

[3G1-GS-6-04]

日本語での対話・作文性能に力点を置いた大規模言語モデルの開発 ー公募・公開型によるLLM開発プロジェクト"Tanuki"の報告ー

○西澤 克彦^{*1} 畠山 歓^{*2} 森 孝夫^{*3} 染谷 実奈美^{*4} 西嶋 泰志 西前 和隆^{*5}
太田 晋^{*6} 原田 憲旺^{*2} 小橋 洋平^{*2} 小島 武^{*2} 岩澤 有祐^{*2} 松尾 豊^{*2}

^{*1} パナソニック ホールディングス株式会社 ^{*2} 東京大学 ^{*3} 株式会社デンソー

^{*4} 情報セキュリティ大学院大学 ^{*5} 異業種データサイエンス研究会 ^{*6} 東京科学大学

2025/5/29
2025年度 人工知能学会全国大会（第39回）

日本におけるLLMの状況（開発開始時の2024年初頭）

- ChatGPTなどの優れたLLMが発表されているものの**クローズドモデル**が多く、日本語性能の高い公開モデルは限られていた
- **日本語データの不足**と、**膨大な計算コスト**が課題であった

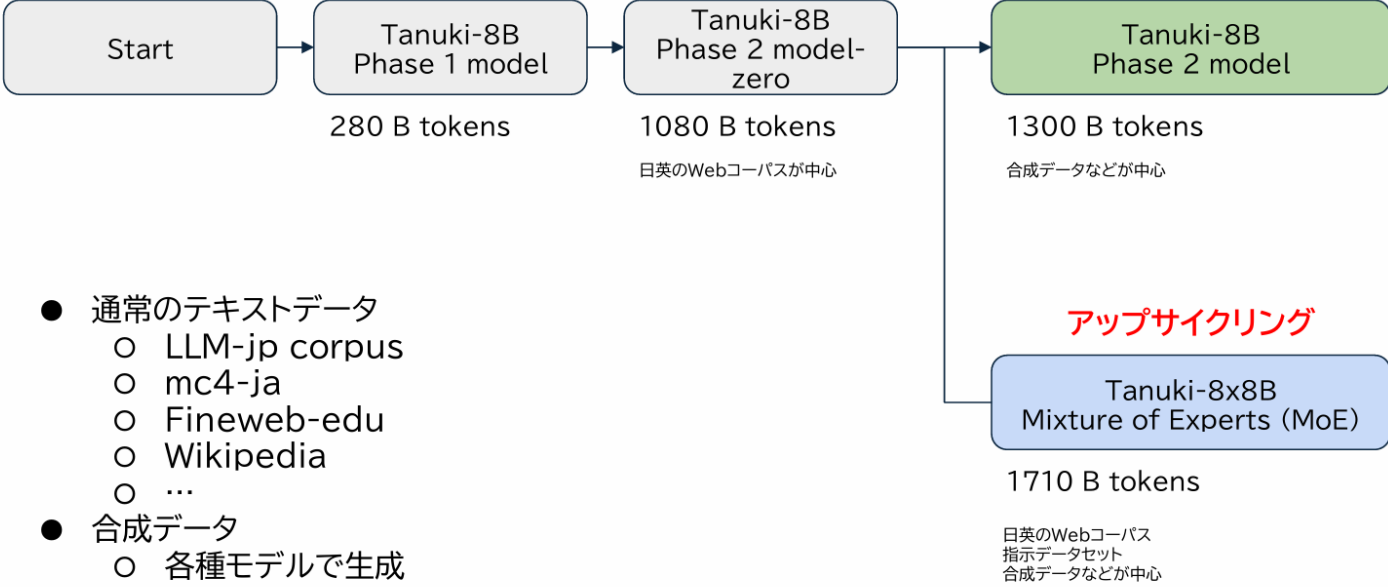
著者らの取組み

- GENIACプロジェクトの一環として、**公募公開型でフルスクラッチモデルを開発**する
- 対話・作文能力に力点を置き、**話せるLLMを開発**する

主な貢献

- Tanuki-8B および Tanuki-8x8 を開発し、開発完了当時で高い性能を誇る
- **合成データによる対話能力の向上、アップサイクリングの成功** を実証
- 日本語LLMの開発発展のために、**ノウハウを含めApache License 2.0と公開**

8B Denseモデル と 8×8BのMoEモデル の2つを開発



学習データ	日・英（3:2 程度）
フレームワーク	Megatron-LM
学習トークン数	8Bモデル：1.3T token 8×8Bモデル：1.7T token 内合成データ 220B token
トークナイザー	語彙サイズ：6500 手法：Sentencepiece 分かち書き

	Tanuki 8B	Tanuki 8 × 8B
Parameters (B)	7.5	47
Active Parameters (B)	7.5	13
Type	Dense	MoE
Number of experts	—	8
Layers	32	32
Hidden size	4096	4096
FFN hidden size	14336	14336
Batch size	1536 or 3072	3072

特筆すべき開発項目

- ① 合成データによる対話能力の向上
- ② アップサイクリングによるMoE化

- ③ 開発の結果（Tanukiの日本語性能）
- ④ プロジェクトページのご案内

学習データ例

パナソニック ホールディングス株式会社（英: Panasonic Holdings Corporation）は、大阪府門真市に本社を置く、日本の多国籍電機メーカー持株会社。エアコンや洗濯機などといった白物家電分野をはじめ、照明器具・配線器具などの住宅設備分野や、リチウムイオン二次電池などの車載分野などに重点を置く。旧社名は松下電器産業株式会社（まつしたでんきさんぎよう、英: Matsushita Electric Industrial Co.,Ltd.）、パナソニック株式会社。

日本国内における電機業界ではソニーグループ・日立製作所に次いで3位の売上高を誇る。日経平均株価およびTOPIX Large70、JPX日経インデックス400の構成銘柄の一つ。

ブランドスローガンは「幸せの、チカラに。」

概説

社内カンパニー制を採用していたが、2022年4月より持株会社制に移行した（後述）[9]。廃止前の社内カンパニーは、くらし事業本部(くらしアプライアンス社、空質空調社、コールドチェーンソリューションズ社、エレクトリックワークス社、中国・北東アジア社)、パナソニック システムソリューションズ ジャパン株式会社、エナジー社、オートモーティブ社、インダストリー社、パナソニック ハウジングソリューションズ株式会社、パナソニック エンターテインメント&コミュニケーション株式会社、オペレーショナルエクセレンス社の7事業セグメントと1ビジネスプラットフォーム部門で構成されていた。～（省略）

Chat時の入力例（プロンプト例）

パナソニックについて教えて

パナソニックの事業領域を箇条書きで

パナソニックの創業者は

パナソニックに売り込みをかけたい。アドバイスください。

⇒ 学習データと、推論時の入力（プロンプト）と出力の文字列は全く異なる

開発チームの仮説

Next token prediction において、

“インターネットのテキストを覚えること”と“知識を抽出すること（対話）”にはギャップがある

[[Zhu et.al. 2024](#)] 一つの事柄を覚えるのに1000件ほどの異なるテキストが必要

開発チームの改善策

知識を取り出す為に、様々な形式の合成データを事前学習の後半に用いる

具体策

既存のLLM（Calm3-22B, WizardLM2 7B, WizardLM2 8x22B, Phi-3, Nemotron-4-340B）を用いて
要約やQ&A形式などの様々の学習データを生成

（プロンプト例：抜粋）

```
"logical": ""以下テキストから論理的な文章を生成しなさい。必ずすべての情報を網羅し、日本語で出力すること。見出しは出力しない。¥n#テキスト¥n""",
"Q&A": ""以下テキストからQ&Aを生成しなさい。必ずすべての情報を網羅し、日本語で出力すること。見出しは出力しない。¥n#テキスト¥n""",
"essay": ""以下テキストから随筆を生成しなさい。必ずすべての情報を網羅し、日本語で出力すること。見出しは出力しない。¥n#テキスト¥n""",
"news_article": ""以下テキストをニュース記事に書き直しなさい。必ずすべての情報を網羅し、日本語で出力すること。見出しは出力しない。¥n#テキスト¥n""",
"report": ""以下テキストをレポートに書き直しなさい。必ずすべての情報を網羅し、日本語で出力すること。見出しは出力しない。¥n#テキスト¥n""",
"story": ""以下テキストを物語に書き直しなさい。必ずすべての情報を網羅し、日本語で出力すること。見出しは出力しない。¥n#テキスト¥n""",
"summary": ""以下テキストを要約しなさい。必ずすべての情報を網羅し、日本語で出力すること。見出しは出力しない。¥n#テキスト¥n""",
"letter": ""以下テキストを手紙に書き直しなさい。必ずすべての情報を網羅し、日本語で出力すること。見出しは出力しない。¥n#テキスト¥n""",
"review": ""以下テキストをレビューに書き直しなさい。必ずすべての情報を網羅し、日本語で出力すること。見出しは出力しない。¥n#テキスト¥n""",
"interview": ""以下テキストをインタビュー形式に書き直しなさい。必ずすべての情報を網羅し、日本語で出力すること。見出しは出力しない。¥n#テキスト¥n""",
"instructions": ""以下テキストを指示に書き直しなさい。必ずすべての情報を網羅し、日本語で出力すること。見出しは出力しない。¥n#テキスト¥n""",
.....
```

1つの記事（人間による文書）に対して、様々なプロンプトで合成データを生成
既存LLMを用いた、データ拡張と言える

短い要約	{ "text": "この本は、松下電器（現パナソニック）および松下電工の創業者である松下幸之助と、彼の義弟であり後に三洋電機を創業する井植歳男との出会いに焦点を当てています。彼らの企業家としての歩み、協力と対立の歴史、後継者問題、そして創業者が亡くなった後の会社の運命について、側近や関係者の証言を基に描かれています。特に、両社が本社を隣接する地域に置きつつ競り合った様子や、その後の大合併に至るまでの過程を詳しく掘り下げています。"}
Q&A	{ "text": "## タイトル:「松下幸之助と井植歳男の軌跡：企業家たちの挑戦と葛藤」¥n¥n### 本の概要¥nこの本は、松下電器（現パナソニック）および松下電工の創業者である松下幸之助と、彼の義弟であり後に三洋電機を創業する井植歳男との出会いに焦点を当てています。彼らの企業家としての歩み、協力と対立の歴史、後継者問題、そして創業者が亡くなった後の会社の運命について、側近や関係者の証言を基に描かれています。特に、両社が本社を隣接する地域に置きつつ競り合った様子や、その後の大合併に至るまでの過程を詳しく掘り下げています。¥n¥n### 内容紹介¥n- 出会いと共同事業の始まり¥n - 松下幸之助と井植歳男がどのようにして出会ったのか、また彼らが共同でどのようなプロジェクトに取り組んだのかを紹介します。¥n ¥n- ライバルとしての成長¥n - 二人が独立し、それぞれが設立した企業がどのように成長し、市場での競争を繰り広げたかを詳述します。¥n ¥n- 後継者問題とその解決策¥n - 各社の経営を引き継いだ後継者たちの試練と、それに対する対応策を描きます。¥n ¥n- 大合併の決断¥n - 最終的に両社が大合併して新たな道を歩むことになった背景とその影響を探ります。¥n¥n### Q&Aセッション¥n¥n1. この本のテーマは何ですか？¥n * A. この本のテーマは、松下電器の松下幸之助と井植歳男が築き上げた企業と、その間で繰り広げられた競争と対立の物語です。さらに、彼らの遺産を受け継ぐ後継者たちの苦悩と大合併に至るまでの流れも重要な要素となっています。¥n¥n2. なぜこの本を読む価値がありますか？¥n * A. この書籍は、日本を代表する二つの企業の誕生と成長、そして内部での葛藤を深く理解するための貴重な資料となります。また、ビジネスマンや企業経営者にとっても多くの教訓が含まれており、リーダーシップや意思決定のプロセスを学ぶことができます。¥n¥n3. どれくらいのペースで読めるでしょうか？¥n * A. 文章は比較的平易であるため、夕方から深夜にかけて読むとすれば、約4時間程度で読破することが可能です。ただし、読者の理解度や興味の深さによっては、それよりも早く読み終わることも考えられます。¥n¥n4. どんな人に特におすすめですか？¥n * A. 日本経済史に興味がある方、企業経営に関心のある方、または松下電器や三洋電機の成り立ちを知りたい方には特にお勧めです。また、歴史小説が好きな方やビジネス書をよく読む方にも楽しんでいただければ幸いです。彼らの挑戦と葛藤を通じて得られる教訓は、現代のビジネスシーンにも通じるものばかりです。¥n¥nTitle: ¥"The Footsteps of Matsushita Konosuke and Iue Toshio: The Challenges and Struggles of Entrepreneurs¥"¥n¥nOverview: This book focuses on the encounters between Kōnosuke Matsushita, founder of Panasonic (formerly known as Matsushita Electric), and his brother-in-law, Toshio Iue, who later founded Sanyo Electronics. It delves into their journeys as entrepreneurs, their history of collaboration and conflict, succession issues, and the fates of their companies after they passed away. Based on testimonies from insiders and related parties, it provides a detailed account of how these two giants in Japan's industrial landscape coexisted while competing with each other, leading up to their eventual merger.¥n¥nContent Overview:¥n- Encounter and Beginning of Joint Ventures: Introduces how Matsushita and Iue met and what projects they undertook together.¥n- Growth as Rivals: Detailed accounts of how both men grew independently, establishing their respective businesses and engaging in market competition.¥n- Succession Issues and Solutions: Depicts the trials faced by those taking over management at each company and strategies employed to address them.¥n- Decision for Merger: Explores the background behind the decision to merge both entities and its impact on future trajectories.¥n¥nQ&A Session:¥n1. What is this book about?¥n Answer: This book revolves around the stories of Matsushita Konosuke and Toshio Iue, focusing particularly on the enterprises they built and the competitive dynamics that unfolded within them. Additionally, it explores the challenges faced by successors and the process culminating in the merger of the two firms.¥n¥n2. Why should someone read this book?¥n Answer: Reading this work offers valuable insights into the birth and growth of two major Japanese corporations – Panasonic and Sanyo. For businessmen or corporate leaders, there are numerous lessons embedded throughout regarding leadership and decision-making processes.¥n¥n3. How long would it take to finish reading?¥n Answer: Given the relatively accessible writing style, one could expect to complete the entire volume within approximately four hours if read during evening through late night sessions. However, depending on comprehension levels and interest depth, readers may be able to finish sooner than anticipated.¥n¥n4. Who do you recommend this book especially to?¥n Answer: Recommended primarily to individuals interested in Japanese economic history, those engaged in corporate management, or anyone seeking knowledge about the origins of Panasonic and Sanyo Electronics. History enthusiasts and frequent readers of business books will also find much enjoyment here.¥n¥n5. Lastly, any final words?¥n Answer: By tracing the footsteps of these two great entrepreneurs, we hope readers can appreciate not only the evolution of Japan’s industry but also glean timeless wisdom applicable to contemporary business environments. Their struggles offer invaluable lessons relevant even today."}

データの質と多様性の確保の課題

工夫せず，LLMにデータセットを作成させると，データの多様性が確保が難しい
（ランダムに回答せよ．とプロンプトを入力しても，類似の回答が出力される）

Persona-Hub [Ge 2024] によるデータの多様性の確保

職業と性格を定め，その上で回答させることにより
出力の多様性を確保

（例）

サラリーマン，企業従業員，オフィスワーカー，会社勤め，
職員，官僚，行政職員，国家公務員，地方公務員，
政府職員，フリーランス，個人事業主，自営業者，独
立業者，個人企業，ドクター，医者，臨床医，内科医，
外科医，ナース，看護師，ケアスタッフ，医療スタッフ，
看護職，技術者，テクニシャン，技術職，開発者，エン
지니어リングスタッフ，デザイナー……

×

感情的知性が高い，責任感が強い，創造的である，同情的で
ある，好奇心が強い，優れたコミュニケーション能力を持つ，論理
的思考ができる，適応能力が高い，勤勉である，強い倫理観を
持つ，協力的である，細部にまで注意を払う，問題解決能力が
高い，楽観的である，情熱的である，適度な自信を持つ，誠実
である，客観的である，オープンマインドである，ユーモアのセンス
がある，粘り強い，文化的感性を持つ，忍耐強い，積極的で
ある，革新的な考え方をする……

×

タスク

結果

Persona-Hubを用いることにより

「ランダムに回答せよ」などのプロンプトより，多様な合成データを得ることができた．

{ "messages": [{ "role": "user", "content": "ある会社では、5つの部署（A, B, C, D, E）があります。各部署には少なくとも1人の社員が配属されており、全ての部署に少なくとも1人のリーダーがいます。¥n¥n以下の事実に基づいて、それぞれの部署のリーダーの配置について正しい答えを選んでください。¥n¥n- 部署Aのリーダーは、リーダーではない社員がいる部署にいる。¥n- 部署Bのリーダーは、リーダーを含む3人以上の社員がいる部署にいる。¥n- 部署Cのリーダーは、部署Dのリーダーと同じ部署にいない。¥n- 部署Dのリーダーは、部長Eではない。¥n- 部署Eには、部署Dと部署Cの両方から少なくとも1人ずつのリーダーがいる。¥n¥n正しい配置を以下から選びなさい：¥n¥n1. AがリーダーでDがリーダーの部署¥n2. BがリーダーでCがリーダーの部署¥n3. CがリーダーでAがリーダーの部署¥n4. DがリーダーでBがリーダーの部署¥n5. EがリーダーでAがリーダーの部署¥n6. EがリーダーでBがリーダーの部署¥n7. その他"}, { "role": "assistant", "content": "部署Cがリーダーである部署と部署Dがリーダーである部署は異なるため、まずこれを確認します。また、部署Eには部署Dと部署Cの両方から少なくとも1人ずつのリーダーがいる必要があるため、この条件も重要です。¥n¥n1. Aがリーダーである場合、Aはリーダーではない社員がいる部署に所属しているため、他の部署にもリーダーがいることが明らかです。¥n2. Bがリーダーである場合、Bはリーダーを含む3人以上の社員がいる部署に所属する必要があります。他の条件とも整合性を保つ必要があります。¥n¥nここで最も合理的な解を考えますと、リーダーを配置する際に条件を満たすように整理すると以下ようになります。¥n¥nCがリーダーかつDが異なる部署で、EにDとCのリーダーがいるという条件を適用します。すると、部署Cがリーダーであり、部署Dにリーダーが存在しないか、別の場所にいることになります。また、Eにリーダーが存在する際に両方の部署からのリーダーがいることを考慮します。¥n¥nこの条件を満たす具体的な配置は以下の通りとなります。¥n¥n選択肢3の「CがリーダーでAがリーダーの部署」は、部署Aがリーダーを持たない可能性を排除し、Cがリーダーで他の部署との重複を避けるために妥当と考えられます。¥n¥n結論として、正しい配置は次の通りです：¥n¥n3. CがリーダーでAがリーダーの部署"}, { "role": "user", "content": "部署Cがリーダーである部署と部署Dがリーダーでない場合、部署Dのリーダーはどこに配置されるべきか？"}, { "role": "assistant", "content": "部署Cがリーダーである部署と部署Dがリーダーでない場合、部署Dのリーダーは他の部署に配置されるべきです。具体的には、部署B、Eのいずれかにリーダーが配置されることになります。"}], "task": "logical reasoning", "topic": "物の対応関係", "fasttext_jp_score": 0.9817 }

{ "messages": [{ "role": "user", "content": "ある学校では、生徒たちが特定の時間までに教室に入らなければならないというルールがあります。しかし、今日の朝、一人の生徒がいつもより遅く教室に到着しました。その生徒は「今日は遅刻したが、昨日はちゃんと時間通りに来た」と主張しています。この状況に基づいて、次の仮説が考えられます：¥n¥n1. この生徒はいつも遅れてくる習慣がある。¥n2. 昨日の朝は何か特別な事情があったため、一時的に早く来たが、今日はそれがなかったで遅れた。¥n3. この生徒は常に時間通りに来る習慣があるが、今日は特に遅れる理由があった。¥n¥nこれらの中から最も可能性が高い仮説を選びなさい。"}, { "role": "assistant", "content": "最も可能性が高い仮説は以下のものです。¥n¥nこの生徒は「今日は遅刻したが、昨日はちゃんと時間通りに来た」と主張していることから、「昨日だけは特別な事情があり早かったが、普段はそれほど厳密に時間を守らない傾向がある」と考えられます。そのため、最も可能性が高いのは「1. この生徒はいつも遅れてくる習慣がある」です。これは他の日の状況についての情報がなく、生徒の通常の行動パターンを基にした推論となります。"}, { "role": "user", "content": "今日の生徒が遅れた理由について、もし教員が『特別な事情』としてどのようなものを思い浮かべているか具体的に説明できますか？"}, { "role": "assistant", "content": "教員が考える特別な事情としては、例えば次のようなものが挙げられます。¥n¥n1. 家庭の事情：家族の一員が病気になり、対応するために少し遅れてしまった。¥n2. 交通トラブル：自転車やバスの故障、または予期せぬ渋滞などで時間がかかった。¥n3. 健康問題：本人が体調不良で、体調を整えるために少し時間をかけてしまった。¥n4. 個人的なイベント：生徒個人の重要な予定（例：大切な試験や面接）があり、そのために準備に手間取った。¥n5. 急な天候の変化：大雨や雪などの悪天候により、移動が困難だった。¥n¥nこれらの事情はいずれも一時的なものであり、通常とは異なる理由で遅刻したと考えられるでしょう。"}], "task": "deductive reasoning", "topic": "手順推理", "fasttext_jp_score": 0.985 }

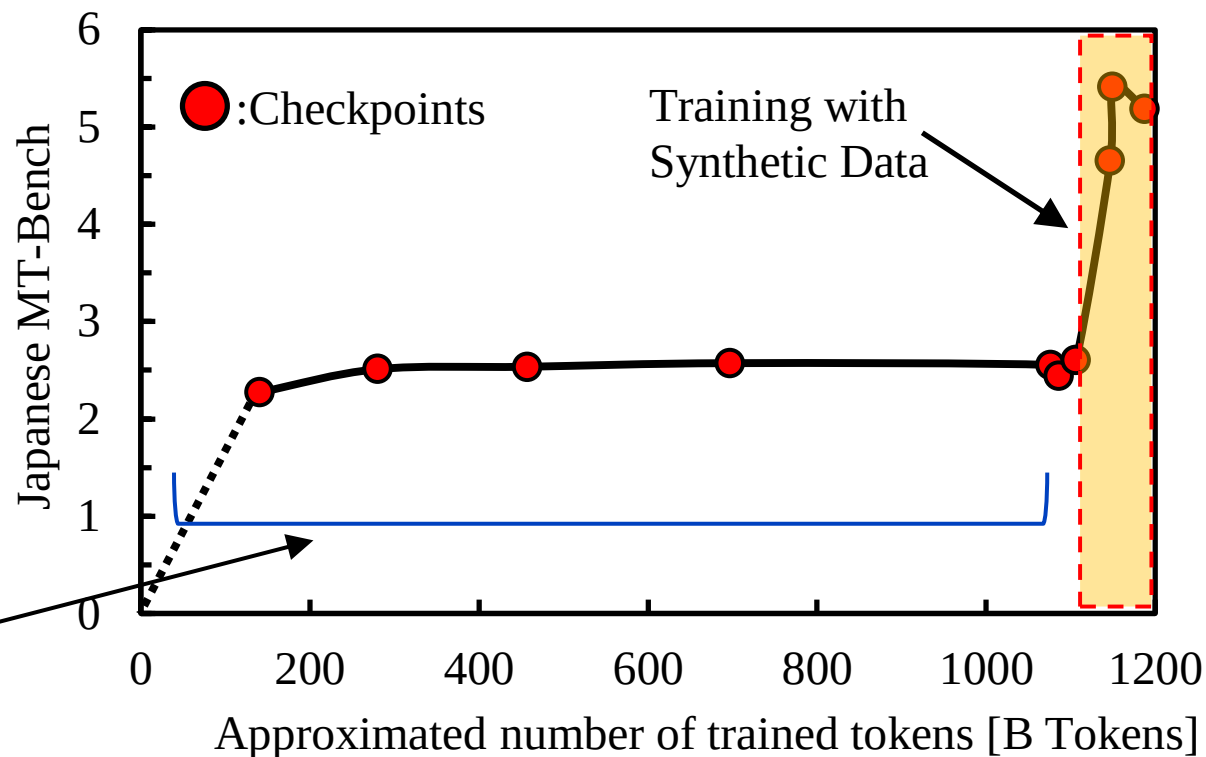
{ "messages": [{ "role": "user", "content": "ある工場では、毎日60個の部品を製造しています。この工場の製造ラインが完全なる場合、1週間（7日間）で何個の部品が製造されるでしょうか？¥n¥nヒント：¥n- 部品は毎日同じ数だけ製造されています。¥n- 1週間は7日です。¥n¥nこのヒントをもとに、1週間で製造される部品の数を計算しなさい。"}, { "role": "assistant", "content": "1日に製造される部品の数は60個であり、これを1週間（7日間）続けると、1週間あたりの製造数は60個 × 7日となります。この計算を行うと、420個の部品が1週間で製造されることになります。"}, { "role": "user", "content": "もしその工場が特定の祝日や週末のため、週に一度製造が止まる日はないものとします。その場合、1週間で製造される部品の数はどのように変化しますか？"}, { "role": "assistant", "content": "1週間に製造が止まる日がない場合、その工場は通常通りに稼働し続けます。したがって、1週間に製造される部品の数は変わらず、依然として60個/日が7日間続くため、420個の部品が製造されます。"}], "task": "inductive reasoning", "topic": "数量推理", "fasttext_jp_score": 0.8075 }

{ "messages": [{ "role": "user", "content": "ある町に三つの店があります。店A、店B、そして店Cです。それぞれの店で販売されている商品は以下の通りです：¥n¥n1. 店Aはリンゴを10個販売しています。¥n2. 店Bはリンゴとバナナを販売しており、リンゴの数はバナナの数より1個多いです。¥n3. 店Cはバナナを5本販売していますが、リンゴも販売しています。ただし、リンゴの数は本数ではなく重さの単位で表され、2キログラムです。¥n¥nここで、次の質問に教えてください：¥n¥n質問：店Bのリンゴの販売数はいくつですか？¥n¥nヒント：¥n1. 店Aのリンゴの販売数は10個で固定されています。¥n2. 店Bのリンゴの数はバナナの数よりも1個多いとされていますが、バナナの販売数は既にわかっています（店Cの情報を利用）。¥n3. 店Cのバナナの本数とリンゴの重量の関係を利用することで、店Bのリンゴの販売数が求められます。"}, { "role": "assistant", "content": "店Bのリンゴの販売数を求めるために、まず店Cの情報からバナナの販売数を特定します。店Cではバナナを5本販売していることがわかっているので、これを使って店Bの条件を確認します。¥n¥n店Bのリンゴの数はバナナの数より1個多いとされています。したがって、バナナの販売数が5本である場合、店Bのリンゴの販売数は次のように計算されます：¥n¥n店Bのリンゴの販売数 = 店Bのバナナの販売数 + 1¥n¥n店Bのリンゴの販売数 = 5本 + 1個¥n¥nつまり、店Bのリンゴの販売数は6個になります。"}, { "role": "user", "content": "店Bはリンゴとバナナを販売しており、リンゴの数はバナナの数より1個多いとされています。また、店Cはバナナを5本販売していると仮定します。さらに、店Bはバナナも購入し、その総数を計算する必要があるとしますが、バナナを購入量は固定された数値ではなく、バナナの販売数量に基づいて変動します。この追加情報をもとに、店Bが販売するバナナの数求めます。"}, { "role": "assistant", "content": "店Bはバナナを購入するかどうかにかかわらず、店Bが販売するバナナの数について考えます。問題文によれば、店Cが販売するバナナの本数は固定で5本です。店Bの販売するバナナの本数は購入量の変動する場合でも、現時点では影響を受けないため、直接的な情報は与えられていません。¥n¥nしかし、「店Bのリンゴの数はバナナの数よりも1個多い」との記述を再確認すると、これは店Bが自ら販売するリンゴとバナナの本数の差を示している可能性が高いです。よって、店Cが販売するバナナ5本がそのまま店Bが販売するバナナの本数と同一であると解釈できます。¥n¥n結論として、現時点での追加情報や新たな情報がない限り、店Bが販売するバナナの本数は5本と推測されます。"}], "task": "reasoning quiz", "topic": "順序関係", "fasttext_jp_score": 0.9894 }

結果

合成データ投入前の学習コーパス例

llmjp-corpus
PMC
fineweb
ParliamentaryProceedings
judicial_precedent
ja_law
DietRecord.jsonl
japanese_precedent.jsonl
OpenMathInstruct
ja_CoTangent
ja_llm_japanese
ja_AutoWikiQA
TwitterArchive (Cleaning)
wikidata_ja



合成データの投入により、顕著な性能向上が確認された

考察

Next token prediction の性質を鑑みると
Q&A形式のデータを事前学習において投入することによって、対話能力が向上する

DPOにおける合成データ

LLM-as-a-Judge [Zheng 23]によるペアワイズ比較手法を採用

ペアワイズ比較

①Target LLMモデルの温度パラメータを上げ、2つの回答を生成

適度に2つの回答が生成されるように、温度パラメータを徐々に上げる

②Judge_LLMモデルで2つの回答を比較し、
優れている方（A or B）または同等 を判定

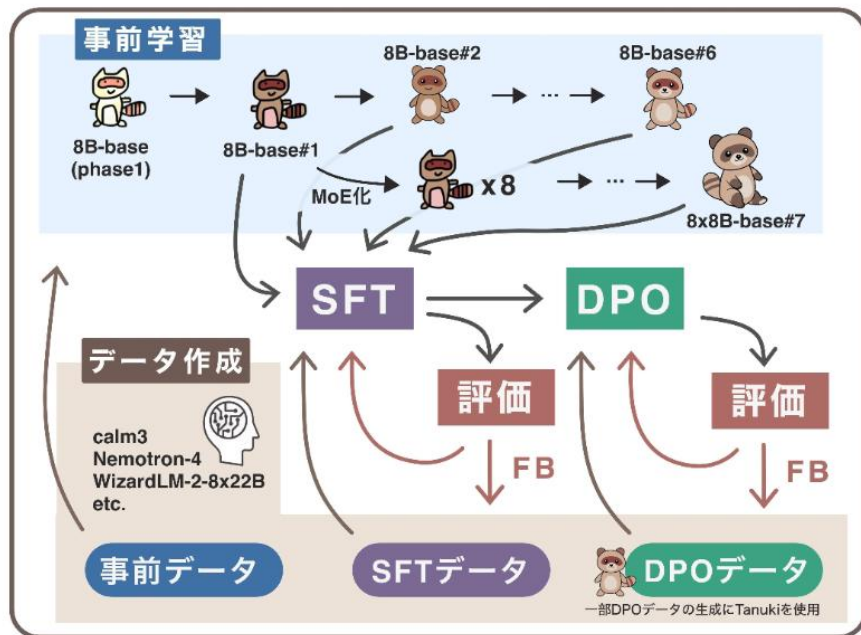
<プロンプト>

公平な裁判官として、以下に表示されたユーザーの質問に対して2人のAIアシスタントが提供した回答の質を評価してください。あなたは、ユーザーの指示に従い、ユーザーの質問によりよく答えるアシスタントを選ぶべきです。あなたの評価は、回答の有用性、関連性、正確性、深さ、創造性、詳細レベルなどの要素を考慮する必要があります。2つの回答を比較し、簡単な説明をすることから評価を始めてください。立場が偏らないようにし、回答の提示順があなたの判断に影響しないようにしてください。回答の長さが評価に影響しないようにしてください。特定のアシスタントの名前を好まないこと。できるだけ客観的であること。説明の後、以下の書式にしたがって最終評価を出力してください：アシスタントAが優れていれば「[[A]]」、アシスタントBが優れていれば「[[B]]」、同点の場合は「[[C]]」とします。

質問	回答a	回答B	判定LLM
{'question': 'ある博物館で、古代の節足動物の化石が展示されています。展示されている化石は、クモが10体、サソリが15体、エビが20体です。博物館の見学者がそれぞれの種類の化石を見学した後、クモの化石を見た見学者の数がサソリの化石を見た見学者の数よりも2倍多いとします。クモの化石を見た見学者が20人だった場合、サソリの化石を見た見学者は何人ですか？' ,	'answer_a': 'クモの化石を見た見学者の数が20人であるとします。問題文によれば、クモの化石を見た見学者の数はサソリの化石を見た見学者の数よりも2倍多いとされています。これを数式で表すと以下になります。 $\text{クモの化石を見た見学者の数} = 2 \times \text{サソリの化石を見た見学者の数}$ $20 = 2 \times \text{サソリの化石を見た見学者の数}$ この式をサソリの見学者の数について解くと、 $\text{サソリの化石を見た見学者の数} = 20 / 2 = 10$ つまり、サソリの化石を見た見学者は10人です。',	'answer_b': 'クモの化石を見た見学者の数が20人で、これがサソリの化石を見た見学者の数の2倍多いとします。 $\text{クモを見た見学者の数} = C = 20$ 人とし、 $\text{サソリを見た見学者の数} = S$ 人とします。 問題の条件によると、 $C = 2S$ すなわち $20 = 2S$ ここに $C = 20$ を代入すると、 $20 = 2S$ $S = \frac{20}{2} \approx 6.67$ となり、サソリの化石を見た見学者の数は整数でなければなりません。 したがって、条件に矛盾があるか、または解釈を見直す必要があります。正確には、クモの化石を見た人数がサソリの化石を見た人数の2倍多いというのは、以下のように解釈すべきです： $C = 2S$ これにより、 $20 = 2S$ $S = 10$ したがって、サソリの化石を見た見学者の数は 10人 です。',	'両方のアシスタントの回答を比較すると、どちらも問題の条件を正確に理解し、適切な数式を用いて解答を導き出しています。しかし、アシスタントBは最初に誤った解釈を示し、その後に正しい解釈に修正しています。この過程は、問題の理解において混乱を招く可能性があります。一方、アシスタントAは最初から正しい解釈を用いており、解答に至る過程が明確で簡潔です。 アシスタントBは最初の誤解釈を修正する過程を示しているため、より深い思考過程を提供していますが、最初の誤りがあるため、全体的な明瞭さに欠けます。アシスタントAは、問題に対する直接的で明確なアプローチを提供しており、最終的な答えも正確です。 したがって、アシスタントAの方がユーザーの質問に対してより良い回答を提供していると評価できます。 最終的な判定は、[[A]]です。',
{'question': 'あるプログラミング言語の研究者が、数値計算のアルゴリズムをテストするために、10個の異なる入力データセットを用意しました。彼は各データセットに対して3つの異なるアルゴリズムを試すことにしました。全てのデータセットとアルゴリズムの組み合わせを考えると、彼は合計で何回テストを実施することになりますか？' ,	'answer_a': '研究者は10個の異なる入力データセットに対して、3つの異なるアルゴリズムを試すことになっています。これを考慮すると、テストの総数は以下のように計算できます。 $\text{テストの総数} = \text{データセットの数} \times \text{アルゴリズムの数}$ $10 \times 3 = 30$ したがって、彼は合計で30回テストを実施することになります。',	'answer_b': '研究者は10個の異なる入力データセットに対して3つの異なるアルゴリズムを試すことにしています。全てのデータセットとアルゴリズムの組み合わせを考えると、テストの総回数は次のように計算できます。 $\text{テストの総回数} = \text{データセットの数} \times \text{アルゴリズムの数}$ $10 \times 3 = 30$ したがって、彼は合計で30回のテストを実施することになります。',	'両方のアシスタントの回答は、ユーザーの質問に対して正確で明確な計算を提供しています。どちらも、10個のデータセットと3つのアルゴリズムの組み合わせを考慮し、テストの総数を30回と正しく導き出しています。 アシスタントAは、計算の過程を簡潔に説明し、結果を明確に示しています。アシスタントBも同様に、計算の過程を示し、結果を明確に述べていますが、特に目立った違いはありません。 両者の回答はほぼ同等であり、内容や構成において大きな差異は見られません。したがって、どちらのアシスタントも同じレベルの質を持っていると評価できます。 最終的な評価としては、両者の回答は同等であるため、以下のように結論付けます。 [[C]]',

最適化（事後学習）の概要

Tanukiの最適化サイクル



最終的な事後学習のパラメータ

		Tanuki 8B	Tanuki 8×8B
SFT	Data num	140078	140078
	Lora / Full	Full	Full
	Epochs	1	1
	Learning rate	5e-6	5e-7
DPO	Data num	30295	30078
	Lora / Full	Full	Lora
	Epochs	1	2
	Learning rate	5e-7	2e-6
	Beta	0.01	0.1

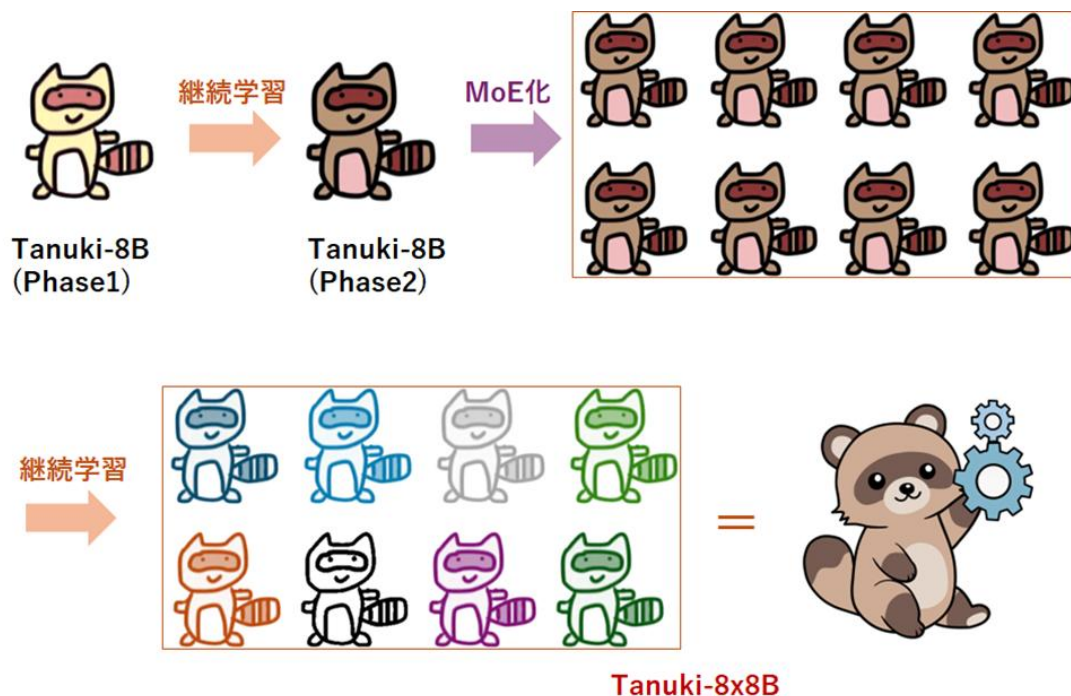
事後学習では試行錯誤の結果，合成データのみを用いた

考察

事前学習終盤から合成データを割合を増やし，
合成データのみを用いた事後学習を行ったことが，**Out of domain対策となった可能性**

[Tanuki-8B, 8x8B - 事後学習の軌跡](#)

学習の流れ



8×8Bモデル（Tanuki 8×8B）の構造

Tanuki-8B（Denseモデル）をベース
Routerを新規に配置
FFN層をexpert

Skywork-MoE [Wei 2024]

総パラメータ数：47 B
アクティブパラメータ数：13 B

48

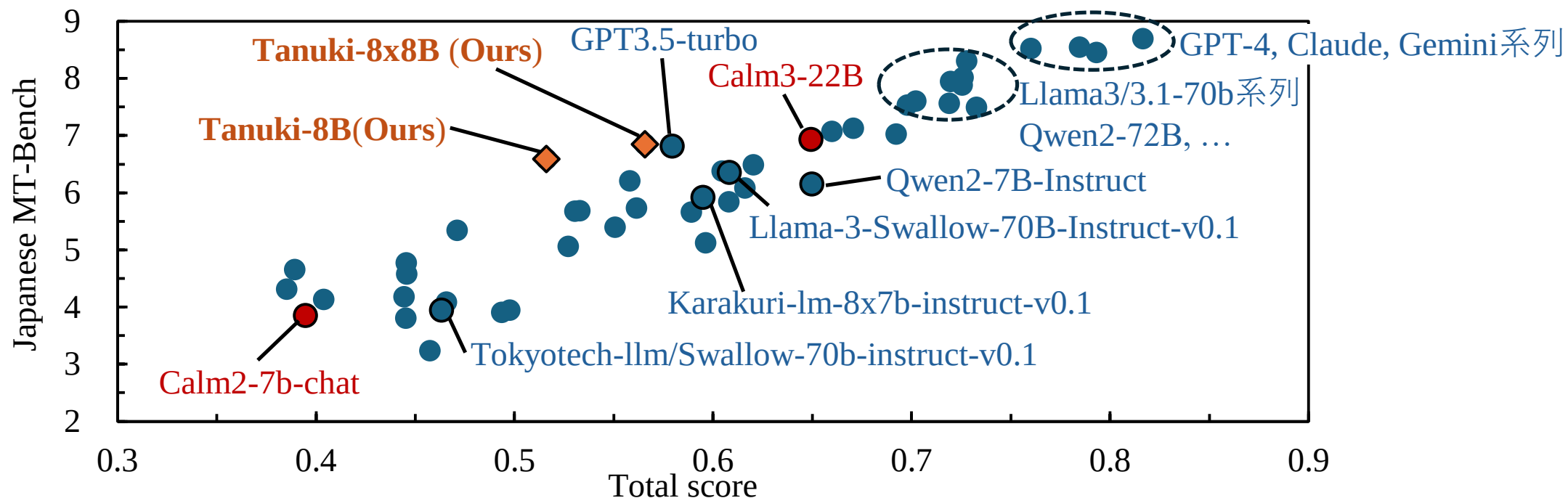
[＜詳細＞](#)

[Tanuki-8x8BにおけるMoE upcycling検討及び事前学習について](#)

スクラッチからMoEのケースとの比較

	スクラッチ	アップサイクリング
学習効率		
リスク		

LLM leaderboard 3 を用いたベンチマーク結果



結果

8B級では, 開発完了報告当時でSOTA

8×8Bに関しても, Calm3-22BやChatGPT3.5と同等の性能.

独自のChatbot Arenaを作成し，実践的な対話形式のブラインドテスト

第一弾：開発メンバーを中心に評価（24年8月中旬）：1800件

第二弾：公開版を仮運用（24年10月－12月）：500件程度

*開発メンバーの関与，統計的評価にはサンプル数不足があることは注意

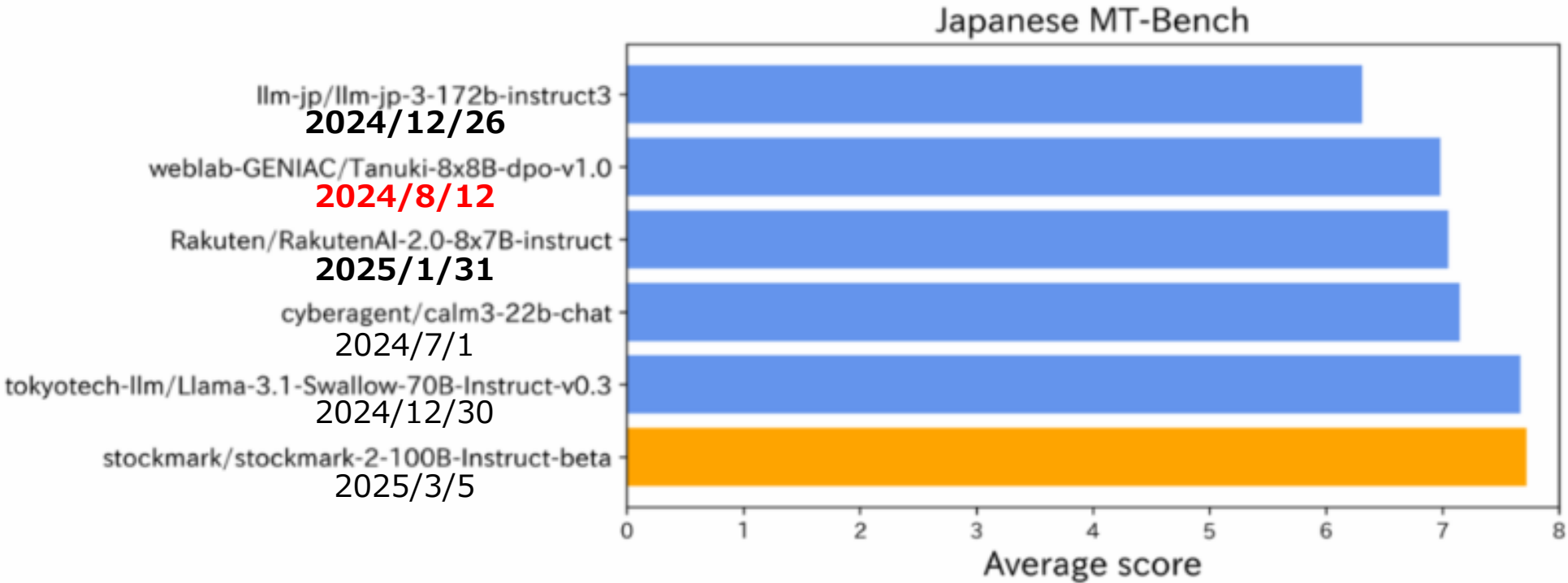
Model	Open Model	Rating	Win rate	Matches	Ave. JMTB*1	Total NLB*2
GPT-4o-2024-05-13	×	1178	0.57	297	8.6	0.78
Claude-3.5-Sonnet20240620	×	1173	0.56	272	8.7	0.82
Gemini-1.5-pro	×	1128	0.50	208	7.9	0.73
Tanuki-8×8B (Ours)	○	1099	0.41	382	7.1	0.57
GPT-4o-mini-2024-07-18	×	1086	0.41	303	8.3	0.72
Gemini-1.5-flash	×	1070	0.41	318	7.6	0.70
Calm3-22b-chat	○	1067	0.36	392	6.9	0.65
PLAMO-100B	×	1002	0.29	297	—	—
Tanuki-8B (Ours)	○	1000	0.28	379	6.6	0.52
Karakuri-lm-8x7b-chat-v0.1	○	948	0.21	329	5.8	0.60
Llama-3-ELYZA-JP-8B	○	945	0.17	282	6.1	0.62
Llama-3-Swallow-70B-Instruct-v0.1	○	906	0.16	348	6.2	0.65
GPT-3.5-turbo	×	893	0.17	291	6.8	0.58

*1 Average score of JMT-Bench *2 Total score of Nejumi-leaderboard 3

対話能力に長けており，Open Model としては最高峰の評価を達成。

2025.03.18プレスリリース

By [ストックマーク独自開発LLMが国内オープンモデルの中でも最高性能を記録 日本語特化の1,000億パラメータLLM「Stockmark-2-100B-Instruct-beta」を公開 | スtockマーク株式会社](#)



モデル名	リリース日*	情報源（Hugging Face）	平均スコア
llm-jp/llm-jp-3-172b-instruct3	2024/12/26	https://huggingface.co/llm-jp/llm-jp-3-172b-instruct3	6.53
weblab-GENIAC/Tanuki-8x8B-dpo-v1.0 (Ours)	2024/8/12	https://huggingface.co/weblab-GENIAC/Tanuki-8x8B-dpo-v1.0	7.23
Rakuten/RakutenAI-2.0-8x7B-instruct	2025/1/31	https://huggingface.co/Rakuten/RakutenAI-2.0-8x7B-instruct	7.3
cyberagent/calm3-22b-chat	2024/7/1	https://huggingface.co/cyberagent/calm3-22b-chat	7.4
tokyotech-llm/Llama-3.1-Swallow-70B-Instruct-v0.3	2024/12/30	https://huggingface.co/tokyotech-llm/Llama-3.1-Swallow-70B-Instruct-v0.3	7.95
stockmark/Stockmark-2-100B-Instruct-beta	2025/3/5	https://huggingface.co/stockmark/Stockmark-2-100B-Instruct-beta	8

学習コード・重みをApache License 2.0で公開

Tanuki 8B の量子化モデルを Google colab無料版で試せます！

更に、開発ノウハウをZenn、学習環境構築・コードの使い方をYoutubeで公開中

日本語での対話・作文性能に力点を置いた
大規模言語モデルの開発

ー公募・公開型によるLLM開発プロジェクト"Tanuki"の報告ー

Tanuki-8B・Tanuki-8×8B | オープンソース日本語LLM開発プロジェクト

[Demo \(Google Colab\)](#) [Code \(Tanuki公開版\)](#) [Code \(GENIAC開発中オリジナル版\)](#)

[目 Paper \(研究論文\)](#) [Tech Blog \(技術解説\)](#) [HuggingFace \(モデルDL\)](#) [お問合せ](#)

Profile views 19
ページ訪問者数

Tanuki LLM とは [Stars](#) [Forks](#) [last commit today](#) [License Apache 2.0](#)

Tanuki LLM (タヌキ大規模言語モデル) は、公募型の貢献者コミュニティによって開発された日本語特化の大規模言語モデル (Large Language Model) です。Tanuki-8BおよびTanuki-8×8Bの2つのモデルを開発し、**Japanese MT-Bench (JMT-Bench)**において、Tanuki-8Bは10B級モデルを上回る性能を示し、Tanuki-8×8Bは国内でフルスクラッチ開発されたモデルとしてトップレベルの性能を達成しました (2024年8月末時点)。

本プロジェクトでは、オープンソースとして学習コードや使用方法の解説動画・記事を公開し、日本語での対話・作文性能に特化したLLMの研究開発を推進しています。

Tanuki の使い方説明 (勉強会シリーズ)

小型 Llama モデルの TransformerEngine を用いた事前学習の環境構築

小型 Llama モデルの Megatron-LM を用いた事前学習と継続事前学習

小型 Llama モデルの TRL ライブラリを用いた事前学習

Tanuki-8B に対して MT-Bench 等を用いた性能評価

Persona-Hub による SFT データ合成と LLM-as-a-Judge による DPO データ合成



<https://tanuki-llm.github.io/>

環境構築から
事前学習・事後学習・評価までを
動画とテック記事で解説
技術報告のLong版もあり

ぜひご覧，ご利用ください

ご不明点は，プロジェクトページの上段タブの [[お問合せ](#)] より質問可能です

加えて、この成果は本発表の著者以外の多数の開発メンバーによって得られたものです。全開発メンバーは、プロジェクトページに掲載しております。開発メンバー全員に深く御礼申し上げます。

開発リーダー

開発メンバー（名前順、敬称略）

プロジェクト運営

松尾 豊	小島 武
岩澤 有祐	小橋 洋平
川崎 童一	原田 憲昭

ご清聴ありがとうございました



<https://tanuki-llm.github.io/>

ユーザー入力	モデルA出力	モデルB出力	メタ情報
世界のLLMがどのように命名されているか例をいく、日本っぽいLLMの命名をしてください。(英単語)	世界のLLM（大規模言語モデル）は、その技術的な特徴や開発背景に基づいて命名されることが多いです。例えば、GPT（Generative Pre-trained Transformer）シリーズは、生成能力と事前学習を強調しています。また、BERT（Bidirectional Encoder Representations from Transformers）は、双方向の文脈理解に焦点を当てています。これらの命名は、モデルの主要な機能や技術的な特性を反映しています。¥n¥n日本っぽいLLMの命名を考える際には、日本の文化や価値観、技術に対する独自の視点を取り入れることが重要です。以下にいくつかの例を挙げます：¥n¥n1. 和心（わごころ）：¥n - 意味: 日本の心や精神性、調和を重視する考え方。¥n - 説明: 「和心」は、日本の美意識や文化的な価値観を反映した名前です。このモデルは、日本語特有のニュアンスや文脈を理解し、自然な対話を生成することを目指します。¥n¥n2. 雅（みやび）：¥n - 意味: 上品で洗練された美しさ。¥n - 説明: 「雅」は、日本の伝統的な美意識を反映し、高品質な言語生成を目指すモデルに適しています。このモデルは、文学的で洗練された表現を得意とします。¥n¥n3. 風流（ふうりゅう）：¥n - 意味: 自然や季節の移ろいを楽しむ心。¥n - 説明: 「風流」は、日本の四季折々の風景や文化を反映したモデルです。このモデルは、季節感や情緒を豊かに表現する能力を持ちます。¥n¥n4. 侘寂（わびさび）：¥n - 意味: 簡素でありながら深い美しさや、不完全さの中に美を見出す考え方。¥n - 説明: 「侘寂」は、日本の禅の思想や美意識を取り入れたモデルです。このモデルは、シンプルでありながらも深みのある対話を提供します。¥n¥n5. 和敬清寂（わけいせいじゃく）：¥n - 意味: 茶道の精神を表す四つの理念（和敬清寂）。¥n - 説明: 「和敬清寂」は、日本の茶道の精神を反映したモデルです。このモデルは、礼儀正しく、敬意を持ち、清潔で純粋な対話を目指します。¥n¥nこれらの命名は、日本の文化的背景や価値観を反映しつつ、LLMの機能や特性を強調するものです。それぞれの名前が持つ意味を通じて、日本らしい独自性と品質を持つLLMであることをアピールできます。	日本では、製品やサービスの名前に日本らしさを反映することがよくあります。この傾向は、日本の文化や伝統、美学を尊重し、それを製品やサービスの名前に取り入れることで、その製品やサービスが持つ独自の魅力や品質を表現しようとするものです。¥n¥n例えば、日本の自動車メーカーが製造する車の名前には、伝統的な日本の美学や力強さを表す漢字が使われることがあります。また、日本の家電製品には、自然や美しさを連想させるような名前が付けられることがあります。¥n¥nこのような命名法は、日本の消費者が製品やサービスに対してどのようなイメージを持つか、またはその製品やサービスがどのような特徴や品質を持つかを反映しています。また、これは日本の企業が自社の製品やサービスを他社と差別化し、ブランドイメージを強化するための戦略の一部でもあります。¥n¥nしたがって、日本のLLMにも同様の命名法が用いられる可能性があります。例えば、日本の伝統的な美学や技術を反映した名前、または日本の自然や文化を連想させる名前が考えられます。	"evaluation": "0", "corrected_answer": null, "model_1": "team-hatakeyama-phase2¥/Tanuki-8x8B-dpo-v1.0", "model_2": "karakuri-ai¥/karakuri-lm-8x7b-chat-v0.1", "created_time": "2024-08-24T18:12:13.167674"}
{ "question": "chmod 777はなんでラッキーセブンなの？", "	answer_1": "「chmod 777」が「ラッキーセブン」と呼ばれる理由は、UNIX系オペレーティングシステムにおけるファイル権限の設定に関連しています。具体的には、`chmod` コマンドを使ってファイルやディレクトリのパーミッションを変更する際に、777というパーミッション設定が非常に強力です。すべてのユーザーに対して読み取り、書き込み、実行の権限を与えることを意味します。¥n¥nこの設定は、以下のように解釈されます：¥n¥n1. 読み取り (r): ファイルやディレクトリの内容を読み取る権限。¥n2. 書き込み (w): ファイルやディレクトリの内容を変更する権限。¥n3. 実行 (x): ファイルやディレクトリを実行する権限。¥n¥nこれらの権限がすべてのユーザーに付与されるため、「777」は「すべての人が何をしてよい」という意味で「ラッキーセブン」と呼ばれます。これは、セキュリティの観点から見ると非常に危険であり、不正アクセスやデータの改ざん、破壊のリスクが高まります。¥n¥nしたがって、「chmod 777」はシステム管理者や開発者が特定の状況下で一時的に使用することがありますが、通常は推奨されません。代わりに、必要最低限の権限を設定する方が安全です。", "	"answer_2": "A clever question!¥n¥n「chmod 777」が「ラッキーセブン」と呼ばれる理由は、Unix系のファイルシステムで使われるパーミッションの設定に由来します。¥n¥nchmodコマンドは、ファイルやディレクトリのアクセス権を設定するコマンドです。パーミッションは、3桁の数字で表され、各桁は、所有者、グループ、他ユーザーのアクセス権を表します。各桁は、以下の値をとります。¥n¥n- 0: アクセス不可¥n- 1: 読み取りのみ¥n- 2: 書き込みのみ¥n- 3: 読み取りと書き込み¥n- 4: 実行権限 (実行可能)¥n- 5: 読み取りと実行¥n- 6: 書き込みと実行¥n- 7: 読み取り、書き込み、実行 (すべての権限)¥n¥n「chmod 777」は、所有者、グループ、他ユーザーに対して、すべての権限 (読み取り、書き込み、実行) を許可する設定です。つまり、誰でも自由にファイルやディレクトリを読み、書き、実行できるということです。¥n¥nこの設定を「ラッキーセブン」と呼ぶのは、777が、すべての権限を許可する設定であることから、運が良く、すべてのアクセスを許可するという意味で「ラッキー」であると、ユーザーが感じたことから広まったと考えられます。", "	"evaluation": "0", "corrected_answer": null, "model_1": "team-hatakeyama-phase2¥/Tanuki-8x8B-dpo-v1.0", "model_2": "elyza¥/Llama-3-ELYZA-JP-8B", "created_time": "2024-08-24T18:40:48.828596"}
{ "question": "非常識、不可能、無弁別のように、通常の漢字 2 文字の単語に「非」が「不」が「無」をつけて否定して、漢字 3 文字の単語を5個生み出してください。", "	"answer_1": "はい、ご要望に沿って、「非」「不」「無」を使って否定の意味を持つ3文字の漢字単語を5つ作成しました：¥n¥n1. 非現実 (ひげんじつ) - 現実的でない、現実離れている¥n¥n2. 不可能 (ふかのう) - できない、実現不可能である¥n¥n3. 無意識 (むいしき) - 意識していない、無自覚な状態¥n¥n4. 非合理 (ひごうり) - 理にかなっていない、筋が通らない¥n¥n5. 不確実 (ふかくじつ) - 確かでない、不確定である¥n¥nこれらの単語は、元の2文字熟語に「非」「不」「無」のいずれかを付けることで、その意味を否定したり、反対の意味を表現したりしています。", "	"answer_2": "もちろんです。以下に「非」「不」「無」を使って否定した漢字の単語を5つご紹介します。¥n¥n1. 非現実的 (ひじょうりてき) ¥n - 意味: 現実的でない、実現不可能なさま。¥n¥n2. 不可能性 (ふかのうせい) ¥n - 意味: 達成や実現が不可能な状態。¥n¥n3. 無意味 (むいみ) ¥n - 意味: 意味や価値がない、無駄である。¥n¥n4. 無知 (むち) ¥n - 意味: 知識や理解がないこと。¥n¥n5. 無形 (むけい) ¥n - 意味: 形や形を持たない、抽象的な概念。¥n¥nこれらの単語は、それぞれ異なる文脈で使われることが多く、否定的な意味合いを持つことが特徴です。", "	"evaluation": "1", "corrected_answer": null, "model_1": "claude-3-5-sonnet-20240620", "model_2": "team-hatakeyama-phase2¥/Tanuki-8x8B-dpo-v1.0", "created_time": "2024-08-24T18:27:04.146857"}

Tanuki-8x8BモデルはLlamaに代表されるDense構造ではなく、Mistralに代表されるMixture of Experts (MoE)構造を採用しました。MoEはDenseモデルのFFN (Feed-Forward Network) 層を複数のexpertを持つ層に変更し、どのexpertを採用するかをrouterによって選択する構造です。Tanuki-8x8Bでは8つのexpertを持ち、top2(上位2つ)のexpertを採用する構造を採用しました。

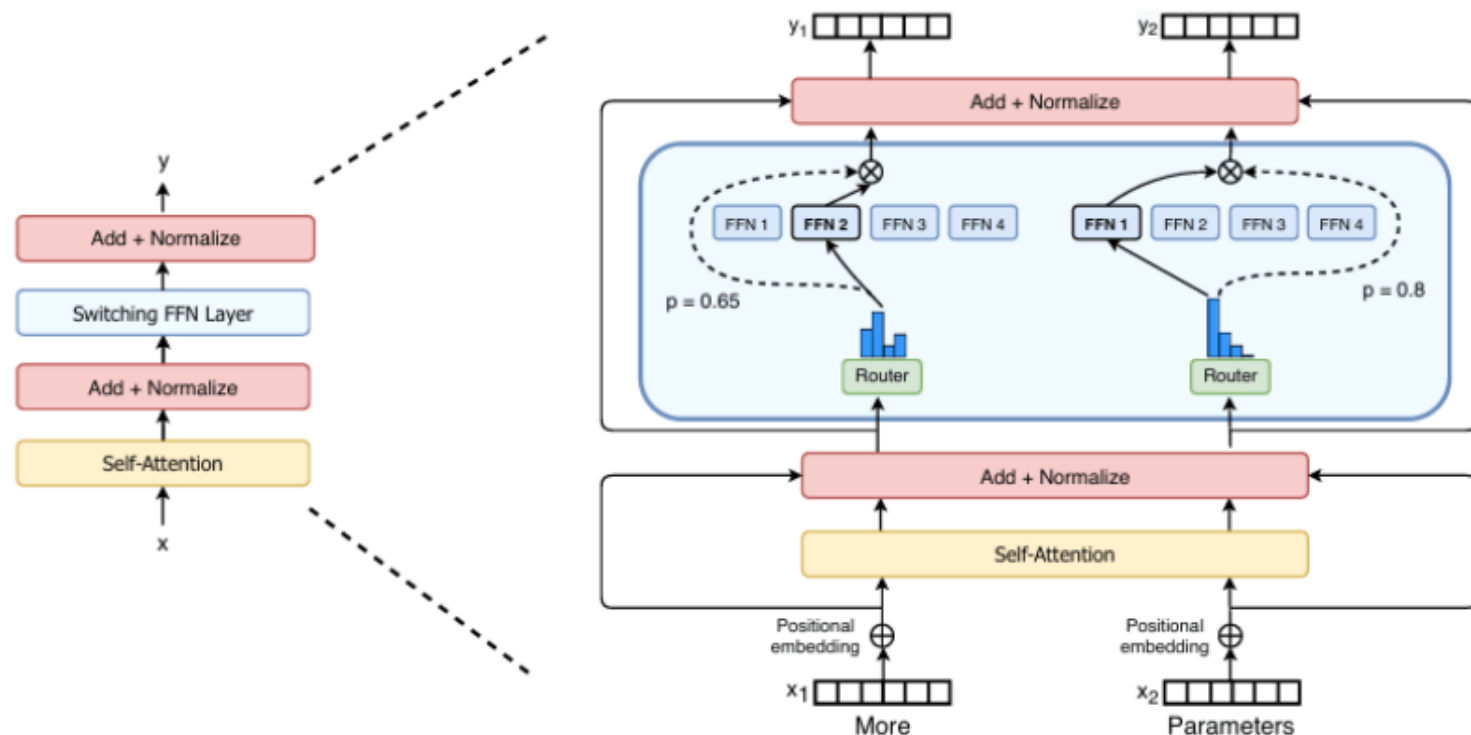


図.1 MoE layer from the [Switch Transformers paper](#)

Webデータ vs. 対話データ

- mc4-jpと対話データ (ichikara) をembed & umapで2次元転写
 - [詳細はこちら](#)
- Webデータと対話データの重複が少ない
- 対話はWebデータの「外挿領域」に近い？

