

判断を、 AIが読める形にする

BtoBマーケティングのAI活用事例から読む
判断基準・業務工程・人間の関わり方

判断基準

業務の実行

人間の確認

AIへ、何を渡しているのか

本レポートは、BtoBマーケティングのAI活用事例から、AIに渡している判断基準・業務工程・人間の関わり方を読み解きます。企業事例の説明と、Prollectの提案を分けて掲載しています。

この号が扱う3つの問い

AIに渡す判断基準は、どのように文書化されているか。

AIは業務のどの工程を担っているか。

人間は、何を・いつ確認しているか。

この資料でわかる範囲

主に英語圏のソフトウェア企業による、限られた公開事例を扱います。企業の自己開示は、その運用の説明として読みます。AI導入の因果効果や、すべての組織に共通する成功法則を示すものではありません。

Web版の全文・全出典：

prollect.jp/reports/2026-ai-native-marketing-org/

持ち帰りたい、3つの視点

判断基準を、繰り返し参照できる形にする

Ahrefsは編集工程、CircleCIはブランドやレビュー基準、Finは指標やデータの定義を外部化しています。これは確認した事例からの整理であり、効果を保証する条件ではありません。

人間の関与は、業務ごとに設計する

公開記事の個別確認、スキルの事前テスト、失注理由分類の例外確認など、関与の対象は異なります。配信前の承認を外すことと、内容を読まなくなることも別です。

一業務から試し、修正理由を残す

Prollectからの提案は、対象業務・判断基準・作業データ・確認方法を決めて試すことです。修正理由を記録し、再発性と重要性に応じて基準を更新します。

出典：[S1] Ahrefs — コンテンツ制作パイプライン / [S3] CircleCI — マーケティングチーム（5名） / [S9] Attio — RevOps / [S10] Fin (fin.ai) — Ground Truth と週次レポート

原典URL・資料名は巻末に記載。

利用率・組織内浸透・方針・成果を、分けて見る

日本の企業のAI活用については4つの調査があります。それぞれ別のことを測っており、対象も設問も時期も違うため、足し合わせて1つの結論にはできません。海外の事例を読む前に、自社を測る物差しを分けておきます。

調査	対象・時期・測っているもの	確認できた数値
帝国データバンク [S14a]	全国23,349社に調査・10,312社回答 (2026年3月17～31日)。全業種・企業 単位の 活用率	生成AIを活用している企業は 34.5% 。活用企業のうち86.7%が「効果あり」。課題として情報の正確性を50.4%が挙げる
才流 [S14b]	BtoBマーケティング従事者 n=600 (2025年7～8月)。組織内の 浸透度	「ときどき活用」41.7%+「日常的に活用」28.8%。ただし「メンバーの80%以上が日常的に活用している組織」は 12.2%
総務省『情報通信白書』 [S14c]	2024年度調査。 方針の有無	AI活用の方針を「明確に定めていない」日本 31.8% (米国9.1%／ドイツ12.9%／中国4.2%)
PwC Japan [S14d]	日本 n=945(売上500億円以上・課長職 以上、2025年2月)。 大企業の実感	活用は56%。「期待を上回る効果」を出せている企業の割合は、米国・英国の およそ4分の1 の水準(原典は比率のみで、日本側の実数は非公表)

「AIを使っているのに、組織の業務プロセスが何も変わっていない」 — 太田翔葵 (シャコウ CEO) / 2026年4月24日 [S15]
1名の見解であり、全国的な傾向を示すものではありません。

Prolectからの提案

自社の状況を見るときは、この4つを分けて確認します。「AIを使っているか」だけでは足りません。個人の利用状況に加えて、**業務手順と判断基準がチーム内で共有されているか、方針が決まっているか、成果をどう見ているか**——を別々に確認します。

出典：[S14a] 帝国データバンク / [S14b] 才流 / [S14c] 総務省 / [S14d] PwC Japan / [S15] 太田翔葵
原典URL・資料名は巻末に記載。

Ahrefs

SEOツールを提供する企業 / コンテンツチーム / Ryan Law (Director of Content Marketing)

編集の問いを、AIの工程へ

SEOツールを提供する Ahrefs のコンテンツチームの事例です。Director of Content Marketing の Ryan Law が、情報系 SEO記事のキーワード選定からほぼ完成原稿までの生成と、既存記事の更新を Claude Code で回しています。

AIに渡したもの	既存の編集プロセスに対応する約23個のスキルファイルと、それらを順に発火させる blog-pipeline。各ステップが自前の出力ファイルを持つ設計です。
AIがする工程	キーワード調査 → トピックギャップ分析 → リサーチ集約 → 構造的アウトライン → ドラフト → HTMLプレビュー → 最終フォーマット。画像挿入は現状手動です。
人間が確認すること	4つのゲート。企画（既存の記事に対して有用な付け加えがあるか）／構成（各節が読者を助け、記事の約束を支えているか）／証拠（重要な主張を裏付けられるか）／原稿（獲得していない確信・埋め草・実例が混ざっていないか）。編集者が全記事の全語を読み、「no」と言う権限と意思を持ちます。

確認できたこと（自己申告） - ドラフト生成は6〜12分 - この新プロセスでの公開実績は、約15本の新規と約30本の更新（原文も概数） - 画像挿入は "not a solved problem, yet" と明記 - 同社の自社調査で、AI検出スコアと検索順位の相関は 0.011（実質ゼロ） [S2]	確認できないこと 6〜12分の測定方法・比較対象・母数 - 約15本・約30本がどの期間のものか
--	---

"My thesis is that small amounts of expert direction provided at the start of the content creation process are vastly more effective than lots of human editing at the end."

工程の開始時に与える少量の専門的方向づけは、最後に大量の人間の編集を加えるよりはるかに効果的だ——本人が「仮説」と明示している主張で、検証結果ではありません。あわせて「数万本規模に拡大することもできるが、そうはしない」とも書いています。

同社が「AIスロップ」と呼ぶもの [S2] : "AI slop is content published without enough human understanding, judgement, evidence, or original contribution to justify the reader's attention."（読者の注意に値するだけの、十分な人間の理解・判断・証拠・独自の貢献を欠いたまま公開されたコンテンツ）。同社の Si Quan Ong は「AI利用が増えるにつれて、検索パフォーマンスは低下する傾向にあった。企業がAIを、困難な作業を飛ばして平均的なコンテンツをより多く公開するために使いがちだからだ」と、条件付きで書いています。

Prolectからの提案

真似する対象は23個のスキルではなく、4つのゲートの問いを自社の言葉に翻訳することです。スキルの本数は、工程を分解した結果として決まります。「6〜12分」は、専門知を先に渡し、読者・製品・競合のデータが揃っている前提の結果値であり、目標にはしません。

出典：[S1] Ahrefs — コンテンツ制作パイプライン / [S2] Ahrefs — AIスロップと検索順位
原典URL・資料名は巻末に記載。

CircleCI

マーケティングチーム（5名） / Ron Powell

生成と公開判断を、分けて置く

CircleCI のマーケティングチーム（5名）の事例です。Ron Powell が、ブログ記事・チュートリアル制作と公開に使う運用基準を、AIが参照するファイル群として整理しました。定量的な成果は原典にありません。

AIに渡したもの	CLAUDE.md と .claude/rules/ の11ファイル（ブランドボイス、禁止フレーズ、レビュー基準の pass/fail、内部リンク規約、会社の行動原則、スケジューリングほか）。エージェント7体、スキル6種（seo-geo-post/tutorial-update/publish-prep/promote/the-guild/workflow-skill-check）。
AIがする工程	生成（Architect/Librarian/Mouthpiece）と、4段階のレビュー（Corrector/Webcrawler/Censor/Focus Group）。
人間が確認すること	プロダクションマネージャーが、公開前のチェックリストを人手で回します。「生成する人」と「公開可否を判断する人」を人的に分離しています。

確認できたこと - チーム5名の構成、ルール11ファイル、エージェント7体、スキル6種 - 公開前チェックリストが人手であること	確認できないこと 時間削減・本数・効率などの定量的な成果は、原典に一切ありません
---	---

"A review process that agrees with everything you write is not a review process, so each of these is built with specific friction."
すべてに同意するレビュープロセスは、レビュープロセスではない。だから、それぞれに固有の摩擦を組み込んである。(Ron Powell)

Prollectからの提案

成果数値がなくても参考になるのは、「何をファイルに書いたか」が真似できる粒度で開示されているからです。5名規模で最初に置くべきは生成者と検査者の人的分離です。1人で担う場合は執筆と確認を別の作業として実行しますが、それが人的分離と同じ効果を持つと実証されたわけではありません。

出典：[S3] CircleCI — マーケティングチーム（5名）
原典URL・資料名は巻末に記載。

Anthropic

社内マーケティング部門の2事例 / Ian Chan（マーケティング・オペレーション）・Adam Ward（フィールドマーケティング）

承認待ちと、内容確認は別のもの

Anthropic 社内のマーケティング部門で、Claude を業務に組み込んだ2つの事例です。いずれも社内向けの業務で、自社製品の宣伝文脈で公開されたものです。数値は〔自己申告〕として読みます。

	マーケティング・オペレーション Ian Chan [S4]	フィールドマーケティング Adam Ward [S5]
業務	週次のマーケティング指標レビューの準備	社内の営業担当者（AE／BDR／CS／アライアンス）への週次 Slack DM 配信。担当アカウントに合わせてイベントや紹介できるコンテンツを知らせる。 顧客への配信ではありません
AIに渡したもの	3つのスキル（prep／proofreading／action-items）	配信プロンプトの 9つのコンテンツルール 。すべて1件ずつのフィードバックに対応
AIがする工程	毎週日曜夜にスケジュール実行。前週のレポートを読み、会議の文字起こしを確認し、Slack を検索し、データにクエリを投げ、指標テーブルと論点候補を用意する	毎週月曜、各担当者のアカウントに合わせて配信内容を生成し Slack DM で送る。 配信前の承認待ちは挟まない
人間が確認すること	月曜朝に担当者が 数値を検証 し、どこにナラティブの焦点を当てるかを定める	本人が 配信内容を読み続けている 。"I still read what goes out, though the system no longer waits for my approval."
確認できたこと （自己申告）	準備工数が「週1〜2日」から「最大2時間」に	承認待ちを外したこと。休暇中も月曜の配信が自動で出たこと
確認できないこと	測定期間、比較方法	内容確認のタイミングと、全件を見ているかどうか

「承認不要」を「人間の確認不要」と読み替えないことが重要です。Adam Ward は非技術者で、記事の小見出しに "You don't need to code, you need to explain"（コードを書く必要はない。説明する必要がある）と置いています。

Prollectからの提案

校正スキルが数値レポートの照合先を持つ設計は、**指標定義を1箇所に揃えている組織**でないと成立しません。配信の承認待ちを外せたのは9つのルールが揃った後のことで、出発点ではありません。これは1社の運用例です。

出典：[S4] Anthropic — マーケティング・オペレーション / [S5] Anthropic — フィールドマーケティング
原典URL・資料名は巻末に記載。

Anthropic（営業組織）

担当者ごとに別の取り組み / Travis Bryant（Head of US Mid-Market GTM）・Jared Sires（AE）・John Albert（BDR）

担当者・業務・判断を分けて読む

Anthropic の営業組織に関する3つの記述です。担当者も業務も異なるため、数値がどれに紐づくかを分けて示します。数値はすべて（自己申告）で、自社製品の宣伝文脈で公開されたものです。

担当者	業務	AIに渡したもの・AIがすること	確認できた数値（自己申告）
Travis Bryant Head of US Mid-Market GTM [S6]	4,000件のアカウントの優先順位づけ	5次元の評価ルーブリック（agent opportunity／internal transformation／AI commitment／white space／industry fit）を渡し、AIが 一晩で4,000件を採点 。人間は「なぜその順序か」を担う	日次90分・週次3時間の削減
Jared Sires AE [S7]	営業向け社内ツール CLAFTS の内製	約4,300行のツールを、ほぼ全て Claude Code が記述。/customer-context が約90秒で顧客のコンテキストをまとめる。本人は入社までコード未経験	週10～15時間の節約。 営業組織の約80%が採用
John Albert BDR [S8]	インバウンド／アウトバウンド	「すべての送信に人を1人置く」を自チームの原則とし、繰り返し答えている質問を 外部向けの1つの文書 に集める	削減後の時間・コンバージョン率・ROI は原典に記載なし

"Claude builds the what; I do the why."（Claudeがwhatを作り、私がwhyをやる） — Travis Bryant / "Keep a person on every send."（すべての送信に、人を1人置け） — John Albert

Prolectからの提案

ルーブリックの5次元は、AIがなくても営業の優先順位づけに使える判断基準です。Albert の "external-facing" は**外部に出せる形のQ&A文書**を指し、社内の万能ナレッジベースではありません。判断基準の整備、準備作業の自動化、顧客への送信確認はそれぞれ別の業務として設計し、他チームの数値を組み合わせ自社の効果予測にはしません。

出典：[S6] Anthropic — 営業（アカウント評価） / [S7] Anthropic — GTMエンジニアリング / [S8] Anthropic — BDR（インバウンド／アウトバウンド）
原典URL・資料名は巻末に記載。

Attio

CRM を提供する企業 / RevOps / Kyle Doherty (寄稿記事)

属性抽出と失注理由分類を、分ける

CRM を提供する Attio の RevOps の事例です。Kyle Doherty の寄稿記事で、商談後の CRM 更新と、失注理由の分類という2つの業務を扱っています。定量的な成果は原典にありません。

業務	AIがすること	人間が確認すること
A. 商談後の CRM 更新	商談の記録から 7属性 を抽出し、CRM のレコードを更新する。 属性は「Update Deal」で定義：stage/close date/value /next steps/people/current solution/competitors	確信度の低いものだけ確認する仕組みが こちら側にあるとは書かれていません
B. 失注理由の分類	別のエージェントが失注理由を分類し、その出力に確信度スコアを付ける	確信度が低い分類だけ を人が見る設計

確認できたこと - CRM 更新で抽出する7属性の内訳 - 失注理由を分類する別エージェントと、その出力への確信度スコア	確認できないこと - 定量的な成果は原典にありません - 業務Aに、確信度による確認があるかどうか
---	--

"Scoring is a prioritization tool, nothing more." (スコアリングは優先順位づけの道具であって、それ以上のものではない)
"Give them good data and let them decide where their hour goes." (良いデータを渡して、その1時間をどこに使うかは本人に決めさせる) —
Kyle Doherty

Prollectからの提案

確信度スコアは「全件を人が見る／全件を自動で通す」の二択を避けるための設計です。ただし、そのまま導入することは勧めません。スコアが誤りの発見に役立つかは、**出力と正誤を突き合わせる検証工程**を経なければ分かりません。高確信のまま誤っている出力があるかを、自社データで先に確かめます。

Fin (fin.ai)

Dee Kapila / 週次レポートの作成と、その基盤「Ground Truth」

指標の意味を、人間が保守する

Fin (fin.ai) の Dee Kapila が、週次レポートの作成に Claude Code を使っている事例です。中心にあるのは、データウェアハウスに対する「Ground Truth」というセマンティックレイヤーです。

AIに渡したもの	Ground Truth＝指標の定義、参照すべきテーブル、既知の罣を書いた、AIファーストのセマンティックレイヤー。
AIがする工程	エージェントは、データに触れる前に必ず Ground Truth を検索し、正しい定義・正しいテーブル・既知の罣を取得する。データ接続は Snowflake MCP と Slack MCP を経由。
人間が確認すること	Research and Data (RAD) チーム が、 Ground Truth を人手で保守する。自動更新ではありません。

確認できたこと（自己申告） - 週次レポートの初版は約45分の対話で構築、以後は約5分の運用 - 週約2時間の手作業を削減 - v67 で3四半期以上稼働	確認できないこと 45分・5分・週2時間の測定方法と比較条件 - RAD チームの人数 - v67 に至るまでに何を変えたか
--	---

"For every 20-25 build experiments you'll have one major breakthrough." (20～25回のビルド実験に対して、大きなブレイクスルーは1回)
— Dee Kapila
本人の経験に基づく見解であり、標準的な成功率や必要な試行回数を示すものではありません。

Prollectからの提案

指標の定義を人間が保守し続ける専任の担い手を置いている点が、この事例の中心です。小規模組織では専任チームは持てませんが、指標定義のファイルを1つ持ち、数値を出すたびにそこを照合先にするとところまでは同じ形で作れます。定義を書くだけでは検算できないので、照合する数値データ・集計期間・集計条件も合わせて渡します。

出典：[S10] Fin (fin.ai) — Ground Truth と週次レポート
原典URL・資料名は巻末に記載。

Mintlify / Buffer / Emily Kramer

少人数チームと個人の3例 / MKT1 のニュースレターと調査レポート制作の記述から

人が通る一点は、置き場所が違う

短い記述しか公開されていない3例です。対象業務、AIがすること、人間が確認することだけを取り出します。Mintlify と Buffer は MKT1 のニュースレター、Emily Kramer は同じく MKT1 の調査レポート制作の記述から引いています。

	Mintlify [S11]	Buffer [S11]	Emily Kramer (MKT1) [S12]
体制	マーケティング7名・エンジニアゼロ	(記載なし)	個人
対象業務	社内ナレッジベースの更新	社内リンクの保守	100社×90以上のデータポイントの調査
AIがすること	Signal agent が30分ごとに動き、更新候補を出す	スキルが内部リンク切れを検出・修正する	企業データと担当者情報を収集・整理する
人間が確認すること	Knowledge Engineer (専任) が更新を全て承認	新しいスキルは、該当分野の社内専門家がテストしてから全社公開	100人分のプロフィールを手作業で目視確認
成果	(記載なし)	内部リンク切れを87%削減 (自己申告)	(記載なし)

Kramer が目視確認をした理由は、AIが「何年も前に退職した役員を、自信満々に現職として出力した」ためです。

"Automation gets you to a starting point; manual verification gets you to the finish line." (自動化はスタート地点まで連れて行く。ゴールまで連れて行くのは手作業の検証だ) — Emily Kramer

Prollectの解釈

3例とも、人が通る一点を置いています。ただし置き場所は同じではありません。Mintlify は更新の承認、Buffer は全社公開前のスキルのテスト、Kramer は公開する調査データの目視確認で、対象も、事前か事後かも違います。Buffer について確認できるのは「新しいスキルを専門家がテストしてから全社公開する」ことだけで、その後の個別出力を確認しているかは原典に書かれていません。

出典：[S11] Mintlify / Buffer — 少人数チームの運用 / [S12] Emily Kramer (MKT1) — 調査制作での手作業検証
原典URL・資料名は巻末に記載。

どの対象を、いつ確認するか

事例を参考に、確認の対象を具体的に決めます。全社共通の一つの承認方法へまとめる必要はありません。

対象	確認できた人間の関与
公開記事 [S1][S3]	編集者による原稿確認、公開前チェックリスト
社内の週次レポート [S4]	数値の検証と、何に焦点を当てるかの判断
社内向けSlack配信 [S5]	配信前の承認待ちはなし。本人は内容を読む。確認タイミングと全件確認の有無は不明。
失注理由の分類 [S9]	確信度が低い分類を人が確認
社内ナレッジベースの更新 [S11]	専任の Knowledge Engineer が更新を全て承認
新しいスキルの全社公開 [S11]	該当分野の社内専門家がテストしてから公開
公開する調査データ [S12]	100人分のプロフィールを手作業で目視確認

「人間は入口と出口、AIは中間工程」と単純化すると、**中間で人間に戻すべき場面が抜けます**。事例から読み取れる、途中でも人間に戻す4つの条件です。

原典どうしが食い違っているとき	同じ施策の成果が、出典によって単位も基準も違う形で書かれていることがある。突き合わせは人間の仕事
顧客固有の情報が足りないとき	Emily Kramer の例のように、AIは欠落を「自信満々の出力」で埋める
商品・提供範囲の判断が要るとき	何が言えて何が言えないかは、判断基準の側の問題
抽出した内容が、以後の前提として蓄積されるとき	検証のないまま社内の正史になると、誤った前提の上に企画が積み上がる

先に決めておきたいこと（提案）

「顧客に出る文面と、数字を含む主張は、必ず人が承認する」を1行で決めます。それとは別に、上の4条件に当たったら工程の途中でも止めると決めておきます。全社共通の一つの承認方法へまとめる必要はなく、「誰が」「何を見るか」「どの条件なら止めるか」を業務ごとに定めます。

出典：[S1] Ahrefs — コンテンツ制作パイプライン / [S3] CircleCI — マーケティングチーム（5名） / [S4] Anthropic — マーケティング・オペレーション / [S5] Anthropic — フィールドマーケティング / [S9] Attio — RevOps / [S11] Mintlify / Buffer / [S12] Emily Kramer

原典URL・資料名は巻末に記載。

一業務に必要なものから始める

以下は Prollect の提案です。判断基準を網羅的に揃えてから着手する必要はありません。1つの業務を選び、その業務に必要なものから始めます。このとき「判断基準」と「作業対象のデータ」は別のもので、判断基準は AI に渡す物差しで、データは物差しを当てる対象です。両方が要ります。

業務	判断基準として要るもの（物差し）	作業対象として要るもの（当てる対象）
数値レポートの校正	指標の定義（何を数え、何を数えないか）	照合先の数値データ、集計期間、集計条件
記事・メールの文体チェック	使わない語・言い回しのリスト	チェック対象の原稿
商談後の CRM 更新	抽出する属性の定義	商談の記録・文字起こし
コンテンツの読者設定	顧客像の定義	企画中のテーマ、既存コンテンツの一覧

指標定義を1ファイル書けば校正が回る、という話ではありません。定義があっても、照合先のデータと集計条件が揃っていないければ検算はできません。Fin の事例 [S10] で「Ground Truth」を人手で保守しているのは、この物差しの側です。

最初を選ぶ業務：そのまま合否を判定できるもの

- 数値が検証済みソースと一致するか
- URL が一字ずつ一致するか
- 文字数制限内か

後に回す業務：判定基準を分解してから

- 「良い記事か」は、そのままでは判定基準が曖昧
- 判定が不可能なのではなく、**読者・主張・証拠**へ分解する必要がある
- Ahrefs の4ゲート [S1] は、その分解の一例（読者を助けているか／主張を裏付けられるか／獲得していない確信が混ざっていないか）

Prollect 自身も、実装中です

顧客像・商品の事実・指標定義・禁止表現の4ファイルを作り、実際の制作で参照を始めたところです。**有効性は検証中**で、この4つで足りるのかも分かっていません。ファイル数を開始条件にはしません。

出典：[S1] Ahrefs — コンテンツ制作パイプライン / [S10] Fin (fin.ai) — Ground Truth と週次レポート
原典URL・資料名は巻末に記載。

試して、直して、判断基準を更新する

以下は Prollect の提案する順序です。実証された処方箋ではなく、期間も目安です。共通しているのは、**小さく試し、結果を見て、続行・修正・停止を判断する**という進め方で、最初から本番規模では走らせません。

順序	実施すること	事例での手がかかり
01 対象を選ぶ	可否を判定できる業務の一つ選ぶ	公開記事 [S1][S3]、社内の週次レポート [S4]、CRM 更新 [S9]、数値レポート [S10] など、事例はいずれも業務単位で始めている
02 基準を書く	その業務に必要な判断基準だけを書く	CircleCI はブランドボイス・禁止フレーズ・レビュー基準の pass/fail を11ファイルに [S3]。Ahrefs は4ゲートの問い [S1]
03 工程を分ける	工程を分解し、各工程の出力をファイルにする	Ahrefs は「各ステップが自前の出力ファイルを持つ」設計 [S1]。CircleCI は生成と4段階レビューを別工程に [S3]
04 確認を決める	人間の承認線を1行で決め、途中で止める条件を決める	「すべての送信に人を1人置く」 [S8]。承認待ちの有無と内容確認の有無は別 [S5]
05 理由を残す	修正した理由を記録し、再発性と重要性に応じ、継続して使う判断基準を更新する	Anthropic の配信プロンプトでは9つのルールがそれぞれ1件のフィードバックに対応 [S5]
06 定期化を判断	繰り返し実行が必要で、確認と停止の条件を決められた業務だけ、スケジュール実行に移す	日曜夜のスケジュール実行と月曜朝の人間の検証 [S4]。すべての業務が対象ではない

05 について：すべての指摘をルールにする、という意味ではありません

Anthropic の事例（9ルール=9件のフィードバック）はこの1社の運用例です。Prollect が提案するのは、**一度きりの修正で済むものと、継続して使う判断基準になるものを、記録の中から選ぶ運用**です。Prollect 自身もその線引きを試している最中で、まだ結論は出ていません。

小さく試し、結果を見て、続行・修正・停止を判断する。

出典：[S1] Ahrefs / [S3] CircleCI / [S4] Anthropic — マーケティング・オペレーション / [S5] Anthropic — フィールドマーケティング / [S8] Anthropic — BDR / [S9] Attio / [S10] Fin
原典URL・資料名は巻末に記載。

最初の一業務を、ここまで具体化する

実装前に確認するためのチェック項目です。点数による診断や、成果を予測する尺度ではありません。各項目の2行目は、本レポートの事例と提案から取った確認の観点です。

01 対象業務

何を AI へ渡し、何が出てくれば完了か。

可否をそのまま判定できるか（数値の一致・URL の一致・文字数）。「良い記事か」のように曖昧なら、読者・主張・証拠へ分解してから始める。

02 判断基準とデータ

参照する基準と、原稿・数値・記録は揃っているか。

判断基準（物差し）と作業対象（当てる対象）の両方があるか。指標定義があっても、照合先データ・集計期間・集計条件がなければ検算できない。

03 人間の確認

誰が、どの出力を、何と照合するか。

「顧客に出る文面と、数字を含む主張は、必ず人が承認する」の1行があるか。承認待ちの有無と、内容確認の有無を分けて決めているか [S5]。

04 停止条件

情報不足や原典の矛盾があったら、どう戻すか。

原典どうしの食い違い／顧客固有の情報の不足／商品・提供範囲の判断／以後の前提として蓄積される——の4条件に当たったら、工程の途中で止めると決めているか。

05 更新の記録

何を修正し、なぜ判断を変えたかを残せるか。

修正の理由を記録し、再発性と重要性で「続けて使う判断基準」を選んでいるか。すべての指摘を機械的にルールへ足していないか。

06 定期実行の条件

繰り返し実行が必要で、確認と停止の条件が決まっているか。

この2つが揃った業務だけをスケジュール実行に移す。すべての業務を定期化の対象にはしない。

自己申告と実測を、混同しない

本レポートの数値は、断りのない限りすべて各社の〔自己申告〕です。測定方法・母数・ベースライン・対照群は開示されていません。ここでは、本調査が言えることの範囲を示します。

事例の説明と、効果の実証は別

掲載した企業のマーケティング実装事例の中に、**対照群を置いた比較として原典で確認できたものではありません**。これは企業が自社で公開した「事例」についての話で、効果を検証した学術研究とは分けて扱います。第2章以降の〔自己申告〕数値は、他施策や時期要因との因果分離がなされていないものとして読み、自社の効果予測へ転用しません。

実感と結果は、逆になることがある

METR の RCT [S16a] (n=16、熟練 OSS コントリビュータ、246課題) では、被験者は事前に「24%速くなる」と予測し、**実際には19%遅くなり、事後も「20%速くなった」と信じていました**。2025年前半の AI ツールと当該の開発作業の条件での結果で、マーケティング業務を調べたものではありません。2026年2月の追跡報告 [S16b] では、被験者の30~50%が「AI なしではやりたくない課題を提出しなかった」と答えるなど選択バイアスが強く、「現在の AI ツールの生産性への効果について信頼できるシグナルを与えていない」と書かれています。**「効果が改善した証明」としても使えません**。ここから言えるのは、自己申告だけでは実際の時間短縮を判断できない、ということです。

本調査で確認できなかったこと	読み方
「マーケティング組織を AI エージェント前提に再設計した」国内の一次事例（数字つき）	存在しないと断定はできない。本調査がアクセスできた公開情報の範囲で見つからなかった。 公開事例が少ないことから、先行者利益や競争上の優位は導けない
「マーケ1人目がフルファネルを1人で回している」という一次証言	確認できなかった
Reddit/Indie Hackers/LinkedIn の長文投稿	未調査。小規模な運用の実装レポートが集まりやすい場で、本調査の空白
主要ツールの、日本語・日本企業データにおける精度	いずれの出典でも検証されていない
文書化なしで成果を出す組織／文書化しても進まなかった組織	比較していない。反証側を見ていないため、「判断基準の外部化」を法則としては書けない

〔解釈〕〔提案〕と付した箇所は、すべて Prolect の判断であり、出典に書かれていることではありません。検証しながら使ってください。

作れる量が増えたとき、 何を作るかを決める

Prollect は、受注から逆算して「どんなコンテンツを、誰に、どの順番で作るか」を設計する会社です。AI ツールの導入支援や AI 研修を提供する会社ではありません。

本調査で AI を正面から扱ったのは、AI が商品だからではありません。**作れる量が増えたときに「何を作るか」の判断がどこまで価値を持つのかを、確かめる必要があったから**です。

Prollect 自身も、実装中です

他社に提案を出しておいて自社では試していない状態を避けるため、Prollect 自身でも判断基準の外部化に着手しています。現時点で報告できるのは次の2点だけです。

- 顧客像・商品の事実・指標定義・禁止表現の**4ファイル**を作成した
- それを実際の制作業務で参照し始めた

これは組織再設計の成果事例ではありません。効果は測っていません。何がうまくいき、何がうまくいかなかったかは、検証できる形になった時点で、本レポートが他社の数値に当てたのと同じ基準——対照群の有無、自己申告か実測か——を自分たちにも当てて公開します。

引用について：本レポートの引用・言及は歓迎します。数値を引用する場合は、〔自己申告〕の表記と「14 LIMITS」を併せてご参照ください。詳細な事例、国内調査の全文、追加の留保は Web 版（全出典つき）にあります。

[Web版の全文・全出典を読む →](#)

[プロレクト株式会社について →](#)

参考文献一覧

本レポートで使用した出典です。番号はWeb版の全文と共通です。タイトルまたはURLから原典を開けます。

[S1] Ahrefs — コンテンツ制作パイプライン

Ryan Law 「How I Do Content Engineering With Claude Code」 Ahrefs Blog／2026-04-28

<https://ahrefs.com/blog/how-i-do-content-engineering-with-claude-code/>

[S2] Ahrefs — AIスロップと検索順位

Si Quan Ong (Reviewed by Ryan Law) 「How We Use AI Without Making AI Slop」 Ahrefs Blog／2026-08-28

<https://ahrefs.com/blog/how-we-use-ai-without-making-ai-slop/>

[S3] CircleCI — マーケティングチーム (5名)

Ron Powell 「Inside CircleCI's Claude marketing folder」 CircleCI Blog／2026-04-14

<https://circleci.com/blog/claude-marketing-folder/>

[S4] Anthropic — マーケティング・オペレーション

「How Anthropic's marketing operations team uses Claude Cowork to automate reporting and campaign builds」 Anthropic／2026-07-08

<https://claude.com/blog/how-anthropics-marketing-operations-team-uses-claude-cowork-to-automate-reporting-and-campaign-builds>

[S5] Anthropic — フィールドマーケティング

「How an Anthropic field marketer uses Claude Code to send weekly personalized updates to every sales rep」 Anthropic／2026-08-24

<https://claude.com/blog/how-an-anthropic-field-marketer-uses-claude-code-to-send-weekly-personalized-updates-to-every-sales-rep>

[S6] Anthropic — 営業 (アカウント評価)

Travis Bryant 「How an Anthropic sales leader uses Claude Cowork to run a 4,000-account book」 Anthropic／2026-05-20

<https://claude.com/blog/how-an-anthropic-sales-leader-uses-claude-cowork-to-run-a-4-000-account-book>

[S7] Anthropic — GTMエンジニアリング

「How Anthropic uses Claude for GTM engineering」 Anthropic／2026-06-05

<https://claude.com/blog/how-anthropic-uses-claude-gtm-engineering>

参考文献一覧

本レポートで使用した出典です。番号はWeb版の全文と共通です。タイトルまたはURLから原典を開けます。

[S8] Anthropic — BDR（インバウンド／アウトバウンド）

John Albert 「How Anthropic's business development team uses Claude to run inbound and outbound at scale」 Anthropic／2026-08-07
<https://claude.com/blog/how-anthropics-business-development-team-uses-claude-to-run-inbound-and-outbound-at-scale>

[S9] Attio — RevOps

Kyle Doherty 「How Attio runs RevOps on Attio」 GTM Strategist／2026-09-04
<https://knowledge.gtmstrategist.com/p/how-attio-runs-revops-on-attio>

[S10] Fin（fin.ai） — Ground Truth と週次レポート

Dee Kapila 「How to use Claude Code as a GTM leader」 Growth Unhinged／2026-08-26
<https://www.growthunhinged.com/p/how-to-use-claude-code-as-a-gtm-leader>

[S11] Mintlify / Buffer — 少人数チームの運用

「Multiplayer Claude」 MKT1 Newsletter／2026-07-27
<https://newsletter.mkt1.co/p/multiplayer-claude-1>

[S12] Emily Kramer（MKT1） — 調査制作での手作業検証

「State of Marketing Report Part 3: How to research in Claude Code」 MKT1 Newsletter／2026-05-06
<https://newsletter.mkt1.co/p/state-of-marketing-report-part-3-how-to-research-in-claude-code>

[S14a] 帝国データバンク — 生成AIの活用に関する企業の意識調査

全国23,349社に調査・10,312社回答（2026年3月17～31日）／2026-05-14
<https://www.tdb.co.jp/report/economic/20260514-genai/>

参考文献一覧

本レポートで使用した出典です。番号はWeb版の全文と共通です。タイトルまたはURLから原典を開けます。

[S14b] 才流 — BtoBマーケティングにおける生成AI活用実態調査

n=600（2025年7～8月）。BtoBマーケティング支援会社による自社調査／PR TIMES／2025-09-19

<https://prtimes.jp/main/html/rd/p/000000023.000022074.html>

[S14c] 総務省 — 『情報通信白書』令和7年版

2024年度調査／2025-07

https://www.soumu.go.jp/main_content/001019264.pdf

[S14d] PwC Japan — 生成AIに関する実態調査2025

日本 n=945（売上500億円以上・課長職以上、2025年2月19～25日）。コンサルティング会社による自社調査／2025-06-23

<https://www.pwc.com/jp/ja/knowledge/thoughtleadership/generative-ai-survey2025.html>

[S15] 太田翔葵（シャコウ CEO） — 個人の見解

note／2026-04-24。1名の見解であり、全国的な傾向を示すものではありません

https://note.com/shoking_shakou/n/n1cfc6a99a2e2

[S16a] METR — 開発者のAI利用に関するRCT

「Measuring the Impact of Early-2025 AI on Experienced Open-Source Developer Productivity」／2025-07-10

<https://metr.org/blog/2025-07-10-early-2025-ai-experienced-os-dev-study/>

[S16b] METR — 追跡報告

「Uplift update」／2026-02-24

<https://metr.org/blog/2026-02-24-uplift-update/>