

WHITE PAPER

非エンジニアのためのClaude実践ガイド

AIの答えを信じていい？ 出力 チェック5つの型

生成AIの誤りを見抜く — 非エンジニアのための出力チェックと安全な使い方
(2026年7月版)

発行：2026年7月 / 対象：非エンジニアの実務担当・管理職・経営者

Claude Works (クロードワークス)

非エンジニアのためのClaude実践メディア claudelab.jp

1. なぜ「AIの答えが不安」なのか（30秒で要点）

生成AIを業務で使うとき、多くの人が二つの極端に振れます。ひとつは、出てきた答えをそのまま信じて貼り付けてしまう。もうひとつは、間違うのが怖くて結局使わない。責任感のある人ほど後者になりがちで、これが現場でAIが使われない大きな理由の一つです。

正解は、その中間にあります。AIは「**下書きを高速で作る優秀な新人**」であって、**内容を保証する専門家ではない**。この距離感さえ持てれば、AIは安全に、しかも大幅な時短で使えます。

この資料は、非エンジニアの現場で「AIの答えをどこまで信じ、どこを自分で確かめるか」を、5つの型と業務別チェックリストに落とし込んだものです。特別な知識は要りません。

2. AIが間違える仕組みを、1分で

生成AIは、次に来る言葉を「確率的にもっともらしい順」で選んでいます。事実を調べて答えているのではなく、それらしい文章を作るのが得意なだけです。だから、実在しない出典や、それらしいが誤った数字を、堂々と自信ありげに書くことがあります。これを「ハルシネーション（もっともらしい作り話）」と呼びます。

大事なのは、AIには**得意なことと苦手なことの地図**を持つことです。

AIの得意・苦手の地図

得意 文章の要約・言い換え・下書き・アイデア出し・翻訳

+

やや苦手 最新情報・固有名詞・正確な計算

+

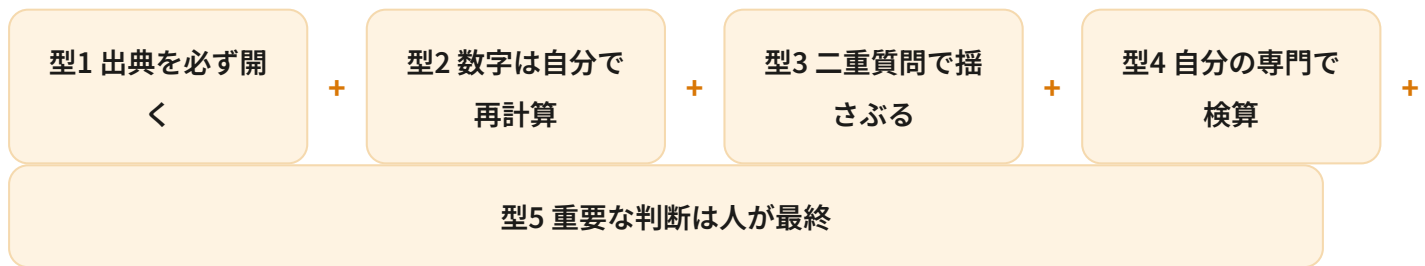
苦手 出典の断定・専門判断・個人や自社だけの事実

要約や下書きは得意なので、ほぼそのまま使えます。一方で、数字・固有名詞・出典・専門判断は、AIが一番間違いやすい領域です。ここだけを人が確かめれば、リスクの大半は消えます。

3. 出力チェック「5つの型」

全部を疑う必要はありません。次の5つの型を、内容に応じて使い分けてください。

出力チェック 5つの型



型1：出典を必ず開く。 AIが「～という調査があります」「～法の第○条」と書いたら、そのリンクや条文を必ず自分で開いて確認します。存在しない出典を作ることがあるためです。開けない・見つからない出典は、無いものとして扱ってください。

型2：数字は自分で再計算。 合計・割合・単価などの計算は、電卓や表計算で1回だけ検算します。AIは計算を間違えることがあり、しかも自信ありげに書きます。

型3：二重質問で揺さぶる。 同じことを角度を変えてもう一度聞く、または「本当に正しいか、根拠とともに批判的に確認して」と追い質問する。答えがブレる箇所は、AIも自信がない箇所です。

型4：自分の専門で検算。 あなたが一番詳しい領域（自社の業務・数字・お客様のこと）については、AIより人が正しい。違和感を覚えたら、あなたの感覚を優先してください。

型5：重要な判断は人が最終。 契約・お金・人の評価・お客様への公開物は、AIの答えを下書きとして使い、最終確認と決定は必ず人が行います。

4. リスクで使い分ける（全部を疑わない）

毎回5つの型を全部使うと疲れて続きません。**間違えたときの影響の大きさ**で、チェックの深さを変えます。

リスク別のチェックの深さ

低リスク（社内・下書き）

会議メモの要約

メールのたたき台

アイデア出し

→ ほぼそのまま使う（型2だけ軽く）



高リスク（社外・お金・法務）

請求・見積の金額

契約・規約の文言

お客様への公開文

→ 型1～5を必ず全部

社内で自分だけが読む下書きは、多少の誤りがあっても影響は小さいので、スピード優先で使います。逆に、お金・契約・社外公開に関わるものは、5つの型をすべて通す。この使い分けができると、安全と速さを両立できます。

5. 業務別チェックリスト

自分の職種で「どこを必ず確かめるか」を1枚にしました。

職種・場面	AIに任せてよい部分	必ず人が確かめる部分
経理（請求・経費）	文面・体裁・分類のたたき台	金額・税率・振込先・取引先名（型2・型5）
法務・契約	条文の要約・論点の洗い出し	条文番号・法解釈・締結判断（型1・型5）
広報・SNS	投稿文・見出しの複数案	事実・固有名詞・数字・公開可否（型1・型5）
営業（提案・見積）	構成・言い回し・下書き	価格・納期・約束事項（型2・型5）
人事	求人票・評価コメントの素案	個人の評価の妥当性・法令（型4・型5）

6. これはAIに丸投げしない（禁止リスト）

次の領域は、AIの答えを最終判断にしないでください。下書きに使うのは構いませんが、決めるのは人です。

- ・契約の締結可否、条項の最終確定
- ・医療・健康・安全に関わる判断
- ・与信・融資・投資の可否
- ・法律・税務の最終的な解釈と申告
- ・個人（社員・候補者）の評価や可否の決定
- ・お客様や社外に出す、事実を含む公開物の最終稿

7. チームで「安全に使う」最小ルール

一人ひとりの心がけに任せると、ばらつきます。次の3つだけ、チームの共通ルールにしてください。

☐ 出典・数字・固有名詞は、貼り付ける前に人が1回確認する

- 社外・お金・契約に関わるものは、必ず別の人がもう一度見る（ダブルチェック）
- 「AIが作った下書きです」と分かる状態で共有し、最終責任は担当者が持つ

このルールは、AIを止めるためではなく、安心して速く使うためのものです。禁止を増やすより、「ここだけ確かめれば大丈夫」を明確にするほうが、現場は前に進みます。

5つの型は身についた。次は「自社版のルール」

ここまでの5つの型・リスク別の使い分け・業務別チェックリストは、どの会社でも使える共通の土台です。ただし、御社のどの業務にどこまでAIを任せてよいか、社内ルールをどう1枚にまとめるかは、業種や体制によって最適解が変わります。

私たちは、非エンジニアの現場が「安全に、かつ止まらずに」AIを使えるよう、社内ルールづくりと運用の定着まで伴走しています。下のリンクから、無料の30分相談でお気軽にご相談ください。御社の業務に合わせた「確認すべき所」の線引きを、その場で一緒に整理します。

無料30分オンライン相談を受け付けています

「自社の場合どう進めればいいのか」を、御社の状況に合わせて具体的にご提案します。売り込みはいたしません。

 ご予約：<https://app.spirinc.com/patterns/availability-sharing/evuvVnwxGC-HC8t6imBtr/confirm>

 support@lexor.jp /  <https://claudelab.jp>