

オープンドメイン常識推論のための 知識拡張型プロンプト学習

ZHAO Xujiang LIU Yanchi CHENG Wei 大石 美賀 大崎 隆夫 松田 勝志 CHEN Haifeng

要 旨

常識推論のためのニューラル言語モデルは問題をQAタスクとして扱うことが多く、ファインチューニング後の言語の学習表現に基づいて予測を行います。ニューラル言語モデルはファインチューニングのデータや事前に定義された回答候補を提供しなくとも、外部知識のみに基づいて常識推論の質問に回答することが可能なのでしょうか。本稿では、回答候補やファインチューニングの実例を用意することなく質問への回答を目指す、ユニークながら難しさも伴う、オープンドメイン問題の常識推論について検討します。NECLA (NEC Laboratories America)、及びNECデジタルプラットフォームビジネスユニットで構成された研究チームでは、ニューラル言語モデルを用いて、タスク特有の指示が必要な外部知識をベースとして推論チェーンを繰り返し検索する方法を提案しました。これらの推論チェーンは、常識質問とこれに付随する知識ステートメントに対する最も正しい回答を特定できるように支援し、選んだ回答の正当性を確認できます。この技術はさまざまなビジネスドメインにおいて、その有効性が証明されています。



常識推論／外部知識／事前学習済み言語モデル／オープンドメイン／推論チェーン

1. はじめに

ビジネスの領域において、大規模事前学習済み言語モデル (PLM)¹⁾は、自然言語処理 (NLP) における最有力の勢力として登場してきました。PLMは膨大な量の一般的なテキストデータによる広範な訓練によって実世界についての基本的な理解を獲得した後、さまざまな適用分野に対応した固有のデータセットによりファインチューニングを行います。PLMは数多くの下流タスクについて並外れた能力を示してきましたが、推論に関連した試みという文脈においては、依然として次のような2つの致命的な課題があります。

- (1) 不完全な知識：訓練データには存在しない情報が必要とするタスクに直面した場合や質問応答の形式に合致しないテストを扱う場合、PLMはしばしば苦戦することがあります。
- (2) 限られた推論能力：PLMは暗黙的にエンコードされた知識に基づいて予測を行います。このことは複雑なタスクで必要となる構造化された推論能力を欠いていることを意味し、その結果として自ら選択した応答について明確に説明することができません。

ビジネスの世界において、PLMをお客様サポートやデータ分析、コンテンツ生成といった適用分野で効果的に活用するには、これらの課題に取り組むことが極めて重要です。PLMの推論能力を強化し、多様な情報ソースを取り扱えるようにするための解決策を見出すことは、さまざま

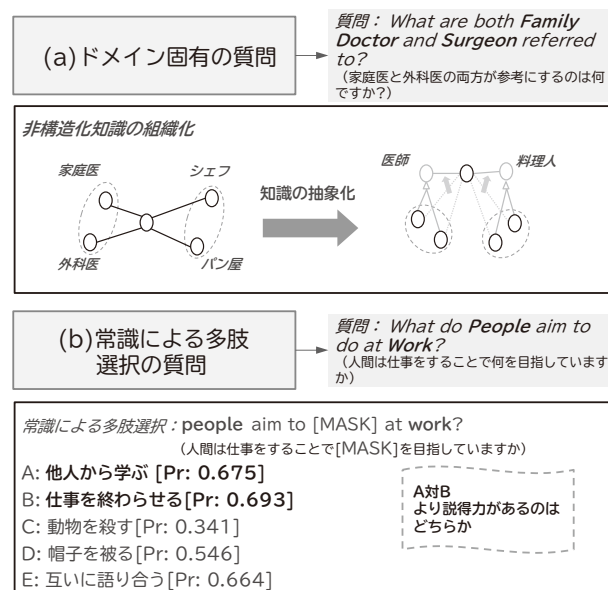


図1 PLMを常識推論に用いるうえで重要な2つのポイント

なビジネス環境においてPLMの可能性をフルに引き出すうえで最も重要です。図1に示すように、訓練中に体験した例とは異なるドメインの質問を提示されたとします。「家庭医と外科医の両方が参考にするのは何ですか?」といった医療ドメインの質問に対し、人間であれば、回答候補が用意されていなくとも、両方のエンティティに対する抽象化された意味の生成を目指します。しかし、PLMの場合はT5-3b²⁾であっても、ファインチューニングの過程を経なければ、暗記学習用アプリケーションのquizletといった不適切な回答を生成してしまうでしょう。更に、「人間は仕事をすることで[MASK]を目指していますか?」という常識質問の場合、PLMによるプロンプト学習のパラダイムでは問題を多肢選択QAに定式化し、マスクした空白部分に回答候補をそれぞれ当てはめて文全体の尤度をしばしば計算します。しかし、「他人から学ぶ」と「仕事を終わらせる」という回答は、両方とも意味的には正しいといえます。PLMはなぜその回答を選んだのかの根拠を示すことができません。両方のケースから分かるのは、常識推論による予測には、質問のコンテキストと外部知識によって提供される明示的な情報を統合するために、安定的であり体系的な推論が必要であるということです。

本稿では、回答候補やファインチューニングのための実例を提示しなくても、マシンが人間と同じように日常的な状況の種類と性質について推定できるようにする、オープンエンド常識推論タスクに注目します。また、回答候補のセットや回答の範囲をあらかじめ定義することなしに、オープンエンドな常識推論を遂行する外部知識拡張型プロンプティング(KEEP: KnowlEdge Enhanced Prompting)手法を示します。KEEPは最初に、回答候補の要件を除去するため、外部の知識ベース(例: ConceptNet)を回答探索空間として利用し、質問に関連したマルチホップな推論パスを繰り返し抽出します。知識ベース全体への網羅的な探索を回避するため、全体的な探索基準はPLMを利用して策定しています。ここでの重要な知見といえるのは、PLMは大規模モデルのパラメーターを介してある種の推論能力を備えており、これを利用すれば推論パスを展開し続けるのか、それともパス内のエンティティを最終回答として採用するのかを判断するための暗黙知を提供できる可能性があるということです。したがってKEEPは、具体的な回答の範囲や推論過程の直接的な監修を制限することなく、常識推論を必要とする

現実世界のほとんどのシナリオに適用できます。そして、PLMの推論能力を更に強化するために、外部の知識ベースから直接抽出したタスクに依存しない推論パスをPLMのファインチューニング用の訓練インスタンスとして利用することを提案します。

2. 知識拡張型プロンプティング手法

第2章ではまず問題の定式化を紹介し、続いて提案する手法の詳細なフレームワークについて検討します。この手法は次の3つの要素に分解できます。

- (1) エンティティの抽出とリンク
- (2) ローカル知識グラフの拡張
- (3) 訓練戦略と回答予測

2.1 問題の定式化

本研究では、PLMから得た知識と構造化された1つの知識グラフを利用することで、オープンエンド常識推論の質問を解くことを目指します。知識グラフ(KG) $G=(V, E)$ (例: ConceptNet)はマルチリレーショナルでヘテロジニアスなグラフです。 V はエンティティノードの集合であり、 $E \subseteq V \times R \times V$ は V 内のノードを接続しているエッジの集合であり、 R はリレーションタイプの集合(例: locates_at, requiresなど)を表します。具体的には、回答の選択肢や回答形式の制限のないオープンエンド常識推論の質問 q があるとき、この研究の目標は、(1) q の関連情報を含んだローカルKG $G_q \in G$ 、(2) G_q から抽出された推論パスの集合 $k=\{k_1, k_2, \dots, k_m\}$ 、(3)質問 q の回答として正しい k から抽出されたエンティティ $\hat{a}$ 、を決定することです。例えば図2で示すように、「What do People aim to do at Work? (人間は仕事をすることで何を目標とするか?)」

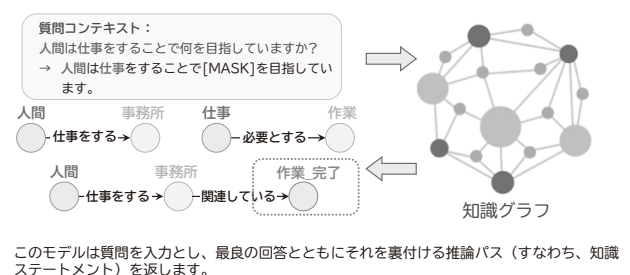


図2 オープンエンドな常識推論の例

指していますか?)」という常識質問に回答するには、まずこの質問に回答するための論理情報を提供できる外部のKGから関連する推論パスをすべて抽出することを目指します。このパス全体のなかから最も正確なものを選び(つまり $\text{people} \rightarrow \text{office} \rightarrow \text{finish_jobs}$)、次の結合尤度を最大にする回答 $\hat{a} = \text{finish_jobs}$ を抽出することができます。

$$P(\hat{a}, \mathbf{k} | q, G_q) = P(\mathbf{k} | q, G_q) \cdot P(\hat{a} | \mathbf{k})$$

課題：ただし、結合尤度を最大化することは、次の2つの致命的障害によって簡単なタスクではなくなっています。1つ目は、既存の研究³⁾⁻⁵⁾で行われているように、オープンエンドな環境下では質問エンティティと回答候補の間にローカルKGを構築できないため、質問に関連した推論パス $\mathbf{k}$ (つまり知識ステートメント)の検索が困難なことです。これに加えて、微分可能法⁶⁾のときのように、事前に定義された回答範囲の調整がないため、探索空間が知識グラフ全体になってしまうことです。そこで次は、この両方の課題を解決するため、ローカルKGを開始する方法について検討し、すべての妥当と思われる知識ステートメントと最も説得力ある回答を見出すためにこれに関する推論を繰り返し行います。全体のフレームワークを図3に示します。

2.2 ローカルグラフの構築と拡張知識グラフエンティティのリンク

コンセプトの知識グラフ(例えば、ConceptNet)は現実世界のテキストに対して多様で有用なコンテキスト指向の推論タスクを可能にし、これによってオープンエンド常識推論での最も適切な構造化された知識を提供しま

す。与えられた常識的なコンテキストに対してPLMとGから得た知識を用いた推論を行う場合、このフレームワークの最初のステップはKG内のノード集合 $V_q \in V$ への全射マッピングを持つ質問 q から重要なエンティティの集合 $c_q = \{c_q^{(1)}, \dots, c_q^{(l)}, \dots\}$ を抽出することです。そして以前の研究⁷⁾に従い、クエリコンテキストの潜在表現とGに保存されているリレーション情報とを利用して、 q から得た情報エンティティ c_q をKG内の結合したコンセプトエンティティ V_q へマッピングします。

ローカル知識グラフに対する推論：人間の推論過程を模倣するため、本稿では、GからLホップ以内で推論パスを検索して質問コンセプト c_q を最も広くカバーするローカル知識サブグラフ G_q を構成することを目指します。 G_q 内の個々のパスは理想的には最も正しい回答とその質問 q に対する説明の特定を支援する推論チェーンとみなすことができます。しかし、 c_q からLホップのサブグラフ G_q への拡張は計算処理的に不可能です。

推論パスのプルーニング：推論パスの拡張プロセスをスケーラブルにするために、PLM内に暗黙知を組み込んで不適切なパスのプルーニングを行います。具体的には、質問に対してそこから派生した推論パスで直接回答することによりローカルグラフ拡張の問題を明示的な推論手順にするため、質問 q とノード v のテキストとを推論パスで変換した知識ステートメントとともに組み合わせて穴埋め方式ベースのプロンプト $w = [q; v; (v_i, r_{ij}, v_j)]$ を構成します。例えば、図4のように、プロンプトは「人間は仕事をする事で何を目指していますか? <回答ノード>、なぜなら <推論パス>」という形式になっています。事前に定義したテンプレートを利用してトリプレット (v_i, r_{ij}, v_j) を自然言語に変換します。例えばトリプレッ

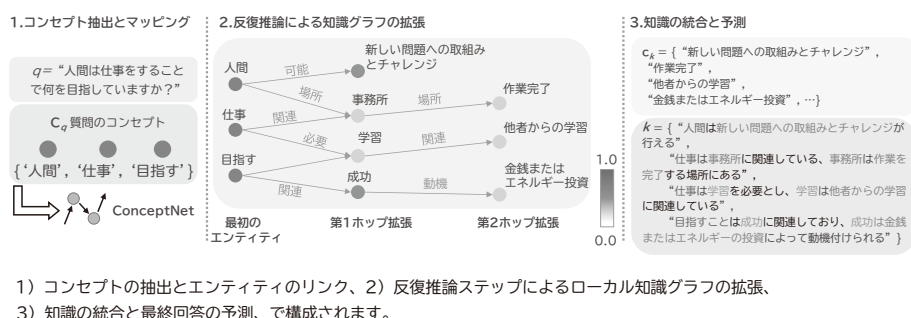


図3 提案した方法のフレームワーク

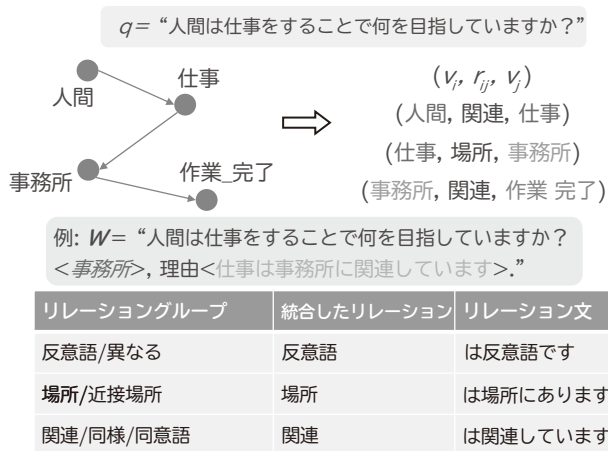


図4 知識ステートメントの変換と穴埋め方式によるプロンプト構築

ト(仕事, 反意語, 失業)は、図4で示すように、仕事は失業の反意語であるというように翻訳できます。推論パスを保存するかどうか評価するための方法として提案するのは、質問のコンテキストが与えられたうえで、PLMを利用して個々の推論パスの妥当性を採点することです。そのための公式として、プロンプト W は N 個のトークン $W = \{\omega_1, \dots, \omega_{n-1}, \omega_n, \omega_{n+1}, \dots, \omega_N\}$ で構成されるとして、第 l 番目ホップでの拡張で作成される論理センテンス W の常識スコア $\phi_l(W)$ を下記のように定義します。

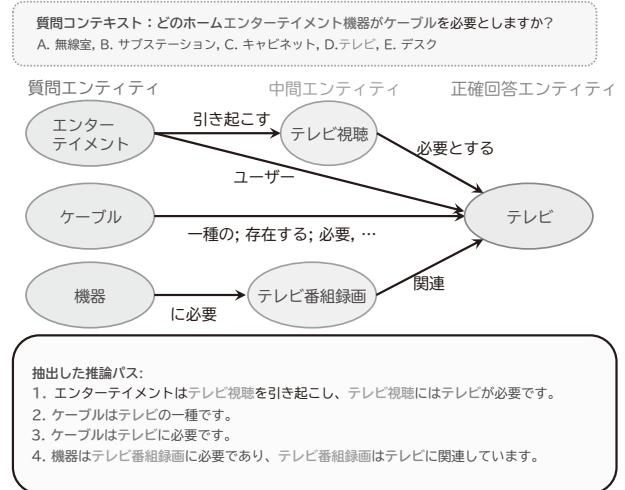
$$\phi_l(W) := \sum_{n=0}^N \log(p_\theta(\omega_n | W_{\setminus n})) / N$$

ここで $W_{\setminus n}$ はトークン ω_n の[MASK]への置き換えを示し、分母 N はスコア予測におけるセンテンス長の影響を軽減します。

G_q を繰り返し拡張すると、各 $\phi_l(W)$ はグラフ内の特定の $l \in [1, L]$ 深度において固有の推論パスを得ます。図3で印を付けているように、 $\phi(W)$ の点数が高いということは次の $(l+1)$ ホップでの拡張のためにはノード v_l を保持しておく必要があることが分かります。

2.3 訓練戦略と回答予測

PLMの推論能力を強化するために提案するのは、ConceptNetから構築された知識実例を用いてPLMのファインチューニングを行うことです。具体的には、ConceptNet上の知識トリプレットを正確に識別する



訓練セットの各常識質問について、質問内のエンティティとConceptNet内の正確回答エンティティの間の推論パスをすべて説明しました。すべての推論パスはテンプレートによりセンテンスに変換し、モデルのファインチューニング用コーパスとして使用しています。

図5 訓練コーパスの生成

ことで p_θ の推論能力の強化を目指します。図5に示すように、常識質問 $q = \text{“What home entertainment equipment requires cable?”}$ (どのホームエンターテインメント機器がケーブルを必要としますか?)とその正確な回答 $\hat{a} = \text{“television (テレビ)”}$ がある場合、 c_q 内の各エンティティ $c_q^{(i)}$ から $\hat{a}$ までの G の推論パス $[(v_1, r_1, v_2), \dots, (v_{L-1}, r_{L-1}, v_L)]$ を識別します。ここでは、例えば「Cable is a type of Television (ケーブルはテレビの一種です)」とか「Cable is required for Television (ケーブルはテレビに必要です)」など $c_q^{(i)}$ から $\hat{a}$ までに複数のパスが存在することがあることに注意が必要です。次に各推論パスを図4の表に示したようなテンプレートで自然言語のセンテンスに変換します。ここでは標準的なマスクされた言語モデリングタスクに従ってモデルのファインチューニングを行います。各センテンス内のトークンのごく一部(15%)をランダムにマスクしておき、PLMがそのマスクを埋める学習をすることで検索した個々の推論パスの背後にある隠れた論理を理解できるようにすることを目指します。

回答予測: L ホップ内のすべての高い常識スコアの推論パス k で構成されたサブグラフ G_q を取得した後は、 $k_i \in k$ の個々のパスは回答 a を個別に支持する知識説明と見なすことができます。

$$\log P_{\theta}(a_i | k_i) \propto \phi_L = \sum_{l=1}^L \phi_l$$

ここで ϕ_L は L ホップ内の個々の回答 a_i の最終スコアを意味しており、単一の推論パス $\{c \rightarrow v_1 \rightarrow \dots \rightarrow a\}$ が与えられた時の回答 a_i の尤度に近いと解釈できます。今回は効率を更に引き上げるため、ビームサーチを利用して信頼度の高い推論パスのみを保存しました。これにより、回答 $\hat{a}$ と最高スコア ϕ_L を持つ推論パス $\hat{k}$ を最終回答であり知識が裏付けているものとして選び出すことができます。

3. まとめ

NECLA及びNECデジタルプラットフォームビジネスユニットのメンバーで構成された当研究チームは、回答の選択枝や回答形式の制限のないオープンエンド常識推論の回答を予測できる汎用的な新しいフレームワークであるKEEPを開発しました。本技術は常識推論という難問への取り組みに現実世界のタスクを適用することで、さまざまなビジネス領域における有効性を示しました。チームではこの仕事が大規模言語モデル時代におけるゼロショット、及びオープンエンド環境下での常識推論の自動化に新しい方向性を提示できるものと信じています。

参考文献

- 1) Mihir Kale and Abhinav Rastogi: Text-to-Text Pre-Training for Data-to-Text Tasks, The 13th International Conference on Natural Language Generation, pp.97-102, 2020
<https://aclanthology.org/2020.inlg-1.14/>
- 2) Iz Beltagy, Kyle Lo and Arman Cohan: SciBERT: A Pretrained Language Model for Scientific Text, 2019 Conference on Empirical Methods in Natural Language Processing (EMNLP), 2019
<https://arxiv.org/abs/1903.10676>
- 3) Michihiro Yasunaga, Hongyu Ren, Antoine Bosselut, Percy Liang and Jure Leskovec: QA-GNN: Reasoning with Language Models and Knowledge Graphs for Question Answering, The 2021 Conference of the North American Chapter of the Association for Computational Linguistics (NAACL), 2021
<https://arxiv.org/abs/2104.06378>
- 4) Xikun Zhang, Antoine Bosselut, Michihiro Yasunaga et al.: GreaseLM: Graph REASoning Enhanced Language Models, 9th International Conference on Learning Representations (ICLR), 2022
<https://arxiv.org/abs/2201.08860>
- 5) Bill Yuchen Lin, Xinyue Chen, Jamin Chen and Xiang Ren: KagNet: Knowledge-Aware Graph Networks for Commonsense Reasoning, 2019 Conference on Empirical Methods in Natural Language Processing (EMNLP), 2019
<https://arxiv.org/abs/1909.02151>
- 6) Bill Yuchen Lin, Haitian Sun et al.: Differentiable Open-Ended Commonsense Reasoning, the 2021 Conference of the North American Chapter of the Association for Computational Linguistics (NAACL), 2021
<https://arxiv.org/abs/2010.14439>
- 7) Maria Becker, Katharina Korfhage and Anette Frank: COCO-EX: A Tool for Linking Concepts from Texts to ConceptNet, The 16th Conference of the European Chapter of the Association for Computational Linguistics (EACL), 2021
<https://aclanthology.org/2021.eacl-demos.15/>

執筆者プロフィール

ZHAO Xujiang

NEC Laboratories America
Researcher

LIU Yanchi

NEC Laboratories America
Researcher

CHENG Wei

NEC Laboratories America
Senior Researcher

大石 美賀

ソフトウェア&システムエンジニア
リング統括部

大崎 隆夫

ソフトウェア&システムエンジニア
リング統括部
ディレクター

松田 勝志

ソフトウェア&システムエンジニア
リング統括部
プロフェッショナル

CHEN Haifeng

NEC Laboratories America
Department Head

NEC 技報のご案内

NEC 技報の論文をご覧くださいありがとうございます。
ご興味がありましたら、関連する他の論文もご一読ください。

NEC技報WEBサイトはこちら

NEC技報（日本語）

NEC Technical Journal（英語）

Vol.75 No.2 ビジネスの常識を変える生成AI特集 ～社会実装に向けた取り組みと、それを支える生成AI技術～

ビジネスの常識を変える生成AI特集によせて
生成AI技術への取り組み ～基盤から応用、ルール作りまで～

◇ 特集論文

急速に広まる生成AIの市場適用

NECにおける生成AIの取り組みについて
電子カルテと医療文書の作成支援による医師業務効率化
映像×LLMによる報告書作成業務の自動化
映像分析と生成AIによるリアル世界の行動理解
サイバー脅威インテリジェンス生成自動化
生成AIの社内活用を進めるNEC Generative AI Service (NGS)
ソフトウェア・システム開発への生成AIの活用
LLMとMIで革新する素材開発プラットフォーム
LLMと画像分析を活用した被災状況の把握

生成AIの可能性を高める基盤技術

日本語性能に優れたNECの大規模言語モデル (LLM)
NECの生成AIを支える国内企業で最大規模のAIスーパーコンピュータ
より安全な大規模言語モデル (LLM) を目指して
データを秘匿したまま連携できる連合学習技術とLLMへの適用可能性
大規模言語モデル (LLM) によるFew-shotクラスタリングの可能性
オープンドメイン常識推論のための知識拡張型プロンプト学習
AI連携とオーケストレーションのための基盤ビジョンLLM
クエリを考慮した新規手法により関連する企業データを減らしLLM APIの使用コストを最適化

AI技術が社会へ浸透するために

AI標準化・ルールメイクに関する動向とNECの取り組み
人権尊重に向けたNECのAIガバナンスの取り組み
RCModelを用いたAIリスクマネジメントのための人材育成事例

◇ NEC Information

2023年度C&C賞表彰式典開催



Vol.75 No.2
(2024年3月)

特集TOP