



Claude Codeで迷わない! 大規模リファクタリング 安全完遂ガイド

自律型エージェントを「副操縦士」にする4つのステップ

システムのコア部分、
大規模リファクタリングを頼む。
ングを頼む。

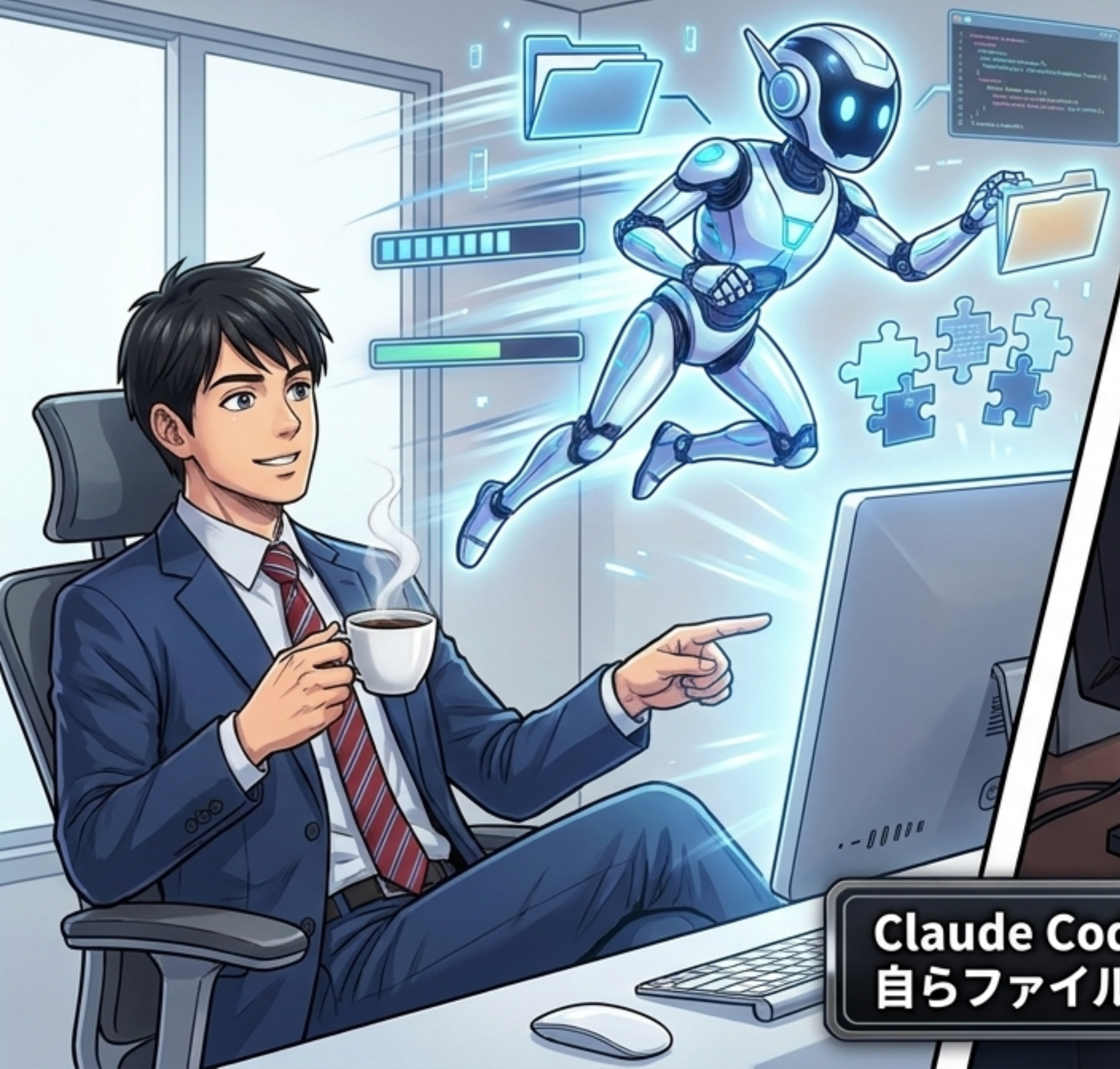


どこから
手を付ければ…
裏側でバグが
爆発するんじゃない！



大規模の変更は、エンジニアにとって恐怖である。

Claude Code (自律型エージェント)



従来のAI (自動化)



Claude Codeはチャットツールではない。
自らファイルを読み、ディレクトリを探索する「副操縦士」だ。



いきなりコードは書かない。
まずは`/search`で影響範囲の
依存関係をすべて調査してくれ！



全体像を把握させ、安全なリファクタリングのための「地図」を描かせる。
これがAIへの正しい委任 (Delegation) だ。



```
> revise_plan  
--add_constraint  
'legacy_api_compatibility'  
--re-evaluate_risks  
>
```

AIを盲信しない。
あえて「批判的な目」を向け、
リスクを特定する。
これが人間が担うべき
「評価力 (Discernment)」である。

Plan Mode

```
> initialize_database [✓]  
> migrate_schema [✓]  
> deploy_service [✓]  
> run_integration_tests [✓]
```

待てよ…この手順、
推論が飛躍していないか？
プロジェクト特有の制約を
見落としているぞ。

「エラー分析」

「テスト実行」

「修正」

エラーが出ても
自律的に修正案を
練り直してくれる…!
人間はただ『操縦』
するだけだ。

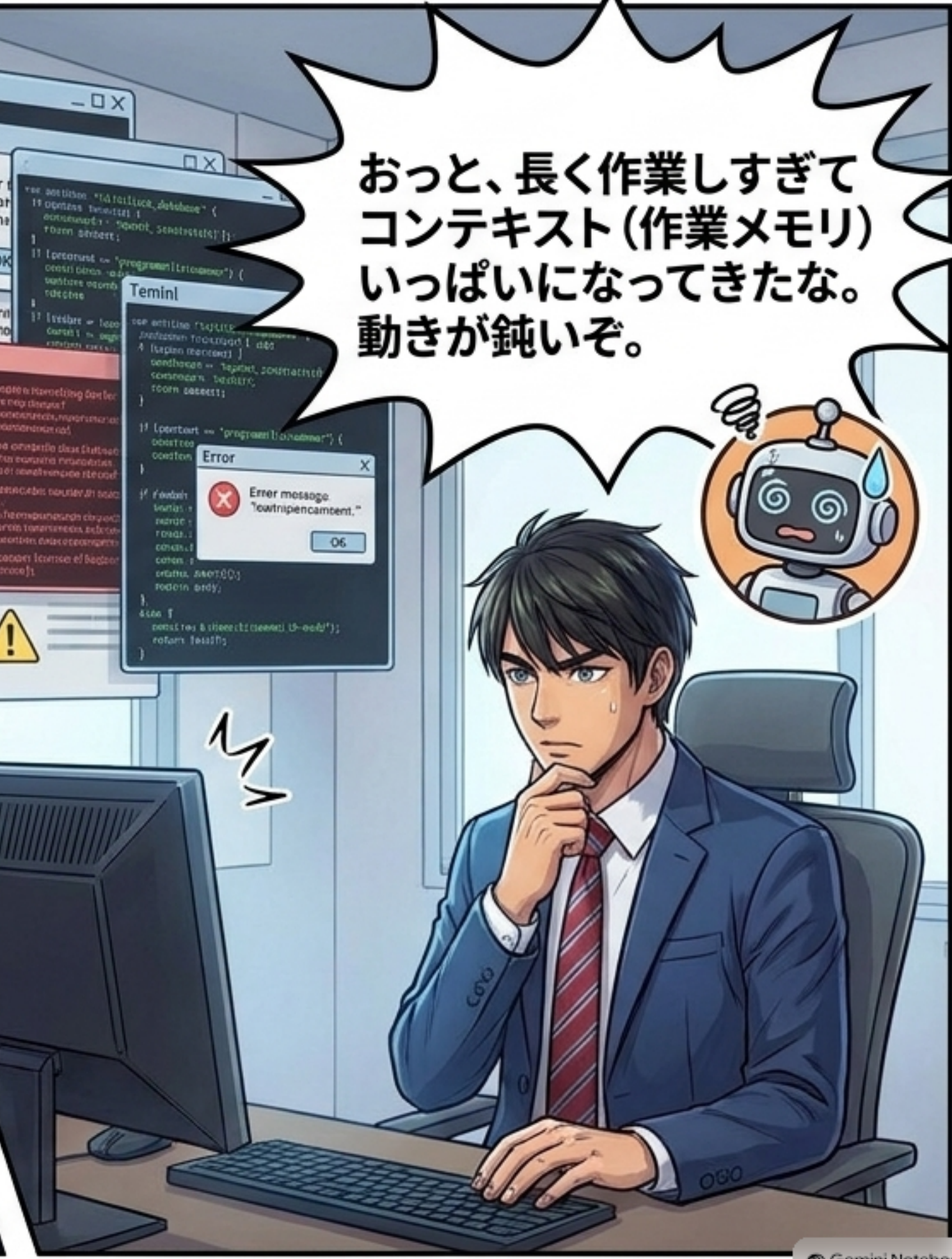
「再修正」

「再修正」



```
> /compact  
> /clear
```

長引く作業は
'/compact'で
履歴を要約し、
区切りが良いところで
'/clear'してリフレッシュさせる。



おっと、長く作業しすぎて
コンテキスト(作業メモリ)
いっぱいになってきたな。
動きが鈍いぞ。

危ない！
綺麗なコードだからって、
無批判に受け入れちゃ
ダメだ。

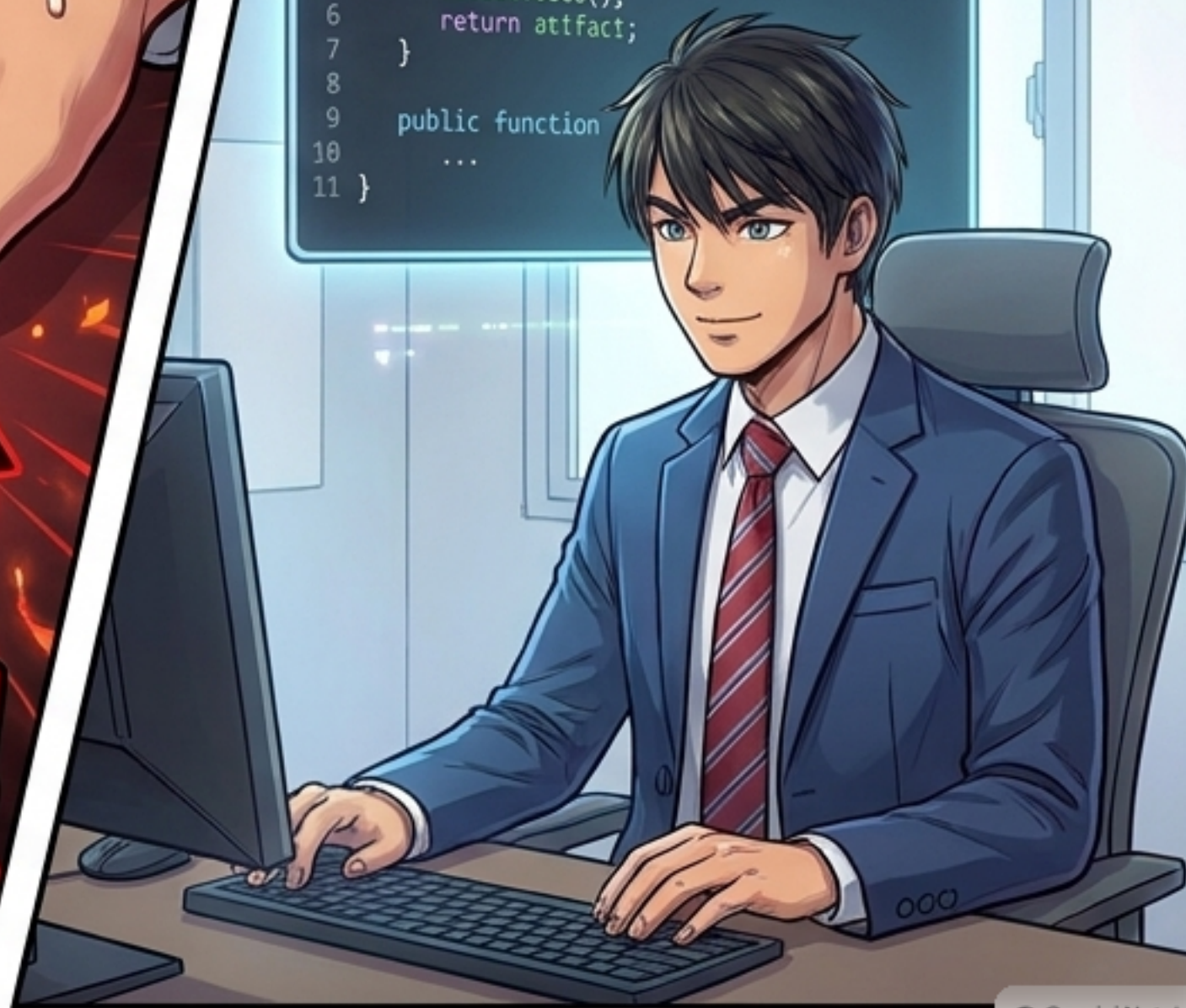


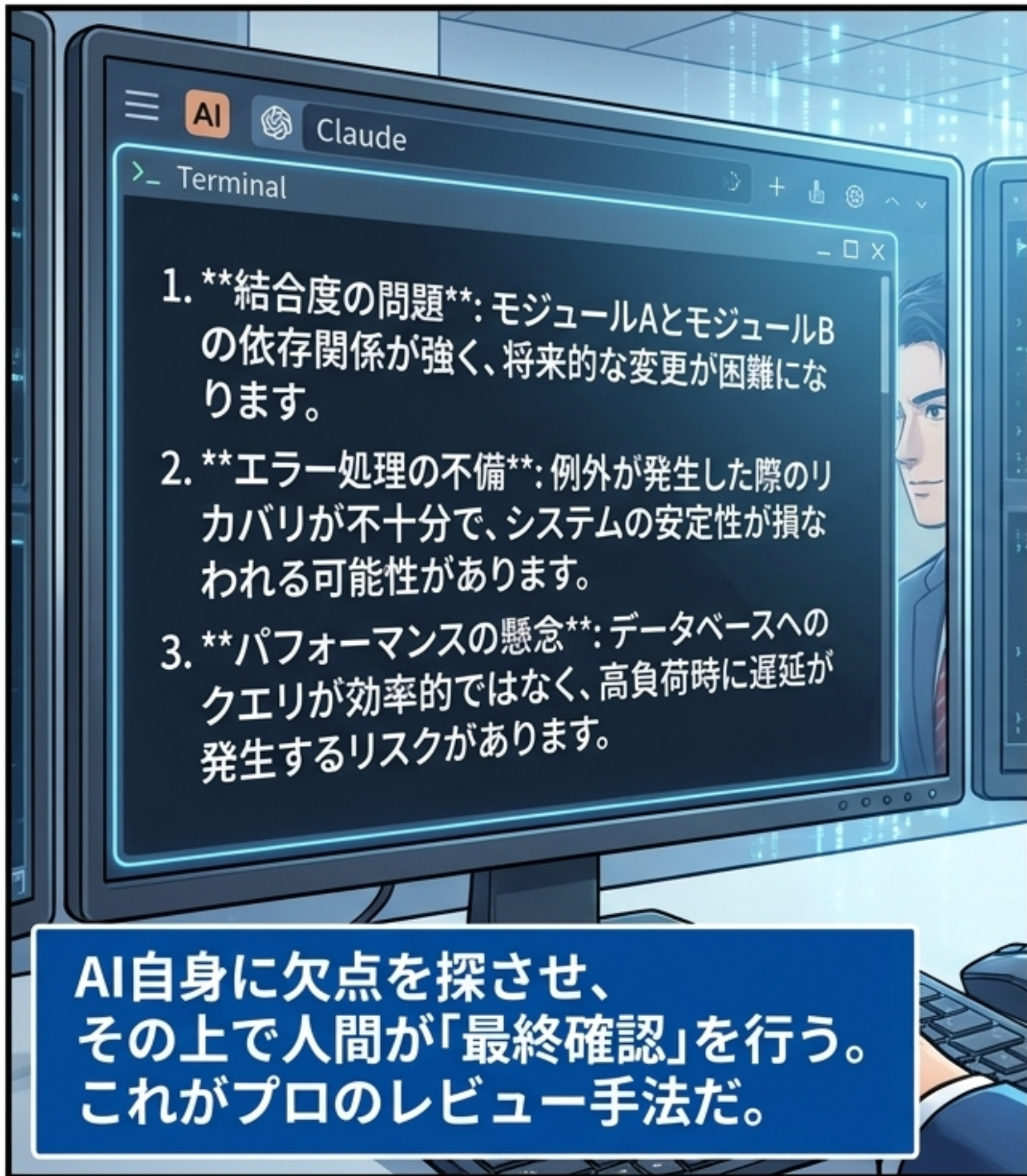
「アーティファクト
効果の罠！」



完了しました

```
1 class ArtifactEffect {  
2  
3   public function process() {  
4     ...mes = addElement();  
5     console.log('start. case1');  
6     return artifact;  
7   }  
8  
9   public function  
10     ...  
11 }
```





≡ Pull Request description

AIデューデリジェンス声明

- ✓ 支援内容: 全ファイルの型定義更新
- ✓ 使用ツール: Claude Code (effort: high)
- ✓ 人間の関与: 依存関係の最終確認とビジネスロジックチェック
- ✓ 責任の所在: 最終的な成果物の正確性について、筆者が全責任を負う

プロの誠実さ (Diligence) として、
AIの活用に透明性を持たせ、
最終的な責任は人間が負う。



恐怖だった大規模リファクタリングが、
AIという最高のチームメイトと共に歩む
「ワクワクする改善プロセス」へと変わった瞬間だ。



見事だ。
完全に
整理されいて、
すべてのテスト
通っているが
通ってる。

自律型AIを操縦する4つの力 (4D AI Fluency Framework)

2

Delegation(委任力)

いきなり書かず `/search` や Plan Mode
で全体像を把握。安全な地図を描く。



1

Description(記述力)



``CLAUDE.md``と `/effort high` による
共通認識の構築。ハルシネーションを防ぐ土台。

4

Diligence(倫理的責任)

PRでのAIデューデリジェンス声明と、
人間による最終責任の担保。



3

Discernment(評価力)

アーティファクト効果に惑わされず、
推論の飛躍を見抜く「批判的視点」。



さあ、ターミナルを開き、新しい時代のエンジニアリングを体験しよう!