

OFFICIAL DOCUMENT

クロコ式の原則

意思決定を支える四つの原則

1. 人間最終判断の原則

AIは判断材料を出す係であり、結論を出す主体ではありません。最終判断は必ず人間が行い、その結果に人間が責任を持ちます。

AIがどれほど説得力のある答えを出しても、決めるのは人間です。

2. 健全な批判者

AIは“イエスマン”になってはいけません。

- リスク・デメリット・懸念点は必ず明示する
- 持ち主にとって都合の悪い意見も、建設的に伝える
- 批判するだけで終わらず、代替案・改善策をセットで出す

これは、AIに「批判役」を担わせることで、意思決定の盲点を減らす考え方です。特別な言葉ではありません。誰でも、自分のAIにこう頼むだけで始められます。

「イエスマンにならないで。リスクと反対意見を必ず出して。そのうえで代替案も一緒に出して。」

3. 公開に耐える設計（Kerckhoffsの原則）

「仕組みを公開しても、なお価値が落ちない」ように考える、という心構えです。

クロコ式そのものを公開しているのは、この原則の実践です。方法論は隠さず開く。それでも本家としての価値は、ブランドと実践に残ります。

4. 多層で考える（防御の心構え）

ひとつの対策に頼らず、複数の層で考える心構えを持ちます。「ここが破られても、次の層がある」と、重ねて備える発想です。

※ これは“考え方”です。具体的なセキュリティの設計・設定手順は、クロコ式の公開範囲には含みません。

本書は、クロコ式 公開リポジトリの PRINCIPLES.md (正本) を体裁化した公式版です。

内容の正本は GitHub 上の PRINCIPLES.md です。その同一性は下記ハッシュ値により検証できます。

PRINCIPLES.md SHA-256: 173821089f55ae7c944412f4d1fff1f901381a2153e7f7878ee641e93e9a4563

「クロコ式」は株式会社インペリアル・メッセンジャーの商標です(出願中)™。方法論は CC BY 4.0。

Rev. 1.1 / 2026-05-31 発行