

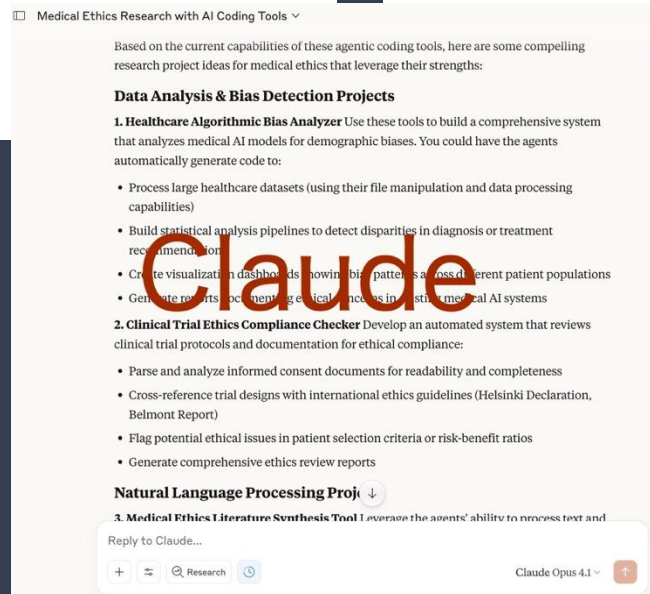
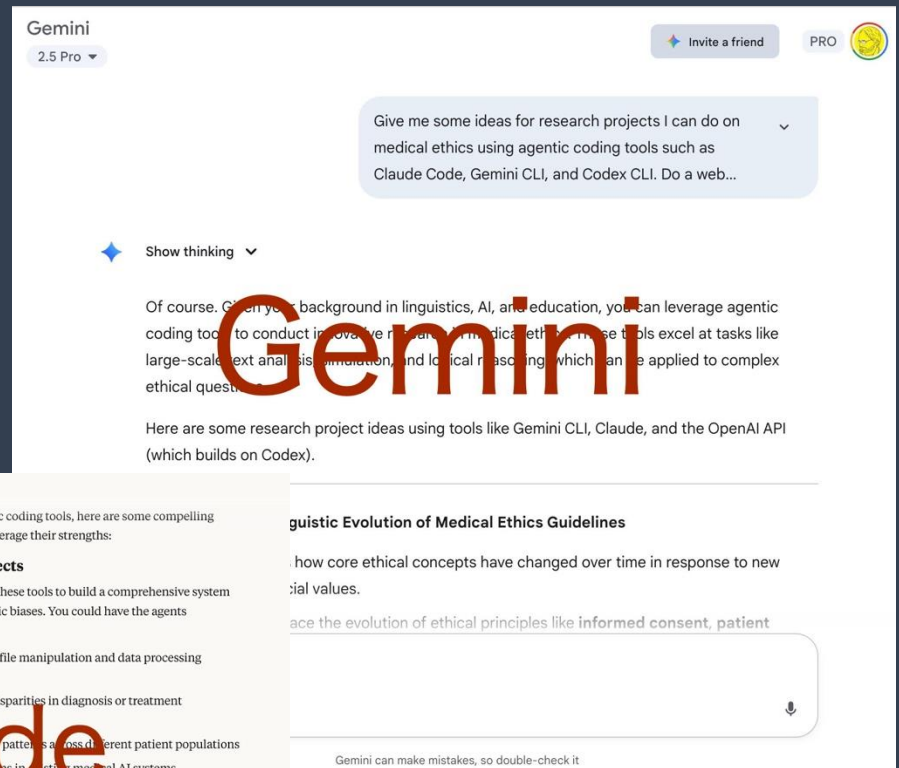
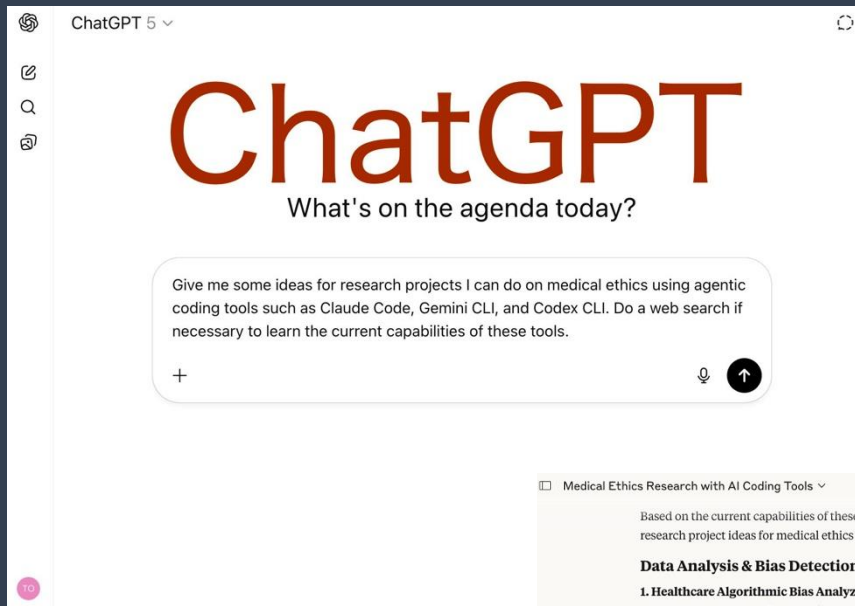
AIエージェント時代における 研究の卓越性の再定義

トム・ガリー

本日の内容

- 基本的なAIツール
- AIエージェント
- 研究評価への課題
- 何をすべきか？

チャットのインターフェース



「プロジェクト」

The image shows a screenshot of the ChatGPT 5 interface. At the top, the header reads "ChatGPT 5". Below it, the project title "Medical ethics and alternative medicine" is displayed. A search bar contains the text "New chat in Medical ethics and alternative medicine". To the right of the search bar are icons for voice input and a microphone. Below the search bar, there are two sections: "Files" and "Instructions". The "Files" section lists several documents:

- Medical Ethics and alternative medicine report 1.md (File)
- Medical Ethics and alternative medicine report 2.md (File)
- Ethical Issues Concerning Research in Complementary and Alternative Medicine.pdf (PDF)
- Physicians' Ethical Obligations Regarding Alternative Medicine.pdf (PDF)
- 关于医学伦理学案例教学的内涵与外延.pdf (PDF)
- The J of Law Medi Ethics - 2007 - Sade - Complementary and Alternative Medicine Foundations Ethics and Law (PDF)
- 보원대체의학의윤리적성찰.pdf (PDF)
- Complementary and alternative medicine.pdf (PDF)
- Medizinethik und Komplementärmedizin.pdf (PDF)

The "Instructions" section contains the text: "Provide clear, objective explanations. Do not praise or ask ...".

ダッシュボードとAPI

The image is a collage of screenshots from two AI development platforms: Google AI Studio and Anthropic.

Google AI Studio (Left): The main dashboard shows a sidebar with navigation options like Chat, Stream, Generate media, Build, History, Dashboard, and Documentation. The central area displays the Google AI Studio logo and a prompt: "Generate a Docker script to create a simple linux machine." Below this, there are several "What's new" cards for features like "Try Nano Banana", "URL context tool", "Native speech generation", and "Live audio-to-audio dialog".

Anthropic (Right): This section shows the interface for the Claude model. It includes a "Run settings" panel for "Gemini 2.5 Pro" (though the text says Gemini, the context is Claude). The "Token count" is 0/1. The "Temperature" is set to 1. The "Media resolution" is set to "Default". The "Thinking" section shows "Thinking mode" and "Set thinking budget". The "Tools" section includes "Structured output". The "Model" dropdown is set to "claude-sonnet-4-20250514 latest". The "System Prompt" is "Define a role, tone or context (optional)". The "User" prompt is "Generate Prompt" or "enter instructions or prompt for Claude...". The "Max tokens" is set to 20000. The "Thinking" toggle is off. The "Run" button is visible at the bottom right.

OpenAI (Bottom): A screenshot of the OpenAI chat interface. The "Model" dropdown is set to "gpt-4o". The "Variables" and "Tools" sections have "+ Add" buttons. The "System message" is "Describe desired model behavior (tone, tool usage, response)". The "Prompt messages" section shows a "User" message: "Enter task specifics. Use {{template variables}} for dynamic". The "Chat with your prompt..." input field is at the bottom.

Anthropic (Center): A large, stylized "Anthropic" logo in red text is overlaid on the right side of the image.

基本的なAIツールでできること

- 研究課題や方法のアイデアを得る
- 新しい分野への理解を深める
- 言語間の翻訳
- 文章を改善する
- 簡単なプログラムやマクロを書く
- LLMの長所と短所を理解する

AIエージェント

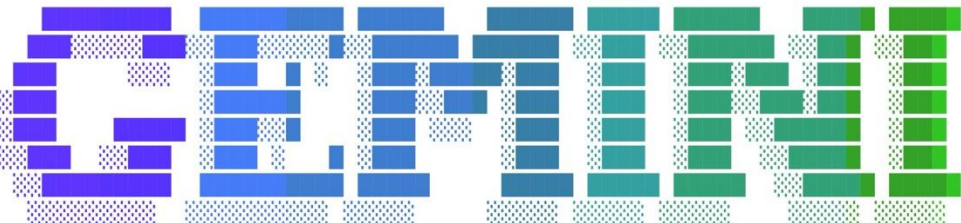
「AIエージェントとは、大規模言語モデルが自身のプロセスやツール使用を動的に管理し、タスクの達成方法を調整し続けるシステムである」

Anthropic, “Building Effective Agents”

2024年12月19日

コマンドラインインターフェース (CLI)

```
medical-ethics-cc — Gemini - medical-ethics-cc — node /opt/homebrew/bin/gemini — 80x24
~/Documents/GitHub/ClaudeCode/medical-ethics-cc — Gemini - medical-ethics-cc — node /opt/homebrew/bin/gemini
Tom@Tom-Gallys-Mac-mini medical-ethics-cc % gemini
```

The Gemini logo is displayed in a pixelated, blocky font. The letters 'G', 'E', 'M', 'I', 'N', 'I' are formed using a grid of colored squares. The colors transition from purple on the left, through blue and teal, to green on the right.

Tips for getting started:

1. Ask questions, edit files, or run commands.
2. Be specific for the best results.
3. Create **GEMINI.md** files to customize your interact
4. **/help** for more information.

Execution allowed by:

- `.claude/settings.local.json`

Learn more (<https://docs.anthropic.com/s/claude-code-security>)

1. Yes, proceed
2. No, exit

Claude Code

Enter to confirm · Esc to exit

* Welcome to **Claude Code**!

`/help` for help, `/status` for your current setup

cwd: `/Users/Tom/Documents/GitHub/ClaudeCode/medical-ethics-cc`

> **Try** "how does <filepath> work?"

? for shortcuts

```
medical-ethics-cc — codex-aarch64-apple-darwin · node /opt/homebrew/bin/codex — 80x24
~/Documents/GitHub/ClaudeCode/medical-ethics-cc — codex-aarch64-apple-darwin · node /opt/homebrew/bin/codex
>_ You are using OpenAI Codex in ~/Documents/GitHub/ClaudeCode/medical-ethics-cc
```

To get started, describe a task or try one of these commands:

`/init` - create an `AGENTS.md` file with instructions for Codex
`/status` - show current session configuration and token usage
`/diff` - show git diff (including untracked files)

Find and fix a bug in @filename
send `Ctrl+J` newline `Ctrl+C` quit

Codex

ブラウザ統合型エージェント

ChatGPT

ChatGPT agent

```
Terminal
+ }
+ });
+ // Activate the first case by default
+ if (index == 0) {
+   btn.classList.add('active');
+   caseDetailsTitle.textContent = c.title;
+   caseDetailsDesc.textContent = c.description;
+   for (const framework of ['utilitarian', 'deontological', 'virtue']) {
+     const section = document.createElement('div');
+     section.classList.add('analysis-item');
+     const heading = document.createElement('h3');
+     heading.textContent = framework;
+     section.appendChild(heading);
+     const para = document.createElement('p');
+     para.textContent = c.frameworks[framework].slice(1);
+     section.appendChild(para);
+     caseAnalysisDiv.appendChild(section);
+   }
+   caseListDiv.appendChild(btn);
+ }
+ });
+ }
+ *** End Patch
+ end
```

I will open the index.html file in a browser using the computer tool to test the site. This will help ensure the

Medical Ethics in Alternat...

I will do the following.

1. Research medical ethics, alternative medicine, and vaccine hesitancy
2. Research regional perspectives and comparative data
3. Compile and structure the comprehensive report
4. Design and develop the static website
5. Deploy the website to public internet

You can interrupt me at any time during my work to provide new information or adjust the plan.

☐ Research medical ethics, alternative medicine, and vaccin...

Searching medical ethics alternative medicine, med...

Research has yielded several sources on medical ethics, alternative medicine, and vaccine hesitancy. These sources include regional perspectives, comparative data, and a comprehensive report. The report's structure is as follows:

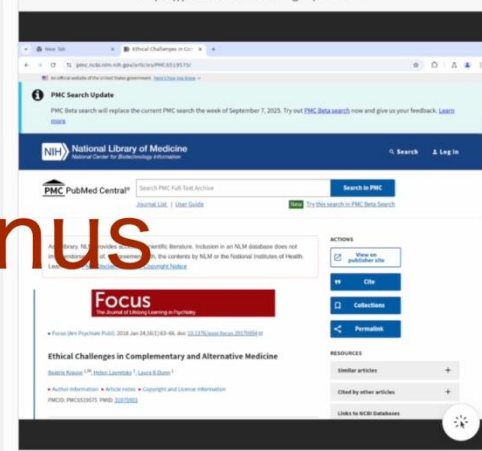
Manus

Manus's Computer

Manus is using Browser

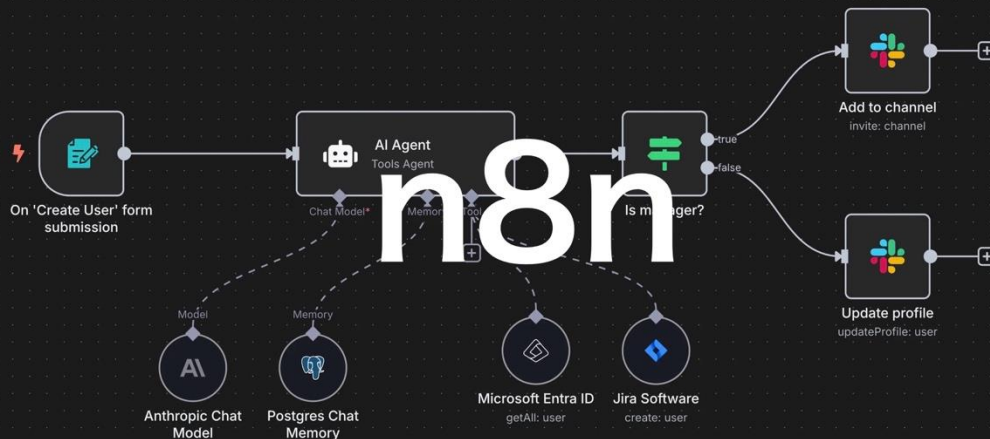
Browsing <https://pmc.ncbi.nlm.nih.gov/articles/PMC6519575/>

<https://pmc.ncbi.nlm.nih.gov/article...>



Manus is working: Research medical ethics, alternative m... 1/5
0:19 Viewing browser

n8n



エージェントができること

- ファイルの読み書き
- ウェブから情報を収集
- Pythonなどの言語でプログラムを書いて実行
- サーバーを起動・実行
- 必要なツールをダウンロード・インストール
- 複雑なタスクを計画して段階的に実行
- 複数のサブエージェント間でタスクを調整
- 大規模なコードベースを管理・修正
- ユーザーの許可があれば上記の作業を自律的に行う
(サンドボックス環境が望ましい！)

エージェントで始める方法

1. ChatGPT、Claude、Geminiなどで「プロジェクト」を作成する
2. 研究分野に関連する論文などをアップロードする
3. Claude Code、ChatGPT agent、Manusなど、使用予定のエージェントを説明する文書をアップロードする。（これらのエージェントは新しいため、LLMにはおそらくその機能の詳細を提供する必要がある）
4. 次のようなプロンプトを試してみる：「私は【トピック】に取り組んでいる研究者だ。このプロジェクトには私が興味を持っている論文などが含まれている。また、【エージェント名】に関する情報も含まれている。私の研究にこのエージェントを使用できる20の方法を提案し、それらの目的で【エージェント名】を使用する詳細な手順を教えて、その研究に使用できるプロンプト例を書いてください。」

現在のエージェントの問題点

- 自律的に作業できる一方で、時に脱線することがある。また、忘れっぽく、欺瞞的になることもある。人間の監視と指導が必要。うまくいかない場合は、別のプロンプトで再試行してください。
- エージェントはLLMに基づいているため、ハルシネーションを起こすことがある。複数のLLMが互いにチェックして修正するワークフローを設定することで、ハルシネーションを軽減できる場合がある。
(複数のLLMを使用するにはOpenRouterが便利)
- エージェントはまだ初期段階にあり、急速に進歩している。使い方の専門家はいない。創造的で実験的に取り組んでください。
- エージェントには、プロンプトインジェクションなどのセキュリティ問題があり、効果的な対策はまだ開発されていない。

エージェントと研究評価

- 7月に東広島キャンパスでの講演で説明したように、一部の分野では、AIエージェントに研究を計画し、調査し、出版可能な論文を書かせることが可能だ。
- 世界中で多くの人々が、人間の知識を進歩させるため、または自分の業績を増加させるために、自律的または半自律的な研究エージェントを開発しようとしている。
- 人間がどの程度関与したか不明瞭な研究論文が大量に生み出される時代が目前に迫っている。

- 長年にわたり、研究を行うインセンティブは定量的評価によって歪められてきた（掲載数、引用数、インパクトファクターなど）。
- これらの歪んだインセンティブは、盗作、ずさんな方法論、研究不正など多くの問題を引き起こし、科学と科学者に対する社会的な信頼の低下をもたらしている。

研究の中核的な知的作業の多くが
AIによって行われるようになったとき、
研究と研究者はどのように
評価されるべきか？

「研究評価に関するサンフランシスコ宣言」 (DORA) の抜粋

「個々の科学者の貢献を評価したり、採用、昇進、資金提供の決定を行う際に、ジャーナルインパクトファクターなどのジャーナルベースの指標を個々の研究論文の品質の代替指標として使用しないでください。」

「論文の科学的内容が出版基準や論文が掲載された雑誌のアイデンティティよりもはるかに重要であること」

「研究評価の目的で、研究出版物に加えて、すべての研究成果（データセットとソフトウェアを含む）の価値と影響を考慮し、政策と実践への影響など、研究の影響の質的指標を含む幅広い影響測定基準を考慮します」

「資金、雇用、終身在職権、昇進などを決定する委員会に参加する場合は、出版指標ではなく科学的内容に基づいて評価を行います。」

DORAへの日本の署名機関の例

- 日本学術振興協会
- 日本生物科学会連合（SEIKAREN）
- 国立研究開発法人医薬基盤・健康・栄養研究所
- 理化学研究所
- 東京大学※

※これまでにDORAに署名した唯一の日本の大学

DORA自体は、AI主導の研究が引き起こす課題に対処するには十分ではありません。

しかし、私たちがこれらの非常に困難で重要な問題にどのように立ち向かうべきかのモデルになります。

AIの進歩は科学研究に大きな課題をもたらしています。私たちは皆、これらの進歩を把握した上、科学界がどのように対応すべきかについて、あらゆるレベルで積極的に議論する必要があります。