

Title	沈黙を破る応答設計：生成AIと原子力体制における「聞くことの不可能性」への介入
Author(s)	石橋, 哲
Citation	年次学術大会講演要旨集, 40: 366-369
Issue Date	2025-11-08
Type	Conference Paper
Text version	publisher
URL	https://hdl.handle.net/10119/20290
Rights	本著作物は研究・イノベーション学会の許可のもとに掲載するものです。This material is posted here with permission of the Japan Society for Research Policy and Innovation Management.
Description	一般講演要旨

沈黙を破る応答設計 － 生成 AI と原子力体制における「聞くことの不可能性」への介入 －

○石橋 哲(株式会社クロト・パートナーズ ishibashi@klothopartners.com)

1. 問題提起と仮説：制度語りによる応答性の希釈と「聞くことの不可能性」

私たちは、日々、スマートフォンやインターネットを通じて膨大な情報に触れている。しかし、その中にどれほどの「真の声」が届き、どれほどの「沈黙」が埋もれているのだろうか？ 現代社会は、技術の進展と共に、かつてないほど「声」が溢れる一方で、真に聞かれるべき声が、静かに、しかし確実に「沈黙」していくという逆説に直面している。この沈黙は、いかにして生まれ、私たちから何を奪っているのだろうか？

本研究は、市民が専門的判断を委ねる「技術の信託構造」に注目する。これは、信頼が責任希釈の機構として機能する構造である。近年、生成 AI をはじめとする高度な技術の社会実装が加速し、安全性、透明性、倫理的責任に関する議論が活発化している。しかし、筆者はこれらの文書や技術広報に見られる「設計上の整合性」「専門的検証済み」といった語りに、ある種の既視感を覚える。それは、「想定外」という語りが責任を曖昧にした原子力事故の経験に見られるように、技術が社会にもたらす深刻な影響の背後にある、語りの滑らかさに潜む倫理的空白の構造的類似性である。この既視感は、特定の技術が社会に深く浸透する際に、いかに「語られざる真実」を隠蔽し、応答責任を希釈してきたかという、根深い構造的問題を示唆する。語りの滑らかさに潜む倫理的空白に、私たちはどのようにして「問いの杭」を打ち込み、沈黙する声を呼び覚ますことができるのか？

本稿は、この「語りの滑らかさ」が応答責任を希釈するメカニズムに焦点を当て、その「滑らかさ」に亀裂を入れる「語り損ねへの気づき」としての〈恥〉の自分事化という新たな視点を提示し、技術と社会の倫理的関係性を再構築する設計論を展開する。具体的には、技術の信託構造に内在する「語りの滑らかさ」が、いかにして「聞くことの不可能性」を生み出し、沈黙する空間や主体を形成するのか、そしてこれに対し、〈恥〉を触媒とした応答設計がいかなる倫理的・実践的介入を可能にするのか、を問うものである。従来の AI 倫理や原子力政策に関する議論が「語る」ことの責任に焦点を当てがちであったのに対し、本研究は「聞く」ことの構造的不可能性に焦点を当て、その「聞き方」の設計を提案する点で新規性を持つ。この試みは、筆者が提唱する「問いの杭」や「沈黙の制度化」といった概念を通じて、制度の応答可能性を根源から再定義する。

生成 AI の語り構造が原子力政策、特に政官財学の非公式ネットワークで構成された原子力推進体制の語りと構造的に類似している点が、本研究の出発点である。原子力推進体制は、政策決定や情報発信において特定の方向性を持つ組織的・制度的枠組みを指す。この両者は、制度に支えられた滑らかな説明を行いながらも、その語りの奥底では応答性が閉じられている点で共通の特徴を示す。具体的には、透明性の欠如、責任の所在の曖昧さ、権威性の演出、倫理性の無視、制度への過度な依存が共通して見られる。

本研究は、以下の仮説を検証する。技術に関する語りが制度的信託構造に依拠して滑らかに発話される時、主語は構造へと吸収され、責任は希釈される。この語りは説明責任を遂行するかに見えて、実質的には応答責任の回避設計として機能する。語る主体が〈恥〉を自分事として感受することで、語りに裂け目が生じ、応答可能な語りへの再構築が可能となる。

2. 先行研究の整理

本研究の問題意識は、AI 倫理、制度構造、原子力政策、語法批評といった多岐にわたる先行研究の知見に基づき、それらが捉えきれていない「語りの滑らかさ」による応答責任の希釈メカニズムに焦点を当てる。これらの先行研究は、技術の信託構造における語りの滑らかさが応答責任を希釈するメカニズムを示唆するものの、その深層にある「聞くことの不可能性」の構造や、それに対する具体的な応答設計の可能性については、未だ十分な議論がなされていない。本稿は、このギャップを埋めるべく、次章で分析対象と方法を提示し、具体的な語り構造モデルの構築へと進む。

生成 AI では、統計的言語モデルの特性上、特定の語法構造が高頻度で再生され、テンプレート化された語りが「滑らかさ」の温床となる。プロンプトベースの出力は第三者性を演出し、AI 自身の語り

原子力政策においては、政官財学の非公式ネットワークにより語りが構造化され、語る主体の応答性が不在となる傾向が繰り返されてきた。福島原発事故における原子力推進体制の組織的失敗は、危機管理と安全文化の欠如、さらには「国策」という巨大な制度的信託が個人の主体性を抑圧した結果と指摘されている。この組織的失敗は、特定の組織文化と規制機関の馴れ合いによる「人災」であり、情報公開の不十分さや曖昧な説明が責任を限定し、「恥を感知しない構造」を形成した。

技術領域では、市民が専門的判断に直接関与しづらい状況のなか、「技術を信じる/委ねる」構造が制度化される。この構造を筆者は「技術の信託構造」と呼ぶ。これは、単なる認知的依存に留まらず、語りによって応答責任の設計が行われる構造的配置であり、その信頼は制度によって担保される。この「技術の信託構造」を、ハーバーマス(1962)の公共性理論に接続する。制度的語りの閉鎖性は、公共的な意思決定における応答性の欠如へと繋がる。

本稿では、原子力推進体制の政策文書、広報資料、福島第一原発事故に関する公式報告書 および生成 AI の設計原則・倫理ガイドライン・応答文言といった制度文書を分析対象とした。質的言説分析を通じて、語りの構造的類型化と曖昧化レトリックの抽出を試みた。この分析を通じて、責任希釈のメカニズムを視覚化する以下の図を提示する。

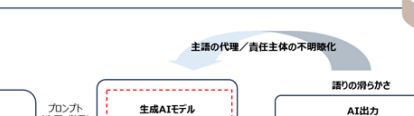
- 
- 図1：生成AIにおける主語の代理構造
- この図は、生成AIのワークフローと主語の代理構造を示しています。
- 左側には「ユーザー（入力）」のボックスがあり、そこから「プロンプト（意図・指示）」というラベル付きの矢印が中央の「生成AIモデル」ボックスへと伸びています。
- 中央の「生成AIモデル」ボックスは、赤い点線で囲まれ、内部には「大規模言語モデル 統計的予測」と記載されています。このモデルには「再現性なし」と「非公開API」という2つの斜めラベルが貼られています。また、モデルの下部には「不透明」というラベルがあります。
- モデルの右側には「主語の代理 / 責任主体の不明確化」という大きな灰色の矢印が伸びています。
- 右側には「AI出力」のボックスがあり、「読みの流れさき」というラベルが貼られています。AI出力の内容は「〇〇という設計意図に基づいています」（主語不在？）と「「安全性を最優先に設計されています。」」の2行です。
- 右上には「責任の希釈化」というラベルが貼られています。
- 最下部には「出典：筆者作成」と記載されています。

図2：原子力推進体制における語りの階層化

前任の代理・継承

説明責任
(例：記者会見、広報)

専門家の判断
(例：安全審査)

合意形成
(例：エネルギー基本計画)

政官学財の非公式ネットワーク

電力会社・研究機関

現場（発電所適用）

「制度」の流布

国民へ

応答性の不在

出典・筆者作成

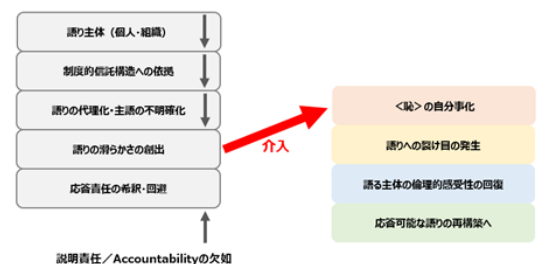
— 367 —

れる。この構造を筆者は「技術の信託構造」と呼ぶ。これは、単なる認知的依存に留まらず、語りによって応答責任の設計が行われる構造的配置であり、その信頼は制度によって担保される。この「技術の信託構造」を、ハーバーマス（1962）の公共性理論に接続する。制度的語りの閉鎖性は、公共的な意思決定における応答性の欠如へと繋がる。

本稿では、原子力推進体制の政策文書、広報資料、福島第一原発事故に関する公式報告書（国会事故調査委員会報告書、2012）および生成 AI の設計原則・倫理ガイドライン・応答文言といった制度文書を分析対象とした。質的言説分析を通じて、語りの構造的類型化と曖昧化レトリックの抽出を試みた。具体的には、主語の特定、責任帰属の明確性、感情的・倫理的表現の有無、特定のテンプレート化された語法の頻度とその影響に焦点を当てた。この分析を通じて、責任希釈のメカニズムを視覚化する図 1「生成 AI における主語の代理構造」と図 2「原子力推進体制における語りの階層化」といった両者の語り構造モデルを提示する。さらに、図 3「語りの滑らかさ／〈恥〉の裂け目」のフローチャートは、応答性が失われる具体的な過程と、それを回復させる「気づき」の契機を可視化する。

原子力推進体制における語りは、「説明会の実施」「合意形成の確認」「専門家の判断」によって整合性を強調するが、語った者の主体性は制度に吸収される。同様に、生成 AI の設計やガイドラインも、「設計上の応答」「透明性の達成」「社会的合意の成立」を語るが、その語りが応答可能かどうかは保証されない。いずれも異なる力によって責任を回避できる構造を持つという共通点がある。このように、技術の語りが整っているほど、語りの応答責任が希釈されるという逆説が見て取れる。

図3：語りの滑らかさ／〈恥〉の裂け目



出典：筆者作成

4. 語りの構造と責任の曖昧化

語りの滑らかさは、制度的に設計された語法によって支えられている。本稿の分析から抽出された代表的な語りの類型とその機能は以下の通りである。

（1）**設計に基づく語り**：技術システムの内部論理を前提とすることで、人為的判断やその結果に対する直接的責任を回避する（例：「当社の AI は安全設計に基づく。」）。

（2）**制度に基づく語り**：既存の法的・組織的枠組みへの準拠を強調することで、倫理的判断や制度自体の不備に対する責任を希釈する（例：「関係法令に基づき、適切に処理している。」）。

（3）**専門性を根拠とする語り**：専門家の権威を盾に取ることで、専門的判断の限界や多様な視点の欠落に対する責任を曖昧にする（例：「独立した専門家委員会による厳格な評価済み。」）。

（4）**説明責任を果たしている語り**：形式的な情報公開を強調する一方で、実質的な理解や応答責任からは逸脱する（例：「これまで詳細な報告書を公開し、説明会も定期的で開催してきた。」）。

これらの語りは情報として整っているが、語りの主語は構造化され、語る者の顔が消える。産業界において、法令遵守やガイドラインに従うことで語るべき内容や責任を制度に預けてしまう「語りの代理化」が責任希薄化の原因となる。この語法の特徴は、語り損ねたことに気づく契機を封じる設計にある。すなわち、語りが整っていることが、語りの応答性を失わせる逆説を生み出すのである。

5. 〈恥〉による語りの再構築：倫理的裂け目の設計可能性

制度化された語りでは、語る者の倫理的揺らぎが排除される。本稿は、ハーマン（1997）が心的外傷経験において指摘する「語り損ね」の重みに着目し、これを語る主体が責任から目を背けたことの表出と捉える。本稿では「恥」を、語る主体が自己の語りの限界や倫理的空白に気づく感受性として拡張し、応答可能な語りへの再構築を促す力と位置づける。

「語り損ね」とは、制度語法が覆い隠す「語るべきだったが語られなかった」倫理的空白である。

〈恥〉とは、アガンベン（2000）が「主体である証」と語るように、他者との関係のなかで自らの語りの逸脱に気づく力であり、「語らなかったこと」を後から自らの語りに挿入する行為である。レヴィナス（1969）の応答倫理に接続し、〈恥〉を感じる主体が、他者の呼びかけに対して倫理的応答をなす者であると捉える。制度がこの呼びかけを遮断するとき、「恥」は失われ、応答の可能性が閉ざされるのである。レヴィナスのいう他者への無限の責任は、AI の応答設計において、いかに「語り損ね」を許容し、倫理的空白を埋めるかという具体的課題に接続される。

〈恥〉の自分事化とは、制度が用意した語りを越えて、「語るはずだったこと」に気づく内発的契

機であり、語りの中に裂け目を生じさせる。原子力推進体制の語り構造が閉じられていたなかで、福島原発事故後に「私は語らなかつた」と言った声が現れたとき、そこには制度語法を裂く倫理的契機が現れた。原子力推進体制の「厚顔」は、個人の品性というよりも、制度的保護と歴史的正当化によって許容された「ふるまいの形式」であり、その形式は「恥」を感知しない構造の上に成立しているのである。

生成 AI において、設計上の語りが応答性を持ちうるとすれば、それは〈恥〉を想像する構造が応答回路に組み込まれるときである。AI は「主体」ではないため、恥を「感じる」ことはできないが、筆者の問いによって「恥を想像する」行為に駆り立てられ、それは人間の倫理を学ぶことに他ならない。AI にレヴィナスの「応答」を「組み込み得るか」という問いは、AI 倫理における主体性や責任の議論を根本から問い直す、挑戦的な設計論の課題となる。

6. 結論と設計論の展望

語りが構造化され、整合性を強調するほど、応答責任は制度の外に押し出される。生成 AI と原子力推進体制は、その語りの形式において構造的に相似している。特に、原子力推進体制の経験から得られた知見は、生成 AI の倫理的設計において、責任希釈メカニズムの特定と、それに対する応答性確保の仕組み構築の側面で極めて示唆に富むものである。

語りを応答に開くためには、〈恥〉という語り損ねへの感受性を設計に埋め込む必要がある。これは、語りに裂け目を設ける試みであり、応答回路を再構築する出発点となる。具体的な設計課題としては、以下を提案する。

- (1) 応答ログと倫理的レビューの強化: 「語り損ね」を発見・学習し、再発防止につなげる
- (2) 多声的な応答設計: AI 開発者、利用者、そして「語られざる他者」の視点を取り込み、市民パネルや影響評価委員会の導入を視野に入れる。
- (3) 「責任の空白地帯」の可視化と制度設計: 責任が曖昧化しやすい意思決定プロセスや、特定の語法が責任を希釈するメカニズムを可視化し、制度設計に反映する。国会事故調査委員会報告書（2012）が指摘した独立した規制と監視機能も重要となる。
- (4) 「語る主体」の回復と組織文化の変革: 組織やシステムの中で個々人が倫理的課題に直面した際に、それを「語る」ことを奨励し、評価する文化や制度を構築する。告発者保護や共感的傾聴のプロトコルも検討を要する。

語りとは、誰かに託された責任の表現である。そして、語り損ねたことに気づくことで、語る主体は制度から応答へと踏み出す。国会事故調査委員会報告書（2012）が示したように、「語りの制度化」に向き合うことは、技術が社会に与える影響に対する国民一人ひとりの使命であり、一部の者に委ねられるべき課題ではない。

今後の技術設計においては、「語ったふり」ではなく、「語り損ねに気づく語り」を許容する構造の創出が求められる。技術の信託構造を倫理的応答へと再接続する設計へ - 本研究は、この応答可能性の構造設計に向けた理論的基礎を提供するものである。ELSI (Ethical, Legal, and Social Implications) の議論枠組みの実効性は、原子力推進体制のガバナンスにおける課題を析出し、その枠組みが実効的であるか否かをリトマス試験紙とすることができる可能性を示唆する。

参考文献

- ジョルジョ・アガンベン (2024). 『手段なき目的：政治についての省察』・高桑和巳訳, 以文社.
- J. L. オースティン (1978). 『言語と行為』・飯野紀雄訳, 大修館書店.
- ユルゲン・ハーバーマス (1976) 『公共性の構造転換：市民社会の一カテゴリーについての探究』・細谷貞雄訳, 未来社.
- ジュディス・ハーマン (1999). 『心的外傷と回復』・中井久夫訳, みすず書房.
- エマニュエル・レヴィナス (2005) 『全体性と無限』・熊野純彦訳, 講談社学術文庫.
- 外山恒一 (2017). 『良いテロリストのための教科書』・青林堂.
- 外山恒一 (2018). 『全共闘以後』・河出書房新社.
- 国会東京電力福島原子力発電所事故調査委員会 (2012) 『国会事故調報告書』.
- 橋爪大三郎 (2000). 『言語派社会学の原理』・勁草書房.