



```
1 import you
2 process config from terminal/com
3 const studia = from terminar/neg
4 @chaccu@od
5 chow r re.exets(
6 na
7 eni
8 arg
9 str
10 o
11 checkr
12 #
13 #
14
```

Terminal

```
tetgr)
update: code");
("esran : updates aspckers"));
```

ターミナルに常駐する自律型エージェントとの協働プロセス

Claude Codeが変える、エンジニアの新しい日常

**「伝わる」Gitコミット
メッセージを自動生成する**

よし、
複雑なロジックの
実装終わった……！

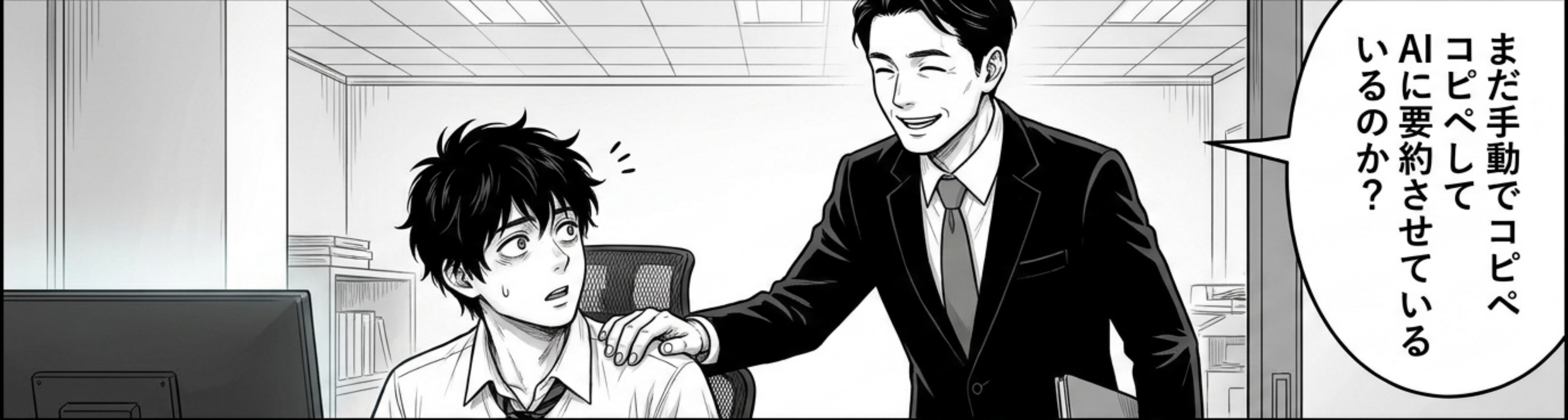


```
git commit -m ""
```

実装直後の「言語化」。
これが意外と重い
認知的負荷になる。

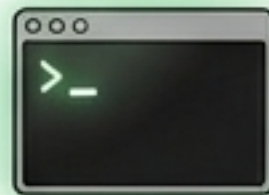
自分が書いたコード
何を変更したか
簡潔にまとめる気力が
気力残っていない……





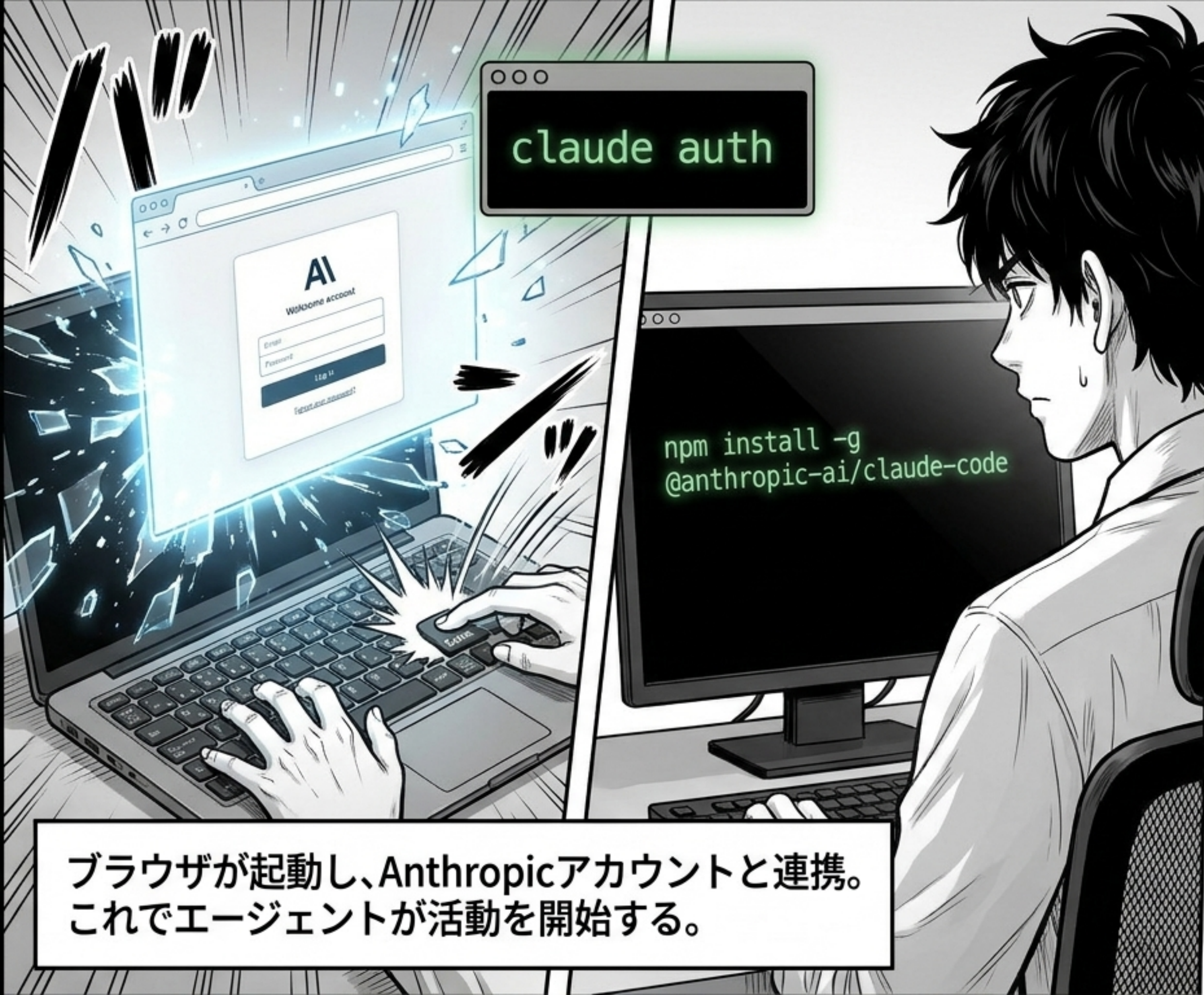
従来のチャットAI

- コンテキスト理解なし
- 変更内容をコピーして依頼
- 単なる「自動化 (Automation)」ツール



Claude Code

- ターミナルに常駐
- ファイル探索・Git操作を自律代行する「エージェント」
- 共に考える「思考の拡張 (Augmentation)」ツール



ブラウザが起動し、Anthropicアカウントと連携。
これでエージェントが活動を開始する。



君のターミナルに極めて
極めて優秀な配属させよう。
Node.js環境があれば
準備は一瞬だ。

コマンドは /effortだ。
コミットメッセージの生成は
生成は見通が良いタスクだから
デフォルトの「medium」で
最高の結果が出る。

深い思考。
大規模リファクタリング用。

思考最小。
単純なタイポ修正用。

【最適】バランス型。
通常の機能追加や
バグ修正に。

思考最小。
単純なタイポ修正用。

全リソース使用。
基盤コードの変更用。



指示と同時に、AIが内部で自律的にツールを使用し、
変更箇所を探索(Explore)している…!



現在のステージングされた
変更内容を確認して、
Conventional Commits形式
メッセージを提案して。
提案して。

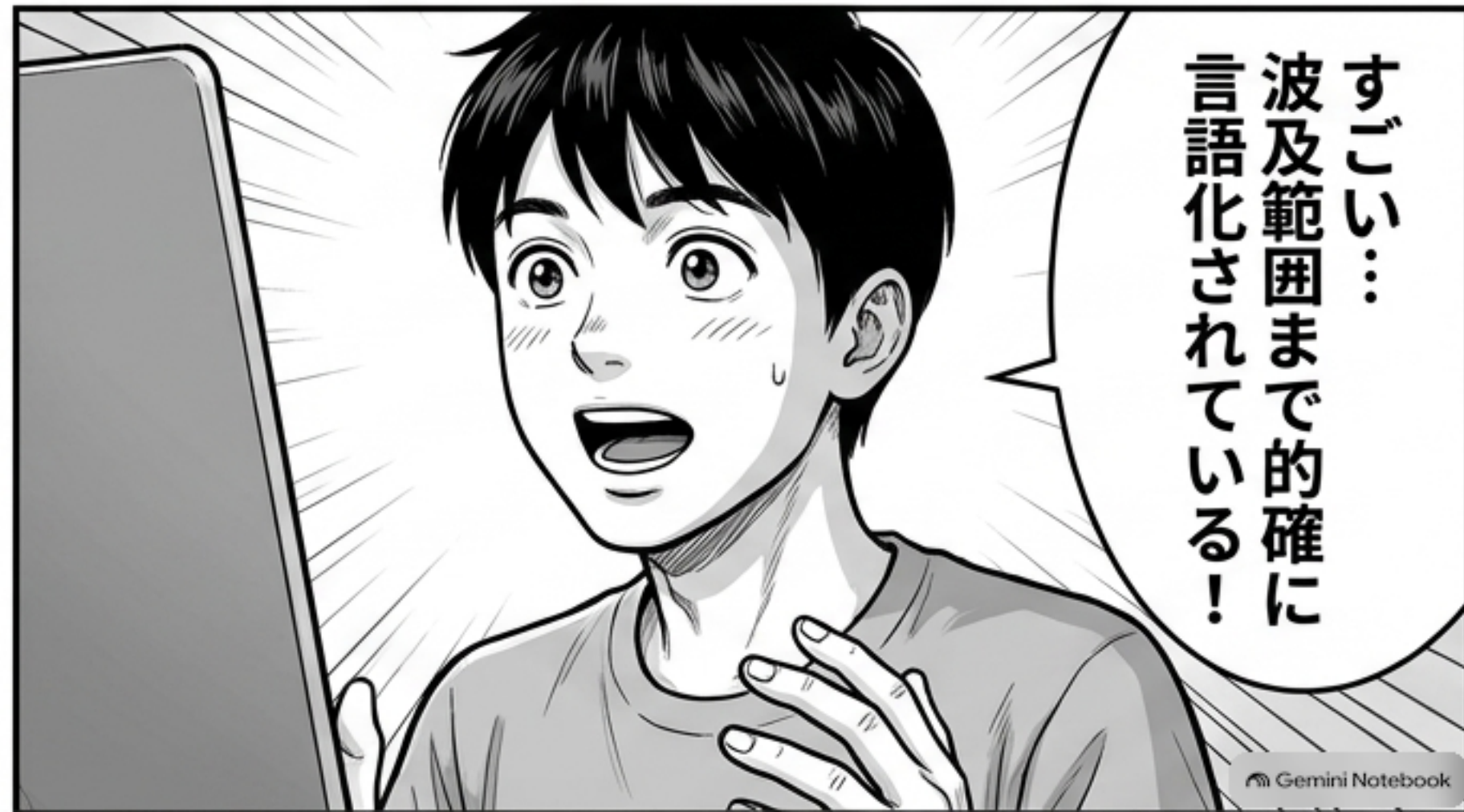
feat: ユーザー認証フローの追加

- ～ 詳細内容は、ユーザー認証フローの追加しております。
- ～ これまでお困りだった場所が購入してコースに実動型を再現した。
- ～ 変更化メッセージを随時確認する事だよ。

納得できれば、
そのまま「commit」を
を指示しろ。



生成された内容を
人間が評価
(Discernment)
するんだ。



すごい…
波及範囲までの確に
言語化されている！

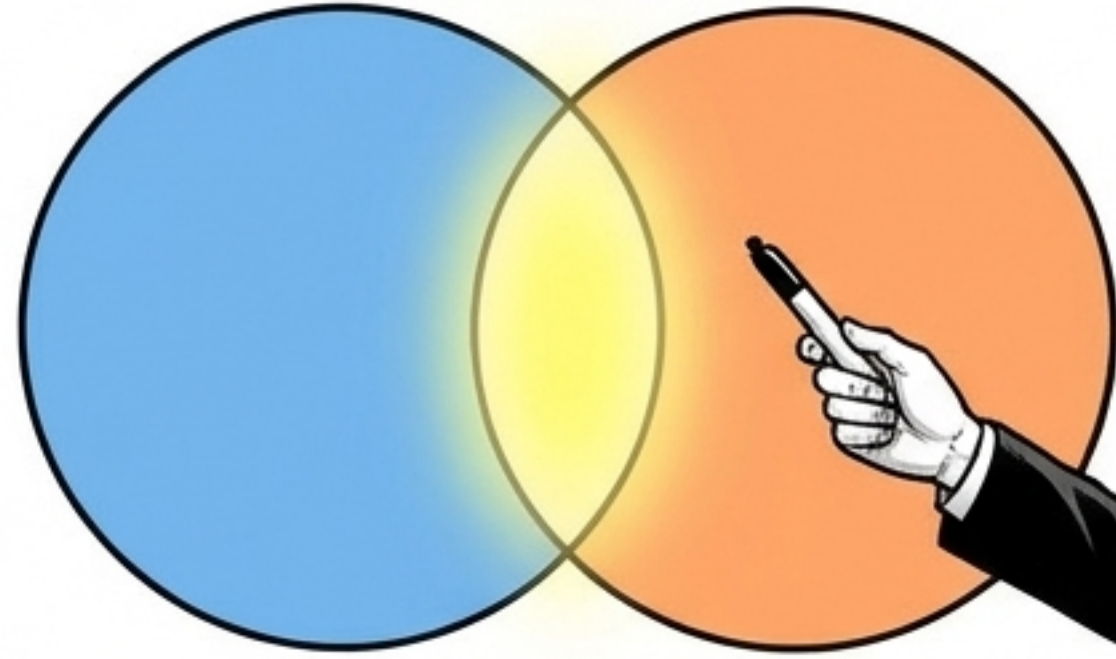
...将来的なUI拡張を見据えた
データ構造の変更

あれ？
頼んでいない
「将来の拡張」のことまで
背景として書かれている…
コードにはそんな情報
ないのに！

次トークン予測の限界。もっともらしい言葉を続けるため、
コードにない『背景』を勝手に推測してしまうことがある。

AIには
「何を変えたか」は
分かってても
「なぜ変えたか」は
は分からない。

背景は人間が
指示するんだ。
違っなて
（
背景の差分）
なぜ変えたか
何かする…。



コードの差分
（AIが見える領域）
＝何を変えたか

開発の背景・意図
（人間の頭の中）
＝なぜ変えたか

「修正内容はコードから取得してください。
ただし、この修正の背景は【将来的な
スケーラビリティの確保】ですので、
その点を含めてください」…と。

BEFORE (CLUTTERED CONTEXT)



AFTER (CLEARED & COMPACTED)



メッセージ生成の直前に、
直前に、
脳内メモリを
クリアにしてやると
精度が落ちないぞ。



/compact: 過去のやり取りを要約し、
不要なトークンを削ぎ落として
作業メモリをリフレッシュ。



/clear: セッションを完全に
リセットし、最新の差分だけに
集中させる。

CLAUDE.md

- コミットメッセージは日本語で記述すること。
- Conventional Commitsのルールに従うこと。
- コミット前に必ずテスト (npm test) をパスしたことを確認すること。

これをルートに置くだけで、チームの『コミットの流儀』をAIが永続的に記憶し、テストまで自律的に実行してくれるのか！

ターゲット読者による出力トーンの調整 (AI Fluency: Description)



技術レビュー向け

指示のポイント:
技術的な詳細と影響範囲を重視。

```
export codes {  
  constructcontext(() => {  
    // URL: "https://src/auth/ai/downioat",  
    indoxn: #leaf "hetp.ai/context="
```

fix(auth): トークン検証時の境界値エラーを修正し、401応答を正確に返すよう変更

```
// output #oot  
$ githut% /--donatmic commits error
```



マネージャー向け

指示のポイント:
ビジネス上の価値と解決した問題を重視。

fix: ログイン時に稀に発生していた接続遮断問題を解消。ユーザー離脱を防止

誰が読むのかを明示するだけで、AIの出力は劇的に変化する。



AIが生成した
したメッセージを
そセージを
使うまま使うのは
効率的だ。
だが、不正確な推論
推論がかを検証し、
最終的な説明責任を
あくまで人間である
間である君自身だ。

ツールに依存
するのではなく、
共に品質に責任を
責任持ちます。

はい…！



「何を変えたか」だけでなく「なぜその設計判断をしたのか」を残す。
これが未来のメンテナンス性を支える礎となる。



よし、完璧な一行だ。
これなら数ヶ月後の
自分も、数ヶ月自分
チームメンバーも
助かるはず。

事務的な要約はAIに。
人間は意図の確認と品質の担保に集中する。
これがAI時代のエンジニアリング。

さあ、ターミナルを開き **claude** と入力しよう。
/usage で確認できる微々たるコストで、
計り知れないプロジェクト履歴の価値が手に入る。