

AI 価値観分析型メール返信アプリ

メール対話から価値仮説・意思決定構造を抽出し、統制された返信戦略へ変換する研究

Value-Oriented AI Mail Reply System

Research Overview Version 1.0

研究・開発	FANDDF
基礎理論資料初出	2026 年 4 月 7 日
研究着手日	2026 年 7 月 28 日
基礎研究正本	価値観起点・顧客意思決定構造推論エンジン 統合正本 V1.0
メール返信アプリ個別仕様初出	2026 年 7 月 30 日
統合正本 V0.2.7	2026 年 8 月 1 日
研究概要初版	2026 年 8 月 18 日

要旨

本研究は、生成 AI によるメール返信を単なる文章生成として扱うのではなく、メールの往復から相手の要求、質問、選択、拒否、保留等を捉え、その背景にある目標、評価、心理反応、ニーズ、行動阻害要因および価値判断を整理し、自社の承認済み情報と照合して返信戦略へ変換する方法を検討するものである。

本研究では、価値観を人物類型として確定するのではなく、原文上の根拠、反証、代替可能性、適用範囲および未確認事項を伴う「価値仮説」として扱う。また、相手側の分析と自社の提供条件を分離し、分析結果から直接送信文を作るのではなく、返信戦略、生成、監査、人間確認を段階的に分ける。これにより、表層的に自然な返信文ではなく、相手の意思決定に必要な情報を整理しながら、自社の判断基準を保持した業務利用を目指す。

キーワード：生成 AI、価値仮説、意思決定構造、メール返信、B2B コミュニケーション、AI 監査、人間承認

1・研究背景・目的

生成 AI は、メールの要約や文章調整、返信案の作成を支援できる。一方、実際の商談や問い合わせでは、同じ「検討します」という表現でも、予算、権限、情報不足、導入負担、時期、リスクなど、行動を左右する条件は異なり得る。表層文だけを基礎に返信すると、相手が実際に必要としている判断材料と返信内容が一致しない可能性がある。

そこで本研究は、メール返信を「入力文から出力文を生成する処理」としてではなく、対話の経過から意思決定に関係する情報を整理し、返信設計へ変換する問題として捉える。研究目的は、相手の内面を断定することではなく、メール上で確認できる証拠から、何を重視し、何を得たい・守りたい・避けたい可能性があるのかを仮説として整理し、実務上の返信判断に利用できる形へ構造化することである。

本研究は、メール要約・返信生成、対話分析、顧客理解、Human-in-the-loop 型 AI 等と隣接するが、本 Research Overview は既存研究に対する厳密な新規性評価を目的とするものではない。

2・研究対象と「価値仮説」

本研究における価値仮説とは、メール上の発言・選択・拒否・保留等の証拠から、相手が何を望ましい状態とし、何を得たい・守りたい・避けたいと判断している可能性があるかを、根拠、反証、代替可能性、適用範囲および未確認状態を伴って表現した仮説である。

したがって、価値仮説は性格診断ではなく、単一の発言から恒常的な人物像を確定するものでもない。複数の価値が同時に存在すること、状況によって優先順位が変化すること、情報不足により判断できないことを前提とする。

分析では、確認可能な事実、目標・意図、評価、心理反応、ニーズ、行動阻害要因、役割上の制約などを同一概念として扱わない。例えば「導入したいが予算がない」という状態と、「提案そのものに価値を感じていない」という状態を区別する。また B2B では、担当者個人、役割、部署、組織、案件等の範囲を分け、一人の発言を組織全体の価値へ安易に一般化しない。

3・研究方法

本研究では、最新の一通だけでなく、可能な範囲でメールの往復全体を分析対象とする。相手の発言に対して自社が何を回答し、その後どのような反応が生じたかを連続した対話として扱い、事実と推定を分離しながら意思決定に係る情報を整理する。

価値仮説は単語一致で確定せず、原文上の証拠を起点として、複数の説明可能性や反対方向の証拠を残す。分析結果が不十分な場合には無理に確定せず、未確認状態を保持する。

また、相手側の分析に自社の商品・サービス都合を先に混入させない。相手側の分析を行った後に、承認済みの自社情報、提供可能範囲、対応不可条件、契約・価格等の条件を照合する。この分離により、「相手にとって理解しやすい返信を作成すること」と「自社として何を回答・提案・約束できるか」を別の判断として扱う。

4・返信生成と統制の基本原則

本研究では、分析結果から直接返信文を確定しない。意思決定構造の整理、返信戦略の設計、文章生成、監査、人間による確認を分けることで、文章として自然であることと、業務上利用してよいことを同一視しない。

返信戦略では、相手の明示的な質問、未回答事項、追加で必要な判断材料、自社が回答できる範囲、確認が必要な事項などを整理する。生成された返信案についても、事実と仮説の混同、過去回答との矛盾、過剰な確約、自社条件との不整合等を確認し、最終利用は人の承認を前提とする。

この統制は、AI へ最終判断を代替させることを目的とするものではない。AI を、対話情報の整理、比較、返信設計の前工程を支援する手段として位置づける。

5・現在の到達点

本研究の基礎理論資料は 2026 年 4 月 7 日に確認でき、2026 年 7 月 28 日付の「価値観起点・顧客意思決定構造推論エンジン 統合正本 V1.0」において、メール対話、価値仮説、意思決定構造、返信戦略、監査および人間確認までを一体の研究仕様として明文化した。本 Research Overview では、この日を研究着手日として扱う。

2026 年 7 月 30 日にメール返信アプリとしての個別仕様が成立し、2026 年 8 月 1 日の V0.2.7 では、単独運用を想定した統合実装仕様と機械可読成果物の整合性を静的に確認する段階まで進展した。V0.2.7 に対する旧版照合・静的監査の総合判定は PASS である。

一方、本番相当のメール接続・送信や公開 API を含む動の実証は未実施である。したがって、2026 年 8 月 18 日時点の位置づけは、実装仕様・静的監査体系を構築した段階／動の実証前であり、本番稼働済みのサービスではない。

6・研究上の限界と今後の実証

本研究は、相手の内面を直接観測するものではない。価値仮説は取得可能なメールや関連証拠から構成されるため、観測されていない価値を低いと判断することはできない。また、B2B では個人の価値判断と、役割・部署・組織上の制約を完全に分離できない場合がある。返信戦略の妥当性が高くても、それだけで契約成立や売上向上を保証するものではない。

今後の動の実証では、価値仮説が原文根拠と整合しているか、誤った断定を抑制できるか、人が返信案をどの程度修正するか、返信作成負担をどの程度軽減できるか、監査が不適切な出力を見逃さないか等を評価する。成約率のみを直接の評価指標とはせず、返信設計そのものの妥当性と、人の判断過程を維持した業務負担削減の両面から検証する。

7・結論

本研究が目指すのは、相手の価値観を AI が「見抜く」ことでも、営業メールを自動的に確定・送信することでもない。メール対話から、相手が何を判断材料とし、何を得たい・守りたい・避けたい可能性があり、何が意思決定を促進または阻害しているのかを、根拠付きの仮説として整理することにある。

その構造を自社の承認済み情報および提供可能条件と照合し、返信戦略へ変換し、生成結果を監査したうえで人が利用可否を判断する。すなわち本研究は、文章生成そのものではなく、意思決定構造を媒介として生成 AI を業務へ組み込む方法を研究するものである。

研究開発ステータス

Research Overview	Version 1.0
研究着手日	2026 年 7 月 28 日
基礎理論資料初出	2026 年 4 月 7 日
個別アプリ仕様初出	2026 年 7 月 30 日
V0.2.7 時点	実装仕様・静的監査完了／動の実証前

更新履歴

Version 1.0 | 2026 年 8 月 18 日

基礎理論資料初出、研究着手日、個別アプリ仕様初出、統合正本 V0.2.7 までの研究・開発経過を明示し、公開研究概要作成基準に基づいて、研究背景、価値仮説、研究方法、生成・統制原則、現在の到達点、限界および今後の実証を公開可能な範囲に再整理した。