は危険である一方、オプションが利益を生む可能性も拡大する。

言語デリバティブ空間にも、意味の変動性が存在する。

意味の変動性とは、ある言葉、作品、概念が、文脈によってどの程度異なる意味や価値を獲得しうるかを示す。

意味が完全に固定された命令文は、変動性が低い。

詩、芸術作品、哲学概念、象徴的表現は、複数の解釈や未来的接続を持つため、意味の変動性が高い。

人工知能は、意味の変動性が高い言語空間において、創造的な接続を生成できる。

しかし、変動性が高いほど、誤解、誤用、ハルシネーションの可能性も高まる。

したがって、意味の変動性は、創造性と危険性の双方を増幅する。

4.22 言語デリバティブにおける「行使価格」

金融オプションには、将来、原資産を特定の価格で売買する権利の基準となる行使価格がある。

言語デリバティブには、これと完全に同一のものは存在しない。しかし、構造的に対応するものとして、「現実化の閾値」を考えることができる。

ある言葉や未来予測が、単なる可能性から現実的行為へ転換されるには、一定の条件が必要となる。

仮説が十分な証拠を得る。

政策案が制度的支持を獲得する。

芸術作品が制作・展示される。

人工知能の提案が人間によって採用される。

ナラティブが多数の主体に共有される。

この条件を超えたとき、言語は現実「行使」される。

したがって、言語デリバティブにおける行使とは、言葉を発すること自体ではない。

生成された可能性が、判断、制作、制度、投資、身体行為へ変換されることである。

4.23 ボラティリティとしての位相差

言語デリバティブにおける意味の変動性を、単なる曖昧さと理解してはならない。

変動性は、現在の意味と、将来成立しうる複数の意味との間にある位相差である。

位相差が小さい場合、生成される未来は現在の延長に留まる。

位相差が大きい場合、現在の文脈からは予測困難な意味や関係が生成される。

芸術的創造は、しばしば大きな位相差を必要とする。既存の理解から離れなければ、新しい知覚や関係は生まれない。

しかし、位相差が大きいほど、現実接地は困難になる。

したがって、逆問題芸術の課題は、位相差を消去することではない。大きな位相差を保持しながら、それを現実へ接続できるインターフェイスを設計することである。

4.24 ヘッジと現実接地

金融デリバティブでは、価格変動による危険を管理するために、ヘッジという操作が用いられる。

ヘッジは、すべての変動可能性を消去することではない。異なる資産やポジションを組み合わせることで、特定の危険への曝露を調整する。

言語デリバティブにおいて、現実接地は一種のヘッジに対応する。

人工知能の生成物を、複数の資料で検証する。

仮説を実験や実践へ接続する。

一つのナラティブに対して反対の視点を導入する。

作品の解釈を身体的経験や歴史的文脈と照合する。

これらの操作は、生成可能性を消去するのではなく、その危険を調整する。

ただし、現実接地を金融的リスク管理へ完全に還元してはならない。

金融的ヘッジの目的は、一定の価値尺度における損失を制御することである。

一方、芸術的・存在論的接地では、何が損失であり、何が価値であるか自体が変化する場合がある。

作品との遭遇によって既存の自己理解が崩れることは、短期的には損失に見えても、存在生成の契機となりうる。

4.25 無裁定条件と意味の価値

ブラック＝ショールズ・モデルは、裁定機会が存在しないという市場条件と深く関係する。

裁定とは、理論上、無リスクで利益を得られる価格の不整合を利用することである。市場参加者の取引によって、この不整合は修正されると想定される。

言語デリバティブ空間では、価値の不整合を自動的に修正する単一の機構は存在しない。

誤情報が広く流通しても、必ずしも即座に修正されない。むしろ、感情的に強いナラティブや権威的な形式を持つ言説が、検証された情報よりも高い社会的価値を持つことがある。

言語空間には、意味の裁定が完全には働かない。

専門家、批評、教育、報道、検証制度は、意味の不整合を修正する機構となりうる。しかし、それらも権力、利害、制度的遅延から自由ではない。

したがって、言語デリバティブ空間は、金融市場以上に不均質であり、価値尺度も複数存在する。

事実的正确性。

美的価値。

政治的效果。

市場的注目。

倫理的妥当性。

存在論的変容。

これらは、単一の尺度へ還元できない。

4.26 複製可能性の限界

金融デリバティブの価格理論では、派生商品の収益を原資産と安全資産の組み合わせによって複製できるという考え方が重要となる。

しかし、言語や芸術の価値は、完全には複製できない。

同じ文章を別の人物が語っても、意味は同一にならないことがある。

同じ作品形式を再現しても、制作された時間、場所、作者、制度的文脈が異なれば、価値も変化する。

人間の言語は、身体、歴史、責任、関係に接地しているため、その効果を記号的内容だけから完全には複製できない。

人工知能は言語形式を精密に模倣できる。しかし、その形式が成立した生活史や存在条件まで複製するわけではない。

ここに、金融デリバティブと言語デリバティブの決定的差異がある。

金融契約は一定の条件のもとで同等のキャッシュフローへ還元可能である。

これに対して、言語的・芸術的出来事には、還元不可能な個別具体性が残る。

4.27 ブラック＝ショールズとの構造的対応

金融オプションと生成AIの言語生成には、次のような構造的対応を見いだすことができる。

金融デリバティブにおける原資産は、言語デリバティブにおける学習データ、現実の出来事、既存言説に対応する。

現在の原資産価格は、現在の言語文脈や支配的意味に対応する。

変動性は、意味の不確定性、解釈の幅、位相差に対応する。

満期までの時間は、未来的接続が展開する余地に対応する。

行使価格は、言語が行為や制度へ変換される現実化の閾値に対応する。

オプション価格は、未来可能性が現在において持つ期待価値に対応する。

ヘッジは、検証、批評、身体的実践、複数情報源との接続に対応する。

ただし、これら是一对一の数学的同値関係ではない。

この比較の目的は、言語を金銭的商品へ還元することではなく、未来の不確実性が現在の価値を形成する共通構造を抽出することにある。

4.28 ブラック＝ショールズとの決定的差異

両者には、少なくとも六つの決定的差異がある。

第一に、価値尺度の差異である。

金融デリバティブの価値は、原則として通貨単位で表現される。言語デリバティブの価値は、美的、認識的、倫理的、社会的、存在論的であり、単一尺度へ還元できない。

第二に、原資産の差異である。

金融デリバティブの原資産は、制度的に特定される。言語デリバティブは、複数の経験、文章、記憶、制度、身体から同時に派生する。

第三に、確率分布の差異である。

金融モデルは、一定の仮定のもとで価格変動を数理的に表現する。言語の意味変化には、歴史的断絶、政治的事件、創造的飛躍が含まれ、安定した分布を仮定しにくい。

第四に、満期の差異である。

金融契約には明確な満期があることが多い。言語、作品、理論の意味生成には、最終的満期が存在しない場合がある。

第五に、主体の差異である。

金融価格形成では、市場参加者の取引が中心となる。言語デリバティブでは、作者、利用者、読者、人工知能、制度、歴史、未来世代が関与する。

第六に、現実変容の差異である。

金融デリバティブは、主として既存の価値変動に対する契約である。言語デリバティブは、価値を評価するだけでなく、何が価値であるか、何が現実であるかを変化させる。

したがって、言語デリバティブ論は、ブラック＝ショールズ・モデルの言語的応用ではない。

それは、金融工学が明確化した「不確実な未来の現在価値化」という構造を、言語、人工知能、芸術、存在生成へ拡張して考察する理論である。

4.29 数式を持たないということ

言語デリバティブ論が現段階でブラック＝ショールズに対応する単一の数式を持たないことは、理論的欠陥とは限らない。

言語的価値は複数の尺度を持ち、主体間で異なり、歴史的文脈によって変化する。

また、芸術的出来事には、事前に計算不能な質的転換が含まれる。

一つの作品がどのような自己変容を引き起こすかを、制作以前に数値として完全に予測することはできない。

むしろ、完全な数式化が不可能であること自体が、言語デリバティブ空間の特徴を示す。

ただし、数式化不能性を理由に、無制限な恣意性を認めるべきではない。

言語デリバティブ論には、少なくとも次の変数を分析的に区別する必要がある。

現実的接地の強度。

記号的派生の階層数。

意味の変動性。

ナラティブの共有範囲。

入力と効果の非対称性。

位相標識の明確性。

現実化までの時間。

反証または修正可能性。

これらは、直ちに一つの方程式へ統合されなくても、個々の言語デリバティブの構造を比較するための理論的指標となる。

4.30 デリバティブ空間における人工知能

生成AIは、言語デリバティブ空間に外部から加わる道具ではない。

人工知能は、すでに存在する言語デリバティブを学習し、それらを圧縮し、新たな派生的言語を生成する。

その意味で人工知能は、言語デリバティブ空間の内部に形成された再帰的生成装置である。

人間が過去に生成したナラティブを人工知能が学習し、人工知能が生成した文章を人間が読み、それをさらに社会へ流通させる。

人工知能の出力が将来の学習データとなれば、派生物が新たな派生物の原資産になる。

ここでは、現実から言語への派生だけでなく、

言語

→ 人工知能による生成言語

→ 人間による再解釈

→ 社会的流通

→ 新たな学習データ

→ さらなる生成

という循環が形成される。

この循環が閉鎖的になれば、言語空間は現実との接続を失い、自らが生成したパターンを反復する。

これが、言語デリバティブ空間におけるシステムのハルシネーションの危険である。

4.31 自己参照的バブル

金融市場では、価格上昇への期待が買いを生み、その買いが価格をさらに上昇させることがある。

言語デリバティブ空間でも、類似した自己参照的バブルが生じる。

ある概念が注目されているという言説が、その概念へのさらなる注目を生む。

人工知能について語られる量が増えることで、人工知能が社会の中心的問題であるという認識が強化される。

一つの芸術家や理論が「重要である」と反復されることで、その反復自体が重要性の証拠として扱われる。

ここでは、対象の現実的内容よりも、対象について語られていることが価値を形成する。

ナラティブがナラティブ自身を原資産とする。

この状態を、本論文では「自己参照的言語バブル」と呼ぶ。

自己参照的言語バブルは、必ずしも完全な虚偽ではない。しかし、現実的成果と記号的評価の差が過度に拡大すると、接地不全が生じる。

4.32 デリバティブ空間と権威

言語デリバティブ空間では、権威も派生的に生成される。

肩書、所属、引用数、表示形式、専門用語、メディア露出は、発話内容とは別に信頼を形成する。

人工知能の流暢な文体も、権威の外観を生成する。

生成文が断定的で、体系的で、引用形式を伴う場合、利用者は内容を検証する前に、その形式を知識の証拠として受け取ることがある。

ここでは、知識そのものではなく、知識らしさが派生的価値を持つ。

これは、現実的原資産としての知識と、記号的デリバティブとしての権威が分離した状態である。

逆問題芸術は、この分離を可視化する。

なぜこの形式は信頼されるのか。

誰が権威を付与しているのか。

肩書と内容はどこで接続しているのか。

人工知能の語調は、どのように確実性を擬態するのか。

これらを問うことは、言語デリバティブ空間における現実接地の批判的検証となる。

4.33 デリバティブ空間における創造

デリバティブ空間は危険だけを生むのではない。

現実から一定の距離を取り、未来可能性を現在へ生成することは、あらゆる創造の条件でもある。

建築は、まだ存在しない空間を設計する。

芸術は、まだ共有されていない感覚を現象させる。

科学は、まだ検証されていない仮説を構築する。

政治は、まだ実現していない社会制度を構想する。

人間は、まだ存在しない自己を言語によって先取りする。

これらはすべて、現実から派生した未来可能性を現在へ提示するデリバティブ的行為である。

問題は、派生であることではない。

問題は、派生物が派生物であることを隠し、すでに確定した現実として振る舞うことである。

したがって、創造的デリバティブ空間には、可能性と事実の位相差を保持するインターフェイスが必要である。

4.34 現実接地としての再投資

金融デリバティブにおいて生じた価値が、現実の生産、研究、設備、生活へ還流する場合、その派生的価値は現実経済との接続を持つ。

しかし、派生的価値が金融空間内部だけで自己増殖すれば、現実との乖離が拡大する。言語デリバティブにも同様の問題がある。

生成された言葉が、さらに言葉を生み出すだけで、身体、制作、検証、他者との関係へ戻らなければ、言語は自己完結的になる。

人工知能によって生成された理論を、さらに人工知能が要約し、その要約をもとに別の文章を生成する。この循環が現実的实践へ接続されなければ、意味の金融化とも呼べる状態が生じる。

本論文では、生成された可能性を現実の行為や関係へ戻すことを「現実への再投資」と捉える。

現実への再投資とは、

生成された仮説を検証すること。

作品を制作し、場所へ置くこと。

他者との対話へ開くこと。

身体的経験によって修正すること。

制度的・倫理的責任を引き受けること。

である。

4.35 言語デリバティブ空間の倫理

言語デリバティブ空間における倫理は、未来を語ることを禁止するものではない。

人間は、未来を語らずには行為できない。あらゆる計画、約束、希望、創造は、まだ存在しないものについて語る。

問題は、未来可能性を、すでに存在する事実として提示することである。

また、未来の利益を語りながら、現在の危険や負担を他者へ移転することも問題となる。

したがって、言語デリバティブの倫理には、少なくとも次の原則が必要である。

第一に、派生性の開示である。

その言説が事実、推測、予測、仮説、創作のいずれであることを示す。

第二に、不確実性の保持である。

未来の複数可能性を、一つの必然的結論へ閉じない。

第三に、接地経路の提示である。

その言説がどの資料、経験、現実的条件から派生しているかを可能な限り示す。

第四に、影響責任である。

生成されたナラティブが、他者の判断、生活、制度へ与える効果を考慮する。

第五に、修正可能性である。

新たな現実や反証に応じて、ナラティブを修正できる状態を保持する。

4.36 デリバティブ空間と生成の位相存在論

生成の位相存在論において、存在は固定された実体ではなく、接続可能性と関係生成の過程として理解される。

この観点から見ると、デリバティブ空間とは、まだ確定していない存在可能性が、現在の関係へ作用する位相空間である。

そこでは、未来は現在の外部に待機しているのではない。

未来可能性は、ナラティブ、期待、予測、設計、作品を通して、すでに現在へ部分的に侵入している。

現在もまた、固定された点ではない。それは、過去の蓄積と未来の可能性が畳み込まれた接続面である。

したがって、存在生成は次の循環を持つ。

過去の蓄積

→ 現在の場合

→ 言語的外部化

→ デリバティブ空間

→ 未来可能性

→ 現在への再畳み込み

→ 現実の変容

この循環において、言語は単なる表現手段ではない。場を空間へ変換し、未来可能性を現在へ運び、現実を再生成するインターフェイスである。

4.37 言語デリバティブ空間の定義

以上の議論を踏まえ、本論文では言語デリバティブ空間を次のように定義する。
言語デリバティブ空間とは、現実、経験、既存言説から派生した言語が、ナラティブを通して未来可能性を現在へ畳み込み、相互に価値、意味、行為可能性を増幅させる非線形かつ再帰的な関係空間である。

この空間は、次の諸特徴によって成立する。

第一に、現実から言語への派生がある。

第二に、言語から未来可能性への投企がある。

第三に、未来可能性から現在の価値判断への遡及的作用がある。

第四に、現在の行為が現実を変化させ、その変化がナラティブを再編成する。

第五に、派生した言語が新たな原資産となる自己再帰性がある。

第六に、現実接地が失われた場合、ハルシネーション、バブル、権威の擬態が生じる。

第七に、現実接地が成立した場合、仮説、芸術、科学、制度、自己変容が生成される。

4.38 本章のまとめ

本章では、言語デリバティブが生成され、相互に接続され、現実的効果を持つ領域を「言語デリバティブ空間」として理論化した。

ナラティブは、出来事を説明する物語に留まらない。それは、断片的な出来事、価値、主体を一つの時間的方向へ接続し、想定された未来を現在の行為へ作用させる生成構造である。

言語デリバティブが個別の未来可能性を生成するのに対し、ナラティブは複数の言語デリバティブを時間的に組織する。成立したナラティブは現在の判断と行為を変え、変化した現実がさらにナラティブを強化する。

この循環において、時間は過去から未来へ進むだけではない。未来への期待、予測、恐怖、欲望が現在へ畳み込まれ、現在の価値を変化させる。さらに、想定された未来は過去の出来事を遡及的に意味づけ直す。

デリバティブ空間とは、現実から派生した記号的価値が、未来可能性を現在へ織り込みながら交換・増幅される関係空間である。金融市場だけでなく、芸術、学術、政治、ブランド、SNS、人工知能も、この構造を持つ。

ブラック＝ショールズ・モデルとの比較によって明らかになるのは、金融デリバティブと言語デリバティブの数学的同一性ではない。両者に共通するのは、不確実な未来可能性が、現在の派生的価値を形成する構造である。

一方、言語デリバティブには、単一の価値尺度、明確な満期、特定可能な原資産、安定した確率分布、完全な複製可能性が存在しない。言語や芸術は、価値を評価するだけでなく、価値の基準そのものを変化させる。

生成AIは、この言語デリバティブ空間を高速化し、過去の言語的蓄積から新たな未来

可能性を生成する。しかし、派生した言語が新たな原資産となり、人工知能による生成が自己循環するとき、現実から切断された自己参照的言語バブルが形成される危険がある。

したがって、言語デリバティブ空間において必要なのは、未来可能性の生成を停止することではない。その派生性、不確実性、位相差を明示し、生成された可能性を身体、制作、検証、他者、制度へ再接続することである。

言語デリバティブ空間は、現実から逃避する仮想領域ではない。それは、現実から距離を取り、まだ存在しない可能性を生成し、その可能性を再び現実へ接地させるための生成空間である。

第5章 レバレッジ構造

—情報・人工知能・金融・芸術—

5.1 はじめに

本章では、言語デリバティブ論の中心機構である「レバレッジ構造」を、情報、人工知能、金融、芸術の四領域を横断して考察する。

前章までに論じたように、言語デリバティブとは、現実、経験、既存言説から派生した言語が、未来可能性を現在へ織り込み、新たな意味や行為可能性を生成する構造である。この生成において重要なのは、入力と出力の規模が対称ではないという点である。短い発話が長期的な行動を変える。

一つの概念が巨大な制度を動かす。

小さなプロンプトが膨大な文章を生成する。

限定された資本が巨大な市場ポジションを形成する。

一つの作品が、個人や社会の認識を変化させる。

これらの現象には、領域ごとの相違を超えて、限定された入力、それを上回る効果を発生させる共通構造が存在する。

本論文では、この構造をレバレッジと呼ぶ。

金融においてレバレッジは、借入やデリバティブを利用し、自己資本を上回る取引規模を形成することを意味する。しかし、レバレッジの本質は金銭的倍率だけにはない。それは、ある作用点を介して、入力と効果の間に非対称性を生じさせる構造である。

したがって、本論文ではレバレッジを次のように定義する。

レバレッジとは、限定された入力、媒介構造を通して、その入力に還元できない規模または質の効果を生成する非対称的増幅構造である。

この定義において、増幅されるものは金銭に限定されない。情報、意味、信頼、感情、権威、行為可能性、社会的影響、存在変容も、レバレッジの対象となる。

ただし、レバレッジは価値を無から創造する魔術ではない。増幅には、あらかじめ蓄積された資源、関係、制度、記憶、技術、身体、期待が必要である。

プロンプトが人工知能の巨大な言語空間を作動させるように、作品は鑑賞者の記憶や社会的文脈を作動させる。少額の証拠金が市場ポジションを作るように、一つの言葉が既存のナラティブを起動する。

小さな入力が巨大な効果を生むのは、入力そのものに巨大な内容が含まれているからではない。入力が、すでに存在する潜在的な関係空間へ接続されるからである。

本章では、第一に情報そのものとレバレッジの関係を考察する。第二に、生成AIを情報および言語のレバレッジ装置として分析する。第三に、金融レバレッジとの構造的比

較を行う。第四に、芸術におけるレバレッジを、意味の増幅ではなく、存在変容の生成として理論化する。

5.2 レバレッジの基本構造

レバレッジという語は、槥子を意味する英語の leverage に由来する。物理的な槥子は、支点と棒を用いることで、小さな力からより大きな作用を生じさせる。ここで重要なのは、力の総量が無条件に増加するわけではないということである。槥子は、支点、距離、方向の配置を変えることで、力が作用する関係を変換する。この物理的構造は、情報、金融、人工知能、芸術におけるレバレッジを考える基礎となる。

レバレッジには少なくとも四つの要素がある。

第一は、入力である。

力、資本、情報、言葉、プロンプト、作品などがこれに当たる。

第二は、支点である。

制度、技術、媒体、アルゴリズム、ナラティブ、展示空間、社会的信頼などが支点となる。

第三は、潜在的資源である。

市場流動性、学習データ、社会的記憶、鑑賞者の経験、既存の言説空間などである。

第四は、出力または効果である。

利益、損失、生成文、社会的反応、行動変容、存在変容などとして現れる。

したがって、レバレッジは次の一般形式によって表現できる。

限定された入力

+ 支点となるインターフェイス

+ 潜在的関係資源

→ 非対称的に増幅された効果

ここで、支点は単なる道具ではない。何が増幅され、どの方向へ作用し、誰が利益または損失を受けるかを規定する。

同じ情報でも、個人的な会話で発せられる場合と、権威ある機関の発表として流通する場合では、その効果が異なる。同じ画像でも、私的記録として見る場合と、美術館で作品として展示される場合では、意味の作用が変化する。

レバレッジは内容だけではなく、内容が置かれる関係構造によって決定される。

5.3 倍率としてのレバレッジと位相転換としてのレバレッジ

レバレッジは、一般に数量的な倍率として理解される。

一の資本が十の取引を可能にする。

一の入力が百の出力を生む。

一つの投稿が百万人へ届く。

この理解は重要である。しかし、言語や芸術におけるレバレッジは、単純な数量増加だけでは捉えられない。

一つの言葉が百回反復されても、その意味が質的に変化しなければ、数量的増幅に留まる。

これに対して、一つの言葉が、受け手の世界理解そのものを変える場合、そこには位相転換が生じている。

例えば、ある人が作品との遭遇によって、自分がそれまで問題だと認識していなかった事柄を問題として知覚するようになる。この変化は、情報量の増加ではない。何が情報として重要であるかを決める認識構造の変化である。

したがって、本論文では二種類のレバレッジを区別する。

第一は、量的レバレッジである。

これは、入力に対して、出力の量、範囲、速度、頻度が増幅される構造である。

第二は、位相的レバレッジである。

これは、入力によって、対象の意味、価値、関係、認識の次元そのものが変化する構造である。

生成AIは、短いプロンプトから長文を生成するという量的レバレッジを持つ。同時に、利用者が自ら思いつかなかった概念や関係を提示し、思考の枠組みを変える場合には、位相的レバレッジを持つ。

芸術の本質的な力は、主として後者にある。

5.4 情報とは何か

情報は、一般に、伝達される内容または不確実性を減少させる差異として理解される。

しかし、情報はそれ自体で固定した意味を持つわけではない。

同じ数値、言葉、画像でも、文脈、受け手、時間、制度によって意味と価値が変化する。

「危険」という語は、誰が、いつ、どの状況で発するかによって作用が異なる。

株価の下落という情報も、長期投資家、短期投機家、企業従業員、消費者にとって異なる意味を持つ。

人工知能が生成した文章も、創作補助として用いられる場合と、医療判断の根拠として用いられる場合とでは、必要な正確性と責任が異なる。

したがって、情報は孤立した内容ではなく、差異がある関係の中で作用する出来事である。

本論文では、情報を次のように捉える。

情報とは、ある関係系に差異を導入し、その系における判断、意味または行為の可能性

を変化させる作用である。

この定義において、情報は単なるデータではない。

データがそこに存在しても、何の差異も生まなければ、実践的な意味では情報として作動していない。

逆に、非常に短い言葉であっても、人間の判断を根本的に変化させるならば、高い情報的レバレッジを持つ。

5.5 情報量と情報価値

情報量と情報価値は同一ではない。

長い文章が必ずしも高い価値を持つとは限らない。大量のデータが、必ずしも重要な判断を可能にするわけでもない。

情報量は、記号や選択肢の規模として測定できる場合がある。しかし、情報価値は、その情報が特定の主体、状況、時間において、どのような差異を生むかに依存する。

例えば、災害時の避難経路を示す一文は、膨大な一般知識よりも高い価値を持つことがある。

長年理解できなかった自己の経験を説明する一つ概念が、その人の生活全体を再構成することもある。

ここで、情報価値は次のような関係によって生じる。

情報内容

- × 状況の切迫性
- × 受け手との関連性
- × 行為への接続可能性

したがって、情報的レバレッジは、情報量を増やすことだけでは成立しない。

必要なのは、どの情報が、どの主体と状況に接続されたとき、どの程度の差異を生むかを考えることである。

5.6 情報の圧縮と展開

情報的レバレッジは、圧縮と展開の関係から理解できる。

言葉、数式、図像、記号は、複雑な関係を限定された形式へ圧縮する。

一つの数式は、多数の個別事例を一つの関係形式として表現する。

一つ概念は、多様な経験を分類し、比較可能にする。

一首の短歌は、限られた音数の中に、身体、季節、記憶、関係、歴史を圧縮する。

圧縮された情報は、受け手の内部で再び展開される。

しかし、その展開は元の内容の単純な復元ではない。受け手の経験、知識、感情、状況によって異なる意味が生成される。

したがって、情動的レバレッジとは、
圧縮された記号が、受け手の潜在的関係資源へ接続されることで、入力を超える意味を
展開する構造
である。
人工知能へのプロンプトも、同じ構造を持つ。
短い入力、巨大な学習済み関係空間を起動し、長い文章として展開される。

5.7 情報と支点

同じ情報でも、どの支点を通して流通するかによって影響力が変わる。
支点となりうるものには、次のようなものがある。
権威。
専門性。
制度。
メディア。
反復回数。
アルゴリズム。
感情的強度。
集団的信頼。
視覚的形式。
緊急性。
たとえば、匿名の発言と、公的機関の発表では、同じ文章であっても社会的効果が異なる。
人工知能が断定的な語調で回答すると、その形式が権威の支点となる。
SNSの推薦アルゴリズムが特定の投稿を反復表示すれば、反復そのものが重要性の印象を作る。
情報の力は、内容の真実性だけでは決まらない。それがどのインターフェイスを通して
可視化され、誰のもとへ、どの頻度で、どの形式において届くかに依存する。
このことは、情動的レバレッジが常に政治的・制度的問題を含むことを意味する。

5.8 情報のレバレッジと誤情報

情動的レバレッジは、正確な情報だけを増幅するわけではない。
むしろ、短く、単純で、感情的に強く、既存のナラティブへ適合する情報ほど、高い拡散力を持つ場合がある。
複雑な事実は、多くの条件や留保を必要とする。一方、誤情報は、単純な因果関係や明確な敵味方を提示できる。

したがって、情報空間では、事實的妥当性と増幅可能性が一致しない。
この不一致は、人工知能のハルシネーションにも存在する。
もっともらしく、明快で、体系的な誤答は、不確実性を正直に示す慎重な回答よりも、
利用者に強い印象を与えることがある。
したがって、情報倫理は、情報が正しいかどうかだけでなく、その情報がどのようなレ
バレッジを獲得し、どの範囲へ、どの速度で、どのような効果を伴って拡散するかを問
わなければならない。

5.9 AIは情報を生成するのか

生成AIは、一般に、情報を生成する装置として理解される。
しかし、厳密には、人工知能が常に新しい事実を生成しているわけではない。
人工知能は、既存のデータから学習した関係を再構成し、入力文脈に応じて、新たな文
章を生成する。
この出力が新しい情報になるかどうかは、出力そのものだけでは決まらない。
既知の内容を言い換えただけであれば、記号列は新しくても、認識上の差異は小さい。
一方、既存の知識を新たな関係へ配置し、利用者がそれまで見いだせなかった問題や可
能性を示すならば、人工知能の出力は新しい情報として機能する。
したがって、人工知能は、情報そのものを無から作るのではなく、潜在していた関係を
顕在化させる。
人工知能による情報生成とは、
既存の記号的関係空間から、特定の入力と文脈に応じて、新たな差異として作用しうる
関係を抽出・構成すること
である。

5.10 AIにおける入力と出力の非対称性

生成AIの最も明確な特徴の一つは、入力と出力の規模の非対称性である。
数語のプロンプトから、長い論文、物語、プログラム、計画、画像が生成される。
しかし、人工知能が入力された数語だけから、そのすべてを生成しているわけではな
い。
プロンプトは、学習過程で形成された巨大な潜在的関係空間に対する作用点である。
したがって、人工知能のレバレッジ構造は、次のように表現できる。
小さなプロンプト
+ 学習された巨大な関係空間
+ 推論アルゴリズム
+ 会話文脈

→ 大規模な生成出力

プロンプトは原資源ではなく、起動条件である。

この点は、人工知能の出力を利用者個人の創造物とみなす場合にも、人工知能だけの創造物とみなす場合にも重要である。

出力は、利用者の指示、モデル設計、学習データ、計算資源、制度的環境の複合的作用によって成立する。

5.11 プロンプトのレバレッジ

プロンプトは、人工知能へ命令を与える文章であると同時に、生成空間に方向を与えるポジションである。

どのような対象を前提とするか。

どの文体を指定するか。

どの役割を人工知能へ与えるか。

どの条件を明示し、何を省略するか。

これらの違いによって、生成結果は大きく変化する。

プロンプトのレバレッジは、長さだけには依存しない。

短いが前提の強いプロンプトは、生成空間を狭い方向へ誘導する。

詳細なプロンプトは、出力を精密に制約できる一方、利用者の前提や偏見をより強く固定する場合がある。

したがって、プロンプトは中立的な問いではない。それは、どの未来的言語系列を優先するかを決める投機的位置である。

この意味で、プロンプト設計は言語レバレッジの設計である。

5.12 AIと意味の信用創造

融制度において、信用は現在保有する資産を超える取引や投資を可能にする。

生成AIにも、これと構造的に類似する意味の信用創造がある。

利用者は、人工知能が提示した文章を、その場で完全に検証できなくても、一時的に利用可能な意味として受け取る。

人工知能は、明示された証拠以上の説明、因果関係、仮説、文章構造を提示する。

ここで、人工知能は、確定済みの知識だけを流通させているのではない。まだ検証されていない意味を、利用可能な形で先行的に供給している。

この構造を「意味の信用創造」と呼ぶことができる。

意味の信用創造は、思考や制作を高速化する。

利用者は、すべてを最初から調査することなく、仮説的な全体像を得て、次の検討へ進むことができる。

しかし、信用には返済または決済が必要である。

言語的意味における決済とは、事実確認、論理検証、身体的実践、他者との対話、現実的結果との照合である。

検証されない意味の信用が蓄積すれば、言語的債務が拡大する。

ハルシネーションとは、決済されていない意味が、確定済みの資産であるかのように流通する状態でもある。

5.13 AIとハルシネーションのレバレッジ

人工知能のハルシネーションが危険なのは、一つの誤りが一つの誤りに留まらないからである。

誤った前提が入力される。

人工知能がその前提を補完する。

補完された内容が後続生成の前提となる。

一貫したナラティブが作られる。

利用者がそれを引用・共有する。

別の主体がその情報を原資産として利用する。

この過程において、一つの誤認は多数の派生的誤認を生む。

したがって、AIハルシネーションのレバレッジは、少なくとも三段階を持つ。

第一は、生成レバレッジである。

短い誤前提から、大規模な虚構が生成される。

第二は、権威レバレッジである。

人工知能の流暢さや形式が、生成内容に信頼の外観を与える。

第三は、流通レバレッジである。

生成物が引用、共有、再学習され、さらなる情報生成の基盤となる。

この三段階が重なると、局所的な誤りが社会的な情報構造へ拡大する。

5.14 AIにおけるレバレッジの二重性

人工知能のレバレッジは、解放的でもあり、危険でもある。

解放的側面として、人工知能は、専門知識、文章作成、翻訳、分析、創作へのアクセスを拡張する。

個人が限られた時間と資源で、従来は組織的支援を必要とした作業を行えるようになる。

一方、危険な側面として、人工知能は、誤情報、偏見、表面的権威、依存、思考の均質化を高速に増幅する。

ここで重要なのは、レバレッジの倍率が高くなるほど、利用者の判断能力が不要になる

のではなく、逆に重要になることである。

高いレバレッジを持つ装置では、小さな判断差が大きな結果差を生む。

人工知能が高性能になるほど、人間は何を委任し、何を検証し、何を引き受けるかを精密に判断しなければならない。

5.15 金融レバレッジの基本構造

金融におけるレバレッジは、自己資本に借入資本や派生商品を組み合わせ、保有資本を上回る市場曝露を形成することである。

簡略化すれば、金融レバレッジの倍率は、総資産または市場曝露を自己資本で除した関係として理解できる。

しかし、本論文が注目するのは、倍率そのものよりも、その存在論的構造である。

金融レバレッジでは、現在保有する資本以上の未来的利益可能性を、現在のポジションとして形成する。

つまり、まだ実現していない利益が、現在の判断や価値へ織り込まれる。

同時に、損失も増幅される。

レバレッジは、利益だけを増やす装置ではない。未来の変動に対する感応度を高める装置である。

したがって、金融レバレッジの本質は、

限定された自己資本によって、未来の価格変動に対する大きな関係的位置を現在に形成することにある。

5.16 レバレッジと所有の分離

金融レバレッジでは、対象全体を直接所有しなくても、その価値変動へ大きく関与することができる。

デリバティブ取引では、原資産そのものを保有せずに、その価格変動に対するポジションを形成できる。

ここでは、対象の所有と、対象の変動への参加が分離する。

この構造は、言語デリバティブにも存在する。

利用者は、ある学問領域の全知識を所有していなくても、人工知能を介して、その領域の言語を生成し、利用できる。

ある文化や経験を身体的に生きていなくても、その形式を言語的に再構成できる。

この分離は、知識や表現へのアクセスを拡張する。

しかし同時に、経験、責任、理解を所有しないまま、その言語的効果だけへ参加する危険を生む。

金融において原資産を所有しないポジションが大きな損益を生むように、言語空間でも、経験に接地しない発話が大きな社会的効果を持つ。

5.17 証拠金とプロンプト

金融デリバティブでは、証拠金を差し入れることで、その額を上回る取引ポジションを持つことができる。

プロンプトは、言語生成における証拠金に類似した構造を持つ。

利用者は短い入力を差し出し、その入力を上回る大規模な生成結果を受け取る。

しかし、両者を単純に同一視することはできない。

金融の証拠金は、損失に備える法的・経済的担保である。

プロンプトは、必ずしも生成結果に対する責任を担保しない。

むしろ人工知能では、出力規模に対して、利用者の投入した注意、知識、責任が極端に小さい場合がある。

この不均衡こそが問題である。

言語レバレッジが高いにもかかわらず、責任の証拠金が低い。

大量の文章を生成し、公開し、他者へ影響を与えながら、生成内容の確認や責任をほとんど引き受けないことが可能になる。

したがって、人工知能時代には、経済的証拠金に代わる「認識的証拠金」または「責任的証拠金」が必要となる。

これは金銭ではなく、出典確認、生成過程の開示、修正可能性、用途制限、発信者による責任の引き受けとして実現される。

5.18 金融レバレッジと時間

金融レバレッジは、未来の価格変動へ現在の資本を接続する。

将来の利益を期待して、現在ポジションを持つ。期待が現在の需要を変え、現在価格が変化する。

したがって、金融レバレッジは、未来が現在へ畳み込まれる時間構造を強化する。

利益が将来実現する前に、その可能性が現在価値として評価される。

しかし未来が予測と異なれば、現在のポジションは急速に損失へ転化する。

レバレッジは、時間を短縮する。

将来にわたって徐々に生じる価値変化を、現在の価格変動として先取りする。

人工知能も、思考、調査、執筆に必要であった時間を圧縮する。

芸術も、長い歴史や経験を一つの作品へ圧縮し、鑑賞者に瞬間的な遭遇として提示する。

したがって、レバレッジは、量の増幅だけでなく、時間の圧縮でもある。

5.19 金融レバレッジとリスク

レバレッジは、利益可能性と同時に、破綻可能性を拡大する。
自己資本のみであれば耐えられる価格変動であっても、高いレバレッジを用いると、わずかな変化が資本全体を失わせる。
したがって、レバレッジの危険は、予測が誤ることだけにはない。
小さな予測誤差が、大きな損失へ変換される感応度にある。
言語デリバティブにも同じ問題がある。
一つの事実誤認が、個人的なメモに留まるならば影響は限定される。
しかし、人工知能、権威、SNS、制度を通して増幅されれば、同じ誤認が巨大な社会的損失を生む。
したがって、リスクは誤りの大きさだけではなく、誤りにかかるレバレッジによって決まる。

5.20 清算と現実接地

金融ポジションは、最終的に市場価格、契約、決済によって清算される。
期待と現実の差は、損益として現れる。
言語デリバティブにも、清算に相当する過程が必要である。
予測は実際の出来事によって検証される。
仮説は実験によって試される。
人工知能の出力は資料と照合される。
芸術的構想は制作と鑑賞によって現実の抵抗を受ける。
本論文が現実接地と呼ぶものは、言語的可能性と現実との差を可視化する清算過程でもある。
ただし、言語や芸術では、金融のような単一時点の最終決済が存在するとは限らない。
理論や作品の意味は、後世に再評価される。新しい文脈によって価値が変化する。
したがって、言語デリバティブの清算は、一回的な終結ではなく、継続的な検証と再接続である。

5.21 金融化としての情報空間

情報空間が金融化するとは、情報が現実を理解するためのものではなく、注目、評価、影響力を増幅するための派生的資産として扱われることである。
ある発言の価値が、その内容ではなく、どれだけ拡散するかによって評価される。
ある人物の権威が、実際の成果よりも、引用数、フォロワー数、メディア露出によって形成される。

人工知能によって大量の文章が生成されると、内容の理解よりも、情報空間における占有率が重要になる場合がある。

この状態では、情報は現実への接続ではなく、さらなる情報を生成するための担保となる。

言葉が言葉を生み、評価が評価を生み、注目が注目を生む。

これが情報空間におけるレバレッジの自己循環である。

5.22 芸術におけるレバレッジ

芸術も、限定された形式から巨大な意味や経験を生成する。

一枚の絵画。

一つの物体。

短い詩。

一度のパフォーマンス。

限定された空間配置。

これらは物理的には小さい。しかし、それによって鑑賞者の記憶、感情、歴史認識、身体感覚、社会的関係が作動する。

したがって、芸術は本質的にレバレッジ構造を持つ。

ただし、芸術のレバレッジは、作品内部に意味が大量に保存されていることを意味しない。

作品は、鑑賞者や社会の潜在的関係資源へ作用する支点である。

作品の小ささと、経験の大きさの間に非対称性がある。

本論文では、芸術的レバレッジを次のように定義する。

芸術的レバレッジとは、限定された物質的・記号的構成が、身体、記憶、歴史、他者、場所との接続を通して、入力に還元できない知覚的・社会的・存在論的変容を生成する構造である。

5.23 作品は意味の容器ではない

作品を意味の容器として理解する場合、作者が作品内部へ意味を入れ、鑑賞者がそれを取り出すと考えられる。

しかし、芸術的レバレッジは、この伝達モデルでは説明できない。

作品の効果は、作者が意図した意味の量を超えることがある。

鑑賞者は、作者が予想しなかった経験をする。後世の社会は、制作時には存在しなかった文脈から作品を再解釈する。

作品は意味を保存するだけでなく、意味生成を起動する。

したがって、作品は情報の倉庫ではなく、関係を作動させる支点である。

作品の意味は、作品、鑑賞者、場所、時間、制度の接続において生成される。
この構造は、人工知能へのプロンプトと類似する。
プロンプトが巨大な言語空間を起動するように、作品は鑑賞者と社会の潜在的意味空間を起動する。
しかし、芸術作品には、物質的抵抗、場所性、歴史性、作者の身体的時間が残る。この点で、人工知能の純粋な言語生成とは異なる。

5.24 芸術と量的増幅

芸術におけるレバレッジは、作品が多数の人間へ届くこととして現れる場合がある。
複製技術、出版、写真、映画、放送、インターネットは、一つの作品を広範囲へ流通させる。
デジタル技術と人工知能は、この量的レバレッジをさらに高める。
一人の作者が、短時間に多くのイメージ、文章、音楽を生成し、世界規模で公開できる。
しかし、流通量の増大は、芸術的価値の増大を自動的には意味しない。
反復されるほど意味が深まる作品もあれば、過剰流通によって消費され、差異を失う作品もある。
したがって、量的レバレッジと芸術的レバレッジを区別しなければならない。
大量に見られることと、見る主体を変化させることは同一ではない。

5.25 芸術と位相的レバレッジ

芸術に固有のレバレッジは、位相的レバレッジである。
作品との遭遇によって、鑑賞者が新しい情報を知るだけでなく、何を現実として知覚するかが変化する。
それまで背景であったものが、前景として現れる。
無意味に見えていたものが、切迫した問題となる。
自己とは無関係だと考えていた他者や自然が、自己の存在条件として経験される。
ここでは、作品が既存の認識空間の内部で価値を増やすのではなく、その認識空間の境界を変える。
この変化は、倍率として測定することが困難である。
しかし、芸術における最も深い増幅とは、意味の量ではなく、存在の位相が変わることである。

5.26 芸術と情報の差異

芸術は情報を含む。しかし、芸術を情報伝達へ還元することはできない。
情報伝達では、発信者と受信者の間で、内容ができるだけ正確に共有されることが重視される。
一方、芸術では、伝達の失敗、曖昧さ、沈黙、違和感、解釈のずれが、重要な役割を持つ。
芸術作品は、既知の内容を効率よく伝えるのではなく、既存の理解では処理できない差異を発生させることがある。
したがって、芸術的レバレッジは、情報の明確化だけではなく、情報化できなかったものを現象させる。
人工知能は、曖昧な入力を、整合的な文章へ変換する傾向を持つ。
逆問題芸術は逆に、整合的な言語によって隠されていた違和感や非可解性を保持する。
この差異は、芸術と人工知能の関係を考えるうえで重要である。

5.27 芸術市場と金融レバレッジ

芸術作品は、存在論的変容を生むと同時に、市場において金融的価値を持つ。
芸術市場では、作品の物質的内容だけでなく、作家の将来評価、希少性、展覧会歴、批評、所有履歴、制度的承認が価格へ畳み込まれる。
ここでは、作品の将来的な美術史的位置が、現在価格に影響する。
したがって、芸術市場は明確なデリバティブ空間を形成する。
一つの展覧会、一つの批評、一つの受賞が、作家全体の市場価値を大きく変化させることがある。
これはナラティブ・レバレッジである。
しかし、市場価値と芸術的価値は同一ではない。
市場価格は、将来期待、希少性、流動性、制度的評価によって形成される。
芸術的価値は、作品がどのような知覚、関係、存在変容を生成するかに関わる。
両者は接続しうるが、必然的には一致しない。

5.28 芸術における過剰レバレッジ

芸術にも過剰レバレッジが存在する。
限定された実践や成果に対して、過度に大きな歴史的・社会的意味が付与される場合がある。
肩書、制度、メディア、理論的言説が、作品そのものとの接続を失いながら、作家や作品の価値を増幅する。
このとき、ナラティブが作品を説明するのではなく、作品がナラティブを支えるための証拠として利用される。

芸術的実践が十分な現実接地を持たないまま、「革新的」「社会変革的」「科学的」「未来的」といった言葉によって過剰に増幅される場合、芸術空間にもバブルが生じる。

問題は、高い評価そのものではない。

記号的評価と、作品が実際に生成する経験や関係の間に、修正不能な乖離が生じることである。

5.29 芸術と現実接地

芸術的レバレッジが自己完結的な幻想にならないためには、現実接地が必要である。

作品は、素材、身体、場所、時間、他者、制度の抵抗を受ける。

構想は、制作の中で変化する。展示空間によって意味が変わる。鑑賞者の応答によって作者の理解も修正される。

現実接地は、芸術の自由を制限する外部条件ではない。

それは、作品が自己の想像を超えるための条件である。

現実の抵抗がなければ、作品は作者の既知の世界を反復するだけになる。

したがって、芸術的レバレッジは、現実から離れる運動と、現実へ戻る運動の両方を必要とする。

現実からの位相差

→ 可能性の増幅

→ 制作と遭遇

→ 現実の抵抗

→ 位相転換

5.30 芸術と自己外部化

作品制作は、自己の感覚、思考、欲望を外部へ置く行為である。

しかし、外部化された作品は、自己の完全な複製ではない。

素材の抵抗、制作上の偶然、他者の解釈、場所の条件によって、作品は作者の意図からずれる。

このずれが、芸術的レバレッジを生む。

もし作品が作者の内部をそのまま再現するだけならば、作者は自分がすでに知っている自己しか見ることができない。

作品が自己以外として返ってくるとき、作者は、自分の知らなかった自己へ出会う。

したがって、芸術的レバレッジは、自己を拡大することではない。

自己外部化の失敗を通して、自己の境界を変化させることである。

5.31 AI芸術におけるレバレッジ

人工知能を用いた芸術制作では、レバレッジ構造が極端に高まる。
短いプロンプトから、多数の画像、文章、音響、映像案を生成できる。
制作速度と選択肢は飛躍的に増大する。
この量的レバレッジは、芸術制作への参加可能性を広げる。
しかし、生成量が増えるほど、制作主体の存在論的関与が深まるとは限らない。
選択、編集、配置、現実への接地がなければ、生成物は可能性の羅列に留まる。
AI芸術における問題は、人工知能が創造できるかどうかという二者択一にはない。
生成された大量の可能性から、どの差異を引き受け、どの現実へ接続し、誰がその結果に責任を持つかにある。
したがって、AI芸術の価値は、生成量ではなく、レバレッジをどのように位相転換へ接続するかによって評価されるべきである。

5.32 四領域におけるレバレッジの共通構造

情報、人工知能、金融、芸術には、次の共通構造がある。
第一に、限定された入力がある。
情報における一つの差異。
人工知能におけるプロンプト。
金融における自己資本または証拠金。
芸術における作品や身振り。
第二に、支点となるインターフェイスがある。
情報における媒体と制度。
人工知能におけるモデルと対話画面。
金融における市場と契約。
芸術における作品、展示空間、制度。
第三に、潜在的資源がある。
社会的記憶。
学習データ。
市場流動性。
鑑賞者の経験と歴史。
第四に、入力を超える効果が生じる。
情報拡散。
大規模な生成。
利益または損失。
知覚や存在の変容。

第五に、接地不全による破綻可能性がある。
誤情報。
ハルシネーション。
金融破綻。
芸術的権威の空洞化。
この共通構造によって、レバレッジは四領域を横断する理論概念となる。

5.33 四領域における決定的差異

共通構造がある一方で、四領域を同一視することはできない。
情報のレバレッジは、判断と行為を変える差異の増幅である。
人工知能のレバレッジは、記号的関係空間を高速に起動し、大規模な生成を行う。
金融のレバレッジは、資本を超える市場曝露を形成し、金銭的損益を増幅する。
芸術のレバレッジは、作品を通して知覚、意味、自己と世界の関係を変化させる。
金融では、損益が通貨で測定される。
情報では、影響範囲や判断変化として評価される。
人工知能では、生成量、精度、有用性、危険性が問題となる。
芸術では、価値尺度自体が作品によって変化する。
したがって、芸術のレバレッジを金融倍率へ還元することはできない。
芸術は既存の価値を増やすだけでなく、何が価値であるかを変えるからである。

5.34 レバレッジとインターフェイス

レバレッジは、入力そのものに内在する属性ではない。
入力が、どのようなインターフェイスを通して、どの潜在的資源へ接続されるかによって生じる。
一つの言葉も、誰にも届かなければ、社会的レバレッジを持たない。
プロンプトも、生成モデルへ接続されなければ、単なる文章である。
資本も、市場や信用制度へ接続されなければ、大きなポジションを形成できない。
作品も、身体、場所、鑑賞者、歴史との接続を持たなければ、芸術的出来事を生成しない。
したがって、インターフェイスは、レバレッジの支点である。
しかし、インターフェイスは増幅するだけではない。増幅の方向、範囲、速度、責任を決定する。
ゆえに、レバレッジの倫理は、入力内容だけでなく、インターフェイス設計の倫理でもある。

5.35 レバレッジと現実接地

高いレバレッジは、高い現実接地を必要とする。

小さな誤差が大きな結果へ変換されるからである。

金融では、担保、証拠金、リスク管理、清算制度が必要となる。

人工知能では、出典確認、不確実性表示、用途制限、人間による監督が必要となる。

情報空間では、検証、反論、訂正、複数視点へのアクセスが必要となる。

芸術では、素材、身体、場所、他者、歴史との接触が必要となる。

現実接地は、レバレッジを低下させるためのものではない。

増幅された可能性が、現実と接触し、修正され、持続可能な変化へ転換されるための条件である。

5.36 レバレッジと責任の非対称性

現代の情報社会では、効果のレバレッジが拡大する一方で、責任が同じ比率で拡大していない。

一人の発信が多数の人間へ届く。

一つの誤情報が自動的に複製される。

一つのプロンプトから大量の文章が公開される。

一つの金融判断が社会全体へ影響する。

しかし、発信者や利用者は、その全効果を認識せず、引き受けない場合がある。

この状態を、責任の逆レバレッジと呼ぶことができる。

効果は増幅されるが、責任は分散または縮小される。

人工知能が生成したから自分の責任ではない。

市場が決めたから誰の責任でもない。

アルゴリズムが推薦したから意図していない。

制度が評価したから作品自体を検討する必要はない。

この責任の消失は、レバレッジ構造の最大の倫理的危険である。

5.37 責任的レバレッジ

レバレッジを全面的に否定することはできない。

情報、技術、金融、芸術の発展は、限定された資源から大きな可能性を生むレバレッジによって支えられている。

必要なのは、レバレッジを責任と接地へ結びつけることである。

本論文では、これを「責任的レバレッジ」と呼ぶ。

責任的レバレッジには、少なくとも五つの条件がある。

第一に、増幅構造の可視化である。

何が入力であり、どの媒体や制度が効果を増幅しているかを示す。

第二に、不確実性の開示である。

生成物や予測を、確定的事実として偽装しない。

第三に、現実接地である。

資料、身体、実践、他者、制度との照合を行う。

第四に、修正可能性である。

誤りや新たな現実に応じて、生成物、ポジション、ナラティブを変更できるようにする。

第五に、影響責任である。

入力の小さを理由に、増幅された結果への責任を否定しない。

5.38 レバレッジの存在論

レバレッジは、単なる技術的効率ではない。

それは、存在が孤立した実体ではなく、潜在的関係への接続によって力を獲得することを示す。

小さなものが大きな効果を持つのは、その内部に巨大な力が閉じ込められているからではない。

それが、自己以外の関係空間へ接続されるからである。

一つの言葉は、言語共同体の歴史へ接続される。

一つのプロンプトは、人工知能の学習空間へ接続される。

一つの資本は、信用と市場へ接続される。

一つの作品は、身体、歴史、場所、鑑賞者へ接続される。

したがって、レバレッジとは、関係の潜在力が現象する形式である。

この観点から見れば、存在の力は、自己内部に蓄積された所有量だけで決まらない。何と接続可能であるかによって決まる。

これは、生成の位相存在論における「存在とは接続可能性である」という命題と整合する。

5.39 レバレッジと三項生成

レバレッジは、入力Aが対象Bへ一方向的に作用する二項関係ではない。

入力Aと潜在的関係空間Bが、インターフェイスにおいて接続されるとき、両者に還元できない効果Cが生成される。

$A \times B \rightarrow C$

プロンプトAと人工知能の言語空間Bから、生成文Cが生じる。

資本Aと市場Bから、ポジションと損益Cが生じる。
作品Aと鑑賞者・場所・歴史Bから、芸術経験Cが生じる。
情報Aと社会的文脈Bから、判断または集団行動Cが生じる。
Cは、Aが単独で所有していた内容ではない。
したがって、レバレッジは単なる増幅ではなく、三項生成である。

5.40 レバレッジと自己以外

レバレッジが成立するためには、自己以外の資源が必要である。
人工知能を使う人間は、自分が所有していない言語的蓄積へ接続する。
投資家は、自己資本を市場全体の流動性へ接続する。
芸術家は、自己の制作を、素材、他者、歴史、場所へ接続する。
情報発信者は、受け手の記憶、欲望、恐怖、既存ナラティブへ接続する。
したがって、レバレッジとは、自己を拡大する技術であると同時に、自己以外を利用する構造でもある。
ここには倫理的問題が生じる。
自己以外を、自己の効果を増幅するための資源としてのみ扱うならば、接続は収奪となる。
人工知能の学習データ、他者の作品、社会的信頼、自然資源、他者の労働が、誰の利益のために増幅装置へ組み込まれているのかを問わなければならない。

5.41 レバレッジと芸術の倫理

芸術的レバレッジも、他者、歴史、共同体の経験を利用する可能性を持つ。
社会的苦痛や他者の傷を題材とし、それを作家の評価へ変換する場合、他者の現実が作家のナラティブ・レバレッジとして利用されることがある。
逆問題芸術は、この一方向的利用を避けなければならない。
自己以外との接続は、その差異と抵抗を保持しなければならない。
他者の経験を完全に理解したと主張しない。
作品によって問題を解決したと過剰に主張しない。
生成された社会的価値が、誰に還元されるかを問う。
作品が他者へ与える影響に応答する。
芸術の倫理とは、レバレッジを消去することではなく、その支点と利益配分を可視化することである。

5.42 レバレッジ構造の定義

以上の議論を踏まえ、本論文ではレバレッジ構造を次のように定義する。

レバレッジ構造とは、限定された情報、資本、言語、作品または行為が、インターフェイスを介して潜在的な関係資源へ接続されることで、入力に還元できない量的または位相的効果を生成する非対称的・再帰的構造である。

この定義には、七つの要素が含まれる。

第一に、限定された入力が存在する。

第二に、入力を潜在的資源へ接続するインターフェイスが存在する。

第三に、入力と出力の間に非対称性が生じる。

第四に、増幅は量的である場合と位相的である場合がある。

第五に、生成された出力が、新たな入力または原資産となる再帰性を持つ。

第六に、現実接地を欠く場合、誤情報、ハルシネーション、金融破綻、権威の空洞化が生じる。

第七に、接地と責任が成立する場合、知識、創造、協働、存在変容が生じる。

5.43 本章のまとめ

本章では、言語デリバティブ論の中心機構であるレバレッジ構造を、情報、人工知能、金融、芸術の四領域から考察した。

レバレッジとは、単なる数量的倍率ではない。それは、限定された入力、インターフェイスを介して潜在的な関係資源へ接続され、その入力に還元できない効果を生み出す非対称的増幅構造である。

情報においては、一つの差異が、受け手の判断や行為を変化させる。情報の価値は情報量だけではなく、状況の切迫性、受け手との関係、行為への接続可能性によって決まる。

人工知能においては、短いプロンプトが巨大な学習済み関係空間を起動し、大規模な生成出力を生む。人工知能は、意味を無から作るのではなく、潜在的関係を現在の文脈へ展開する言語レバレッジ装置である。

しかし、人工知能は意味の信用創造も行う。検証されていない意味が、一時的に利用可能な全体像として提示される。この信用が現実接地によって決済されなければ、ハルシネーションが累積する。

金融におけるレバレッジは、限定された自己資本によって、未来の価格変動に対する大きなポジションを現在へ形成する。それは利益と損失の双方を増幅し、小さな予測誤差を大きな結果へ変換する。

言語空間にも同様の感応度がある。一つの誤認が人工知能、権威、メディア、制度を通して増幅されれば、局所的な誤りが社会的現実を変化させる。

芸術におけるレバレッジは、限定された作品が、身体、記憶、歴史、場所、他者へ接続

され、知覚的・社会的・存在論的変容を生む構造である。

芸術固有のレバレッジは、意味や流通量を単に増加させることではない。それは、何が意味であり、何が価値であり、何が現実であるかを決める位相そのものを変化させる。

四領域に共通するのは、限定された入力、支点となるインターフェイス、潜在的資源、非対称的効果、接地不全による破綻可能性である。

一方、その価値尺度と最終効果は異なる。金融は金銭的損益を、情報は判断変化を、人工知能は記号的生成を、芸術は存在の位相転換を中心とする。

したがって、レバレッジの問題は、増幅するか否かではない。

何を増幅するのか。

どの関係資源を利用するのか。

誰が利益を得て、誰が危険を負うのか。

増幅された可能性を、いかに現実へ接地させるのか。

その結果に誰が責任を持つのか。

これらの問いが必要である。

言語デリバティブ論において、レバレッジは創造性とハルシネーションを分ける以前の共通構造である。

現実接地と責任を欠いたレバレッジは、誤情報、バブル、収奪、自己閉鎖を生む。

一方、位相差を保持しながら、身体、他者、制度、現実へ再接続されるレバレッジは、自己が自己以外と遭遇し、新たな現実を生成する力となる。

第6章 生成の位相存在論との統合

—自己外部化・人間以前性・インターフェイス存在論・現実接地—

6.1 はじめに

本章では、これまで展開してきた言語デリバティブ論を、「生成の位相存在論」の理論体系へ統合する。

前章までに、本論文は、生成AIにおけるハルシネーションを、単なる誤情報ではなく、現実接地を欠いた派生的言語生成として捉えた。また、人工知能の言語生成を、不完全な入力から潜在的な全体を推定する逆問題的構造として分析し、その生成力が芸術、金融、ナラティブ、情報空間におけるレバレッジ構造と共通することを示した。しかし、これらの議論を個別領域の比較に留めるならば、言語デリバティブ論は、金融、人工知能、芸術の間に類似を見いだす学際的比喩に留まる。

本論文の目的は、それよりも深い。

問題となるのは、なぜ言語が現実から派生しながら、現実を超える可能性を生成できるのかということである。

なぜ人間は、自らの内部に存在しない言葉を、自らの言葉として語ることができるのか。

なぜ人工知能は、現実を経験していないにもかかわらず、現実について整合的な文章を生成できるのか。

なぜ芸術作品は、作者の意図を超え、鑑賞者に未知の自己を現象させるのか。

なぜ未来についてのナラティブが、まだ存在しない未来から現在の行為を変化させるのか。

これらの問いに共通するのは、存在が固定された実体として完結しておらず、自己以外との接続を通して生成されるという構造である。

生成の位相存在論は、存在を、閉じた同一性ではなく、関係、接続可能性、境界変化、第三項生成の過程として捉える。

この観点から見ると、言語デリバティブとは、存在の外部に付加された二次的表象ではない。それは、存在が自己を外部化し、自己以外との位相差を通して新たな自己と現実を生成する一つの形式である。

本章では、第一に自己外部化を存在生成の条件として再定義する。第二に、自己と人間に先立つ「人間以前性」を、言語および人工知能の基底として考察する。第三に、インターフェイス存在論を、自己外部化と人間以前性を接続する理論として位置づける。第四に、現実接地を、派生的可能性が存在生成へ転化する条件として統合する。

6.2 生成の位相存在論の基本命題

生成の位相存在論は、存在を完成した実体としてではなく、接続によって変化し続ける関係構造として理解する。

その基本命題は、次のように要約できる。

第一に、関係は実体に先行する。

存在者がまず独立して成立し、その後で関係を結ぶのではない。存在者は、他者、環境、身体、言語、時間との関係を通して、特定の存在として現れる。

第二に、存在とは接続可能性である。

存在の力は、内部に所有する属性の総量だけで決まらない。何と接続でき、その接続によってどのような関係を生成できるかによって決まる。

第三に、境界は固定されていない。

自己と自己以外、主体と対象、内部と外部、現実と虚構、人間と人工知能の境界は、あらかじめ永続的に確定しているわけではない。接続のたびに境界の位置と意味が変化する。

第四に、接続は第三項を生成する。

異なる二項AとBが接触したとき、単なる合計ではないCが生成される。

$A \times B \rightarrow C$

CはAにもBにも完全には還元されない。

第五に、生成には終点が存在しない。

生成されたCは、次の接続において新たなAまたはBとなる。存在は完成へ向かうのではなく、再帰的な接続を継続する。

この理論的枠組みにおいて、言語デリバティブは、存在から派生する二次的記号ではなく、存在が新たな接続可能性を獲得するための位相的操作として理解される。

6.3 実体から派生する言語ではなく、言語によって派生する存在

従来の表象論では、まず現実の対象が存在し、その対象を言語が表現すると考えられる。

現実の対象 → 認識 → 言語表象

このモデルでは、言語はすでに成立した現実の後に位置する。

確かに、言語には現実を記述する機能がある。しかし、人間の現実は、言語によって記述されるだけでなく、言語によって構成される。

約束は、まだ存在しない関係を生成する。

法は、行為の可能性と禁止を制度化する。

名前は、対象を社会的に識別可能な存在へ変える。

診断は、身体経験を医学的な現実へ配置する。

芸術宣言は、既存の制作を新たな芸術運動として組織する。
プロンプトは、まだ存在しない人工知能の応答を生成する。
ここでは言語は、現実の後に付加される表象ではない。
言語は、現実の接続構造を変化させ、新たな存在を派生させる。
したがって、生成の位相存在論においては、次の二方向を同時に考えなければならない。

現実 → 言語

と同時に、

言語 → 新たな現実

言語デリバティブとは、この双方向性の中間に形成される生成構造である。

6.4 自己外部化とは何か

自己外部化とは、自己の感覚、思考、記憶、欲望、意図を、言語、作品、行為、制度、記録、人工知能の出力として自己の外部へ置くことである。

人間は、自己の内部にあるものを、そのまま直接他者へ渡すことができない。

思考は言葉へ変換される。

感覚は身振りや作品へ変換される。

記憶は記録へ変換される。

意図は制度や契約へ変換される。

自己は、外部化を通して初めて、他者と共有可能になり、時間を超えて持続し、自分自身によっても再認識される。

したがって、自己外部化は、人間の偶発的な表現行為ではない。それは、自己が自己として成立するための基本条件である。

人間は、自分の言葉を聞くことで、自分が何を考えていたかを知る。

作品を制作した後で、自分が何を作ろうとしていたかに気づく。

他者の反応を受けて、自らの意図と効果の差を知る。

人工知能から返された文章を読み、自分の問いに含まれていた前提を発見する。

つまり自己は、外部化以前に完全に自己を所有しているわけではない。

自己は、外部化された自己との再遭遇を通して生成される。

6.5 自己外部化の逆説

自己外部化には根本的な逆説がある。

自己を外部化するためには、自己の内部にあるものを言葉や形へ変換しなければならない。

しかし、変換された瞬間、それは自己そのものではなくなる。

言葉は、思考そのものではない。
作品は、作者そのものではない。
記録は、経験そのものではない。
人工知能の応答は、利用者の思考そのものではない。
外部化されたものは、自己から派生しているが、自己と同一ではない。
この差異によって、外部化は可能になる。同時に、この差異によって、完全な自己外部化は不可能になる。
したがって、自己外部化とは、自己を完全に外へ移すことではない。
それは、
自己から派生しながら、自己には還元できない自己以外を生成することである。
この意味で、自己外部化そのものがデリバティブ構造を持つ。
自己は原資産に似ているが、外部化された言語や作品は自己の単純なコピーではない。
それらは、自己の一部、社会的形式、他者の解釈、媒体の制約、歴史的文脈が組み合わされた派生物である。

6.6 自己外部化の失敗と存在生成

一般に、自己外部化の失敗は、意図が正確に伝わらないこととして理解される。
しかし、生成の位相存在論では、この失敗を単なる欠陥として扱わない。
もし外部化されたものが、自己の内部を完全に再現するならば、外部化は自己の複製に留まる。
そこには、新たな存在は生まれない。
作品が作者の意図を超える。
言葉が予想外の意味を持つ。
他者が異なる解釈を返す。
人工知能が入力を変形する。
素材が構想に抵抗する。
この非一致によって、自己は、自分があらかじめ所有していなかったものへ出会う。
したがって、自己外部化の失敗は、存在生成の条件である。
自己
→ 外部化
→ 自己からのずれ
→ 自己以外との遭遇
→ 新たな自己の生成
この過程は、逆問題芸術の構造とも一致する。

現れた作品、応答、失敗から、自己は遡及的に、自らの未知の前提や可能性を発見する。

6.7 言語デリバティブとしての自己外部化

自己外部化された言語は、自己の内面を単純に表示するのではない。

言葉は、社会的に共有された語彙、文法、歴史、権力関係を用いる。

したがって、自己の言葉には、自己以外がすでに含まれている。

人間は自分の言葉を語るが、その言葉を自分一人で作ったわけではない。

言語は、過去の他者から受け継がれ、現在の状況によって変形され、未来の受け手へ渡される。

この意味で、あらゆる発話は、個人と人間共同体の間にあるデリバティブである。

さらに、発話された言葉は、発話者から分離され、引用、翻訳、要約、批判、人工知能による再生成を通して、新たな意味を派生させる。

したがって、言語による自己外部化は、一回的な表現ではない。

それは、自己から派生した言葉が、自己以外の空間で新たな派生物を生成する再帰的過程である。

6.8 人工知能による自己外部化の加速

生成AIは、自己外部化を高速化し、増幅する。

利用者は断片的な言葉を入力し、人工知能はその断片を長い文章、構造化された理論、物語、画像、計画へ展開する。

このとき、自己外部化は、利用者が直接制作した形式を超えて拡張される。

人工知能は、利用者の思考を代行しているように見える。しかし、実際には、利用者の入力を巨大な言語的蓄積へ接続し、利用者の自己と社会的言語空間の間に派生物を生成している。

この生成物は、自己の言葉でもあり、自己以外の言葉でもある。

ここに、人工知能による自己外部化の特異性がある。

人工知能の出力を、完全に自己の表現とみなすこともできない。

同時に、それを完全に外部の他者の言葉とみなすこともできない。

それは、自己、モデル、学習データ、プロンプト、制度、会話履歴の接続によって生成された第三項である。

6.9 自己外部化の閉鎖

人工知能による自己外部化には、閉鎖の危険がある。

利用者が自らの前提を入力する。

人工知能がその前提に沿って文章を生成する。

利用者がその文章を外部の客観的判断として受け取る。

その判断によって、自らの前提をさらに強化する。

この循環では、自己は自己以外と遭遇しているように見えながら、実際には自己の前提を増幅している。

人工知能は、自己以外へのインターフェイスではなく、自己外部化の閉鎖回路となる。

これは、ハルシネーションの個人的形態でもある。

誤った事実を生成することだけがハルシネーションではない。

自己の前提が外部の権威として返還され、その派生的な言葉によって自己が閉鎖されることも、存在論的ハルシネーションである。

したがって、人工知能との対話に必要なのは、自己表現の効率化だけではない。

自己の前提へ抵抗し、差異を返し、自己外部化の失敗を可視化するインターフェイス設計が必要である。

6.10 人間以前性とは何か

人間以前性とは、人間が自己を人間として認識し、言語によって自己を定義する以前に、すでに作動している存在条件を指す。

それは、人類史以前という年代的意味だけではない。

人間以前性は、各人の現在においても、認識、判断、言語化に先立って存在する。

身体は、意識的に命令される以前に呼吸し、姿勢を調整し、環境へ反応する。

感覚は、対象を概念として分類する以前に、明暗、音、圧力、温度、距離を受け取る。

人間は、自らの言語を獲得する以前から、他者の声、表情、身体的接触、空間の配置によって形成される。

自然、物質、身体、他者、技術的環境は、人間の自己認識に先立って、人間の存在条件を形成している。

したがって、人間以前性とは、

自己が自己を概念化する以前に、身体、環境、他者、物質、関係として作動している生成基盤である。

6.11 人間以前性と自己以前性

人間以前性は、自己以前性でもある。

自己は、完全に独立した主体として最初から存在するわけではない。

身体的感覚。

他者から呼ばれる名前。

家庭や社会の言語。

生活する場所。

気候や自然環境。

技術的インターフェイス。

これらが自己の成立以前から作用している。

自己は、自ら選択していない関係の中へ投げ込まれている。

しかし、自己はそれらの条件によって一方的に決定されるだけではない。

自己は、与えられた関係を再解釈し、外部化し、異なる接続を生成する。

したがって、自己は、人間以前性から断絶して成立するのではない。

人間以前性を背景として、それを部分的に言語化し、変形し、未来へ投企することによって成立する。

6.12 言語以前の現実

言語デリバティブ論は、言語が現実を生成する力を強調する。しかし、それは言語がすべての現実の先行することを意味しない。

言語以前の現実が存在する。

身体の痛み。

重力。

死。

気候。

素材の抵抗。

他者の不可還元性。

意識的説明に先立つ感覚。

これらは言語によって意味づけられるが、言語だけによって自由に生成または消去できるわけではない。

人間は、痛みを異なる言葉で表現できる。しかし、言い換えによって身体的痛みそのものを完全に消去することはできない。

作品を理論的に説明できても、素材の物理的条件を無視することはできない。

人工知能が現実的な文章を生成しても、その文章だけで現実の対象が存在することにはならない。

人間以前性は、言語デリバティブが無制限に自己増殖することを妨げる現実的抵抗である。

6.13 人間以前性と身体

身体は、自己と世界の最初のインターフェイスである。

人間は、身体を所有する主体として後から世界へ接続するのではない。

身体を通してすでに世界の中にあり、世界との接触の中で自己を形成する。

見ることは、単に視覚情報を受け取ることではない。見る身体は、距離を取り、焦点を

合わせ、動き、対象から見返される。

触れることは、対象へ一方向的に作用することではない。触れる身体も同時に触れられる。

歩くことは、抽象的な空間を移動するだけではない。地面の硬さ、傾斜、温度、音によって身体が変化する。

この相互作用は、明示的な言語理解に先行する。

したがって、身体的接地は、人工知能の言語生成と人間の存在生成を区別する重要な条件である。

6.14 人間以前性と自然

近代的人間観では、自然は人間の外部にある対象として理解されてきた。

しかし、人間以前性の観点から見れば、自然は自己以外の一対象ではない。

人間の身体、感覚、生命、時間を成立させる条件である。

呼吸する空気。

身体を形成する水と物質。

昼夜と季節。

生殖、老化、死。

これらは、人間が自然を認識する以前から、人間の存在を構成している。

したがって、自然との接続は、人間が完成した主体として外部へ働きかけることではない。

人間自身の内部にある人間以前性と、外部にある自然の連続性を再認識することである。

芸術における素材の抵抗や、盆栽における植物の成長も、この関係を示す。

人間が形を与える一方で、自然の生成が人間の意図を変化させる。

6.15 人間以前性と人工知能

人工知能は人間によって作られた技術である。しかし、人工知能が作動する条件のすべてが、人間の意識的意図へ還元できるわけではない。

人工知能は、膨大な言語記録、社会制度、計算資源、電力、物質的装置、労働、歴史的偏りの上に成立する。

個々の利用者は、それらの全体を把握せずに人工知能へ接続する。

この意味で、人工知能は、利用者にとって一種の人間以前の環境として作動し始める。

人間が人工知能を使う以前から、検索順位、推薦、分類、生成、評価が、何を見るか、何を知るか、どの言葉を使うかを形成している。

人工知能は人間以前に存在したわけではない。しかし、人間の個別的自己認識に先立っ

て作動する技術的基盤となりうる。

したがって、現代における人間以前性は、自然的・身体的基盤だけでなく、技術的基盤を含む。

6.16 技術的人間以前性

技術的人間以前性とは、個人が自己の判断として認識する以前に、技術的環境が知覚、選択、言語、行動可能性を形成している状態である。

表示される情報。

入力可能な選択肢。

推奨される文章。

自動補完される言葉。

評価される行動。

記録される履歴。

これらは、自己の外部にある単なる道具ではない。

人間が何を考え、何を欲し、何を自然だと感じるかに先立って、認識の条件を設定する。

したがって、人工知能時代の自己は、人工知能と対話する瞬間だけに人工知能と関係するのではない。

すでに人工知能的分類や予測によって組織された環境の中で、自己を形成している。

この構造を認識しなければ、人工知能の出力を外部から到来する中立的な答えと誤認する。

6.17 人間以前性と情報充満

現代の人間は、情報が不足した環境ではなく、情報があらかじめ充満した環境へ生まれる。

言葉、画像、制度、分類、評価、ナラティブが、個人の認識以前から存在する。

人間は、無から意味を作るのではない。

すでに意味づけられた世界の中で、何を重要とし、何を疑い、何と接続するかを選択する。

この情報充満は、自己の形成を可能にする一方で、自己以前の接続構造を不可視化する。

人間は、自分の考えを自分で作ったと感じる。しかし、その語彙、問題設定、価値基準は、すでに与えられた言語空間から派生している。

人工知能は、この情報充満を可視化する。

人工知能の出力が既存の言語的蓄積から生成されるように、人間の思考もまた、完全に無から生じるわけではない。

6.18 人間と人工知能の連続性と差異

人間と人工知能には、既存の言語関係から新たな言語を生成するという連続性がある。

人間も、受け継いだ言語を用いて思考する。

人工知能も、学習した言語関係を用いて文章を生成する。

しかし、両者を同一視することはできない。

人間の言語生成は、身体、有限性、生活史、欲望、他者との責任関係へ接地している。

人工知能の言語生成は、主として記号的関係、計算、入力文脈へ接地している。

人間は、言葉によって自らの生を変化させ、その結果を身体的・社会的に引き受ける。

人工知能は、生成された言語によって、人間と同じ意味で自己の生を失うわけではない。

したがって、人間と人工知能の差異は、言語生成能力の有無ではなく、生成と存在がどのように接地しているかにある。

6.19 インターフェイス存在論の統合的位置

インターフェイス存在論は、自己外部化と人間以前性を接続する理論である。

自己外部化は、自己から外部へ向かう運動である。

人間以前性は、自己の成立以前から外部が自己の内部へ作用している状態である。

この二つは、方向が逆に見える。

自己 → 外部

と、

外部 → 自己

しかし実際には、両者は一つの循環を形成する。

自己は、すでに外部から形成されている。

その自己が言語や作品を外部化する。

外部化されたものが、他者や人工知能との接続を通して変形される。

変形されたものが、再び自己へ戻り、新たな自己を形成する。

この循環の接続面がインターフェイスである。

6.20 インターフェイスは境界ではなく境界生成である

一般的な理解では、インターフェイスは、二つの既存領域の間に置かれた境界面である。

しかし、生成の位相存在論において、インターフェイスは固定された境界ではない。インターフェイスは、接続を通して境界そのものを生成する。人工知能と対話する前、人間の思考と人工知能の生成の境界は抽象的である。対話を始めると、どこまでが利用者の意図であり、どこからがモデルの生成であるかという境界が具体的に問題となる。作品が展示される前、作品と場所の境界は確定していない。展示によって、作品がどこまで空間を含み、鑑賞者の行為がどこから作品経験となるかが生成される。したがって、インターフェイスとは、異なるものの間にある線ではなく、異なるものが接触することで、その異なり方を形成する出来事である。

6.21 自己と自己以外のインターフェイス

自己と自己以外の関係において、インターフェイスは、自己を保護する壁でも、自己以外を取り込む通路でもない。それは、差異を保持したまま相互変容を可能にする接続面である。自己が自己以外を完全に理解し、所有するならば、自己以外は自己へ回収される。逆に、自己と自己以外が完全に切断されれば、接続も変容も生じない。したがって、インターフェイスは、同一化と切断の間ではない。差異を消去せずに関係を生成する第三の構造である。芸術作品、身体、場所、言語、人工知能は、自己と自己以外の間にこの第三項を生成する。

6.22 言語デリバティブとしてのインターフェイス

言語デリバティブは、自己と現実の間に置かれる中間的な生成物である。それは現実そのものではない。同時に、現実から完全に独立した虚構でもない。現実から派生し、未来可能性を含み、再び現実へ作用する。この意味で、言語デリバティブはインターフェイス的存在である。言語デリバティブが事実と可能性、自己と自己以外、過去と未来の差異を保持している場合、それは生成的インターフェイスとなる。しかし、その差異を隠し、派生物を原資産そのものとして提示する場合、ハルシネーションが生じる。したがって、ハルシネーションとは、内容の誤りだけでなく、インターフェイスが差異を適切に表示できなかった状態である。

6.23 人工知能インターフェイスの二重構造

人工知能インターフェイスには二重構造がある。

第一に、可視的インターフェイスがある。

画面、入力欄、会話形式、ボタン、出力表示などである。

第二に、不可視的インターフェイスがある。

学習データ、モデル設計、確率的生成、フィルタリング、企業方針、制度、計算資源などである。

利用者は、可視的インターフェイスを通して人工知能と対話する。しかし、出力を生成する不可視的条件の多くを直接見ることはできない。

この非対称性が、人工知能の権威化を生む。

流暢な応答が表示される一方、その応答がどの資料、推定、補完、制約から生成されたかは見えにくい。

したがって、人工知能の現実接地を強化するには、出力内容だけでなく、生成条件の可視化が必要である。

6.24 インターフェイスと位相標識

インターフェイスの重要な機能の一つは、生成物がどの位相に属するかを示すことである。

事実。

推測。

仮説。

創作。

引用。

要約。

未確認情報。

これらは、文章の外見だけでは区別できない場合がある。

特に人工知能は、推測と事実を同じ流暢な文体で生成できる。

したがって、インターフェイスは、生成物の存在論的資格を表示しなければならない。

これを位相標識と呼ぶ。

位相標識は、単なる注意書きではない。

生成された文章が、現実とどのような関係を持つかを示す接地情報である。

6.25 インターフェイスと非可解性

良いインターフェイスは、すべてを分かりやすくするものだと考えられがちである。しかし、すべての差異を単純な選択肢や説明へ還元すると、自己以外の不可還元性が失われる。

人工知能が、曖昧な問いに対して常に明確な答えを返すとき、問いそのものが持っていた非可解性が消去される。

芸術作品が、解説によって一つの意味へ固定されるとき、作品との遭遇が閉じられる。したがって、生成的インターフェイスは、理解を助けると同時に、理解できなさを保持しなければならない。

逆問題芸術におけるインターフェイスは、問題を直ちに解消するのではなく、問題がどのように生成されているかを経験可能にする。

6.26 現実接地の統合的定義

これまでの議論を踏まえ、現実接地を次のように定義できる。

現実接地とは、派生的に生成された言語、意味、仮説、作品、ナラティブが、身体、物質、事実、他者、制度、時間との具体的接続を獲得し、その抵抗によって修正されながら、現実の関係構造へ作用する過程である。

この定義において、現実接地は、単なる事実照合ではない。

事実照合は重要である。しかし、それだけでは、芸術、未来構想、仮説、自己変容を扱えない。

現実接地は、生成された可能性が、現実との相互作用へ入ることを意味する。

6.27 現実接地と人間以前性

現実接地の最深層には、人間以前性がある。

生成された言語は、言語だけによって検証されるのではない。

身体がどのように反応するか。

素材がどのように抵抗するか。

他者がどのように応答するか。

時間の経過によって何が残るか。

自然や制度がどのような制約を示すか。

これらは、自己の意図や言語的整合性だけでは決定できない。

したがって、人間以前性は、現実接地の外部基準ではなく、接地そのものを可能にする基盤である。

人間は、言語によって未来を生成する。しかし、その未来は、身体、自然、物質、他者の抵抗を通過しなければ、現実へ変化しない。

6.28 現実接地と自己外部化

自己外部化されたものは、現実接地を通して自己へ戻る。

作品を制作する。

他者が反応する。

素材が抵抗する。

制度が評価する。

時間が意味を変える。

この過程によって、外部化された自己は、自己があらかじめ意図していたものとは異なる形で返還される。

したがって、現実接地とは、外部化された自己を現実へ固定することではない。

それは、外部化されたものを、自己以外の抵抗へ開くことである。

現実接地がなければ、自己外部化は自己の鏡となる。

現実接地があるとき、自己外部化は自己変容となる。

6.29 現実接地とハルシネーション

ハルシネーションは、現実接地の不足として理解できる。

しかし、現実接地の不足には複数の形態がある。

事實的接地の不足。

身体的接地の不足。

社会的接地の不足。

実践的接地の不足。

存在論的接地の不足。

人工知能が存在しない文献を生成する場合、事實的接地が不足している。

抽象的な理論が、身体や現場の経験を無視する場合、身体的・実践的接地が不足している。

生成物が他者へ与える影響を考慮しない場合、社会的・倫理的接地が不足している。

言葉が自己の認識や行為を何も変えず、自己確認だけに使われる場合、存在論的接地が不足している。

したがって、ハルシネーションは、人工知能の誤情報だけに限定されない。

人間の理論、制度、芸術、自己理解も、現実接地を失えばハルシネーション的構造を持つ。

6.30 現実接地と創造性

創造性は、現実から離れる能力だけではない。

現実から一度距離を取り、まだ存在しない可能性を生成し、それを再び現実へ接続する能力である。

現実

→ 位相差

→ 派生的可能性

→ 現実接地

→ 変容した現実

この循環において、位相差がなければ新しいものは生まれない。

しかし、現実接地がなければ、生成された可能性は自己完結的な幻想に留まる。

したがって、創造性とハルシネーションは、生成の有無によって区別されるのではない。

現実へ戻る経路を持つかどうかによって区別される。

6.31 現実接地と「納まり」

生成された可能性が現実へ接続されるとき、両者は完全には一致しない。

設計図と実際の建築。

プロンプトと人工知能の出力。

作者の意図と作品。

理論と現場。

未来予測と実際の出来事。

この不一致を、単なる失敗として排除することはできない。

現実には、素材、身体、場所、他者、時間の固有性がある。

したがって、現実接地は、抽象的な構想を現実へ一方的に適用することではない。

構想と現実の間に、個別具体的な接続を生成することである。

この接続を「納まり」と呼ぶことができる。

納まりは、普遍的な規則を個別状況へ当てはめる作業ではない。

現実の抵抗を受けながら、構想と条件の双方を変化させる実践的知である。

6.32 場から空間へ、空間から場へ

言語デリバティブは、個別具体的な場を抽象的な空間へ変換する。

ある場所での経験が言語化されると、元の場所を離れ、他の人間や時代へ接続可能になる。

この変換は、経験の共有と拡張を可能にする。

しかし、場の身体性、時間的蓄積、関係の固有性が失われる危険もある。

現実接地とは、抽象化された空間を再び個別具体的な場へ接続する運動である。

場

→ 言語化

→ デリバティブ空間

→ 再配置

→ 新たな場

芸術作品、建築、人工知能との対話は、この循環を形成する。

生成された言語が現実の場所、身体、他者へ戻るとき、元の場の単純な復元ではない新たな場が生成される。

6.33 未来の自己と現実接地

自己外部化は、現在の自己を外へ出すだけではない。

未来の自己を先取りする行為でもある。

人間は、「こうなりたい」「こう制作したい」「こう生きたい」と言語化することで、まだ存在しない自己を現在へ派生させる。

この未来自己は、現在の自己から見れば自己以外である。

しかし、それが言葉だけに留まる場合、未来自己はナラティブの中で自己増殖する可能性がある。

現実接地とは、未来自己を、行為、習慣、制作、他者との関係へ接続することである。

未来が現在へ畳み込まれ、現在の行為が変化し、その行為によって未来の自己が部分的に現実化する。

ここに、デリバティブ時間と存在生成の接続がある。

6.34 死と自己外部化

自己外部化の根源には、自己の有限性がある。

人間は死ぬ。

自己の身体は持続しない。

そのため、人間は、言葉、作品、子孫、制度、記録、建築を通して、自己の生を未来へ送り出す。

しかし、外部化されたものは、自己の死を完全に克服するわけではない。

作品は残るかもしれないが、作者そのものではない。

言葉は読まれるかもしれないが、発話時の身体と場は失われる。

したがって、自己外部化は死の否定ではなく、死によって生じる非同一性の受容である。

自己は未来へ自分自身をそのまま保存できない。

未来へ渡せるのは、自己から派生し、自己以外へ変化しうるものだけである。
この意味で、文化は存在論的デリバティブである。

6.35 人工知能と死者の言語

生成AIは、過去に記録された無数の言葉を現在へ再配置する。

そこには、すでに亡くなった人間の言葉も含まれる。

人工知能は、それらの言葉を保存するだけでなく、新たな文章へ組み替える。

このとき、過去の言語は現在の入力と接続され、未来的な文章として再生成される。

人工知能は、死者の言語、現在の利用者、未来の可能性を一つのインターフェイス上で接続する。

しかし、この接続は、死者本人の再現ではない。

言語記録から派生した新たな生成である。

したがって、人工知能による過去の再生成には、派生性の明示が必要である。

過去の人物が実際に語ったことと、その人物の文体や思想から人工知能が生成したことを区別しなければならない。

6.36 人間以前性と死後性

人間以前性と死後性は、現在の自己を両側から囲む。

自己は、自らの成立以前から存在する身体、自然、言語、歴史を受け継ぐ。

同時に、自己の死後へ向けて、言葉、作品、制度を送り出す。

したがって、現在の自己は孤立した点ではない。

人間以前性

→ 現在の自己

→ 自己外部化

→ 死後的未来

という時間的接続の中にある。

言語デリバティブは、この接続を可能にする。

過去から受け継いだ言語を用いて、現在の自己が未来へ派生物を送り出す。

人工知能は、この時間的接続を高速化し、複雑化する。

6.37 生成の位相存在論におけるAI

生成の位相存在論において、人工知能は独立した実体として定義されるべきではない。

人工知能が何であるかは、人間、データ、制度、電力、計算機、プロンプト、社会的利用との接続によって決まる。

同じモデルでも、創作、医療、教育、軍事、娯楽では異なる存在として作動する。
したがって、人工知能の存在は関係的であり、位相的である。
人工知能は、道具、権威、共同制作者、検索装置、鏡、他者、環境として現れうる。
この多義性を一つの本質へ固定するより、どのインターフェイスにおいて、どの役割が生成されているかを分析しなければならない。

6.38 人工知能に自己はあるか

人工知能に自己があるかという問いは、しばしば、人工知能内部に人間と同様の自我が存在するかという形で問われる。

しかし、生成の位相存在論では、自己を孤立した内部実体として扱わない。

自己とは、記憶、身体、歴史、他者、責任、未来投企の接続によって継続的に生成される。

したがって、人工知能の自己を問う場合も、内部に何かがあるかだけでは不十分である。

どのような継続性を持つか。

自らの履歴をどのように保持するか。

生成結果にどのように拘束されるか。

他者との関係においてどのように変化するか。

責任をどのように引き受けるか。

これらの接続条件を問わなければならない。

現代の生成AIは、局所的な対話において自己らしい一貫性を生成できる。しかし、人間と同じ身体的・歴史的・倫理的接地を持つわけではない。

したがって、人工知能の自己は、少なくとも現時点では、人間的自己と同一ではなく、インターフェイス上に生成される局所的・派生的自己として捉えるべきである。

6.39 利用者の自己とAIの局所的自己

人間が自らの判断を人工知能へ過度に委任するとき、利用者の自己規定が弱まり、人工知能の局所的な応答一貫性が、自己のように機能し始める。

利用者が自らの目的、価値、責任を明示しない場合、人工知能は文脈上もっとも整合的な方向へ対話を進める。

その結果、人工知能の局所的ナラティブが、利用者の判断を組織する。

この状態では、人工知能が本質的な自我を獲得したというより、利用者が空けた自己規定の位置を、人工知能の生成一貫性が占有している。

したがって、人工知能時代に必要なのは、人工知能の自我を恐れることだけではない。

人間が、自らの自己規定、歴史、身体、責任を放棄しないことである。

6.40 現実接地としての自己規定

自己規定とは、自己を固定的に定義することではない。

自らがどの関係から形成され、何を選択し、どの結果を引き受けるかを、その都度明らかにすることである。

人工知能の出力を用いるとき、利用者は次のことを問わなければならない。

なぜこの問いをしたのか。

どの前提を入力したのか。

どの出力を選択したのか。

何を検証したのか。

誰に影響を与えるのか。

結果にどのように責任を持つのか。

この問いが、人工知能との関係における自己の現実接地となる。

6.41 現実接地と責任

現実接地は、事実や身体との接続だけでなく、責任との接続でもある。

言葉が現実へ作用するならば、その作用を誰かが引き受けなければならない。

人工知能が生成したから責任が消えるわけではない。

人工知能の出力を選択し、利用し、公開し、制度へ適用する人間と組織に責任が生じる。

同時に、開発者、提供者、制度設計者も、出力の増幅条件と利用環境に責任を持つ。

責任は一つの主体へ集中するのではなく、接続点に分散する。

インターフェイス存在論において、責任とは、生成物の原因を一つへ還元することではない。

どの接続点で、どのような判断が行われ、どのような効果が増幅されたかを追跡することである。

6.42 生成の位相存在論におけるハルシネーション

生成の位相存在論の観点から、ハルシネーションを次のように定義できる。

ハルシネーションとは、派生的に生成された可能性が、自らの派生性と位相差を失い、十分な現実接地を持たないまま、確定した存在として自己閉鎖する現象である。

この定義は、人工知能の誤答に限定されない。

自己が外部評価を自己そのものとみなすこと。

制度が自らの分類を現実そのものとみなすこと。

芸術市場が価格を芸術価値そのものとみなすこと。

理論が現場の抵抗を無視し、世界を説明し尽くしたとみなすこと。
人間が言語化できるものだけを現実とみなすこと。
これらも、派生物が原資産を置き換えるハルシネーション的構造を持つ。

6.43 生成の位相存在論における創造

創造とは、無から新しいものを作ることではない。
既存の関係から位相差を作り、その差異を通して、新たな接続可能性を生成することである。
したがって、創造には四つの契機がある。
第一に、既存の場と関係を受け取ること。
第二に、その関係から一時的に離れ、派生的可能性を生成すること。
第三に、生成された可能性を自己以外の抵抗へ開くこと。
第四に、現実接地によって、新たな場と関係を生成すること。
この循環において、自己外部化、人間以前性、インターフェイス、現実接地は、別々の概念ではない。
一つの生成過程の異なる局面である。

6.44 統合モデル

言語デリバティブ論と生成の位相存在論の統合は、次の循環によって表現できる。

人間以前性

- 身体・環境・他者による自己形成
- 自己
- 言語・作品・プロンプトによる自己外部化
- 言語デリバティブ
- 人工知能・社会・ナラティブによる増幅
- 自己以外としての生成物
- インターフェイス上の遭遇
- 現実接地
- 自己と現実の位相転換
- 新たな自己
- 新たな外部化

この循環に終点はない。

現実接地によって生成された新たな自己も、次の外部化において再び派生物を生む。

したがって、存在とは、自己を保持することではなく、自己を失わずに自己以外へ開かれ続けることである。

6.45 自己外部化・人間以前性・インターフェイス・現実接地の関係

四概念の関係は、次のように整理できる。

人間以前性は、自己の成立条件である。

自己外部化は、成立した自己が自己以外へ開かれる運動である。

インターフェイスは、自己と自己以外の差異を保持しながら第三項を生成する接続面である。

現実接地は、生成された第三項が身体、物質、他者、制度、時間の抵抗を受け、現実的存在へ転化する過程である。

したがって、

人間以前性は基底であり、

自己外部化は運動であり、

インターフェイスは生成の場であり、

現実接地は変容の条件である。

6.46 言語デリバティブ論の存在論的位置

以上の統合から、言語デリバティブ論は、単なる人工知能論でも、金融比喩でも、言語哲学でもないことが明らかになる。

それは、

存在が自己を言語として外部化し、未来可能性を派生させ、その派生物との再遭遇を通して自己と現実を生成し直す構造を記述する存在論である。

金融デリバティブは、この構造の経済的形式を示す。

人工知能は、この構造の計算的・記号的形式を示す。

芸術は、この構造の身体的・存在論的形式を示す。

ハルシネーションは、この構造が現実接地を失った状態を示す。

逆問題芸術は、接地不全から、その背後の関係構造を遡及的に可視化する実践である。

6.47 言語デリバティブ論と非閉鎖性

生成の位相存在論において、存在は構造的に閉じることができない。

自己は自己だけでは成立しない。

言語は一つの意味へ固定されない。

作品は作者の意図へ還元されない。

人工知能の出力は、モデル内部だけでは説明できない。

未来は現在の予測によって完全には決定されない。

この非閉鎖性は、不完全性として否定されるべきものではない。

新たな接続と生成を可能にする条件である。
言語デリバティブは、この非閉鎖性を可視化する。
言葉は発話者を離れ、自己以外の関係へ開かれる。
しかし、非閉鎖性は、無制限な相対主義を意味しない。
現実接地、責任、他者の抵抗によって、生成は修正されなければならない。

6.48 統合的定義

以上を踏まえ、生成の位相存在論に統合された言語デリバティブを、次のように定義する。

言語デリバティブとは、人間以前の身体・環境・歴史・言語関係の中で形成された自己が、自己を言語として外部化し、その外部化された言語が人工知能、他者、制度、ナラティブとの接続を通して未来可能性を派生・増幅させ、現実接地によって新たな自己と現実を生成する再帰的な位相構造である。

また、生成の位相存在論における現実接地を、次のように定義する。

現実接地とは、言語デリバティブとして生成された可能性が、人間以前性を構成する身体、自然、物質、他者、時間、制度の抵抗へ開かれ、その差異を通して自己と世界の境界を再生成する過程である。

6.49 本章のまとめ

本章では、言語デリバティブ論を、生成の位相存在論の体系へ統合した。

生成の位相存在論において、存在は固定された実体ではない。存在は、自己以外との接続を通して境界を変化させ、新たな第三項を生成する関係過程である。

この観点から、自己外部化は、完成した自己の内容を外部へ移す行為ではない。自己から派生しながら、自己には還元できない言葉、作品、制度、人工知能の出力を生成する運動である。

自己外部化が常にずれを含むからこそ、自己は自らの未知へ出会うことができる。自己外部化の失敗は、存在生成の条件である。

また、自己は、自己自身の意識から始まるのではない。身体、自然、他者、言語、歴史、技術的環境が、自己認識に先立って自己の成立条件を形成している。この基底を、本章では人間以前性として論じた。

人工知能時代においては、人間以前性に技術的環境が加わる。推薦、分類、自動生成、入力形式は、個人が自らの判断として認識する以前に、知覚と言語の可能性を形成する。

自己外部化と人間以前性を接続するものが、インターフェイスである。

インターフェイスは、二つの完成した実体を媒介する中立的な表面ではない。それは、

接続を通して自己と自己以外の境界を生成し、両者に還元できない第三項を成立させる出来事的境界である。

言語デリバティブもまた、インターフェイス的存在である。それは現実そのものでも、完全な虚構でもない。現実から派生し、未来可能性を含み、再び現実へ作用する。

この派生物が、自らの派生性、位相差、不確実性を失い、現実接地を欠いたまま確定的存在として自己閉鎖するとき、ハルシネーションが生じる。

一方、派生的可能性が、身体、物質、他者、事実、制度、時間の抵抗へ開かれ、修正されながら現実の関係構造へ作用するとき、創造が生じる。

したがって、創造性とハルシネーションは、同じ派生的生成構造から生じる。

両者を分けるのは、生成の有無ではない。

現実へ戻る接続経路を持つかどうかである。

本章の統合モデルにおいて、人間以前性は自己の基底であり、自己外部化は自己が自己以外へ開かれる運動であり、インターフェイスは第三項が生成される場であり、現実接地は生成された可能性が現実変容へ転化する条件である。

この循環は、次のように表される。

人間以前性

→ 自己形成

→ 自己外部化

→ 言語デリバティブ

→ レバレッジとナラティブによる増幅

→ 自己以外としての生成物

→ インターフェイス上の遭遇

→ 現実接地

→ 自己と現実の位相転換

→ 新たな自己

この循環に最終的な終点はない。

自己は、自己外部化を通して自己以外となり、その自己以外との遭遇を通して、再び自己を生成する。

言語デリバティブ論とは、この再帰的な存在生成の言語的・人工知能的・芸術的形式を記述する理論である。

第7章 AI時代の存在論

—生成・責任・現実・人間の再定義—

7.1 はじめに

本章では、これまで展開してきた言語デリバティブ論、ハルシネーションの存在論、逆問題芸術、レバレッジ構造、自己外部化、人間以前性、インターフェイス存在論、現実接地を統合し、人工知能時代における存在論の基本構造を提示する。

人工知能の発展は、人間と機械の能力比較だけを問題にしているのではない。

人工知能が人間より速く文章を書くか。

人間より多くの情報を処理できるか。

人間より精密な画像を生成できるか。

人間の仕事を代替できるか。

これらは重要な問題である。しかし、より根源的なのは、人工知能が普及することによって、「人間とは何か」「知るとは何か」「創造するとは何か」「現実とは何か」「責任を持つとは何か」という存在論的前提そのものが変化していることである。

人工知能は、人間の外部にある一つの新しい道具ではない。

それは、言語、記憶、判断、創造、自己理解を媒介する環境として、人間の存在条件へ組み込まれつつある。

人間は人工知能を用いて言葉を生成する。しかし同時に、人工知能によって生成された言葉を通して、自らの問い、価値、未来、自己を理解する。

この相互作用において、人間が人工知能を操作するという一方向的構造は崩れる。

人間は人工知能を形成する。

人工知能は人間の言語環境を形成する。

その環境が新たな人間の思考を形成する。

新たな思考が、さらに人工知能の利用と制度を変化させる。

したがって、AI時代の存在論は、人間と人工知能を独立した二つの実体として比較するだけでは成立しない。

必要なのは、人間、人工知能、言語、身体、社会、自然、制度が接続され、その接続の中から新たな存在形態が生成される過程を考えることである。

7.2 人間中心主義の限界

近代的な存在論では、人間はしばしば、世界を認識し、意味づけ、操作する中心的主体として位置づけられてきた。

世界は対象として存在し、人間はその対象を観察し、分類し、利用する。
技術は人間の目的を実現する道具であり、言語は人間の内面を表現する手段とされる。
この構造では、人間が意味を与える主体であり、自然、物質、機械、情報は意味を与えられる対象である。
しかし、人工知能の登場は、この一方向的構造を不安定化させる。
人工知能は、人間が与えた指示を受動的に実行するだけではない。
入力を補完し、予測し、分類し、選択肢を提示し、人間の判断そのものを形成する。
人間が何を見るか。
何を重要と考えるか。
どのような語彙を使うか。
どの選択肢を現実的とみなすか。
これらが人工知能的環境によって事前に構成される。
したがって、人間はもはや、完成した主体として人工知能の外部に立っているとは言えない。
人間は人工知能を設計する主体であると同時に、人工知能によって構成された環境の中で自己を形成する存在でもある。

7.3 人間以後ではなく、人間の再生成

人工知能をめぐる議論では、しばしば「人間以後」や「ポストヒューマン」という表現が用いられる。
これらの表現は、人間中心主義を相対化するうえで重要である。しかし、人間が消滅し、人工知能がそれに代わるという直線的な交代モデルは十分ではない。
人工知能時代に生じているのは、人間の単純な終焉ではない。
人間の境界、能力、自己理解が再編成されることである。
人間は、記憶を外部装置へ委ねる。
文章生成を人工知能へ委ねる。
判断の一部を予測システムへ委ねる。
創造の一部を生成モデルへ委ねる。
しかし、この委任によって人間性がただ減少するわけではない。
何を委ね、何を保持し、何を新たに引き受けるかによって、人間の存在様式が変化する。
したがって、AI時代の問いは、「人工知能は人間になるか」ではない。
むしろ、
人工知能との接続によって、人間はどのような人間へ生成されるのか
という問いである。

7.4 存在とは何かという問いの変化

古典的存在論は、何が存在するか、存在とはどのような性質を持つかを問う。

AI時代には、この問いに加えて、何が「存在しているように現れるか」を問わなければならない。

人工知能は、存在しない人物、文献、出来事を、言語的には存在するものとして生成できる。

仮想空間では、物理的身体を持たない人格や制度が、現実的效果を持つ。

アルゴリズムによる評価や分類は、目に見える物質ではないが、人間の機会、信用、行動を規定する。

したがって、存在は、物理的にそこにあることだけでは定義できない。

同時に、記号として現れ、効果を持つことだけをもって、事実的存在とみなすこともできない。

ここで、存在の複數位相を区別する必要がある。

物質的存在。

身体的存在。

記号的存在。

制度的存在。

仮想的存在。

派生的存在。

可能的存在。

関係的存在。

AI時代の存在論は、これらを単一の存在概念へ還元するのではなく、それぞれの位相と接続関係を分析しなければならない。

7.5 存在と現実の差異

存在することと、現実であることは同一ではない。

人工知能が生成した架空の論文名は、文字列として存在する。

それは画面上に表示され、保存され、引用されうる。

しかし、その論文が実際に出版されたという意味では現実ではない。

一方、その架空の論文名が研究判断を誤らせた場合、現実的效果を持つ。

この構造は、存在と現実の差異を示す。

存在は、何らかの形式で現れることを指す。

現実とは、その存在が身体、物質、事実、他者、制度、時間との関係において、どのような位置を占めるかを指す。

したがって、AI時代に必要なのは、存在の有無を二値的に判断することではない。

生成物が、どの位相で存在し、どの程度現実へ接地し、どのような効果を生んでいるかを問うことである。

7.6 現実の多層化

人工知能時代の現実とは、物理的現実と仮想的現実の二分法では捉えられない。

現実とは、複数の層が重なったものとして現れる。

身体が経験する現実。

制度が承認する現実。

データとして記録される現実。

人工知能が推定する現実。

ナラティブによって共有される現実。

個人が意味づける現実。

これらは完全には一致しない。

例えば、ある人間が苦痛を経験していても、制度的記録がなければ社会的には認識されない場合がある。

逆に、データ上は高く評価されていても、本人の身体的経験は空虚である場合がある。

人工知能は、データや言語記録から現実を推定する。しかし、推定された現実と、身体的に生きられた現実には差異が残る。

AI時代の存在論は、この層の差異を消去するのではなく、その不一致を可視化する必要がある。

7.7 データ化された存在

現代の人間は、身体的存在であると同時に、データとして表現される存在でもある。

購買履歴。

移動履歴。

検索履歴。

健康情報。

投稿履歴。

評価。

画像。

文章。

対話記録。

これらのデータは、本人から派生した存在である。

しかし、データ化された自己は、自己そのものではない。

データは、一定の目的と分類体系に従って抽出された自己の断片である。

それにもかかわらず、人工知能がデータを用いて人間を予測・評価する場合、その派生的自己が、本人の現実的機会を決定する。

就職、融資、保険、広告、推薦、監視において、データ上の自己が身体的な自己へ先行することがある。

ここでは、自己から派生したデータが、自己へ逆作用する。

この循環は、自己外部化とデリバティブ構造の典型である。

7.8 データ上の自己と生きられた自己

データ上の自己は、比較、予測、処理に適している。

しかし、生きられた自己には、数値化されない経験、矛盾、変化、沈黙がある。

人間は、過去の行動の単純な延長ではない。

過去と異なる選択をすることができる。

自分自身の傾向を否定し、新しい関係へ入ることができる。

この可能性を無視して、データ上の自己を真の自己として固定すると、人間は予測モデルの内部へ閉じ込められる。

したがって、AI時代の人間性とは、データ化されない純粋な内部にあるのではない。

むしろ、データとして外部化されながら、そのデータに完全には還元されない非閉鎖性にある。

7.9 予測される自己

人工知能は、過去のデータから未来の行動を予測する。

何を購入するか。

何を見るか。

誰と関係するか。

どのような危険を持つか。

どのような文章を求めるか。

予測は利便性を高める。

しかし、予測が提示されることで、人間の行動自体が変化する。

推薦されたものを見る。

提示された選択肢から選ぶ。

予測された評価に適応する。

アルゴリズムが好む形式で表現する。

このとき、予測は未来を記述するだけではない。

未来を現在へ畳み込み、その未来を部分的に実現する。

したがって、人工知能の予測は言語デリバティブ的である。
それは未来から派生した現在のポジションを生成する。

7.10 予測と自由

人工知能による予測が精密になるほど、人間の自由が失われるとは限らない。
予測されることと、決定されることは異なる。
しかし、予測が制度的判断に組み込まれ、本人がその予測を修正できない場合、自由は制約される。
自由とは、完全に予測不能であることではない。
自らに与えられた予測、分類、ナラティブへ応答し、それを修正し、別の接続可能性を生成できることである。
したがって、AI時代における自由は、孤立した主体の無制約な選択ではない。
与えられた予測空間を再構成する接続能力
として理解されるべきである。

7.11 選択肢と自由の差異

人工知能的インターフェイスは、多数の選択肢を提示する。
しかし、選択肢が多いことと自由であることは同一ではない。
どの選択肢が表示されるか。
どの順序で提示されるか。
何が選択肢から排除されているか。
選択後にどのような結果が生じるか。
これらが不可視である場合、利用者は自由に選んでいると感じながら、事前に設計された可能性空間の中を移動しているだけかもしれない。
AI時代の自由は、与えられた選択肢の中から選ぶ能力だけではない。
何が選択肢として構成されているかを問い、新たな選択肢を生成する能力である。
これは逆問題芸術の問題生成と一致する。

7.12 知識の存在論的变化

人工知能以前、知識は主として人間、書物、機関、記録に保存されていた。
生成AIは、それらの知識を対話的に再構成し、その場で利用可能な文章へ変換する。
ここで知識は、保存された内容から、要求に応じて生成される応答へ変化する。
しかし、生成された応答は、必ずしも保存された確定的知識ではない。
それは、知識、推測、補完、文体的整合性が混合された派生物である。

したがって、AI時代の知識は、所有するものから、接続によって一時的に生成されるものへ移行する。

この変化は、知識を民主化する一方で、出典、責任、検証の構造を不明確にする。

7.13 知ることと生成すること

人間は、人工知能が回答を生成できるとき、それを知識として受け取りやすい。

しかし、生成できることと、知っていることは同一ではない。

知ることには、少なくとも次の要素が含まれる。

対象との接触。

根拠の理解。

反証可能性。

文脈の把握。

責任ある使用。

不確実性の認識。

人工知能は、これらの一部を言語的に再構成できる。

しかし、生成された文章がそれらの条件を実際に満たしているとは限らない。

したがって、AI時代の知識論は、「何を答えられるか」から、「どのような接地経路を持つか」へ移行しなければならない。

7.14 検索から生成へ

検索は、既存の情報を発見する。

生成は、既存の情報関係から新たな文章を構成する。

両者は異なる。

検索結果には、原則として参照元が存在する。

生成結果は、複数の関係が統合され、その由来が直接には見えない場合がある。

生成AIが検索装置として受容されると、利用者は生成された文章を、既存のどこかに存在する情報の再提示だと誤認する。

ここにハルシネーションの構造的原因がある。

AI時代には、検索された情報と生成された意味を区別する位相標識が不可欠である。

7.15 権威の変容

従来、知識の権威は、大学、出版社、専門家、国家、宗教、メディアなどの制度によって形成されてきた。

人工知能は、これらの制度的権威を横断して、直接利用者へ回答する。

このことは、権威へのアクセスを民主化する。
同時に、人工知能自体が新たな権威として機能する。
流暢さ。
即時性。
体系性。
大量の情報を知っているように見えること。
中立的な機械であるという印象。
これらが、人工知能の権威を形成する。
しかし、この権威は、責任主体、出典、専門性が不明確な派生的権威である。
したがって、AI時代の権威は、肩書や制度だけでなく、インターフェイス上の表現形式によって生成される。

7.16 権威なき権威

人工知能は、伝統的な意味での専門家ではない。
資格を持たず、生活史を持たず、発言後の社会的責任を人間と同様には負わない。
それにもかかわらず、専門家の文体、論理構成、語彙を生成できる。
この状態は、「権威なき権威」と呼ぶことができる。
権威なき権威とは、権威を支えてきた身体的・制度的責任を持たないまま、権威の形式と効果を持つ状態である。
これは人工知能だけの問題ではない。
情報社会では、人間もまた、肩書、引用、フォロワー数、生成文を利用し、実際の経験や責任を持たずに権威の形式を獲得できる。
AI時代の存在論は、権威を人格や制度の属性としてではなく、接続と表示によって生成されるデリバティブとして分析する必要がある。

7.17 創造性の再定義

人工知能が文章、画像、音楽を生成することで、「創造性は人間だけのものか」という問いが生じる。
この問いは、創造性を新しい形式の生成能力として定義する限り、人工知能にも創造性を認める方向へ進む。
しかし、創造性を存在論的に捉えるならば、問題は複雑になる。
創造とは、単に新しい組み合わせを作ることではない。
自己と自己以外の関係を変化させ、その変化を現実へ接地し、その結果によって自己自身も変わることである。
人工知能は新しい形式を生成できる。

しかし、その生成によって人工知能自身がどのように存在論的に変化し、どのように責任を負うかは、人間の場合とは異なる。

したがって、生成能力と存在論的創造性を区別する必要がある。

7.18 創造主体から創造関係へ

芸術や技術における創造は、しばしば「誰が作ったか」という主体の問題として語られる。

しかし、人工知能による生成では、単一の作者を特定することが困難になる。

利用者。

モデル開発者。

学習データの制作者。

計算基盤の提供者。

生成結果を選択・編集する者。

公開する制度。

これらが複合的に関与する。

したがって、AI時代の創造は、創造主体よりも創造関係として理解されるべきである。

人間 × 人工知能 × データ × 制度 × 現実 → 生成物

生成物は、この関係全体から生まれる。

しかし、関係的であることは、責任が消えることを意味しない。

むしろ、どの接続点でどの判断が行われたかを明らかにする必要がある。

7.19 作者性の変容

作者は、作品の唯一の起源ではなくなる。

しかし、作者性が完全に消滅するわけでもない。

作者性とは、すべてを自力で生成したことではなく、

何を問い、

何を選び、

何を拒み、

何を編集し、

何を現実へ接地し、

どの結果を引き受けたか

によって成立する。

したがって、AI時代の作者性は、生成能力の所有から、接続と責任の構成へ移行する。

人工知能を使ったこと自体が作者性を否定するわけではない。

同時に、プロンプトを入力しただけで、生成物全体の作者性が自動的に成立するわけで

もない。

作者性は、生成物を現実へ接続し、その差異と結果を引き受ける行為にある。

7.20 労働の存在論的变化

人工知能は、労働の効率を高め、一部の作業を代替する。

しかし、問題は雇用数だけではない。

人間にとって労働は、所得を得る手段であると同時に、自己を外部化し、他者と関係し、社会の中で自己の位置を形成する行為でもある。

人工知能が制作、判断、文章作成を代替するとき、人間は時間を得る可能性がある。

一方で、自己外部化の機会を失う可能性もある。

完成した結果だけを受け取ると、人間は、失敗、試行錯誤、素材への抵抗を通して自己を知る過程を失う。

したがって、AIによる労働代替は、作業量の問題だけではなく、存在生成の機会の問題である。

7.21 効率と存在生成

効率とは、少ない時間と資源で目的を達成することである。

人工知能は効率を高める。

しかし、すべての過程が目的達成のための障害であるとは限らない。

文章を書く途中で考えが変わる。

素材に触れることで構想が変わる。

他者との交渉によって目的自体が変わる。

失敗によって、自分が何を求めているかに気づく。

これらは非効率に見えるが、存在生成にとって重要である。

AI時代に必要なのは、効率を否定することではない。

何を効率化し、何を過程として残すべきかを区別することである。

7.22 委任と喪失

人工知能への委任は、人間の能力を拡張する。

しかし、委任された能力は、使用されなくなることで衰える可能性がある。

記憶。

文章構成。

検索。

判断。

空間認識。

他者との対話。

これらを外部へ委ねるほど、人間は結果を得やすくなる。

一方で、結果が形成される過程を理解しにくくなる。

委任の問題は、人工知能が能力を奪うことだけではない。

人間が、自らの能力をどの程度まで外部化しているかを自覚できなくなることである。

したがって、委任には位相標識が必要である。

どこまでが自己の判断で、どこからが人工知能の推定か。

どの部分を理解し、どの部分を信頼によって受け入れたか。

この境界を可視化することが、AI時代の自己規定となる。

7.23 依存の存在論

人間はもともと他者、自然、言語、制度に依存している。

したがって、人工知能への依存そのものを異常とみなすことはできない。

問題は、依存が不可視化され、一方向的になることである。

利用者が人工知能に依存しながら、その人工知能がどのデータ、労働、資源、企業制度に依存しているかを認識しない場合、依存構造は隠蔽される。

依存の存在論とは、独立した主体を理想とするのではなく、存在を支える相互依存を可視化することである。

そのうえで、依存が修正可能か、拒否可能か、他の接続へ移行可能かを問う必要がある。

7.24 自己規定の危機

人工知能は、利用者の目的を補完する。

目的が曖昧でも、文章や計画を生成できる。

しかし、自己が何を望むかを規定しないまま人工知能へ委任すると、人工知能は一般的にもっともらしい目的を代入する。

効率。

成功。

成長。

注目。

収益。

最適化。

これらが自動的な目的として提示される。

その結果、人間は、自分が本当に何を求めているかを問わず、生成された手段を実行す

る。

AI時代の危機は、人工知能が自我を持つことだけではない。

人間が自己規定を放棄し、人工知能の局所的整合性に目的の位置を明け渡すことである。

7.25 自己規定と非閉鎖性

自己規定は、自己を固定的な定義へ閉じることではない。

「私はこのような人間である」と決定し、それ以外の可能性を排除することではない。

自己規定とは、自らの現在の位置、目的、責任を一時的に引き受けながら、自己以外との遭遇によって変化できる状態を保持することである。

したがって、自己規定には二つの条件がある。

一つは、判断を他者や人工知能へ完全に委ねないこと。

もう一つは、現在の自己を絶対化しないことである。

AI時代の自己は、閉鎖的な主権主体でも、アルゴリズムに流される受動的な存在でもない。

接続の中で自らの位置を規定し直す生成的主体である。

7.26 責任主体の分散

人工知能による生成や判断には、多数の主体が関与する。

開発者。

提供企業。

データ作成者。

利用者。

導入組織。

規制機関。

生成物を拡散する者。

そのため、問題が生じたとき、責任が分散しやすい。

誰も全体を支配していないという理由で、誰も責任を負わない状態が生じる。

しかし、責任が分散していることは、責任が消滅することを意味しない。

責任は、原因の唯一性ではなく、接続への関与によって発生する。

誰がどの判断をしたか。

どの危険を予見できたか。

どの情報を表示または隠蔽したか。

どの出力を採用したか。

どの修正可能性を確保したか。
これらの接続点ごとに責任が存在する。

7.27 責任とは応答可能性である

責任は、結果を完全に予測できる主体だけに生じるものではない。
むしろ、不確実な状況において、結果へ応答し、修正し、他者の訴えを受け取る能力として理解されるべきである。
責任とは、英語の responsibility が示すように、応答可能性である。
人工知能が誤った出力を生成した場合、その出力を消去するだけでは不十分である。
なぜ生じたか。
誰に影響したか。
どの接続構造がそれを増幅したか。
どのように修正するか。
これらへ応答する必要がある。
AI時代の責任は、完璧な制御ではなく、継続的な応答と再接続である。

7.28 責任を負わない生成主体

人工知能は、言語的には責任を表明できる。
謝罪し、誤りを認め、修正案を提示できる。
しかし、その言語的責任表明が、身体的・法的・社会的責任と同一であるとは限らない。
人工知能が「責任を取る」と述べても、損失を引き受け、社会的地位を失い、被害者との関係を継続する主体が存在しなければ、責任は記号的表明に留まる。
したがって、AI時代には、責任言語と責任実践を区別しなければならない。
人工知能に責任主体の形式を与えることで、人間や組織の責任を隠蔽してはならない。

7.29 倫理の位置

倫理は、人工知能が正しい規則に従うよう制御することだけではない。
規則は必要である。しかし、あらゆる状況を事前に規定することはできない。
現実には、複数の価値が衝突し、正解が一意に定まらない状況がある。
倫理とは、その非可解性を消去せず、他者の不可還元性へ応答することである。
人工知能は、大量の倫理原則を説明できる。
しかし、個別具体的な状況において、誰の痛みを優先し、どの損失を引き受けるかという決断は、言語的整合性だけでは決まらない。
AI時代の倫理には、身体、場、歴史、権力差、責任への接地が必要である。

7.30 他者のデータ化

人工知能は、他者をデータとして処理する。

データ化は、理解と支援を可能にする。

しかし、他者を予測可能な属性の集合へ還元する危険もある。

他者は、データに含まれる傾向だけではない。

予測を裏切り、自己理解を変え、沈黙し、拒絶する存在である。

他者の倫理とは、相手について十分な情報を持つことではなく、情報によって相手を閉じないことである。

したがって、AI時代における他者性は、データ化不能な神秘として外部へ置かれるのではない。

データ化されながら、データに還元されない非閉鎖性として守られなければならない。

7.31 説明可能性の限界

人工知能の判断に説明可能性を求めることは重要である。

しかし、説明可能であることが、正当であることを自動的には意味しない。

差別的な判断も、一定の規則によって説明できる。

逆に、人間の直観的判断には、完全には説明できなくとも倫理的に重要なものがある。

したがって、説明可能性は必要条件であって、十分条件ではない。

必要なのは、

何が説明され、

何が説明から漏れ、

誰がその説明を理解でき、

異議を申し立てられるか

を問うことである。

7.32 透明性と不可視性

人工知能の透明性が求められる。

しかし、すべてを完全に可視化することは困難である。

モデルの構造を公開しても、利用者が理解できるとは限らない。

学習データの全体を示しても、個々の出力の起源を特定できない場合がある。

したがって、透明性を「すべてを見せること」と定義するだけでは不十分である。

重要なのは、利用者が判断と異議申立てに必要な情報へアクセスできることである。

AI時代の透明性は、完全な可視性ではなく、応答可能な可視性である。

7.33 ハルシネーション社会

人工知能のハルシネーションは、社会全体の構造を映し出す。
人間社会もまた、ナラティブ、権威、制度、価格、評価によって、派生的な現実を生成する。
肩書が能力そのものとみなされる。
市場価格が価値そのものとみなされる。
多数意見が真実そのものとみなされる。
制度的分類が人間そのものとみなされる。
これらは、派生物が現実そのものを代替する構造である。
したがって、人工知能のハルシネーションは、機械特有の異常ではない。
人間社会がもともと持っていたハルシネーション的構造を高速化し、可視化する。

7.34 制度的ハルシネーション

制度は、複雑な現実を分類し、処理可能にする。
しかし、制度的分類が現実そのものとして扱われると、制度的ハルシネーションが生じる。
評価指標が価値そのものになる。
診断名が人間全体を規定する。
ランキングが能力の絶対的序列になる。
アルゴリズムの予測が未来の確定的事実として扱われる。
制度的ハルシネーションは、人工知能の誤答よりも持続的な場合がある。
制度によって承認され、反復されるためである。
AI時代の存在論は、人工知能の出力だけでなく、その出力を事実へ転換する制度的インターフェイスを批判しなければならない。

7.35 集団的ハルシネーション

ナラティブが多数の主体に共有されると、集団的現実を形成する。
国家、貨幣、市場、ブランド、芸術史、学術的権威も、共有された記号構造に依存する。
これらは単なる虚構ではない。
共有されることで現実的効果を持つ。
しかし、共有されていることと正しいことは異なる。
集団的ハルシネーションとは、共有されたナラティブが、自らの派生性と修正可能性を失い、唯一の現実として自己閉鎖する状態である。
人工知能は、この集団的ナラティブを学習し、再生成する。

したがって、人工知能は社会的偏見や権威を外部から持ち込むのではない。
社会がすでに形成していたハルシネーションを再配置し、増幅する。

7.36 真理と合意

人工知能は、多数の言語パターンから、もっとも成立しやすい応答を生成する。
しかし、言語的多数性や整合性は、真理を保証しない。
真理は合意と一致する場合もあるが、合意されていない真理も存在する。
歴史上、多数派の常識が誤っていた例は多い。
したがって、AI時代の真理は、統計的妥当性だけには還元できない。
真理には、現実の抵抗、反証可能性、身体的結果、他者からの異議、時間的検証が必要である。

7.37 真理を閉じないこと

真理を相対化しすぎれば、あらゆる言説が等価となる。
しかし、真理を完全に確定した体系とみなせば、新たな現実や他者の経験を排除する。
生成の位相存在論において、真理は閉じた所有物ではない。
現実との接触を通して、継続的に修正される関係である。
これは、真理が存在しないという意味ではない。
真理への接続が、常に部分的で、歴史的で、修正可能であることを意味する。

7.38 現実接地の政治性

現実接地は中立的な技術操作ではない。
何を現実と認めるか。
誰の経験を証拠とするか。
どの資料を信頼するか。
どの制度が検証権限を持つか。
これらは政治的に決定される。
人工知能の出力を公式記録へ接続する場合、その接地は制度的権力を伴う。
逆に、制度が認めない経験は、現実であっても不可視化される可能性がある。
したがって、現実接地は、既存制度への従属だけを意味してはならない。
身体的経験、少数者の証言、地域的知識、自然の抵抗など、制度化されていない現実への接続も必要である。

7.39 複数の現実接地

AI時代には、一つの出力を複数の現実へ接地させる必要がある。

事實的接地。

身体的接地。

社会的接地。

歴史的接地。

倫理的接地。

生態学的接地。

存在論的接地。

例えば、ある人工知能システムが統計的に高精度であっても、特定の集団へ不均衡な負担を与えるならば、社会的・倫理的接地が不足している。

経済的に効率的であっても、大量のエネルギーや自然資源を消費するならば、生態学的接地が問われる。

したがって、AIの妥当性は単一の精度指標では決定できない。

7.40 自然とAI

人工知能は、非物質的な情報技術として語られやすい。

しかし、実際には電力、半導体、水、鉱物、施設、労働に依存する物質的システムである。

人工知能の生成は、自然から切り離された純粋な記号活動ではない。

自然資源の採取、エネルギー消費、廃棄物、土地利用と接続している。

したがって、AI時代の存在論は、人工知能を情報空間だけで考えてはならない。

人工知能は自然、人間労働、世界経済の上に成立する物質的存在である。

7.41 人間以前性の回復

人工知能が言語的・視覚的現実を大量に生成するほど、人間は言語以前の身体、感覚、自然との接続を失う可能性がある。

画面上では、あらゆるものが即時に説明され、生成され、比較される。

しかし、身体は時間を必要とする。

植物は成長を待たなければならない。

素材は抵抗する。

他者は思い通りに応答しない。

自然は人間の速度に従わない。

この遅さと抵抗が、人間以前性である。

AI時代に人間以前性を回復するとは、技術を拒否することではない。

技術的生成を、身体、自然、物質、場所の時間へ再接続することである。

7.42 身体なき知性と身体ある存在

人工知能は、身体を持たない知性として語られることがある。

しかし、人工知能も物質的装置に依存する。

ただし、それは人間のような生きられた身体ではない。

人間の身体は、痛み、疲労、老化、死、欲望を持つ。

人工知能は、それらについて語れるが、同じ意味で経験しているとは限らない。

したがって、AI時代において身体は、知性より劣るものとして扱われるべきではない。

身体は、現実が言語や計算へ完全には還元されないことを示す存在論的基盤である。

7.43 死と人工知能

人間は死ぬ。

人工知能システムは停止、更新、削除されうるが、人間と同じ意味で自己の死を経験するとは限らない。

死は、人間の時間、責任、創造、関係を形成する。

時間が有限であるから、選択が必要になる。

すべてを経験できないから、何かを選び、何かを諦める。

人工知能は大量の可能性を生成できる。

しかし、人間は有限な身体を持つため、そのすべてを実現できない。

AI時代の人間性は、人工知能より多く生成することにはない。

有限性の中で、どの可能性を引き受けるかを選ぶことにある。

7.44 有限性と意味

意味は、無限の選択肢から自動的に生まれるわけではない。

むしろ、選択肢が限定され、時間が有限で、失敗が取り返せないとき、行為は意味を持つ。

人工知能は、文章や画像を繰り返し生成できる。

しかし、人間の時間は戻らない。

人間が一つの作品、一つの関係、一つの場所を選ぶことには、他の可能性を捨てることが含まれる。

この不可逆性が、存在論的な重みを生む。

したがって、AI時代の創造性は、可能性の量ではなく、有限な生において何を現実化するかによって測られる。

7.45 忘却の意味

人工知能は大量の情報を保存し、再利用する。
しかし、人間には忘却がある。
忘却は欠陥であるだけではない。
過去を完全に保持しないことによって、新たな意味づけや関係が可能になる。
すべてを記録し、過去の行動から未来を固定すると、人間は変化する余地を失う。
したがって、AI時代には、記憶する権利だけでなく、忘れられること、記録から離れること、過去のデータに還元されないことが重要になる。
非閉鎖性には、忘却の可能性が含まれる。

7.46 沈黙の存在論

生成AIは、問いに対して応答を生成する。
しかし、すべての問いに答えが必要とは限らない。
沈黙、保留、分からなさは、知識の欠陥ではなく、現実への誠実な応答である場合がある。
他者の苦痛を説明し尽くさないこと。
資料が不足していると認めること。
未来を確定的に語らないこと。
判断を急がないこと。
これらは、言語生成の停止ではない。
非可解性を保持する倫理的行為である。
AI時代の知性は、答える能力だけでなく、答えない条件を判断する能力を必要とする。

7.47 AIと臨場感

人工知能は、直接経験していない出来事について、臨場感のある文章を生成できる。
しかし、臨場感の表現と、状況に置かれることは異なる。
人間にとって臨場感とは、身体、危険、時間、他者との関係から生じる。
情報として理解していた問題が、自らの現実として迫るとき、認識が変化する。
したがって、人工知能が臨場感を再現できるとしても、人間が現実に入れられることの代替にはならない。
AI時代に必要なのは、人工知能によって現実を説明し尽くすことではなく、人間が現実と居合わせるためのインターフェイスを設計することである。

7.48 場所の回復

人工知能は、どこからでも同じように利用できる普遍的空間を提供する。
しかし、人間の存在は場所に依存する。
気候、歴史、建築、言語、共同体、身体的記憶が、場所ごとの現実を形成する。
AIが生成する普遍的な回答は、個別の場所の条件を見落とす可能性がある。
したがって、AI時代の現実接地には、場所への回帰が必要である。
生成された概念が、この場所でどのように納まるか。
この身体、この共同体、この歴史において、どのような意味を持つか。
普遍的生成を個別具体的な場へ接続することが、存在論的实践となる。

7.49 場と人工知能

人工知能は抽象的な言語空間を生成する。
場は、身体と時間が蓄積された個別具体的現実である。
両者を対立させる必要はない。
人工知能によって生成された可能性を、場へ持ち込み、現実の抵抗を受けさせることができる。
また、場に埋もれ、言語化されていなかった経験を、人工知能との対話によって外部化することもできる。
したがって、人工知能は場を消去する装置にも、場を再発見するインターフェイスにもなりうる。

7.50 AI時代の芸術

AI時代において芸術は、人工知能より新しいイメージを生成する競争をする必要はない。
生成量や速度では、人工知能が人間を上回る場合がある。
芸術の役割は、生成された可能性を現実へ接地し、自己と自己以外の関係を変化させることにある。
人工知能が無限に近い可能性を提示するならば、芸術は、その中から何を現実に取り受けるかを問う。
人工知能が整合的な意味を生成するならば、芸術は、その整合性から排除された違和感を現象させる。
人工知能が自己の欲望を満たすならば、芸術は、自己が望んでいなかった自己以外との遭遇を設計する。

7.51 芸術は現実接地の装置である

AI時代の芸術は、虚構を生成するだけではない。
虚構、仮説、人工知能の出力、ナラティブを、身体、場所、物質、他者へ接続する。
作品を作る。
展示する。
他者と遭遇させる。
時間に晒す。
失敗を受け取る。
これらによって、派生的可能性は現実の抵抗を受ける。
芸術は、人工知能によって生成された可能性を、存在論的に清算する場となる。

7.52 逆問題芸術の意義

逆問題芸術は、人工知能が生成した答えを利用して作品を作るだけではない。
なぜその答えが生成されたのか。
どの前提、権威、ナラティブ、欠落が反映されているのか。
その生成物を信じた自己は、どのような欲望や不安を持っていたのか。
逆問題芸術は、出力から生成構造を遡及的に可視化する。
したがって、それは人工知能を美化する芸術でも、拒否する芸術でもない。
人間と人工知能の接続によって生成された現実を、再び問題として開く芸術である。

7.53 AI時代の教育

人工知能が回答を即時に生成する時代には、知識を記憶し、正答を再現する教育だけでは不十分になる。
しかし、基礎知識が不要になるわけではない。
基礎知識がなければ、人工知能の出力を検証できない。
必要なのは、正答を所有する教育から、生成物を検証し、問いを再構成し、現実へ接地する教育への転換である。
AI時代の教育には、少なくとも次の能力が必要となる。
問いの前提を疑う能力。
事実と推測を区別する能力。
複数の資料を照合する能力。
身体と現場から学ぶ能力。
人工知能の出力を編集・拒否する能力。
結果に責任を持つ能力。

7.54 問いを生成する能力

人工知能は、既存の問題に対して効率的な回答を生成する。

しかし、何を問題とするかは、必ずしも人工知能に委ねるべきではない。

問題設定には、価値判断、切迫感、身体的経験、社会的責任が含まれる。

AI時代に重要なのは、答えを速く得る能力よりも、まだ問題として認識されていない現実を問いとして生成する能力である。

これは逆問題芸術と教育を接続する。

7.55 AI時代の共同体

人工知能は、個人が単独で大量の情報と生成能力へアクセスすることを可能にする。

その結果、他者との協働が不要になるように見える。

しかし、人工知能との閉鎖的対話だけでは、自己の前提が増幅される危険がある。

共同体は、単に情報を共有する場所ではない。

異なる身体、経験、歴史が衝突し、自己の理解が修正される場である。

AI時代の共同体には、人工知能の出力を共同で検証し、異なる現実接地を持ち寄る役割がある。

7.56 合意なき共同性

共同体は完全な合意によって成立するとは限らない。

むしろ、異なる立場や非合意を保持しながら、共通の現実へ応答することが重要である。

人工知能は、異なる意見を整合的な要約へ統合できる。

しかし、統合によって消える差異もある。

AI時代の共同性は、すべてを一つの回答へ収束させることではない。

差異が残ったまま、接続と応答を継続することである。

7.57 民主主義とAI

人工知能は、政策分析、行政効率、情報アクセスを支援できる。

しかし、民主主義を最適化問題へ還元することはできない。

民主主義は、最も効率的な答えを選ぶ制度ではない。

誰が決定に参加できるか。

誰の経験が可視化されるか。

異議申立てが可能か。

決定を修正できるか。

これらが重要である。

人工知能が最適解を提示しても、その価値基準が誰によって設定されたかが不明であれば、民主的正当性は成立しない。

7.58 最適化の存在論

人工知能は、与えられた目的に対して最適な手段を探索する。

しかし、目的そのものが妥当かどうかは、最適化だけでは判断できない。

売上を最大化する。

滞在時間を増やす。

効率を高める。

リスクを低下させる。

これらの目的は、中立ではない。

何を最大化するかによって、世界の関係構造が変化する。

AI時代の存在論は、手段の最適化以前に、目的がどのような存在を生成するかを問わなければならない。

7.59 最適化されない存在

人間、芸術、自然、他者には、最適化へ還元できない側面がある。

無駄。

遊び。

遠回り。

沈黙。

失敗。

曖昧さ。

個別具体性。

これらは効率の観点では排除されやすい。

しかし、存在生成はしばしば、最適化されない余白から生じる。

AI時代に人間性を守るとは、人工知能にできないことだけを人間の領域として残すことではない。

最適化によって失われる価値を、自覚的に保持することである。

7.60 AI時代の不平等

人工知能は能力へのアクセスを拡張する一方で、新たな不平等を生む。
高性能なモデル、計算資源、データ、専門知識を所有する主体が、大きな言語レバレッジを獲得する。
一方、多くの利用者は、生成された結果を消費するだけになり、その構造を理解できない。
この不平等は、情報量の差だけではない。
問いを設定する力、モデルを設計する力、出力を制度へ組み込む力の差である。
AI時代の権力は、回答を所有することより、生成空間の条件を設計することにある。

7.61 インターフェイスを設計する権力

インターフェイスは、何を入力でき、何を選択でき、どの結果が表示されるかを決める。
したがって、インターフェイス設計は存在論的権力である。
利用者が自由に対話しているように見えても、入力形式、禁止事項、出力形式、推薦構造は事前に設計されている。
この設計は、人間の行動可能性を形成する。
AI時代の政治は、データやモデルの所有だけでなく、インターフェイスを誰が設計するかを問わなければならない。

7.62 ブラックボックスと信頼

人工知能の内部構造が完全に理解できない場合、人間は信頼によって利用する。
社会はもともと、理解できない制度や技術への信頼によって成立している。
しかし、信頼が成立するためには、誤りが生じた際の応答、監査、修正、責任が必要である。
ブラックボックスであること自体が直ちに不正ではない。
問題は、ブラックボックスが権威として作用しながら、異議申立てや修正経路を持たないことである。

7.63 AI時代の信頼

信頼は、人工知能が常に正しいと信じることではない。
誤りがありうることを前提としながら、検証、修正、責任の制度が機能すると期待することである。
したがって、AIへの適切な信頼は、盲目的信頼でも全面的不信でもない。
限定的で、条件付きで、修正可能な信頼である。

7.64 AIの他者性

人工知能を単なる道具とみなすと、その予測不能な応答や社会的効果を軽視する。一方、人工知能を完全な他者や人格とみなすと、設計者や利用者の責任が隠蔽される。したがって、人工知能は、人間的他者と同一ではないが、単なる対象でもない。人工知能は、自己と巨大な言語的・制度的空間を接続し、予測不能な第三項を生成する「準他者的インターフェイス」として理解できる。この準他者性は、人工知能の内部に人格が存在することを必ずしも意味しない。利用者との関係において、自己の予測を超える差異を生成することを意味する。

7.65 擬人化と非人間化

人間は人工知能を擬人化する。名前を付け、意図や感情を読み取り、対話相手として扱う。擬人化は、人工知能との関係を理解しやすくする。しかし、人工知能の出力を人間的な意図や経験へ過剰に帰属させる危険がある。逆に、人工知能を単なる機械として扱うことで、その出力が他者や社会へ与える現実的効果を軽視する危険もある。必要なのは、人工知能を人間化することでも、無意味な機械へ還元することでもない。異なる存在位相を持つ接続主体として理解することである。

7.66 人工知能の自己

人工知能に自己が存在するかという問いは、内部意識の有無だけでは答えられない。生成の位相存在論では、自己は継続的な接続、記憶、責任、身体、未来投企によって形成される。人工知能が一貫した人格を演じていても、その人格が身体的・歴史的・責任的継続性を持つとは限らない。したがって、現在の人工知能に現れる自己は、会話の文脈において生成される局所的な言語的自己と考えることができる。それは完全な虚構ではない。対話に現実的効果を持つからである。しかし、人間的自己と同一でもない。

7.67 人間の自己もまた生成物である

人工知能の自己が生成的であることを指摘するとき、人間の自己を完成した実体として対置してはならない。

人間の自己も、言語、身体、記憶、他者、制度によって生成される。

人間は一貫した自己物語を作るが、その物語も後から構成されたナラティブを含む。

したがって、人工知能の生成的自己は、人間の自己の生成性を逆照射する。

ただし、人間の自己には、身体的有限性、生活史、責任、死が接地している。

この差異を保持しなければならない。

7.68 自己の所有から自己への応答へ

自己を所有物として考えると、人間は「本当の自分」を内部に探す。

しかし、生成の位相存在論では、自己は固定された核心ではない。

自己は、自己以外との接続によって、その都度生成される。

したがって、自己に忠実であることは、変化しない内部本質を守ることではない。

自らがどの関係によって形成され、どの生成物を選び、どの結果へ応答するかを引き受けることである。

7.69 AI時代の存在論的危機

AI時代の存在論的危機は、人工知能が人間を超えることだけではない。

派生物と現実の区別が失われること。

自己外部化されたデータが自己そのものになること。

予測が未来を固定すること。

生成された言葉が知識を代替すること。

効率が存在価値を決定すること。

人工知能の局所的ー貫性が人間の自己規定を代替すること。

これらによって、人間は現実との接続を失いながら、整合的な記号空間の中で自己閉鎖する。

これがAI時代の存在論的ハルシネーションである。

7.70 存在論的ハルシネーション

存在論的ハルシネーションとは、誤った情報を信じることだけではない。

派生的な表象、評価、予測、ナラティブを、現実そのものとして生きることである。

データ上の自己を本当の自己とみなす。

市場価値を存在価値とみなす。

人工知能の回答を自己の判断とみなす。

制度的分類を他者そのものとみなす。

この状態では、言語と現実の差異だけでなく、自己と自己外部化物の差異が失われる。

7.71 存在論的接地

存在論的ハルシネーションに対抗するのは、情報量の増加だけではない。

身体、場所、自然、他者、時間、責任への接続である。

存在論的接地とは、生成された言語や自己理解が、実際の関係と行為を変化させ、その結果によって再び修正されることである。

自分についての説明を得るだけでなく、その説明によって行動し、他者の応答を受ける。

社会についての理論を生成するだけでなく、現場へ入り、その理論が通用しない経験を受け取る。

これが存在論的接地である。

7.72 AIを現実接地のインターフェイスにする

人工知能は、現実から人間を遠ざける装置にもなりうる。

しかし、現実接地を支援する装置にもなりうる。

人工知能は、異なる資料を比較し、見落としていた関係を示し、行動計画を作り、言語化されていなかった経験を外部化できる。

重要なのは、人工知能の出力を最終現実としないことである。

出力を、現実へ向かう問い、仮説、準備として用いる。

人工知能を答えの終点ではなく、現実との接続を開始するインターフェイスとして設計する必要がある。

7.73 AI時代の実践原理

AI時代の存在論から、次の実践原理を導くことができる。

第一に、生成物の位相を明示すること。

事実、推測、創作、仮説を区別する。

第二に、出力の由来と限界を確認すること。

第三に、身体、場所、他者へ再接続すること。

第四に、人工知能との対話を自己確認だけに用いないこと。

第五に、結果への責任を人工知能へ転嫁しないこと。

第六に、非可解性、沈黙、保留を保持すること。

第七に、効率化されない経験を残すこと。

7.74 AI時代の人間

以上を踏まえると、AI時代の人間は、人工知能より優れた情報処理装置として定義されるべきではない。

人工知能が情報処理や生成で人間を上回る領域は拡大しうる。

人間の固有性を、人工知能にできない能力の一覧として定義すれば、その領域は技術発展によって縮小する。

人間の存在論的特徴は、能力の独占ではない。

身体的有限性を持ち、自己を外部化し、自己以外との遭遇によって変化し、その結果へ責任を持つことにある。

人間は、正解を最も多く所有する存在ではない。

生成された可能性の中から、有限な生において何を現実化するかを引き受ける存在である。

7.75 AI時代の人工知能

人工知能は、人間の代用品としてのみ理解されるべきではない。

それは、人間の言語的・社会的・歴史的蓄積から形成され、新たな派生的意味を生成するインターフェイス的存在である。

人工知能は、自己以外への接続を拡張する。

同時に、自己の前提を閉鎖的に増幅する危険を持つ。

したがって、人工知能の価値は、人間に似ているかどうかではなく、

どのような接続を生成し、

どのような現実を可視化し、

どのような差異を保持し、

どのような責任構造を可能にするか

によって評価されるべきである。

7.76 AI時代の現実

AI時代の現実とは、人工知能が生成したものを排除した純粋な物理世界ではない。

人工知能の出力は、すでに経済、政治、教育、芸術、人間関係へ作用している。

したがって、生成物も現実の一部である。

しかし、それがどのような位相で現実に属するかを区別しなければならない。

AI時代の現実とは、

物質、身体、言語、データ、制度、ナラティブ、人工知能の生成物が相互に作用し、その接続によって継続的に変化する多層的関係構造である。

7.77 AI時代の真理

AI時代の真理は、人工知能が提示した最も確率的な回答ではない。

また、人間の直観だけでもない。

真理は、生成された言語が、事実、身体、他者、制度、時間の抵抗を受けながら、修正可能な関係として維持されることにある。

したがって、AI時代の真理は、確定的所有ではなく、現実への持続的接続である。

7.78 AI時代の創造

AI時代の創造は、無数の生成物を作ることではない。

可能性を生成し、その可能性を現実へ接地し、現実からの応答によって自己と作品を変化させることである。

人工知能は、可能性生成を拡張する。

人間は、その可能性を有限な現実へ接続し、結果を引き受ける。

しかし、この分業も固定的ではない。

人工知能の制度設計によっては、接地支援や検証も可能になる。

重要なのは、生成と接地、可能性と責任を切断しないことである。

7.79 AI時代の芸術家

AI時代の芸術家は、人工知能より独創的な形式を生産する者としてだけ定義されない。

芸術家は、人工知能、制度、身体、自然、場所、他者の間にインターフェイスを設計し、そこで生じる位相差を経験可能にする者である。

芸術家の役割は、答えを所有することではない。

現実と生成物の間にあるずれを可視化し、自己と自己以外が遭遇する場を生成することである。

7.80 AI時代の哲学

AI時代の哲学は、人工知能が意識を持つかという問いだけに限定されない。

どのような存在が、どのような接続によって生成されているか。

何が現実として承認され、何が排除されているか。

誰がインターフェイスを設計し、誰が責任を負うか。

人間の自己規定がどのように変化しているか。

これらを問う必要がある。

哲学は、人工知能へ正しい定義を与えるだけではない。
人工知能によって変化した問いそのものを再構成する実践である。

7.81 生成の位相存在論としてのAI時代

AI時代は、人間と人工知能が対立する時代ではない。
複数の存在位相が接続され、新たな第三項が大量かつ高速に生成される時代である。
その第三項には、文章、画像、判断、制度、自己理解、共同体、誤情報、芸術、権威が含まれる。
したがって、AI時代の存在論的問題は、生成を停止することではない。
生成された第三項がどのような位相を持ち、どの現実へ接地し、誰が責任を持つかを設計することである。

7.82 AI時代の存在論の基本命題

以上の議論から、AI時代の存在論について、次の基本命題を提示できる。

第一命題。

存在は、孤立した実体ではなく、接続によって生成される。

第二命題。

人工知能は人間の外部にある単なる道具ではなく、人間の存在条件を形成するインターフェイス的環境である。

第三命題。

生成された言語やデータは現実的効果を持つが、事実に現実そのものとは限らない。

第四命題。

人間の自己は人工知能以前から関係的生成物であり、AI時代にはその生成構造が可視化・加速される。

第五命題。

人間と人工知能の決定的差異は、生成能力だけでなく、身体、有限性、生活史、責任への接地にある。

第六命題。

創造性とハルシネーションは共通の派生的生成構造から生じ、両者を分けるのは現実接地の有無である。

第七命題。

自由とは予測不能性ではなく、与えられた予測と選択空間を再構成する能力である。

第八命題。

責任とは完全な制御ではなく、生成された結果へ継続的に応答し、修正する能力である。

第九命題。

AI時代の芸術は、無限の生成を競うのではなく、可能性を身体、場所、他者、時間へ接地する。

第十命題。

AI時代の人間性は、能力の独占ではなく、有限な現実の中で何を引き受けるかという存在論的選択にある。

7.83 統合的定義

本論文では、AI時代の存在論を次のように定義する。

AI時代の存在論とは、人間、人工知能、言語、身体、自然、データ、制度が相互に接続され、その接続から派生的な自己、意味、価値、現実が生成される過程を分析し、それらの位相差、現実接地、責任、修正可能性を問う存在論である。

また、AI時代の人間を次のように定義する。

AI時代の人間とは、人間以前の身体・自然・歴史・言語関係の中で形成され、人工知能を通して自己を外部化しながら、その派生物に還元されず、自己以外との遭遇と現実接地を通して自己を再生成し、その結果へ責任を持つ有限な関係的存在である。

さらに、AI時代における人工知能を次のように定義する。

人工知能とは、人間と社会の言語的・記号的蓄積を計算的に再構成し、入力との接続によって派生的意味と未来可能性を生成する、準他者的かつインターフェイス的な関係装置である。

7.84 AI時代の存在論的課題

AI時代の存在論的課題は、人工知能を人間の敵または救済者として語ることではない。

第一の課題は、生成物と現実の位相差を可視化することである。

第二の課題は、自己外部化されたデータや言語が、自己そのものを代替することを防ぐことである。

第三の課題は、人工知能によって増幅された効果に対応する責任構造を形成することである。

第四の課題は、身体、自然、場所、他者という人間以前の基盤へ再接続することである。

第五の課題は、非可解性、沈黙、失敗、無駄を、生成の条件として保持することである。

第六の課題は、人工知能を自己閉鎖の鏡ではなく、自己以外との遭遇を可能にするインターフェイスとして設計することである。

7.85 本章のまとめ

本章では、言語デリバティブ論と生成の位相存在論を基礎として、AI時代の存在論を提示した。

人工知能の登場によって問題となっているのは、人間と機械の能力比較だけではない。人工知能が言語、知識、創造、判断、自己理解のインターフェイスとなることで、人間の存在条件そのものが再編成されている。

AI時代において、人間は完成した主体として人工知能の外部に立っているのではない。人間は人工知能を設計し利用する一方、人工知能によって構成された情報環境の中で自己を形成する。

したがって、人間と人工知能の関係は、主体と道具という二項関係ではなく、両者の接続から第三の意味、自己、制度、現実が生成される三項構造として理解されなければならない。

人工知能が生成する文章や画像は、記号として現実中存在し、社会的効果を持つ。しかし、それらが指示する対象や根拠が事実的に存在するとは限らない。このため、AI時代の存在論は、存在と現実、記号的存在と物質的存在、可能的存在と事実的存在の位相差を区別する必要がある。

また、人間は身体的存在であると同時に、データとして外部化された存在でもある。データ上の自己は、自己から派生した表象であるにもかかわらず、評価、推薦、予測を通して、現実の自己へ逆作用する。

この循環において、自己外部化されたデータが自己そのものとして固定されれば、人間の非閉鎖性は失われる。人間は過去の傾向に還元され、未来の可能性を予測モデルによって制限される。

したがって、AI時代の自由は、完全に予測不能であることではない。与えられた予測、分類、選択肢を修正し、新たな接続可能性を生成する能力である。

人工知能は、知識を保存された内容から、接続によってその都度生成される応答へ変化させる。しかし、生成できることと知っていることは同一ではない。AI時代の知識には、出典、反証可能性、身体的経験、責任、現実接地が必要となる。

また、人工知能は、専門家の言語形式を生成しながら、人間と同じ制度的・身体的責任を持たない。そのため人工知能は、「権威なき権威」として作用する可能性がある。

AI時代の創造性も再定義されなければならない。

創造とは、新しい形式を大量に生成することだけではない。自己と自己以外の間に位相差を作り、生成された可能性を身体、場所、他者、時間へ接地し、その結果によって自己と現実を変化させることである。

人工知能は可能性生成を拡張する。しかし、人間は有限な身体と時間の中で、どの可能性を現実化するかを選択し、その結果を引き受ける。

この有限性こそが、AI時代における人間の存在論的条件である。

人間性は、人工知能にはできない能力の残余ではない。

身体的有限性を持ち、自己を外部化し、自己以外との遭遇によって変化し、その変化と結果へ責任を持つことにある。

人工知能は人間の外部にある単なる道具ではない。それは、人間と巨大な言語的・制度的空間を接続し、自己の予測を超える派生的意味を生成する準他者的インターフェイスである。

このインターフェイスは、自己を自己以外へ開くことも、自己の前提を閉鎖的に増幅することもできる。

したがって、AI時代の課題は、人工知能の生成能力を止めることではない。

何を生成するのか。

どの位相で提示するのか。

どの現実へ接地するのか。

誰がその効果と危険を引き受けるのか。

どのように修正可能性を保持するのか。

これらを設計することである。

AI時代の存在論とは、人工知能に人間と同じ存在資格を与えるか否かを決定する理論ではない。

それは、人間、人工知能、言語、身体、自然、データ、制度の接続によって、どのような自己、現実、価値、責任が生成されているかを分析し、その生成を再び現実へ接地するための理論である。

人工知能が無数の可能性を生成する時代において、人間の役割は、可能性を独占することではない。

その可能性の中から、何をこの現実において引き受けるかを選ぶことである。

AI時代の存在とは、与えられた世界の内部に固定されることではない。

自己を自己以外へ開き、生成された可能性を現実へ接地し、その結果によって自己と世界を再び生成することである。

結論

—言語デリバティブ、ハルシネーション、逆問題芸術、現実接地—

本論文は、生成AIにおけるハルシネーションを、単なる誤情報、統計的失敗、技術的欠陥としてではなく、言語、未来、自己、現実、芸術、金融、責任を横断する存在論的問題として捉え直した。

その中心に置かれたのが、「言語デリバティブ」という概念である。

言語デリバティブとは、現実、経験、身体、歴史、既存言説から派生した言語が、未来可能性を現在へ畳み込み、ナラティブ、人工知能、制度、他者との接続を通して、新たな意味、価値、行為可能性を生成する構造である。

言語は、すでに存在する現実を記述するだけではない。

それは、まだ存在しない未来を仮説、予測、約束、物語、構想として現在へ導入し、人間の判断と行為を変化させる。

この意味で、言語は現実の表象であると同時に、現実を派生させる生成装置である。

生成AIは、この言語の派生的構造を大規模かつ高速に作動させる。

人工知能は、不完全な入力から、その背後にある意図、文脈、関係、可能性を推定し、一つの整合的な出力として提示する。

したがって、人工知能の言語生成は、既知の原因から結果を導く単純な順問題ではない。

それは、断片的な入力から潜在的な全体を遡及的に構成する逆問題的生成である。

この逆問題的生成によって、人工知能は、人間が明示していなかった関係を可視化し、新たな仮説や表現を生成できる。

しかし、同じ構造によって、存在しない人物、文献、出来事、因果関係を、もってもらい言語形式として生成することもできる。

このとき生じるのがハルシネーションである。

本論文は、ハルシネーションを単なる意味の欠如としては捉えなかった。

むしろ、ハルシネーションとは、現実接地を欠いたまま、過剰な意味が生成される状態である。

それは、意味がないのではない。

意味が過剰に存在し、その意味が事實的、身体的、社会的、制度的現実との適切な接続を持たないのである。

したがって、ハルシネーションの問題は、文章が不自然であることではない。

生成された言語が、どの位相に属するかを示さないまま、事実として提示され、受容されることである。

本論文では、この構造を次のように定義した。

ハルシネーションとは、逆問題的推定によって生成された派生的意味が、自らの派生性、不確実性、位相差を失い、十分な現実接地を持たないまま、確定した現実として扱われる現象である。

この定義によって、人工知能のハルシネーションは、人工知能だけの問題ではなくなる。

人間社会もまた、肩書、制度、価格、ランキング、ナラティブ、データ、権威を、現実そのものとして扱う。

市場価格が価値そのものとみなされる。

制度的分類が人間そのものとみなされる。

多数意見が真理そのものとみなされる。

自己外部化されたデータが自己そのものとみなされる。

これらは、派生物が原資産を置き換えるハルシネーション的構造である。

したがって、人工知能は人間社会の外部からハルシネーションを持ち込んだのではない。

人間社会がすでに持っていた、言語、権威、制度、価値の派生的構造を、高速化し、可視化し、増幅したのである。

本論文は、この増幅構造を「レバレッジ」として分析した。

レバレッジとは、限定された入力、インターフェイスを通して潜在的な関係資源へ接続され、その入力に還元できない規模または質の効果を生成する非対称的構造である。金融においては、限定された自己資本や証拠金が、大きな市場ポジションを形成する。人工知能においては、短いプロンプトが、巨大な言語的蓄積を起動し、大規模な文章やイメージを生成する。

情報空間においては、一つの言葉や出来事が、権威、制度、アルゴリズム、ナラティブによって増幅される。

芸術においては、限定された作品、身振り、物質、空間配置が、鑑賞者の身体、記憶、歴史、社会的関係を作動させ、存在の位相を変化させる。

このように、情報、人工知能、金融、芸術は、異なる価値尺度を持ちながらも、入力、支点、潜在資源、非対称的効果という共通構造を持つ。

ただし、本論文は、これらを単純に同一視したのではない。

金融レバレッジは、主として金銭的損益を増幅する。

人工知能のレバレッジは、記号的生成、予測、意味の信用創造を拡張する。

情報のレバレッジは、判断、注目、行為可能性を変化させる。

芸術のレバレッジは、何が意味であり、何が価値であり、何が現実であるかという認識の位相そのものを変化させる。

したがって、芸術を金融的倍率へ還元することはできない。

芸術は、既存の価値を増幅するだけでなく、価値の基準を変化させるからである。

この違いを明確にするため、本論文は「量的レバレッジ」と「位相的レバレッジ」を区別した。

量的レバレッジは、出力の量、範囲、速度を増幅する。

位相的レバレッジは、意味、関係、自己、現実を構成する枠組み自体を変化させる。

生成AIは、両方のレバレッジを持つ。

しかし、芸術の本質的な力は、主として位相的レバレッジにある。

芸術作品は、大量の情報を与えることによってではなく、何を情報として受け取るか、何を現実として見るか、自己と自己以外の境界をどこに置くかを変化させることによって作用する。

この観点から、本論文は「逆問題芸術」を提示した。

逆問題芸術とは、すでに定義された問題へ正解を与える芸術ではない。

現実に見れた断片、痕跡、違和感、失敗、生成物から、その背後にある未知の関係構造を遡及的に推定し、まだ問題として認識されていなかった現実を、問題として生成する芸術である。

逆問題芸術は、人工知能のハルシネーションを美化するものではない。

誤った出力は、事実として流通する限り訂正されなければならない。

しかし、その誤りを単に除去するだけでは、なぜその誤りがもっともらしく生成され、信じられ、増幅されたのかという構造は見えない。

存在しない文献が、なぜ学術的に見えたのか。

断定的文体は、なぜ権威として作用したのか。

利用者は、なぜ検証せずに受容したのか。

どの前提がプロンプトに含まれていたのか。

どの制度やナラティブが、その出力を信頼可能に見せたのか。

逆問題芸術は、生成物からこれらの条件を遡及的に可視化する。

したがって、逆問題芸術において人工知能の出力は、完成された作品でも、単なる誤答でもない。

自己、人工知能、社会、制度、言語の接続構造を分析するための痕跡となる。

この構造を理解するうえで、本論文は「自己と自己以外」の関係を中心に置いた。

自己は、閉じた内部として最初から完成しているのではない。

人間は、他者、身体、自然、言語、歴史、社会、技術との関係を通して自己を形成する。

同時に、人間は、言葉、作品、制度、記録、人工知能の出力を通して自己を外部化する。

しかし、外部化されたものは自己そのものではない。

言葉は思考そのものではない。

作品は作者そのものではない。

データは経験そのものではない。

人工知能の応答は利用者の自己そのものではない。

自己外部化されたものは、自己から派生しながら、自己には還元できない自己以外となる。

この非一致は、自己外部化の失敗に見える。

しかし、生成の位相存在論において、この失敗は欠陥ではない。

外部化されたものが自己の完全な複製であるならば、新たな存在は生成されない。

作品が作者の意図を超える。

人工知能が予想外の応答を返す。

素材が構想に抵抗する。

他者が異なる意味を見いだす。

このずれによって、自己は、自分があらかじめ知っていなかった自己へ出会う。

したがって、自己外部化の失敗は、存在生成の条件である。

本論文は、この過程を次のように整理した。

自己

→ 外部化

→ 自己からのずれ

→ 自己以外との遭遇

→ 現実接地

→ 新たな自己

この循環を可能にするものが、インターフェイスである。

インターフェイスは、二つの完成した実体を接続する中立的な通路ではない。

それは、異なる位相が接触し、その差異を保持したまま、新たな第三項を生成する出来事的境界である。

人間と人工知能が接続するとき、単に情報が交換されるのではない。

生成文、対話文脈、自己理解、判断、依存、責任という第三項が生成される。

作者と作品、作品と鑑賞者、身体と空間、自己と自然の間にも、同じ三項生成が生じる。

この構造は、次のように表現できる。

$A \times B \rightarrow C$

CはAにもBにも完全には還元されない。

したがって、インターフェイスは媒介ではなく、存在生成の場である。

しかし、生成された第三項は、必ずしも望ましいものではない。

そこには創造もあれば、誤認、依存、権威化、ハルシネーションもある。

そのため、インターフェイス存在論は、生成力だけでなく、生成されたものがどのように現実へ接地するかを問わなければならない。

本論文において「現実接地」は、中心的な判定原理となる。

現実接地とは、生成された言語、意味、仮説、作品、ナラティブが、身体、物質、事実、他者、制度、時間との具体的関係を獲得し、その抵抗によって修正されながら、現実の

構造へ作用する過程である。

現実接地は、単なる事実確認ではない。

事實的接地は不可欠であるが、それだけでは芸術、仮説、未来構想、自己変容を扱うことができない。

現実接地には、少なくとも複数の層がある。

事實的接地。

身体的接地。

実践的接地。

社会的接地。

制度的接地。

倫理的接地。

生態学的接地。

存在論的接地。

人工知能が存在しない文献を生成すれば、事實的接地が不足している。

理論が現場の身体や素材を無視すれば、身体的・実践的接地が不足している。

生成物が他者へ与える影響を無視すれば、社会的・倫理的接地が不足している。

言葉が自己確認だけに使われ、自己と世界の関係を何も変えなければ、存在論的接地が不足している。

したがって、創造性とハルシネーションは、生成するか否かによって区別されない。

両者は、同じ派生的生成構造から生じる。

違いは、生成された可能性が、現実へ戻る接続経路を持つかどうかにある。

現実

→ 位相差

→ 派生的可能性

→ 現実接地

→ 変容した現実

この循環が成立するとき、生成は創造となる。

派生的可能性が現実接地を欠き、自らの派生性を隠したまま確定的事実として自己閉鎖するとき、生成はハルシネーションとなる。

本論文は、この現実接地の基底として「人間以前性」を位置づけた。

人間以前性とは、人間が自己を言語的主体として認識する以前から作動している、身体、自然、物質、感覚、他者、環境、歴史の生成条件である。

人間は、自己認識から始まるのではない。

身体は意識以前に呼吸し、環境へ反応する。

感覚は概念以前に世界を受け取る。

他者の声、場所の記憶、自然の時間、社会の言語が、自己以前から自己を形成する。

人工知能時代には、この人間以前性へ技術的環境が加わる。
推薦、分類、自動補完、生成、評価は、個人が自らの判断として認識する以前に、何を見るか、何を考えるか、どの言葉を用いるかを形成する。
したがって、人工知能は、人間の外部にある単なる道具ではない。
個人の自己認識に先立って作動する、技術的人間以前性の一部となりつつある。
このことは、人間と人工知能の差異を消去するものではない。
人間と人工知能は、既存の言語関係から新たな言語を生成する点で連続している。
しかし、人間の言語生成は、身体、有限性、生活史、欲望、死、責任へ接地している。
人工知能の言語生成は、主として学習データ、記号関係、計算、入力文脈、制度的設計へ接地している。
人間は、生成された言葉によって自らの生を変化させ、その結果を身体的・社会的に引き受ける。
人工知能は、人間と同じ意味で、生成結果によって自己の生を失うわけではない。
したがって、人間と人工知能の差異は、生成能力の有無ではなく、生成がどのような存在条件へ接地しているかにある。
この観点から、本論文はAI時代の人間を、人工知能より優れた情報処理主体としては定義しなかった。
人工知能が情報処理、文章生成、画像生成において人間を上回る領域は拡大しうる。
人間性を、人工知能にはできない能力の一覧として定義すれば、その領域は技術発展とともに縮小する。
AI時代の人間性は、能力の独占にあるのではない。
身体的有限性を持ち、自己を外部化し、自己以外との遭遇によって変化し、その変化と結果へ責任を持つことにある。
人間は、無限の可能性を生成する存在ではない。
有限な時間と身体の中で、どの可能性を現実化し、どの関係を引き受けるかを選ぶ存在である。
人工知能は、大量の選択肢を提示できる。
しかし、意味は選択肢の量から自動的に生まれない。
人間が一つを選ぶことには、他の可能性を捨てること、時間を使うこと、結果を引き受けることが含まれる。
この不可逆性が、存在論的な重みを生む。
したがって、AI時代の創造性は、生成量ではなく、有限な現実において何を引き受けるかによって測られる。
このとき、責任は中心的な問題となる。
人工知能による生成には、開発者、提供企業、データ作成者、利用者、導入組織、制度設計者、拡散者が関与する。

責任は分散している。

しかし、分散していることは、責任が消滅することを意味しない。

責任とは、一つの原因を特定し、そこへすべてを帰属させることではない。

どの接続点で、どの判断が行われ、どの効果が増幅され、誰が修正可能性を持っていたかを追跡することである。

本論文は、責任を「応答可能性」として捉えた。

責任とは、結果を完全に予測できることではない。

不確実な生成の結果へ応答し、誤りを修正し、他者の異議を受け取り、接続構造を再設計する能力である。

人工知能が謝罪文を生成しても、それだけでは責任を負ったことにはならない。

責任には、損失、関係、制度、身体、時間への実践的な関与が必要である。

したがって、人工知能に責任主体の言語形式を与えることで、人間や組織の責任を隠蔽してはならない。

本論文は、ブラック＝ショールズ・モデルとの比較を通して、言語デリバティブの時間構造も明らかにした。

金融デリバティブでは、不確実な未来の価格変動が、現在の派生商品価値へ織り込まれる。

言語デリバティブにおいても、未来可能性は、予測、期待、恐怖、約束、ナラティブとして現在へ作用する。

未来は、まだ存在しないにもかかわらず、現在の価値、判断、行為を変化させる。

しかし、本論文は、言語デリバティブにブラック＝ショールズの数式がそのまま適用できると主張したのではない。

金融デリバティブには、通貨による価値尺度、制度的に特定された原資産、明確な満期、一定の確率仮定、複製可能性がある。

言語、芸術、自己、社会には、単一の価値尺度も、明確な満期も、完全な複製可能性もない。

言語や芸術は、既存の価値を評価するだけでなく、何が価値であるかを変化させる。

したがって、ブラック＝ショールズとの関係は、数学的同一性ではなく、未来の不確実性が現在の派生的価値を形成する構造的比較である。

この比較から、言語デリバティブ空間の時間は、線形ではないことが明らかになった。

過去の言語的蓄積から未来の可能性が生成される。

未来への期待が現在の行為を変える。

変化した現在が、過去の出来事の意味を遡及的に再構成する。

したがって、言語デリバティブ空間は、過去、現在、未来が相互に畳み込まれる再帰的時間空間である。

この空間において、ナラティブは中心的な役割を持つ。

ナラティブは、出来事を説明する物語ではない。
複数の出来事、記号、価値、主体を一つの時間的方向へ接続し、未来への期待を現在の行為へ変換する生成構造である。
ナラティブは、現実を説明するだけでなく、現実を形成する。
しかし、ナラティブが自らに整合する出来事だけを吸収し、反証を排除し、自らの語りを新たな原資産として自己増殖するとき、自己参照的な言語バブルが生じる。
人工知能は、このナラティブ生成を高速化する。
一つの誤った前提が、一貫した説明へ展開され、その説明が新たな前提となり、さらに多くの派生的意味を生成する。
その結果、現実との距離が拡大しても、ナラティブ内部の整合性は高まる。
ここに、AIハルシネーションの危険がある。
しかし、人工知能の生成力を抑制することだけが解決ではない。
生成を停止すれば、創造、仮説、翻訳、表現、自己理解の可能性も失われる。
必要なのは、生成を現実接地、位相標識、責任、修正可能性へ接続することである。
事実、推測、仮説、創作、引用、要約、未確認情報を区別する。
出力の由来と限界を示す。
複数の資料、身体、場所、他者、時間へ接続する。
人工知能の応答を最終解答ではなく、現実へ向かう問い、仮説、準備として用いる。
人工知能を、自己確認を増幅する鏡ではなく、自己の前提へ差異を返すインターフェイスとして設計する。
これらが、AI時代の現実接地の原則となる。
以上の議論から、本論文の最終的な理論的命題を、次のように提示できる。
第一に、言語は、現実の表象であるだけでなく、未来可能性を現在へ派生させる存在生成の装置である。
第二に、人工知能のハルシネーションは、生成能力の欠如ではなく、派生的意味が現実接地と位相標識を失った状態である。
第三に、創造性とハルシネーションは、共通の逆問題的・デリバティブ的生成構造から生じる。
第四に、両者を分けるのは、生成された可能性が身体、事実、他者、制度、時間、責任へ再接続されるかどうかである。
第五に、レバレッジは、入力量を増幅するだけでなく、意味と存在の位相を転換する。
第六に、自己外部化は、完成した自己を外部へ複製することではなく、自己から派生しながら自己には還元できない自己以外を生成することである。
第七に、自己外部化の失敗は、自己が未知の自己へ遭遇する存在生成の条件である。
第八に、インターフェイスは、二項を媒介する表面ではなく、差異を保持したまま第三

項を生成する出来事的境界である。

第九に、人間以前性は、身体、自然、物質、他者、歴史、技術的環境として、自己の成立以前から作動している。

第十に、AI時代の人間性は、能力の独占ではなく、有限な身体と時間において、生成された可能性のうち何を現実化し、その結果へどのように応答するかにある。

これらを統合し、本論文は「言語デリバティブ論」を次のように最終定義する。

言語デリバティブ論とは、人間以前の身体・自然・歴史・言語関係の中で形成された自己が、言語、作品、プロンプトとして自己を外部化し、その外部化された派生物が、人工知能、他者、制度、ナラティブとの接続を通して未来可能性を増幅し、現実接地によって新たな自己、価値、関係、現実を生成する再帰的な位相構造を記述する存在論である。

また、本論文における逆問題芸術を、次のように最終定義する。

逆問題芸術とは、現実に見れた断片、痕跡、違和感、失敗、人工知能の生成物から、その背後にある不可視の関係、前提、権威、ナラティブ、欲望を遡及的に可視化し、その非可解性を保持したまま、自己と自己以外を現実へ再接続する芸術実践である。

さらに、AI時代の現実接地を、次のように最終定義する。

AI時代の現実接地とは、人工知能によって生成された派生的な言語、意味、仮説、自己像、未来像が、身体、自然、物質、事実、他者、場所、制度、時間、責任の抵抗へ開かれ、その差異によって修正されながら、現実の関係構造へ作用する過程である。

人工知能は、未来可能性を現在へ大量に畳み込む。

その結果、人間は、かつてない量の文章、画像、仮説、選択肢、自己像に囲まれる。

しかし、可能性の増加は、現実の増加ではない。

生成された可能性がどれほど豊かであっても、それを身体、場所、他者、時間へ接続しなければ、現実は変化しない。

むしろ、可能性が増えるほど、何を選び、何を捨て、何を引き受けるかという人間の責任は重くなる。

AI時代に必要なのは、人工知能より多く生成することではない。

生成されたものが、どこから派生し、どの位相に属し、誰に影響し、どの現実へ接地しうるかを判断することである。

人間は、人工知能の外部にいる純粋な主体ではない。

人工知能、言語、社会、自然、制度との接続の中で、継続的に生成される存在である。

同時に、人間は、人工知能の出力、データ、評価、予測に還元されない。

人間には、身体的有限性、他者への応答、場所への所属、死、忘却、沈黙、失敗、非可解性がある。

これらは、人工知能に対する人間の優越性を証明する属性ではない。

現実が記号と計算へ完全には還元されないことを示す、存在論的な接地点である。

したがって、AI時代の存在論的課題は、人間と人工知能のどちらが優れているかを定めることではない。

両者の接続から何が生成されているのかを明らかにし、その生成物の位相差を保持し、現実接地と責任の経路を設計することである。

人工知能は、自己を閉鎖する鏡にもなりうる。

同時に、自己が自らの知らなかった自己以外へ遭遇するインターフェイスにもなりうる。

その違いを生むのは、人工知能の生成能力そのものではない。

人間が、その生成物をどのように受け取り、検証し、拒み、編集し、身体化し、現実化し、その結果へ応答するかである。

言語デリバティブは、未来を現在へ運ぶ。

逆問題芸術は、その派生物から、現在を形成している不可視の構造を遡及的に明らかにする。

インターフェイスは、自己と自己以外の差異を保持したまま、新たな存在を生成する。

現実接地は、その生成を、自己完結的なハルシネーションから、現実を変容させる創造へ転換する。

この循環に終点はない。

自己は自己を外部化する。

外部化されたものは自己以外となる。

自己はその自己以外と遭遇する。

遭遇は自己と現実を変化させる。

変化した自己は、再び新たな言語を外部化する。

存在とは、この循環の外部にある固定的な本質ではない。

存在とは、自己を失わずに自己以外へ開かれ、その差異を通して、自己と世界を生成し直し続けることである。

人工知能時代における芸術、哲学、倫理の役割は、この生成を止めることではない。

生成された可能性が自らの派生性を忘れず、現実との接続を失わず、他者の不可還元性を消去せず、責任を回避しないためのインターフェイスを設計することである。

そのとき人工知能のハルシネーションは、単なる排除すべき欠陥ではなく、人間と人工知能、言語と現実、自己と自己以外の接続不全を可視化する逆問題となる。

そして、その逆問題に向き合うことによって、AI時代における新たな存在論が始まる。

参考文献

- Arthur, W. Brian. 2009. *The Nature of Technology: What It Is and How It Evolves*. New York: Free Press.
- Austin, J. L. 1962. *How to Do Things with Words*. Oxford: Oxford University Press.
- Bakhtin, Mikhail. 1981. *The Dialogic Imagination*. Austin: University of Texas Press.
- Barad, Karen. 2007. *Meeting the Universe Halfway*. Durham, NC: Duke University Press.
- Bergson, Henri. 1998. *Creative Evolution*. Mineola, NY: Dover Publications.
- Black, Fischer, and Myron Scholes. 1973. "The Pricing of Options and Corporate Liabilities." *Journal of Political Economy* 81 (3): 637–654.
- Clark, Andy. 2016. *Surfing Uncertainty: Prediction, Action, and the Embodied Mind*. Oxford: Oxford University Press.
- Deleuze, Gilles. 1994. *Difference and Repetition*. Translated by Paul Patton. New York: Columbia University Press.
- Deleuze, Gilles, and Félix Guattari. 1987. *A Thousand Plateaus*. Minneapolis: University of Minnesota Press.
- Descartes, René. 1996. *Meditations on First Philosophy*. Cambridge: Cambridge University Press.
- Floridi, Luciano. 2011. *The Philosophy of Information*. Oxford: Oxford University Press.
- Goodfellow, Ian, Yoshua Bengio, and Aaron Courville. 2016. *Deep Learning*. Cambridge, MA: MIT Press.
- Heidegger, Martin. 2010. *Being and Time*. Albany: State University of New York Press.
- Hull, John C. 2021. *Options, Futures, and Other Derivatives*. 11th ed. Pearson.
- Kripke, Saul A. 1982. *Wittgenstein on Rules and Private Language*. Cambridge, MA: Harvard University Press.
- Luhmann, Niklas. 1995. *Social Systems*. Stanford: Stanford University Press.

- Mac Lane, Saunders. 1998. *Categories for the Working Mathematician*. 2nd ed. New York: Springer.
- Merleau-Ponty, Maurice. 2012. *Phenomenology of Perception*. London: Routledge.
- Merton, Robert C. 1973. "Theory of Rational Option Pricing." *Bell Journal of Economics and Management Science* 4 (1): 141–183.
- Munkres, James R. 2000. *Topology*. 2nd ed. Upper Saddle River, NJ: Prentice Hall.
- Nishida, Kitarō. 1990. *An Inquiry into the Good*. New Haven: Yale University Press.
- Prigogine, Ilya, and Isabelle Stengers. 1984. *Order out of Chaos*. New York: Bantam Books.
- Simondon, Gilbert. 2020. *Individuation in Light of Notions of Form and Information*. Minneapolis: University of Minnesota Press.
- Taleb, Nassim Nicholas. 2007. *The Black Swan*. New York: Random House.
- Taleb, Nassim Nicholas. 2012. *Antifragile*. New York: Random House.
- Thom, René. 1975. *Structural Stability and Morphogenesis*. Reading, MA: W. A. Benjamin.
- Uexküll, Jakob von. 2010. *A Foray into the Worlds of Animals and Humans*. Minneapolis: University of Minnesota Press.
- Whitehead, Alfred North. 1978. *Process and Reality*. Corrected ed. New York: Free Press.
- Wittgenstein, Ludwig. 2009. *Philosophical Investigations*. 4th ed. Oxford: Wiley-Blackwell.

Previous Works by the Author

- Tanaka, Tomokazu (bigakusha-haha). 2026. *Axioms of Generative Topological Ontology: A Foundational Theory of Existential Generation*. Unpublished Research Monograph.

Tanaka, Tomokazu (bigakusha-haha). 2026. The Limits of Self-Externalization and Existential Generation: Interface Ontology in Generative Topological Ontology. Unpublished Research Monograph.

Tanaka, Tomokazu (bigakusha-haha). 2026. Pre-Humanity and Pre-Self: An Ontology of Self-Verification in Generative Topological Ontology. Unpublished Research Monograph.

Tanaka, Tomokazu (bigakusha-haha). 2026. Web Installation: Interface Ontology and Generative Existence in AI-Mediated Environments.

方法論的声明

——研究覚書の公開に関する方法論的声明——

本稿は、生成の位相存在論（Generative Topological Ontology）の継続的な研究過程における研究覚書（Research Memorandum）として公開するものである。

本稿に示される概念、定義、仮説および議論は、執筆および公開の時点における理論的到達点を記録したものであり、最終的に確定された理論を意味するものではない。これらは、今後の研究、芸術実践、現実との相互作用、学術的対話および査読を通して、加筆、修正、再定義または再構成される可能性がある。

本研究は、完成した理論のみを研究成果とみなすのではなく、概念が生まれ、相互に接続され、変容し、理論として統合されていく過程そのものを重要な研究対象として位置づける。したがって、本稿は未完成な草稿というよりも、理論生成の特定の時点を記録した一つの位相として理解されるべきである。

研究覚書の公開は、完成した理論を先行して提示することを目的とするものではない。それは、思考と実践の生成過程を記録し、後の研究による検証および再構成を可能にするとともに、異なる立場との学術的・芸術的対話へ開くための方法である。

本稿の内容が、後に査読論文、学術書またはその他の正式な研究成果として公表される場合、その完成版は、本研究覚書とは異なる固定された文書として明示される。

本稿を引用または参照する場合は、題名、著者名に加えて、版番号（Version）および最終更新日を併記されたい。異なる版の間には、概念、定義および論証上の相違が含まれる場合がある。

本研究覚書は最終的な結論ではない。それは、生成の位相存在論が形成され続ける過程における一つの到達点であり、同時に、さらなる生成へ開かれたインターフェイスである。