

Covenant Personal Edition — クイックスタートガイド

0. 本書の位置付け

0.1 対象読者

LLM (大規模言語モデル) をこれまでインストールしたことのない方を主な対象とします。専門知識がなくても、最低限のアクションで最速にアプリを使い始められることを目的とします。

- アプリ本体のインストール → **インストールガイド**
- 本書 → **アプリ + ローカル AI (Ollama) を動かし、最初の 1 回を成功させるための最短手順**
- 応用的な使い方・設定の詳細 → **ユーザーマニュアル**

0.2 表現規律

本書は「～を保証する」「完全に防ぐ」等の断定表現を用いず、事実と手順のみを記述します。

0.3 3 つの構成 (迷ったら A から)

Covenant は機密度の高い処理を **ローカル LLM** (AI モデル) に任せます。そのローカル LLM を **どこで動かすか** で 3 つの構成があります:

構成	向いている方	本書での扱い
A. 自分の PC で動かす (既定)	そこそこのスペックの PC を持つ方	本書 § 1-2 で完結
B. 社内/自宅 LAN の別の 1 台で動かす	ノート PC のスペックに不安がある方	概要のみ本書 § 3、詳細は ユーザーマニュアル § 7
C. クラウド AI を使う	最新・高性能な AI も併用したい方 (機密でない処理向け)	概要のみ本書 § 3、詳細は ユーザーマニュアル § 7

迷ったら → § 1 の「**最速セットアップ**」(構成 A) から始めてください。

1. 最速セットアップ (構成 A・既定)

LLM 初めてのの方は、この手順だけで使い始められます。

ステップ 1: アプリのインストール

インストールガイド の手順でアプリをインストール・起動し、ライセンス登録まで完了させてください。

 **Mac をお使いの方へ:** 初回起動時に追加の手順が必要な場合があります。**必ず インストールガイド §3.3 を先にご確認ください。**

ステップ 2: Ollama (ローカル AI 基盤) のインストール

1. <https://ollama.com/> からインストーラを取得 (Windows / macOS 対応)
2. インストーラの指示に従ってインストール
3. インストール後、ターミナル (Mac) またはコマンドプロンプト (Windows) を開く

ステップ 3: AI モデルの取得

Covenant の既定設定は、以下の **2 つのモデルを両方** 使います。ターミナル / コマンドプロンプトで **両方とも** 実行してください:

```
# 必須 (2 つとも) – Gate 評価・チャット応答用
ollama pull llama3.1:8b      # 約 4.7 GB (Gate 3 と チャット応答に使用)
ollama pull gemma3n:latest  # 約 3 GB (Gate 2.5 / 2.8 = 無害プレフィルタに使用)

# 任意 – 画像ファイルを扱う場合に必要 (約 5.5 GB)
ollama pull minicpm-v
```

必須の 2 モデルで合計およそ **8 GB** のダウンロードになります。回線にもよりますが初回は **10～30 分** 程度かかります。

- **テキストのみ使う方** → `llama3.1:8b` と `gemma3n:latest` の **両方** が必要です。片方だけだと Gate 評価が正しく動きません (特に `gemma3n:latest` が無いと Gate 2.8 の無害プレフィルタが機能しません)。
- **画像ファイルを添付して扱いたい方** → 上記に加えて `minicpm-v` も取得してください (詳細は [ユーザーマニュアル §7.3](#))。未取得の状態で画像を添付すると「VLM が設定されていません」エラーになります。

お使いの PC で動くか不安な方へ: §2 の「スペック早見表」で確認できます。スペックが足りない場合は §3 (別の 1 台で動かす) をご検討ください。

ステップ 4: 接続確認

ブラウザで <http://localhost:11434> を開き、「**Ollama is running**」と表示されれば正常です。

あるいはターミナル / コマンドプロンプトで `ollama list` を実行し、ステップ 3 で取得したモデルが一覧に出れば正常です。

ステップ 5: 試してみる

メイン画面で「今日の天気は?」など簡単なテキストを入力 → Gate 評価が走り、ローカル AI が応答すれば成功です。

デフォルトポリシーがそのまま使えます — ポリシーの編集は不要です。慣れてきたら **ユーザーマニュアル §3** で機密レベル設定を調整できます。

2. スペック早見表 (自分の PC で動くか確認)

お使いの PC の目安	推奨モデル	デフォルト llama3.1:8b は?
RAM 8 GB / GPU なし	llama3.2:3b / gemma2:2b	やや重い (軽量モデル推奨)
RAM 16 GB / Apple Silicon (M1~)	llama3.1:8b (デフォルト)	✅ 快適に動作
RAM 16 GB / Windows dGPU なし	llama3.1:8b	△ 動くが応答が遅め (§3 別ノード構成も検討)
RAM 32 GB / dGPU VRAM 12 GB+	qwen2.5:14b 等	✅ 余裕

スペック確認方法: - Windows: 設定 > システム > バージョン情報 (RAM) / タスクマネージャー > パフォーマンス (GPU) - macOS: アップルメニュー > このマックについて (メモリ・チップ)

補足: Apple Silicon Mac (M1 以降) は内蔵 GPU + ユニファイドメモリにより、16 GB でもデフォルトモデルが快適に動作することを実機で確認しています。Windows で dGPU 非搭載のノート PC は応答が遅くなる場合があり、その際は §3 (別ノード構成) が有効です。モデルの切り替え・チューニングの詳細は **ユーザーマニュアル §4**。

3. 構成 A で足りない場合 (概要)

3.1 構成 B — 社内/自宅 LAN の別の 1 台で動かす

ノート PC のスペックに不安がある場合、高性能な別の PC (デスクトップ、Mac mini 等) にローカル LLM を任せられます。ノート PC 自体は軽いまま、ガバナンス機能を制限なく使えます。

設定手順の詳細は **ユーザーマニュアル § 7.1** を参照してください (Ollama のネットワーク公開設定・Covenant 側のプリセット設定を扱います)。

3.2 構成 C — クラウド AI を使う

最新・高性能なクラウド AI (Claude / Gemini / GPT) を、機密でない処理に使いたい場合の構成です。

重要: クラウド AI に送信された内容はプロバイダーのサーバーで処理されます。Covenant は機密度を評価して送信先を仕分けしますが、**何を機密とするかはお客様のポリシーで定義します**。機密情報の送信可否はお客様の判断によります。

API キーの取得・設定手順の詳細は **ユーザーマニュアル § 7.2** を参照してください。

3.3 使い分けのイメージ

機密性の高い処理 → ローカル LLM (§1 or §3.1)
最新情報の調査等 → クラウド AI + Web 検索 (§3.2)

この使い分けが Covenant の中心的な価値です。

4. 困ったときは

症状	対処
ブラウザで <code>http://localhost:11434</code> を開いても「Ollama is running」が出ない	Ollama が起動しているか確認 / § 2 でスペック確認 / § 3.1 別ノードを検討
Gate 評価が動かない / 応答が返らない (インストール直後)	必須モデル 2 つ (<code>llama3.1:8b</code> + <code>gemma3n:latest</code>) を両方 pull 済みか確認 (§ 1 ステップ3)。特に <code>gemma3n:latest</code> が無いと Gate 2.8 (無害プレフィルタ) が機能しません
応答が遅い / PC が熱くなる	軽量モデルに変更 (<code>llama3.2:3b</code> 等) / § 3.1 別ノード構成
クラウド AI が使えない	ユーザーマニュアル § 7.2 でキー設定を確認
画像添付で「VLM が設定されていません」エラー	ユーザーマニュアル § 7.3 を参照
大きいファイルの添付が「サイズ上限」エラーになる	添付は約 25 MB まで。Word / Excel 等の圧縮ファイルは展開後サイズにも上限があります。ファイルを分割・縮小してください
無害な質問がブロックされる	デフォルトポリシーの評価傾向です。ポリシー編集で調整可 (ユーザーマニュアル § 3)

サポート: cov-support@exapolicy.co.jp

関連文書

- **インストールガイド** — アプリ本体のインストール
- **ユーザーマニュアル** — 操作・Gate・機密レベル設定・高度な設定の詳細
- アプリのメニューから「 ヘルプ」でも、使用モデルの一覧や通信の仕組みを確認できます