

Bring Your Own Knowledge: A Survey of Methods for LLM Knowledge Expansion

Anonymous ACL submission

Abstract

Adapting large language models (LLMs) to new and diverse knowledge is essential for their lasting effectiveness in real-world applications. This survey provides an overview of state-of-the-art methods for expanding the knowledge of LLMs, focusing on integrating various knowledge types, including factual information, domain expertise, language proficiency, and user preferences. We explore techniques, such as continual learning, model editing, and retrieval-based explicit adaptation, while discussing challenges like knowledge consistency and scalability. Designed as a guide for researchers and practitioners, this survey sheds light on opportunities for advancing LLMs as adaptable and robust knowledge systems.

1 Introduction

As large language models (LLMs) are increasingly deployed in real-world applications, their ability to adapt to evolving knowledge becomes crucial for maintaining relevance and accuracy. However, LLMs are typically trained once and thus only have knowledge up to a certain cutoff date, limiting their ability to stay updated with new information. This survey provides a comprehensive overview of methods that enable LLMs to incorporate various types of new knowledge, including factual, domain-specific, language, and user preference knowledge. We survey adaptation strategies, including continual learning, model editing, and retrieval-based approaches, and aim at providing guidelines for researchers and practitioners.

To remain effective, LLMs require updates across multiple dimensions. Factual knowledge consists of general truths and real-time information, while domain knowledge pertains to specialized fields, such as medicine or law. Language knowledge enhances multilingual capabilities, and preference knowledge aligns model behavior with user expectations and values. Ensuring that LLMs

	Continual Learning (§4)	Model Editing (§5)	Retrieval (§6)
Knowledge Type			
Fact	✓	✓	✓
Domain	✓	✗	✓
Language	✓	✗	✗
Preference	✓	✓	✗
Applicability			
Large-scale data	✓	✗	✓
Precise control	✗	✓	✓
Computational cost	✗	✓	✓
Black-box applicable	✗	✗	✓

Table 1: We compare three key approaches for adapting LLMs — continual learning, model editing, and retrieval — based on their supported knowledge types and applicability across different criteria.

can integrate updates across these dimensions is essential for their sustained utility.

Existing LLM adaptation methods differ in approach and application. Continual learning enables incremental updates to models’ parametric knowledge, mitigating catastrophic forgetting (McCloskey and Cohen, 1989) while ensuring long-term performance. Model editing allows for precise modifications of learned knowledge, providing controlled updates without requiring full retraining. Unlike these *implicit* knowledge expansion methods, which modify the model’s internal parameters, retrieval-based approaches *explicitly* access external information dynamically during inference, reducing dependency on static parametric knowledge. The suitability of these methods for different knowledge types and their general applicability are summarized in Table 1. By leveraging these strategies, LLMs can maintain accuracy, contextual awareness, and adaptability to new information.

After placing our work into context (Section 2) and defining knowledge types covered in this paper (Section 3), we provide an overview of different knowledge expansion methods as detailed in Fig-

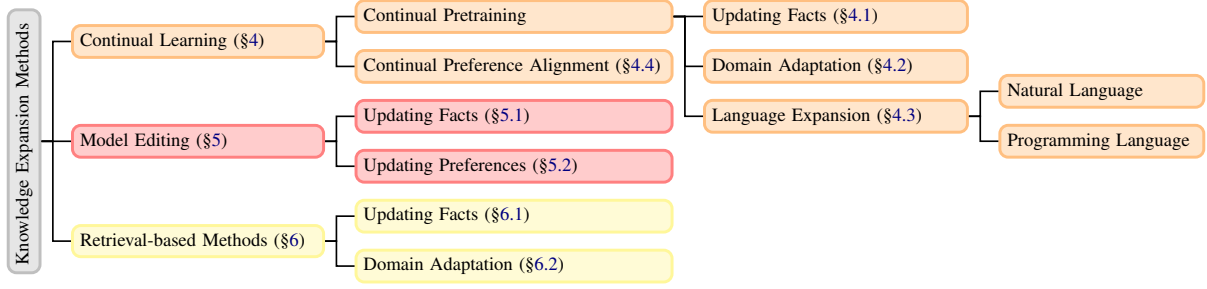


Figure 1: Taxonomy of current methods for expanding LLM knowledge. Due to space constraints, please refer to Appendix A.1 for a comprehensive review of methods and their corresponding citations.

ure 1. This work thus surveys diverse research efforts and may serve as a guide for researchers and practitioners aiming to develop and apply adaptable and robust LLMs. We highlight research opportunities and provide insights into optimizing adaptation techniques for various real-world applications.

2 Related Surveys

The main goal of our work is to provide researchers and practitioners a broad overview of various types of methods to adapt LLMs to diverse types of new knowledge. In this section, we explain how other more specialized surveys relate to our paper.

To the best of our knowledge, there is limited prior work that specifically focuses on continuous knowledge expansion for LLMs. Closest to our work, Zhang et al. (2023c) describe temporal factual knowledge updates, while we take a broader perspective by examining methods for adapting LLMs to unseen domain knowledge, expanding language coverage, and incorporating user preferences. Yao et al. (2023) and Zhang et al. (2024c) provide overviews of knowledge editing methodologies, categorizing approaches of knowledge editing. Similarly, Ke and Liu (2022), Wu et al. (2024b) and Wang et al. (2024b) offer a comprehensive overview of continual learning. In contrast, our survey shifts the focus towards a task-oriented perspective on knowledge expansion, detailing how various types of knowledge — including factual, domain-specific, language, and user preference knowledge — can be seamlessly integrated to ensure LLMs remain relevant and effective.

3 Knowledge Types

Integrating diverse types of knowledge into LLMs is essential to enhance their versatility and effectiveness. Depending on the use case, the type of knowledge that an LLM shall be adapted to, might

differ. In this paper, we distinguish four key types of knowledge, which cover a broad range of use cases of researchers and practitioners: **factual, domain, language, and preference knowledge**.

(1) We define **factual knowledge** as general truths or contextualized information about the world that can be expressed in factual statements. We adopt a broad, high-level definition, encompassing finer-grained categorizations, such as commonsense knowledge, cultural knowledge, temporal knowledge, and entity knowledge as subsets of factual knowledge, in contrast to prior works (Cao et al., 2024; Wu et al., 2024b) using more granular classifications. This inclusive perspective enables a comprehensive exploration of knowledge expansion techniques for LLMs, providing flexibility beyond predefined categories and taxonomies.

(2) We define **domain knowledge** as specialized information relevant to specific fields, such as medicine, law, or engineering, enabling LLMs to perform well in targeted applications. Since LLMs typically excel in general-domain tasks but struggle with specialized content, incorporating domain knowledge is crucial for bridging this gap and improving performance in specific fields.

(3) We define **language knowledge** as the ability of an LLM to understand, generate, and reason in specific natural or programming languages.¹ Its integration focuses on adapting models to new languages and enhancing performance in underrepresented ones for broader applicability.

(4) Finally, we define **preference knowledge** as the capability of LLMs to tailor their behavior to align with user-specific needs, preferences, or values. Preference knowledge integration involves

¹We distinguish language knowledge from linguistic knowledge as defined by Hernandez et al. (2024). Language knowledge refers to the multilingual capabilities of an LLM, whereas linguistic knowledge falls under factual knowledge, encompassing statements about syntax and grammar.

adapting LLM behavior to meet diverse and dynamic user expectations.

In the next sections, we survey knowledge expansion methods and explain for which of these four knowledge types they are suitable.

4 Continual Learning

Continual learning (CL) is a machine learning paradigm that mimics the human ability to continuously learn and accumulate knowledge without forgetting previously acquired information (Chen and Liu, 2018). In the context of knowledge expansion, CL allows LLMs to integrate new corpora and incrementally update the knowledge stored in their parameters. This ensures that LLMs remain adaptable, relevant, and effective as facts, domains, languages, and user preferences evolve over time.

In the era of LLMs, the training of language models usually includes multiple stages: pretraining, instruction tuning, preference alignment, and potentially fine-tuning on a downstream task (Shi et al., 2024a). Depending on the stage, continual learning can be categorized into continual pretraining, continual instruction tuning, continual preference alignment, and continual end-task learning (Ke et al., 2023; Shi et al., 2024a). For knowledge expansion, the focus lies on continual pretraining (CPT) and continual preference alignment (CPA). In contrast, continual instruction tuning and continual end-task learning primarily aim to sequentially fine-tune pretrained LLMs for acquiring new skills and solving new tasks, which fall outside the scope of this survey.

In the following sections, we review existing studies that leverage continual pretraining for updating facts, adapting domains, and expanding languages, and continual alignment for updating user preferences.

4.1 Continual Pretraining for Updating Facts

This line of research focuses on updating a language model’s outdated internal factual knowledge by incrementally integrating up-to-date world knowledge (Jang et al., 2022; Ke et al., 2023).

Early studies (Sun et al., 2020; Röttger and Pierrehumbert, 2021; Lazaridou et al., 2021; Dhingra et al., 2022) empirically analyze continual pretraining on temporal data, demonstrating its potential for integrating new factual knowledge. Jin et al. (2022) and Jang et al. (2022) apply traditional continual learning methods to factual knowledge up-

dates in LLMs, evaluating their effectiveness in continual knowledge acquisition. Similarly, Jang et al. (2022) and Kim et al. (2024) classify world knowledge into time-invariant, outdated, and new categories — requiring knowledge retention, removal, and acquisition, respectively — and benchmark existing continual pretraining methods for knowledge updates.

Additionally, Hu et al. (2023) introduce a meta-trained importance-weighting model to adjust per-token loss dynamically, enabling LLMs to rapidly adapt to new knowledge. Yu and Ji (2024) investigate self-information updating in LLMs through continual learning, addressing exposure bias by incorporating fact selection into training losses.

4.2 Continual Pretraining for Domain Adaptation

Continual domain adaptive pretraining (Ke et al., 2022, 2023; Wu et al., 2024b) focuses on incrementally adapting an LLM using a sequence of unlabeled, domain-specific corpora. The objective is to enable the LLM to accumulate knowledge across multiple domains while mitigating catastrophic forgetting (McCloskey and Cohen, 1989) of previously acquired domain knowledge or general language understanding.

Gururangan et al. (2020) introduced the term of domain-adaptive pretraining, demonstrating that a second phase of pretraining on target domains can effectively update an LLM with new domain knowledge. It is important to note that further pretraining can lead to catastrophic forgetting of general concepts by overwriting essential parameters. To mitigate this, recent works utilize *parameter-isolation* methods which allocate different parameter subsets to distinct tasks or domains and keep the majority of parameters frozen (Razdaibiedina et al., 2023; Wang et al., 2024d,e). DEMix-DAPT (Gururangan et al., 2022) replaces every feed-forward network layer in the Transformer model with a domain expert mixture layer, containing one expert per domain. When acquiring new knowledge, only the newly added expert is trained while all others remain fixed. Qin et al. (2022) propose ELLE for efficient lifelong pretraining on various domains. ELLE starts with a randomly initialized LLM and expands the PLM’s width and depth to acquire new knowledge more efficiently. Ke et al. (2022) introduce a continual pretraining system which inserts continual learning plugins to the frozen pretrained language models that mitigate catastrophic forget-

ting while effectively learn new domain knowledge. Similarly, Lifelong-MoE (Chen et al., 2023) expands expert capacity progressively, freezing previously trained experts and applying output-level regularization to prevent forgetting.

In a later work, Ke et al. (2023) apply regularization to penalize changes to critical parameters learned from previous data, preventing catastrophic forgetting. Their approach computes the importance of LLM components, such as attention heads and neurons, in preserving general knowledge, applying soft-masking and contrastive loss during continual pretraining to maintain learned knowledge while promoting knowledge transfer.

4.3 Continual Pretraining for Language Expansion

Continual pretraining (CPT) has emerged as a pivotal strategy for adapting LLMs to new languages, or enhancing performance in underrepresented languages without full retraining (Wu et al., 2024b). Below, we discuss two major areas of expansion enabled by CPT: *natural language expansion* and *programming language expansion*.

Natural Language Expansion. Several recent studies have demonstrated the effectiveness of CPT in expanding language coverage. Glot500 (Imani et al., 2023) and EMMA-500 (Ji et al., 2024) enhance multilingual capabilities using CPT and vocabulary extension. Glot500, based on XLM-R (Ruder et al., 2019), and EMMA-500, built on LLaMA 2 (Touvron et al., 2023), expand language support up to 500 languages using extensive multilingual corpora. Similarly, Aya (Üstün et al., 2024) applies continual pretraining to the mT5 model (Xue et al., 2021) using a carefully constructed instruction dataset, achieving improved performance across 101 languages. Furthermore, LLaMAX (Lu et al., 2024) enhances multilingual translation by applying continual pretraining to the LLaMA model family. Supporting over 100 languages, it improves translation quality and promotes language inclusivity.

While covering many languages, many multilingual models exhibit suboptimal performance on medium- to low-resource languages (Ruder et al., 2019; Touvron et al., 2023; Imani et al., 2023). To bridge this performance gap, researchers have focused on expanding training corpora and strategically applying continual pretraining to enhance the multilingual capabilities of LLMs. Alabi et al.

(2022), Wang et al. (2023a), Fujii et al. (2024), and Zhang et al. (2024b) show that continual pretraining on one or more specific languages significantly improves performance across related languages. Blevins et al. (2024) extend this approach to the MoE paradigm for better parameter efficiency, while Zheng et al. (2024) investigate scaling laws for continual pretraining by training LLMs of varying sizes under different language distributions and conditions. Additionally, Tran (2020), Minixhofer et al. (2022), Dobler and de Melo (2023), Liu et al. (2024b), and Minixhofer et al. (2024) explore advanced tokenization and word embedding techniques to further improve LLMs’ multilingual performance in low-resource settings.

Programming Language Expansion. Going beyond natural languages, continual pretraining has demonstrated significant potential in enhancing the capabilities of LLMs for understanding and generating programming languages.

CERT, proposed by Zan et al. (2022), addresses the challenges of library-oriented code generation using unlabeled code corpora. It employs a two-stage framework to enable LLMs to effectively capture patterns in library-based code snippets. CodeTask-CL (Yadav et al., 2023) offers a benchmark for continual code learning, encompassing a diverse set of tasks such as code generation, summarization, translation, and refinement across multiple programming languages. Furthermore, continual pretrained models specifically for code understanding and programming from natural language prompts emerged with LLMs, such as CodeLLaMA (Grattafiori et al., 2023), Llama Pro (Wu et al., 2024a), CodeGemma (Team et al., 2024) and StarCoder 2 (Lozhkov et al., 2024), consistently outperform general-purpose LLMs of comparable or larger size on code benchmarks.

4.4 Continual Preference Alignment

Preference alignment ensures that large language models generate responses consistent with human values, improving usability, safety, and ethical behavior. While techniques like Reinforcement Learning from Human Feedback (RLHF) (Ziegler et al., 2019; Lambert et al., 2022) align LLMs with static preferences, societal values evolve, requiring continual preference alignment (CPA). It enables LLMs to adapt to emerging preferences while preserving previously learned values, ensuring relevance, inclusivity, and responsiveness to shifting

societal expectations. Despite its importance, CPA remains a relatively underexplored area. Below, we briefly discuss two representative works that highlight the potential of this approach:

Zhang et al. (2023b) propose a non-reinforcement learning approach for continual preference alignment in LLMs. Their method uses function regularization by computing an optimal policy distribution for each task and applying it to regularize future tasks, preventing catastrophic forgetting while adapting to new domains. This provides a single-phase, reinforcement learning-free solution for maintaining alignment across diverse tasks. Zhang et al. (2024a) introduce Continual Proximal Policy Optimization (CPPO), integrating continual learning into the RLHF framework to accommodate evolving human preferences. CPPO employs a sample-wise weighting strategy to balance policy learning and knowledge retention, consolidating high-reward behaviors while mitigating overfitting and noise.

As the demand for responsive and inclusive AI grows, CPA is key to keeping LLMs ethical and aligned with evolving user needs, requiring further research to reach its full potential.

4.5 Applicability and Limitations

Continual learning is a versatile framework for expanding LLM knowledge across facts, domains, languages, and preferences. It excels in large-scale knowledge integration, retaining previously learned knowledge, making it well-suited for tasks like domain adaptation and language expansion (Bu et al., 2021; Jin et al., 2022; Cossu et al., 2024).

However, CL has notable limitations, including a lack of precise control compared to model editing (cf. Section 5) and retrieval-based methods (cf. Section 6), inefficiency due to the computational demands of retraining, and limited applicability in black-box models. These challenges highlight the need for alternative approaches like model editing and retrieval, which offer more targeted and efficient updates.

5 Model Editing

Model editing offers a controllable and efficient solution to update factual knowledge and user preferences in LLMs. Introduced by Zhu et al. (2020), De Cao et al. (2021) and Mitchell et al. (2022a), it aims at modifying the model’s predictions for specific inputs without affecting unrelated ones.

Yao et al. (2023) and Zhang et al. (2024c) define four key evaluation metrics for model editing: **(1) reliability**, ensuring the edited model produces the target prediction for the target input; **(2) generalization**, requiring the edited knowledge to apply to all in-scope inputs — inputs that are directly related to the target input, including rephrasings and semantically similar variations; **(3) locality**, preserving original outputs for unrelated out-of-scope inputs; and **(4) portability**, extending the generalization metric by assessing how well updated knowledge transfers to complex rephrasings, reasoning chains, and related facts.

While recent works (Mitchell et al., 2022b; Madaan et al., 2022; Zhong et al., 2023; Zheng et al., 2023) use *model editing* and *knowledge editing* interchangeably for updating factual knowledge, we distinguish between them: model editing is a subset of knowledge editing that modifies model parameters, whereas retrieval-based methods update knowledge dynamically without altering the model’s parameters (see Section 6).

5.1 Model Editing for Updating Facts

To address outdated or incorrect information (Lazaridou et al., 2021), model editing research focuses on selectively modifying this knowledge. Below, we highlight key works in this area.

KnowledgeEditor (De Cao et al., 2021) uses a hypernetwork to predict parameter shifts for modifying a fact, trained via constrained optimization for locality. Similarly, MEND (Mitchell et al., 2022a) trains a hypernetwork per LLM layer and decomposes the fine-tuning gradient into a precise one-step parameter update. Given the findings that feed-forward layers in transformers function as key-value memories (Geva et al., 2021), Dai et al. (2022) introduce a knowledge attribution method to identify these neurons and directly modify their values via knowledge surgery.

Recent works employ a locate-and-edit strategy for precise model editing. Using causal tracing, Meng et al. (2022) identify middle-layer feed-forward networks as key to factual predictions and propose ROME, which updates facts by solving a constrained least-squares problem in the MLP weight matrix. MEMIT (Meng et al., 2023) extends ROME to modify thousands of facts simultaneously across critical layers while preserving generalization and locality. BIRD (Ma et al., 2023) introduces bidirectional inverse relationship modeling to mitigate the reverse curse (Berglund et al., 2003)

in model editing. While editing FFN layers has proven effective, PMET (Li et al., 2024d) extends editing to attention heads, achieving improved performance. Wang et al. (2024g) further shift the focus to conceptual knowledge, using ROME and MEMIT to alter concept definitions, finding that concept-level edits are reliable but have limited influence on concrete examples.

5.2 Model Editing for Updating Preferences

Recent works expand model editing beyond factual corrections to aligning LLMs with user preferences, such as ensuring safety, reducing bias, and preserving privacy .

Wang et al. (2024c) use model editing to detoxify LLMs, ensuring safe responses to adversarial inputs and preserving general LLM capabilities, such as fluency, knowledge question answering, and content summarization. Their results show that model editing is promising for detoxification but slightly affects general capabilities. Since LLMs can exhibit social biases (Gallegos et al., 2024), Chen et al. (2024a) propose fine-grained bias mitigation via model editing. Inspired by Meng et al. (2022), they identify key layers responsible for biased knowledge and insert a feed-forward network to adjust outputs with minimal parameter changes, ensuring generalization, locality, and scalability. For privacy protection, Wu et al. (2023) extend Dai et al. (2022)’s work by identifying privacy neurons that store sensitive information. Using gradient attribution, they deactivate these neurons, reducing private data leakage while preserving model performance. Moreover, Mao et al. (2024) apply model editing techniques like MEND to modify personality traits in LLMs, aligning responses to opinion-based questions with target personalities. While effective, this approach can degrade text generation quality.

5.3 Applicability and Limitations

Model editing complements continual learning by allowing fine-grained knowledge updates with lower computational costs. However, research has primarily focused on structured, relational, and instance-level knowledge, with limited exploration of other knowledge types, multilingual generalization, and cross-lingual transfer (Nie et al., 2024; Wei et al., 2025).

Additionally, model editing faces several technical challenges, including limited locality and gradual forgetting in large-scale edits (Bu et al.,

2019; Mitchell et al., 2022b; Gupta et al., 2024; Li et al., 2024b), making it more suitable for minor updates. Additionally, it can impact general LLM capabilities (Gu et al., 2024b; Wang et al., 2024f) and downstream performance (Gupta et al., 2024), potentially causing model collapse (Yang et al., 2024b). Addressing these issues will enhance model editing’s role alongside continual learning and retrieval, ensuring greater precision in dynamic knowledge adaptation.

6 Retrieval-based Methods

Continual learning and model editing modify a model’s parameters to update its internal knowledge, making them implicit knowledge expansion methods (Zhang et al., 2023c). In contrast, retrieval-based methods (Lewis et al., 2020) explicitly integrate external knowledge, allowing models to overwrite outdated or undesired information without parameter modifications. These methods leverage external sources, such as databases, off-the-shelf retriever systems, or the Internet, and thus provide up-to-date or domain-specific knowledge (Zhang et al., 2023c), making them effective for factual updates and domain adaptation.

6.1 Retrieval-based Methods for Updating Facts

Retrieval-based methods enhance LLMs by pairing them with an updatable datastore, ensuring access to current factual information. An early approach, retrieval-augmented generation (RAG) (Lewis et al., 2020), fine-tunes a pre-trained retriever end-to-end with the LLM to improve knowledge retrieval. Similarly, kNN-LM (Khandelwal et al., 2020) interpolates the LLM’s output distribution with k-nearest neighbor search results from the datastore, with later works optimizing efficiency (He et al., 2021; Alon et al., 2022) and adapting it for continual learning (Peng et al., 2023b).

For factual knowledge editing, Tandon et al. (2022) store user feedback for post-hoc corrections, while Mitchell et al. (2022b), Madaan et al. (2022), and Dalvi Mishra et al. (2022) retrieve stored edits to guide responses. Chen et al. (2024b) introduce relevance filtering to efficiently handle multiple edits. Retrieval-based in-context learning (Zheng et al., 2023; Ram et al., 2023; Mallen et al., 2023; Yu et al., 2023; Shi et al., 2024b; Bi et al., 2024) enables dynamic factual updates.

For complex reasoning, retrieval supports multi-

hop question answering and iterative prompting: [Zhong et al. \(2023\)](#) propose iterative prompting for multi-hop knowledge editing, while [Gu et al. \(2024a\)](#) use a scope detector to retrieve relevant edits and improve question decomposition via entity extraction and knowledge prompts. Similarly, [Shi et al. \(2024c\)](#) enhance multi-hop question answering by retrieving fact chains from a knowledge graph with mutual information maximization and redundant fact pruning.

In multi-step decision-making, retrieval is combined with Chain-of-Thought (CoT) reasoning ([Trivedi et al., 2023](#); [Press et al., 2023](#)). Retrieval also aids post-generation fact-checking and refinement ([Gao et al., 2023](#); [Peng et al., 2023a](#); [Song et al., 2024](#)) by revising generated text or prompts based on retrieved facts.

For a more comprehensive review of retrieval-based factual knowledge updates, we refer to [Zhang et al. \(2023c\)](#).

6.2 Retrieval-based Methods for Domain Adaptation

Retrieval-based methods have been widely adopted for various domain-specific tasks, e.g., in science and finance. By integrating retrieved external knowledge, these models enhance their adaptability to specialized domains, improving decision-making, analysis, and information synthesis.

In the biomedical domain, retrieval-based approaches facilitate tasks, such as molecular property identification and drug discovery by integrating structured molecular data and information about biomedical entities like proteins, molecules, and diseases ([Wang et al., 2023b](#); [Liu et al., 2023](#); [Wang et al., 2024h](#); [Yang et al., 2024a](#)). For instance, [Wang et al. \(2023b\)](#) and [Li et al. \(2024a\)](#) introduce retrieval-based frameworks that extract relevant molecular data from databases to guide molecule generation. In protein research, retrieval-based approaches enhance protein representation and generation tasks ([Ma et al., 2024](#); [Sun et al., 2023](#)). Additionally, [Lozano et al. \(2023\)](#) develop a clinical question-answering system that retrieves relevant biomedical literature to provide more accurate responses in medical contexts.

The finance domain, characterized by its data-driven nature, also benefits from retrieval-based methods ([Li et al., 2024f,g](#)). [Zhang et al. \(2023a\)](#) enhance financial sentiment analysis by retrieving real-time financial data from external sources. Furthermore, financial question-answering also bene-

fits from retrieval-based methods, which involves extracting knowledge from professional financial documents. [Lin \(2024\)](#) introduces a PDF parsing method integrated with retrieval-augmented LLMs to retrieve relevant financial insights.

6.3 Applicability and Limitations

Despite their advantages, retrieval-based methods also come with several limitations. A major challenge is their reliance on external knowledge sources, which can introduce inconsistencies or outdated information if not properly curated ([Jin et al., 2024](#); [Xu et al., 2024](#)). Their effectiveness also depends on the quality and scope of the retrieval system ([Bai et al., 2024](#); [Liu et al., 2024a](#)); poor indexing or noisy retrieval may lead to irrelevant or misleading information. Another key issue is maintaining knowledge consistency across queries. Since retrieval-based methods do not update model parameters, contradictions can arise between retrieved facts and previously generated responses, affecting coherence ([Njeh et al., 2024](#); [Zhao et al., 2024](#); [Li et al., 2024c](#)).

Addressing these challenges is essential to improving retrieval-based approaches and ensuring their seamless integration with other LLM adaptation techniques.

7 Challenges, Opportunities, Guidelines

Solving Knowledge Conflicts. An inherent challenge of expanding a model’s knowledge is the emergence of knowledge conflicts, which can undermine the consistency and trustworthiness of LLMs ([Xu et al., 2024](#)). Studies have identified various types of conflicts following knowledge updates, including (i) temporal misalignment ([Luu et al., 2022](#)), where outdated and newly learned facts coexist inconsistently, (ii) model inconsistencies ([Huang et al., 2021](#)), where responses to similar queries vary unpredictably, and (iii) hallucinations ([Ji et al., 2023](#)), where the model generates fabricated or contradictory information. While some efforts have been made to address these issues ([Zhang and Choi, 2023](#); [Mallen et al., 2023](#); [Zhou et al., 2023](#); [Xie et al., 2024](#)), they remain an open challenge that requires further research and more robust solutions.

Minimizing Side Effects. Continual learning and model editing, both of which involve modifying model parameters, inevitably introduce side effects. A major challenge in continual learning

is catastrophic forgetting (McCloskey and Cohen, 1989), where newly acquired knowledge overwrites previously learned information. In LLMs, the multi-stage nature of training exacerbates this issue, leading to cross-stage forgetting (Wu et al., 2024b), where knowledge acquired in earlier stages is lost as new training phases are introduced. For model editing, recent studies have shown that large-scale edits, particularly mass edits, can significantly degrade the model’s general capabilities, such as its language modeling performance (Wang et al., 2024f) or accuracy on general NLP tasks (Li et al., 2024e,b; Wang et al., 2024a). Effectively addressing these challenges is crucial for maximizing the potential of these methods for large-scale knowledge expansion while maintaining model stability and overall performance.

Comprehensive Benchmarks. Although this paper explores the properties, strengths, and weaknesses of various methods for knowledge expansion, the discussion remains largely theoretical due to the lack of a comprehensive benchmark datasets for a uniform evaluation and a proper comparison. Existing works, such as Jang et al. (2022), Liska et al. (2022), and Kim et al. (2024), provide factual knowledge-based datasets and evaluate continual pretraining and/or retrieval-based methods. However, their experiments are limited in scale and fail to offer a comprehensive assessment. Developing benchmarks that encompass a variety of knowledge types and enable the evaluation of all methods would provide a more holistic and systematic understanding of their relative effectiveness.

General Guideline. Selecting the appropriate method for knowledge expansion in LLM depends on the application context and type of knowledge that needs to be updated.

(i) For **factual knowledge**, model editing is ideal for precise, targeted updates, such as correcting specific facts, due to its efficiency and high level of control. Retrieval-based methods are effective for integrating dynamic or frequently changing facts, as they allow updates without modifying the model’s parameters, making them suitable for black-box applications. For large-scale factual updates, continual learning is preferred as it enables the incremental integration of new knowledge while preserving previously learned information.

(ii) For **domain knowledge**, both continual learning and retrieval-based methods are applicable. Continual learning excels in large-scale adap-

tation, using domain-specific corpora to ensure the model retains general knowledge while adapting to specialized contexts. Retrieval-based methods complement this by dynamically providing domain-specific information without requiring model modifications, making them valuable in scenarios where static updates are impractical.

(iii) For **language knowledge**, continual learning is the only method capable of supporting large-scale language expansion. It facilitates the integration of multilingual corpora and provides the foundational updates necessary for underrepresented or low-resource languages.

(iv) For **preference updates**, such as aligning models with evolving user values or ethical norms, continual alignment is typically achieved by combining continual learning techniques with preference optimization methods, such as reinforcement learning from human feedback. These approaches enable models to dynamically adapt to changing preferences while retaining alignment with previously learned values.

Summary. Continual learning is indispensable for large-scale updates like domain adaptation and language expansion, where foundational and incremental updates are required. **Model editing** excels at precise factual updates, while **retrieval-based methods** offer dynamic access to factual and domain knowledge without altering the model. A well-informed selection or combination of these methods ensures efficient and effective knowledge expansion tailored to specific use cases.

8 Conclusions

Adapting large language models to evolving knowledge is essential for maintaining their relevance and effectiveness. This survey explores three key adaptation methods — continual learning for large-scale updates, model editing for precise modifications, and retrieval-based approaches for external knowledge access without altering model parameters. We examine how these methods support updates across factual, domain-specific, language, and user preference knowledge while addressing challenges like scalability, controllability, and efficiency. By consolidating research and presenting a structured taxonomy, this survey provides insights into current strategies and future directions, promoting the development of more adaptable and efficient large language models.

Limitations

This survey provides a comprehensive overview of knowledge expansion techniques for LLMs. However, due to page constraints, we had to limit its scope and prioritize certain aspects:

First, the paper only provides a high-level overview of each method rather than an in-depth analysis. This can limit the understanding of the nuances and specific applications of each technique, as well as implementation details.

Second, our work is a literature review of adaptation methods rather than an empirical study evaluating their actual performance. While we analyze existing strategies, we do not benchmark or experimentally compare their effectiveness, leaving room for future studies to assess their practical impact under real-world conditions.

Third, we focus solely on text-based models and do not cover vision-language models, which integrate multi-modal learning for textual and visual understanding. While the methods covered in this survey could be used to adapt the language encoders of such models in theory, extending these adaptation methods to vision-language models remains an open research direction.

Finally, this survey reflecting the current state of research might become outdated as new research is published, as the field of LLMs is rapidly evolving and new methods for knowledge expansion are continuously being developed.

References

Jesujoba O. Alabi, David Ifeoluwa Adelani, Marius Mosbach, and Dietrich Klakow. 2022. [Adapting pre-trained language models to African languages via multilingual adaptive fine-tuning](#). In *Proceedings of the 29th International Conference on Computational Linguistics*, pages 4336–4349, Gyeongju, Republic of Korea. International Committee on Computational Linguistics.

Uri Alon, Frank Xu, Junxian He, Sudipta Sengupta, Dan Roth, and Graham Neubig. 2022. [Neuro-symbolic language modeling with automaton-augmented retrieval](#). In *Proceedings of the 39th International Conference on Machine Learning*, volume 162 of *Proceedings of Machine Learning Research*, pages 468–485. PMLR.

Ge Bai, Jie Liu, Xingyuan Bu, Yancheng He, Jiaheng Liu, Zhanhui Zhou, Zhuoran Lin, Wenbo Su, Tiezheng Ge, Bo Zheng, and Wanli Ouyang. 2024. [MT-bench-101: A fine-grained benchmark for evaluating large language models in multi-turn dialogues](#). In *Proceedings of the 62nd Annual Meeting of the*

Association for Computational Linguistics (Volume 1: Long Papers), pages 7421–7454, Bangkok, Thailand. Association for Computational Linguistics.

Lukas Berglund, Meg Tong, Maximilian Kaufmann, Mikita Balesni, Asa Cooper Stickland, Tomasz Korbak, and Owain Evans. 2003. The reversal curse: Llms trained on “a is b” fail to learn “b is a”. In *The Twelfth International Conference on Learning Representations*.

Baolong Bi, Shenghua Liu, Lingrui Mei, Yiwei Wang, Pengliang Ji, and Xueqi Cheng. 2024. [Decoding by contrasting knowledge: Enhancing llms’ confidence on edited facts](#). *arXiv preprint arXiv:2405.11613*.

Terra Blevins, Tomasz Limisiewicz, Suchin Gururangan, Margaret Li, Hila Gonen, Noah A. Smith, and Luke Zettlemoyer. 2024. [Breaking the curse of multilinguality with cross-lingual expert language models](#). In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 10822–10837, Miami, Florida, USA. Association for Computational Linguistics.

Xingyuan Bu, Junran Peng, Junjie Yan, Tieniu Tan, and Zhaoxiang Zhang. 2021. [GAIA: A Transfer Learning System of Object Detection that Fits Your Needs](#). In *2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 274–283. IEEE Computer Society.

Xingyuan Bu, Yuwei Wu, Zhi Gao, and Yunde Jia. 2019. [Deep convolutional network with locality and sparsity constraints for texture classification](#). *Pattern Recognition*, 91:34–46.

Boxi Cao, Hongyu Lin, Xianpei Han, and Le Sun. 2024. [The life cycle of knowledge in big language models: A survey](#). *Machine Intelligence Research*, 21(2):217–238.

Ruizhe Chen, Yichen Li, Zikai Xiao, and Zuozhu Liu. 2024a. [Large language model bias mitigation from the perspective of knowledge editing](#). *arXiv preprint arXiv:2405.09341*.

Wuyang Chen, Yanqi Zhou, Nan Du, Yanping Huang, James Laudon, Zhifeng Chen, and Claire Cui. 2023. [Lifelong language pretraining with distribution-specialized experts](#). In *Proceedings of the 40th International Conference on Machine Learning*, volume 202 of *Proceedings of Machine Learning Research*, pages 5383–5395. PMLR.

Yingfa Chen, Zhengyan Zhang, Xu Han, Chaojun Xiao, Zhiyuan Liu, Chen Chen, Kuai Li, Tao Yang, and Maosong Sun. 2024b. [Robust and scalable model editing for large language models](#). In *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, pages 14157–14172, Torino, Italia. ELRA and ICCL.

Zhiyuan Chen and Bing Liu. 2018. [Continual learning and catastrophic forgetting](#). In *Lifelong Machine Learning*, pages 55–75. Springer.

845	Andrea Cossu, Antonio Carta, Lucia Passaro, Vincenzo	<i>Long Papers</i>), pages 16477–16508, Toronto, Canada.	902
846	Lomonaco, Tinne Tuytelaars, and Davide Bacciu.	Association for Computational Linguistics.	903
847	2024. Continual pre-training mitigates forgetting in		
848	language and vision . <i>Neural Networks</i> , 179:106492.		
849	Damai Dai, Li Dong, Yaru Hao, Zhifang Sui, Baobao	Mor Geva, Roei Schuster, Jonathan Berant, and Omer	904
850	Chang, and Furu Wei. 2022. Knowledge neurons in	Levy. 2021. Transformer feed-forward layers are key-	905
851	pretrained transformers . In <i>Proceedings of the 60th</i>	value memories . In <i>Proceedings of the 2021 Confer-</i>	906
852	<i>Annual Meeting of the Association for Computational</i>	<i>ence on Empirical Methods in Natural Language Pro-</i>	907
853	<i>Linguistics (Volume 1: Long Papers)</i> , pages 8493–	<i>cessing</i> , pages 5484–5495, Online and Punta Cana,	908
854	8502, Dublin, Ireland. Association for Computational	Dominican Republic. Association for Computational	909
855	Linguistics.	Linguistics.	910
856	Bhavana Dalvi Mishra, Oyvind Taffjord, and Peter Clark.	Wenhan Xiong Grattafiori, Alexandre Défossez, Jade	911
857	2022. Towards teachable reasoning systems: Using a	Copet, Faisal Azhar, Hugo Touvron, Louis Martin,	912
858	dynamic memory of user feedback for continual sys-	Nicolas Usunier, Thomas Scialom, and Gabriel Syn-	913
859	tem improvement . In <i>Proceedings of the 2022 Confer-</i>	naeve. 2023. Code llama: Open foundation models	914
860	<i>ference on Empirical Methods in Natural Language</i>	for code . <i>arXiv preprint arXiv:2308.12950</i> .	915
861	<i>Processing</i> , pages 9465–9480, Abu Dhabi, United		
862	Arab Emirates. Association for Computational Lin-	Hengrui Gu, Kaixiong Zhou, Xiaotian Han, Ning-	916
863	guistics.	hao Liu, Ruobing Wang, and Xin Wang. 2024a.	917
864	Nicola De Cao, Wilker Aziz, and Ivan Titov. 2021. Edit-	PokeMQA: Programmable knowledge editing for	918
865	ing factual knowledge in language models . In <i>Pro-</i>	multi-hop question answering . In <i>Proceedings of</i>	919
866	<i>ceedings of the 2021 Conference on Empirical Meth-</i>	<i>the 62nd Annual Meeting of the Association for</i>	920
867	<i>ods in Natural Language Processing</i> , pages 6491–	<i>Computational Linguistics (Volume 1: Long Papers)</i> ,	921
868	6506, Online and Punta Cana, Dominican Republic.	pages 8069–8083, Bangkok, Thailand. Association	922
869	Association for Computational Linguistics.	for Computational Linguistics.	923
870	Bhuwan Dhingra, Jeremy R. Cole, Julian Martin	Jia-Chen Gu, Hao-Xiang Xu, Jun-Yu Ma, Pan Lu, Zhen-	924
871	Eisenschlos, Daniel Gillick, Jacob Eisenstein, and	Hua Ling, Kai-Wei Chang, and Nanyun Peng. 2024b.	925
872	William W. Cohen. 2022. Time-aware language mod-	Model editing harms general abilities of large lan-	926
873	els as temporal knowledge bases . <i>Transactions of the</i>	guage models: Regularization to the rescue . In <i>Pro-</i>	927
874	<i>Association for Computational Linguistics</i> , 10:257–	<i>ceedings of the 2024 Conference on Empirical Meth-</i>	928
875	273.	<i>ods in Natural Language Processing</i> , pages 16801–	929
876	Konstantin Dobler and Gerard de Melo. 2023. FOCUS:	16819, Miami, Florida, USA. Association for Com-	930
877	Effective embedding initialization for monolingual	putational Linguistics.	931
878	specialization of multilingual models . In <i>Proceed-</i>	Akshat Gupta, Anurag Rao, and Gopala Anu-	932
879	<i>ings of the 2023 Conference on Empirical Methods in</i>	manchipalli. 2024. Model editing at scale leads to	933
880	<i>Natural Language Processing</i> , pages 13440–13454,	gradual and catastrophic forgetting . In <i>Findings of</i>	934
881	Singapore. Association for Computational Linguis-	<i>the Association for Computational Linguistics: ACL</i>	935
882	tics.	2024, pages 15202–15232, Bangkok, Thailand. As-	936
883	Kazuki Fujii, Taishi Nakamura, Mengsay Loem, Hi-	sociation for Computational Linguistics.	937
884	roki Iida, Masanari Ohi, Kakeru Hattori, Hirai Shota,	Suchin Gururangan, Mike Lewis, Ari Holtzman,	938
885	Sakae Mizuki, Rio Yokota, and Naoaki Okazaki.	Noah A. Smith, and Luke Zettlemoyer. 2022. DEMix	939
886	2024. Continual pre-training for cross-lingual LLM	layers: Disentangling domains for modular language	940
887	adaptation: Enhancing japanese language capabili-	modeling . In <i>Proceedings of the 2022 Conference of</i>	941
888	ties . In <i>First Conference on Language Modeling</i> .	<i>the North American Chapter of the Association for</i>	942
889	Isabel O. Gallegos, Ryan A. Rossi, Joe Barrow,	<i>Computational Linguistics: Human Language Tech-</i>	943
890	Md Mehrab Tanjim, Sungchul Kim, Franck Dernon-	<i>nologies</i> , pages 5557–5576, Seattle, United States.	944
891	court, Tong Yu, Ruiyi Zhang, and Nesreen K. Ahmed.	Association for Computational Linguistics.	945
892	2024. Bias and fairness in large language models:	Suchin Gururangan, Ana Marasović, Swabha	946
893	A survey . <i>Computational Linguistics</i> , 50(3):1097–	Swayamdipta, Kyle Lo, Iz Beltagy, Doug Downey,	947
894	1179.	and Noah A. Smith. 2020. Don’t stop pretraining:	948
895	Luyu Gao, Zhuyun Dai, Panupong Pasupat, Anthony	Adapt language models to domains and tasks . In	949
896	Chen, Arun Tejasvi Chaganty, Yicheng Fan, Vincent	<i>Proceedings of the 58th Annual Meeting of the</i>	950
897	Zhao, Ni Lao, Hongrae Lee, Da-Cheng Juan, and	<i>Association for Computational Linguistics</i> , pages	951
898	Kelvin Guu. 2023. RARR: Researching and revising	8342–8360, Online. Association for Computational	952
899	what language models say, using language models .	Linguistics.	953
900	In <i>Proceedings of the 61st Annual Meeting of the</i>	Junxian He, Graham Neubig, and Taylor Berg-	954
901	<i>Association for Computational Linguistics (Volume 1:</i>	Kirkpatrick. 2021. Efficient nearest neighbor lan-	955
		guage models . In <i>Proceedings of the 2021 Confer-</i>	956
		<i>ence on Empirical Methods in Natural Language Pro-</i>	957
		<i>cessing</i> , pages 5703–5714, Online and Punta Cana,	958

959	Dominican Republic. Association for Computational Linguistics.	1015
960		1016
961	Evan Hernandez, Arnab Sen Sharma, Tal Haklay, Kevin Meng, Martin Wattenberg, Jacob Andreas, Yonatan Belinkov, and David Bau. 2024. Linearity of relation decoding in transformer language models . In <i>The Twelfth International Conference on Learning Representations</i> .	1017
962		1018
963		1019
964		1020
965		1021
966		1022
967	Nathan Hu, Eric Mitchell, Christopher Manning, and Chelsea Finn. 2023. Meta-learning online adaptation of language models . In <i>Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing</i> , pages 4418–4432, Singapore. Association for Computational Linguistics.	1023
968		1024
969		1025
970		1026
971		1027
972		1028
973	Yichong Huang, Xiachong Feng, Xiaocheng Feng, and Bing Qin. 2021. The factual inconsistency problem in abstractive text summarization: A survey . <i>arXiv preprint arXiv:2104.14839</i> .	1029
974		1030
975		1031
976		1032
977	Zeyu Huang, Yikang Shen, Xiaofeng Zhang, Jie Zhou, Wenge Rong, and Zhang Xiong. 2023. Transformer-patcher: One mistake worth one neuron . In <i>The Eleventh International Conference on Learning Representations</i> .	1033
978		1034
979		1035
980		1036
981		1037
982	Ayyoob Imani, Peiqin Lin, Amir Hossein Kargaran, Silvia Severini, Masoud Jalili Sabet, Nora Kassner, Chunlan Ma, Helmut Schmid, André Martins, François Yvon, and Hinrich Schütze. 2023. Glot500: Scaling multilingual corpora and language models to 500 languages . In <i>Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 1082–1117, Toronto, Canada. Association for Computational Linguistics.	1038
983		1039
984		1040
985		1041
986		1042
987		1043
988		1044
989		1045
990		1046
991		1047
992	Joel Jang, Seonghyeon Ye, Sohee Yang, Joongbo Shin, Janghoon Han, Gyeonghun Kim, Stanley Jungkyu Choi, and Minjoon Seo. 2022. Towards continual knowledge learning of language models . In <i>International Conference on Learning Representations</i> .	1048
993		1049
994		1050
995		1051
996		1052
997	Shaoxiong Ji, Zihao Li, Indraneil Paul, Jaakko Paavola, Peiqin Lin, Pinzhen Chen, Dayyán O’Brien, Hengyu Luo, Hinrich Schütze, Jörg Tiedemann, et al. 2024. Emma-500: Enhancing massively multilingual adaptation of large language models . <i>arXiv preprint arXiv:2409.17892</i> .	1053
998		1054
999		1055
1000		1056
1001		1057
1002		1058
1003	Ziwei Ji, Nayeon Lee, Rita Frieske, Tiezheng Yu, Dan Su, Yan Xu, Etsuko Ishii, Ye Jin Bang, Andrea Madotto, and Pascale Fung. 2023. Survey of hallucination in natural language generation . <i>ACM Comput. Surv.</i> , 55(12).	1059
1004		1060
1005		1061
1006		1062
1007		1063
1008	Jiajie Jin, Yutao Zhu, Yujia Zhou, and Zhicheng Dou. 2024. BIDER: Bridging knowledge inconsistency for efficient retrieval-augmented LLMs via key supporting evidence . In <i>Findings of the Association for Computational Linguistics: ACL 2024</i> , pages 750–761, Bangkok, Thailand. Association for Computational Linguistics.	1064
1009		1065
1010		1066
1011		1067
1012		1068
1013		1069
1014		1070
		1071
		1072
	Xisen Jin, Dejiao Zhang, Henghui Zhu, Wei Xiao, Shang-Wen Li, Xiaokai Wei, Andrew Arnold, and Xiang Ren. 2022. Lifelong pretraining: Continually adapting language models to emerging corpora . In <i>Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies</i> , pages 4764–4780, Seattle, United States. Association for Computational Linguistics.	
	Zixuan Ke, Haowei Lin, Yijia Shao, Hu Xu, Lei Shu, and Bing Liu. 2022. Continual training of language models for few-shot learning . In <i>Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing</i> , pages 10205–10216, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.	
	Zixuan Ke and Bing Liu. 2022. Continual learning of natural language processing tasks: A survey . <i>arXiv preprint arXiv:2211.12701</i> .	
	Zixuan Ke, Yijia Shao, Haowei Lin, Tatsuya Konishi, Gyuhak Kim, and Bing Liu. 2023. Continual pre-training of language models . In <i>The Eleventh International Conference on Learning Representations</i> .	
	Urvashi Khandelwal, Omer Levy, Dan Jurafsky, Luke Zettlemoyer, and Mike Lewis. 2020. Generalization through memorization: Nearest neighbor language models . In <i>International Conference on Learning Representations</i> .	
	Yujin Kim, Jaehong Yoon, Seonghyeon Ye, Sangmin Bae, Namgyu Ho, Sung Ju Hwang, and Se-Young Yun. 2024. Carpe diem: On the evaluation of world knowledge in lifelong language models . In <i>Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)</i> , pages 5401–5415, Mexico City, Mexico. Association for Computational Linguistics.	
	Nathan Lambert, Louis Castricato, Leandro von Werra, and Alex Havrilla. 2022. Illustrating reinforcement learning from human feedback (rlhf). <i>Hugging Face Blog</i> . https://huggingface.co/blog/rlhf .	
	Angeliki Lazaridou, Adhiguna Kuncoro, Elena Gribovskaya, Devang Agrawal, Adam Liška, Tayfun Terzi, Mai Gimenez, Cyprien de Masson d’Autume, Tomas Kocisky, Sebastian Ruder, Dani Yogatama, Kris Cao, Susannah Young, and Phil Blunsom. 2021. Mind the gap: Assessing temporal generalization in neural language models. In <i>Proceedings of the 35th International Conference on Neural Information Processing Systems, NIPS ’21</i> . Curran Associates Inc.	
	Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, Sebastian Riedel, and Douwe Kiela. 2020. Retrieval-augmented generation for knowledge-intensive nlp tasks. In <i>Proceedings of the 34th International Conference on Neural Information Processing Systems, NIPS ’20</i> . Curran Associates Inc.	

1073	Jiatong Li, Yunqing Liu, Wenqi Fan, Xiao-Yong Wei,	Dong Liu, Roger Waleffe, Meng Jiang, and Shivaram	1128
1074	Hui Liu, Jiliang Tang, and Qing Li. 2024a. Em-	Venkataraman. 2024a. Graphsnapshot: Graph ma-	1129
1075	powering Molecule Discovery for Molecule-Caption	chine learning acceleration with fast storage and re-	1130
1076	Translation With Large Language Models: A Chat-	trieval. <i>arXiv preprint arXiv:2406.17918</i> .	1131
1077	GPT Perspective . <i>IEEE Transactions on Knowledge</i>		
1078	& Data Engineering , 36(11):6071–6083.		
1079	Qi Li, Xiang Liu, Zhenheng Tang, Peijie Dong, Zeyu	Shengchao Liu, Weili Nie, Chengpeng Wang, Jiarui	1132
1080	Li, Xinglin Pan, and Xiaowen Chu. 2024b. Should	Lu, Zhuoran Qiao, Ling Liu, Jian Tang, Chaowei	1133
1081	we really edit language models? on the evaluation	Xiao, and Animashree Anandkumar. 2023. Multi-	1134
1082	of edited language models. In <i>The Thirty-eighth An-</i>	modal molecule structure–text model for text-based	1135
1083	<i>annual Conference on Neural Information Processing</i>	retrieval and editing. <i>Nature Machine Intelligence</i> ,	1136
1084	<i>Systems</i> .	5(12):1447–1457.	1137
1085	Shilong Li, Yancheng He, Hangyu Guo, Xingyuan Bu,	Yihong Liu, Peiqin Lin, Mingyang Wang, and Hinrich	1138
1086	Ge Bai, Jie Liu, Jiaheng Liu, Xingwei Qu, Yang-	Schuetze. 2024b. OFA: A framework of initializing	1139
1087	guang Li, Wanli Ouyang, Wenbo Su, and Bo Zheng.	unseen subword embeddings for efficient large-scale	1140
1088	2024c. GraphReader: Building graph-based agent to	multilingual continued pretraining. In <i>Findings of the</i>	1141
1089	enhance long-context abilities of large language mod-	<i>Association for Computational Linguistics: NAACL</i>	1142
1090	els. In <i>Findings of the Association for Computational</i>	2024, pages 1067–1097, Mexico City, Mexico. Asso-	1143
1091	<i>Linguistics: EMNLP 2024</i> , pages 12758–12786, Mi-	ciation for Computational Linguistics.	1144
1092	ami, Florida, USA. Association for Computational		
1093	<i>Linguistics</i> .	Alejandro Lozano, Scott L Fleming, Chia-Chun Chiang,	1145
1094	Xiaopeng Li, Shasha Li, Shezheng Song, Jing Yang, Jun	and Nigam Shah. 2023. Clinfo. ai: An open-source	1146
1095	Ma, and Jie Yu. 2024d. Pmet: Precise model editing	retrieval-augmented large language model system for	1147
1096	in a transformer. In <i>Proceedings of the AAAI Con-</i>	answering medical questions using scientific litera-	1148
1097	<i>ference on Artificial Intelligence</i> , volume 38, pages	ture. In <i>PACIFIC SYMPOSIUM ON BIOCOMPUT-</i>	1149
1098	18564–18572.	<i>ING 2024</i> , pages 8–23. World Scientific.	1150
1099	Zhoubo Li, Ningyu Zhang, Yunzhi Yao, Mengru Wang,	Anton Lozhkov, Raymond Li, Loubna Ben Allal, Fed-	1151
1100	Xi Chen, and Huajun Chen. 2024e. Unveiling the pit-	erico Cassano, Joel Lamy-Poirier, Nouamane Tazi,	1152
1101	falls of knowledge editing for large language models.	Ao Tang, Dmytro Pykhtar, Jiawei Liu, Yuxiang Wei,	1153
1102	In <i>The Twelfth International Conference on Learning</i>	et al. 2024. Starcoder 2 and the stack v2: The next	1154
1103	<i>Representations</i> .	generation. <i>arXiv preprint arXiv:2402.19173</i> .	1155
1104	Zichao Li, Bingyang Wang, and Ying Chen. 2024f. In-	Yinquan Lu, Wenhao Zhu, Lei Li, Yu Qiao, and Fei	1156
1105	corporating economic indicators and market senti-	Yuan. 2024. LLaMAX: Scaling linguistic horizons	1157
1106	ment effect into us treasury bond yield prediction	of LLM by enhancing translation capabilities beyond	1158
1107	with machine learning. <i>Journal of Infrastructure,</i>	100 languages. In <i>Findings of the Association for</i>	1159
1108	<i>Policy and Development</i> , 8(9):7671.	<i>Computational Linguistics: EMNLP 2024</i> , pages	1160
1109	Zichao Li, Bingyang Wang, and Ying Chen. 2024g.	10748–10772, Miami, Florida, USA. Association for	1161
1110	Knowledge graph embedding and few-shot relational	<i>Computational Linguistics</i> .	1162
1111	learning methods for digital assets in usa. <i>Journal of</i>	Kelvin Luu, Daniel Khashabi, Suchin Gururangan, Kar-	1163
1112	<i>Industrial Engineering and Applied Science</i> , 2(5):10–	ishma Mandyam, and Noah A. Smith. 2022. Time	1164
1113	18.	waits for no one! analysis and challenges of tem-	1165
1114	Demiao Lin. 2024. Revolutionizing retrieval-	poral misalignment. In <i>Proceedings of the 2022</i>	1166
1115	augmented generation with enhanced pdf structure	<i>Conference of the North American Chapter of the</i>	1167
1116	recognition. <i>arXiv preprint arXiv:2401.12599</i> .	<i>Association for Computational Linguistics: Human</i>	1168
1117	Adam Liska, Tomas Kocisky, Elena Gribovskaya, Tay-	<i>Language Technologies</i> , pages 5944–5958, Seattle,	1169
1118	fun Terzi, Eren Sezener, Devang Agrawal, Cyprien	United States. Association for Computational Lin-	1170
1119	De Masson D’Autume, Tim Scholtes, Manzil Zaheer,	guistics.	1171
1120	Susannah Young, Ellen Gilsonan-Mcmahon, Sophia	Chang Ma, Haiteng Zhao, Lin Zheng, Jiayi Xin, Qin-	1172
1121	Austin, Phil Blunsom, and Angeliki Lazaridou. 2022.	tong Li, Lijun Wu, Zhihong Deng, Yang Young Lu,	1173
1122	StreamingQA: A benchmark for adaptation to new	Qi Liu, Sheng Wang, and Lingpeng Kong. 2024. Re-	1174
1123	knowledge over time in question answering models.	trieved sequence augmentation for protein represen-	1175
1124	In <i>Proceedings of the 39th International Conference</i>	tation learning. In <i>Proceedings of the 2024 Con-</i>	1176
1125	<i>on Machine Learning</i> , volume 162 of <i>Proceedings</i>	<i>ference on Empirical Methods in Natural Language</i>	1177
1126	<i>of Machine Learning Research</i> , pages 13604–13622.	<i>Processing</i> , pages 1738–1767, Miami, Florida, USA.	1178
1127	PMLR.	Association for Computational Linguistics.	1179
		Jun-Yu Ma, Jia-Chen Gu, Zhen-Hua Ling, Quan Liu,	1180
		and Cong Liu. 2023. Untying the reversal curse via	1181
		bidirectional language model editing. <i>arXiv preprint</i>	1182
		<i>arXiv:2310.10322</i> .	1183

1184	Aman Madaan, Niket Tandon, Peter Clark, and Yiming Yang. 2022. Memory-assisted prompt editing to improve GPT-3 after deployment . In <i>Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing</i> , pages 2833–2861, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.	1241
1185		1242
1186		1243
1187		
1188		1244
1189		1245
1190		1246
1191	Alex Mallen, Akari Asai, Victor Zhong, Rajarshi Das, Daniel Khashabi, and Hannaneh Hajishirzi. 2023. When not to trust language models: Investigating effectiveness of parametric and non-parametric memories . In <i>Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 9802–9822, Toronto, Canada. Association for Computational Linguistics.	1247
1192		1248
1193		
1194		1249
1195		1250
1196		1251
1197		1252
1198		1253
1199	Shengyu Mao, Xiaohan Wang, Mengru Wang, Yong Jiang, Pengjun Xie, Fei Huang, and Ningyu Zhang. 2024. Editing personality for large language models . In <i>Natural Language Processing and Chinese Computing: 13th National CCF Conference, NLPCC 2024, Hangzhou, China, November 1–3, 2024, Proceedings, Part II</i> , page 241–254. Springer-Verlag.	1254
1200		1255
1201		1256
1202		1257
1203		1258
1204		
1205		
1206	Michael McCloskey and Neal J Cohen. 1989. Catastrophic interference in connectionist networks: The sequential learning problem . In <i>Psychology of Learning and Motivation</i> , volume 24, pages 109–165. Academic Press.	1259
1207		1260
1208		1261
1209		1262
1210		1263
1211	Kevin Meng, David Bau, Alex J Andonian, and Yonatan Belinkov. 2022. Locating and editing factual associations in GPT . In <i>Advances in Neural Information Processing Systems</i> , volume 35, pages 17359–17372. Curran Associates, Inc.	1264
1212		1265
1213		1266
1214		1267
1215		1268
1216	Kevin Meng, Arnab Sen Sharma, Alex J Andonian, Yonatan Belinkov, and David Bau. 2023. Mass-editing memory in a transformer . In <i>The Eleventh International Conference on Learning Representations</i> .	1269
1217		1270
1218		1271
1219		1272
1220		1273
1221	Benjamin Minixhofer, Fabian Paischer, and Navid Rekasaz. 2022. WECHSEL: Effective initialization of subword embeddings for cross-lingual transfer of monolingual language models . In <i>Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies</i> , pages 3992–4006, Seattle, United States. Association for Computational Linguistics.	1274
1222		1275
1223		1276
1224		1277
1225		1278
1226		1279
1227		1280
1228		1281
1229		1282
1230	Benjamin Minixhofer, Edoardo Maria Ponti, and Ivan Vulić. 2024. Zero-shot tokenizer transfer . In <i>The Thirty-eighth Annual Conference on Neural Information Processing Systems</i> .	1283
1231		1284
1232		1285
1233		
1234	Eric Mitchell, Charles Lin, Antoine Bosselut, Chelsea Finn, and Christopher D Manning. 2022a. Fast model editing at scale . In <i>International Conference on Learning Representations</i> .	1286
1235		1287
1236		1288
1237		1289
1238	Eric Mitchell, Charles Lin, Antoine Bosselut, Christopher D Manning, and Chelsea Finn. 2022b. Memory-based model editing at scale . In <i>Proceedings of the 39th International Conference on Machine Learning</i> , volume 162 of <i>Proceedings of Machine Learning Research</i> , pages 15817–15831. PMLR.	1290
1239		1291
1240		1292
		1293
		1294
		1295
		1296
		1297
	Ercong Nie, Bo Shao, Zifeng Ding, Mingyang Wang, Helmut Schmid, and Hinrich Schütze. 2024. Bmike-53: Investigating cross-lingual knowledge editing with in-context learning . <i>arXiv preprint arXiv:2406.17764</i> .	
	Chaima Njeh, Haïfa Nakouri, and Fehmi Jaafar. 2024. Enhancing RAG-retrieval to improve LLMs robustness and resilience to hallucinations. In <i>Hybrid Artificial Intelligent Systems</i> , pages 201–213. Springer Nature Switzerland.	
	Vaidehi Patil, Peter Hase, and Mohit Bansal. 2024. Can sensitive information be deleted from LLMs? objectives for defending against extraction attacks . In <i>The Twelfth International Conference on Learning Representations</i> .	
	Baolin Peng, Michel Galley, Pengcheng He, Hao Cheng, Yujia Xie, Yu Hu, Qiuyuan Huang, Lars Liden, Zhou Yu, Weizhu Chen, and Jianfeng Gao. 2023a. Check your facts and try again: Improving large language models with external knowledge and automated feedback . <i>arXiv preprint arXiv:2302.12813</i> .	
	Guangyue Peng, Tao Ge, Si-Qing Chen, Furu Wei, and Houfeng Wang. 2023b. Semiparametric language models are scalable continual learners . <i>arXiv preprint arXiv:2303.01421</i> .	
	Ofir Press, Muru Zhang, Sewon Min, Ludwig Schmidt, Noah Smith, and Mike Lewis. 2023. Measuring and narrowing the compositionality gap in language models . In <i>Findings of the Association for Computational Linguistics: EMNLP 2023</i> , pages 5687–5711, Singapore. Association for Computational Linguistics.	
	Yujia Qin, Jiajie Zhang, Yankai Lin, Zhiyuan Liu, Peng Li, Maosong Sun, and Jie Zhou. 2022. ELLE: Efficient lifelong pre-training for emerging data . In <i>Findings of the Association for Computational Linguistics: ACL 2022</i> , pages 2789–2810, Dublin, Ireland. Association for Computational Linguistics.	
	Ori Ram, Yoav Levine, Itay Dalmedigos, Dor Muhlgay, Amnon Shashua, Kevin Leyton-Brown, and Yoav Shoham. 2023. In-context retrieval-augmented language models . <i>Transactions of the Association for Computational Linguistics</i> , 11:1316–1331.	
	Anastasia Razdaibiedina, Yuning Mao, Rui Hou, Madihan Khabsa, Mike Lewis, and Amjad Almahairi. 2023. Progressive prompts: Continual learning for language models . In <i>The Eleventh International Conference on Learning Representations</i> .	
	Paul Röttger and Janet Pierrehumbert. 2021. Temporal adaptation of BERT and performance on downstream document classification: Insights from social media . In <i>Findings of the Association for Computational Linguistics: EMNLP 2021</i> , pages 2400–2412, Punta Cana, Dominican Republic. Association for Computational Linguistics.	

1298	Sebastian Ruder, Anders Søgaard, and Ivan Vulić. 2019.	CodeGemma Team, Heri Zhao, Jeffrey Hui, Joshua	1356
1299	Unsupervised cross-lingual representation learning .	Howland, Nam Nguyen, Siqi Zuo, Andrea Hu,	1357
1300	In <i>Proceedings of the 57th Annual Meeting of the</i>	Christopher A Choquette-Choo, Jingyue Shen, Joe	1358
1301	<i>Association for Computational Linguistics: Tutorial</i>	Kelley, et al. 2024. Codegemma: Open code models	1359
1302	<i>Abstracts</i> , pages 31–38, Florence, Italy. Association	based on gemma . <i>arXiv preprint arXiv:2406.11409</i> .	1360
1303	for Computational Linguistics.		
1304	Haizhou Shi, Zihao Xu, Hengyi Wang, Weiyi Qin,	Hugo Touvron, Louis Martin, Kevin Stone, Peter Al-	1361
1305	Wenyuan Wang, Yibin Wang, Zifeng Wang, Sayna	bert, Amjad Almahairi, Yasmine Babaei, Nikolay	1362
1306	Ebrahimi, and Hao Wang. 2024a. Continual learning	Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti	1363
1307	of large language models: A comprehensive survey .	Bhosale, et al. 2023. Llama 2: Open founda-	1364
1308		tion and fine-tuned chat models . <i>arXiv preprint</i>	1365
1309	Weijia Shi, Julian Michael, Suchin Gururangan, and	<i>arXiv:2307.09288</i> .	1366
1310	Luke Zettlemoyer. 2022. Nearest neighbor zero-shot	Ke Tran. 2020. From english to foreign languages:	1367
1311	inference . In <i>Proceedings of the 2022 Conference on</i>	Transferring pre-trained language models . <i>arXiv</i>	1368
1312	<i>Empirical Methods in Natural Language Processing</i> ,	<i>preprint arXiv:2002.07306</i> .	1369
1313	pages 3254–3265, Abu Dhabi, United Arab Emirates.		
1314	Association for Computational Linguistics.	Harsh Trivedi, Niranjan Balasubramanian, Tushar Khot,	1370
1315	Weijia Shi, Sewon Min, Michihiro Yasunaga, Minjoon	and Ashish Sabharwal. 2023. Interleaving retrieval	1371
1316	Seo, Richard James, Mike Lewis, Luke Zettlemoyer,	with chain-of-thought reasoning for knowledge-	1372
1317	and Wen-tau Yih. 2024b. REPLUG: Retrieval-	intensive multi-step questions . In <i>Proceedings of</i>	1373
1318	augmented black-box language models . In <i>Proceed-</i>	<i>ings of the 61st Annual Meeting of the Association for Com-</i>	1374
1319	<i>ings of the 2024 Conference of the North American</i>	<i>putational Linguistics (Volume 1: Long Papers)</i> ,	1375
1320	<i>Chapter of the Association for Computational Lin-</i>	pages 10014–10037, Toronto, Canada. Association	1376
1321	<i>guistics: Human Language Technologies (Volume</i>	for Computational Linguistics.	1377
1322	<i>1: Long Papers)</i> , pages 8371–8384, Mexico City,		
1323	Mexico. Association for Computational Linguistics.	Ahmet Üstün, Viraat Aryabumi, Zheng Yong, Wei-Yin	1378
1324	Yucheng Shi, Qiaoyu Tan, Xuansheng Wu, Shaochen	Ko, Daniel D’souza, Gbemileke Onilude, Neel Bhan-	1379
1325	Zhong, Kaixiong Zhou, and Ninghao Liu. 2024c.	dari, Shivalika Singh, Hui-Lee Ooi, Amr Kayid, Fred-	1380
1326	Retrieval-enhanced knowledge editing in language	die Vargus, Phil Blunsom, Shayne Longpre, Niklas	1381
1327	models for multi-hop question answering . In <i>Pro-</i>	Muennighoff, Marzieh Fadaee, Julia Kreutzer, and	1382
1328	<i>ceedings of the 33rd ACM International Conference</i>	Sara Hooker. 2024. Aya model: An instruction fine-	1383
1329	<i>on Information and Knowledge Management, CIKM</i>	tuned open-access multilingual language model . In	1384
1330	<i>’24</i> , page 2056–2066. Association for Computing	<i>Proceedings of the 62nd Annual Meeting of the As-</i>	1385
1331	Machinery.	<i>sociation for Computational Linguistics (Volume 1:</i>	1386
1332	Chenglei Si, Zhe Gan, Zhengyuan Yang, Shuohang	<i>Long Papers)</i> , pages 15894–15939, Bangkok, Thai-	1387
1333	Wang, Jianfeng Wang, Jordan Lee Boyd-Graber, and	land. Association for Computational Linguistics.	1388
1334	Lijuan Wang. 2023. Prompting GPT-3 to be reli-	Jianchen Wang, Zhouhong Gu, Zhuozhi Xiong, Hong-	1389
1335	able . In <i>The Eleventh International Conference on</i>	wei Feng, and Yanghua Xiao. 2024a. The missing	1390
1336	<i>Learning Representations</i> .	piece in model editing: A deep dive into the hidden	1391
1337	Xiaoshuai Song, Zhengyang Wang, Keqing He, Guant-	damage brought by model editing . <i>arXiv preprint</i>	1392
1338	ing Dong, Yutao Mou, Jinxu Zhao, and Weiran Xu.	<i>arXiv:2403.07825</i> .	1393
1339	2024. Knowledge editing on black-box large lan-	Liyuan Wang, Xingxing Zhang, Hang Su, and Jun Zhu.	1394
1340	guage models . <i>arXiv preprint arXiv:2402.08631</i> .	2024b. A comprehensive survey of continual learn-	1395
1341	Fang Sun, Zhihao Zhan, Hongyu Guo, Ming Zhang,	ing: theory, method and application . <i>IEEE Transac-</i>	1396
1342	and Jian Tang. 2023. Graphvf: Controllable protein-	<i>tions on Pattern Analysis and Machine Intelligence</i> .	1397
1343	specific 3d molecule generation with variational flow .	Mengru Wang, Ningyu Zhang, Ziwen Xu, Zekun Xi,	1398
1344	<i>arXiv preprint arXiv:2304.12825</i> .	Shumin Deng, Yunzhi Yao, Qishen Zhang, Linyi	1399
1345	Yu Sun, Shuohuan Wang, Yukun Li, Shikun Feng, Hao	Yang, Jindong Wang, and Huajun Chen. 2024c.	1400
1346	Tian, Hua Wu, and Haifeng Wang. 2020. Ernie 2.0:	Detoxifying large language models via knowledge	1401
1347	A continual pre-training framework for language un-	editing . In <i>Proceedings of the 62nd Annual Meeting</i>	1402
1348	derstanding . <i>Proceedings of the AAAI Conference on</i>	<i>of the Association for Computational Linguistics (Vol-</i>	1403
1349	<i>Artificial Intelligence</i> , 34(05):8968–8975.	<i>ume 1: Long Papers)</i> , pages 3093–3118, Bangkok,	1404
1350	Niket Tandon, Aman Madaan, Peter Clark, and Yiming	Thailand. Association for Computational Linguistics.	1405
1351	Yang. 2022. Learning to repair: Repairing model out-	Mingyang Wang, Heike Adel, Lukas Lange, Jannik	1406
1352	put errors after deployment using a dynamic memory	Strötgen, and Hinrich Schuetze. 2024d. Learn it or	1407
1353	of feedback . In <i>Findings of the Association for Com-</i>	leave it: Module composition and pruning for contin-	1408
1354	<i>putational Linguistics: NAACL 2022</i> , pages 339–352,	ual learning . In <i>Proceedings of the 9th Workshop on</i>	1409
1355	Seattle, United States. Association for Computational	<i>Representation Learning for NLP (RepL4NLP-2024)</i> ,	1410
	Linguistics.	pages 163–176, Bangkok, Thailand. Association for	1411
		Computational Linguistics.	1412

1413	Mingyang Wang, Heike Adel, Lukas Lange, Jannik Strötgen, and Hinrich Schuetze. 2024e. Rehearsal-free modular and compositional continual learning for language models . In <i>Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 2: Short Papers)</i> , pages 469–480, Mexico City, Mexico. Association for Computational Linguistics.	1471
1414		1472
1415		
1416		
1417		
1418		
1419		
1420		
1421		
1422	Mingyang Wang, Heike Adel, Lukas Lange, Jannik Strötgen, and Hinrich Schütze. 2023a. NLNDE at SemEval-2023 task 12: Adaptive pretraining and source language selection for low-resource multilingual sentiment analysis . In <i>Proceedings of the 17th International Workshop on Semantic Evaluation (SemEval-2023)</i> , pages 488–497, Toronto, Canada. Association for Computational Linguistics.	
1423		
1424		
1425		
1426		
1427		
1428		
1429		
1430	Mingyang Wang, Lukas Lange, Heike Adel, Jannik Strötgen, and Hinrich Schuetze. 2024f. Better call SAUL: Fluent and consistent language model editing with generation regularization . In <i>Findings of the Association for Computational Linguistics: EMNLP 2024</i> , pages 7990–8000, Miami, Florida, USA. Association for Computational Linguistics.	
1431		
1432		
1433		
1434		
1435		
1436		
1437	Xiaohan Wang, Shengyu Mao, Shumin Deng, Yunzhi Yao, Yue Shen, Lei Liang, Jinjie Gu, Huajun Chen, and Ningyu Zhang. 2024g. Editing conceptual knowledge for large language models . In <i>Findings of the Association for Computational Linguistics: EMNLP 2024</i> , pages 706–724, Miami, Florida, USA. Association for Computational Linguistics.	
1438		
1439		
1440		
1441		
1442		
1443		
1444	Zichao Wang, Weili Nie, Zhuoran Qiao, Chaowei Xiao, Richard Baraniuk, and Anima Anandkumar. 2023b. Retrieval-based controllable molecule generation . In <i>The Eleventh International Conference on Learning Representations</i> .	
1445		
1446		
1447		
1448		
1449	Zifeng Wang, Zichen Wang, Balasubramaniam Srinivasan, Vassilis N. Ioannidis, Huzefa Rangwala, and RISHITA ANUBHAI. 2024h. Biobridge: Bridging biomedical foundation models via knowledge graphs . In <i>The Twelfth International Conference on Learning Representations</i> .	
1450		
1451		
1452		
1453		
1454		
1455	Zihao Wei, Jingcheng Deng, Liang Pang, Hanxing Ding, Huawei Shen, and Xueqi Cheng. 2025. MLaKE: Multilingual knowledge editing benchmark for large language models . In <i>Proceedings of the 31st International Conference on Computational Linguistics</i> , pages 4457–4473, Abu Dhabi, UAE. Association for Computational Linguistics.	
1456		
1457		
1458		
1459		
1460		
1461		
1462	Chengyue Wu, Yukang Gan, Yixiao Ge, Zeyu Lu, Jiahao Wang, Ye Feng, Ying Shan, and Ping Luo. 2024a. LLaMA pro: Progressive LLaMA with block expansion . In <i>Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 6518–6537, Bangkok, Thailand. Association for Computational Linguistics.	
1463		
1464		
1465		
1466		
1467		
1468		
1469	Tongtong Wu, Linhao Luo, Yuan-Fang Li, Shirui Pan, Thuy-Trang Vu, and Gholamreza Haffari. 2024b. Continual learning for large language models: A survey . <i>arXiv preprint arXiv:2402.01364</i> .	1471
1470		1472
	Xinwei Wu, Junzhuo Li, Minghui Xu, Weilong Dong, Shuangzhi Wu, Chao Bian, and Deyi Xiong. 2023. DEPN: Detecting and editing privacy neurons in pre-trained language models . In <i>Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing</i> , pages 2875–2886, Singapore. Association for Computational Linguistics.	1473
		1474
		1475
		1476
		1477
		1478
		1479
	Jian Xie, Kai Zhang, Jiangjie Chen, Renze Lou, and Yu Su. 2024. Adaptive chameleon or stubborn sloth: Revealing the behavior of large language models in knowledge conflicts . In <i>The Twelfth International Conference on Learning Representations</i> .	1480
		1481
		1482
		1483
		1484
	Rongwu Xu, Zehan Qi, Zhijiang Guo, Cunxiang Wang, Hongru Wang, Yue Zhang, and Wei Xu. 2024. Knowledge conflicts for LLMs: A survey . In <i>Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing</i> , pages 8541–8565, Miami, Florida, USA. Association for Computational Linguistics.	1485
		1486
		1487
		1488
		1489
		1490
		1491
	Linting Xue, Noah Constant, Adam Roberts, Mihir Kale, Rami Al-Rfou, Aditya Siddhant, Aditya Barua, and Colin Raffel. 2021. mT5: A massively multilingual pre-trained text-to-text transformer . In <i>Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies</i> , pages 483–498, Online. Association for Computational Linguistics.	1492
		1493
		1494
		1495
		1496
		1497
		1498
		1499
	Prateek Yadav, Qing Sun, Hantian Ding, Xiaopeng Li, Dejiao Zhang, Ming Tan, Parminder Bhatia, Xiaofei Ma, Ramesh Nallapati, Murali Krishna Ramanathan, Mohit Bansal, and Bing Xiang. 2023. Exploring continual learning for code generation models . In <i>Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)</i> , pages 782–792, Toronto, Canada. Association for Computational Linguistics.	1500
		1501
		1502
		1503
		1504
		1505
		1506
		1507
		1508
	Ling Yang, Zhilin Huang, Xiangxin Zhou, Minkai Xu, Wentao Zhang, Yu Wang, Xiawu Zheng, Wenming Yang, Ron O. Dror, Shenda Hong, and Bin CUI. 2024a. Prompt-based 3d molecular diffusion models for structure-based drug design .	1509
		1510
		1511
		1512
		1513
	Wanli Yang, Fei Sun, Jiajun Tan, Xinyu Ma, Du Su, Dawei Yin, and Huawei Shen. 2024b. The fall of ROME: Understanding the collapse of LLMs in model editing . In <i>Findings of the Association for Computational Linguistics: EMNLP 2024</i> , pages 4079–4087, Miami, Florida, USA. Association for Computational Linguistics.	1514
		1515
		1516
		1517
		1518
		1519
		1520
	Yunzhi Yao, Peng Wang, Bozhong Tian, Siyuan Cheng, Zhoubo Li, Shumin Deng, Huajun Chen, and Ningyu Zhang. 2023. Editing large language models: Problems, methods, and opportunities . In <i>Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing</i> , pages 10222–10240, Singapore. Association for Computational Linguistics.	1521
		1522
		1523
		1524
		1525
		1526
		1527
		1528

comprehensive study of knowledge editing for large language models. *arXiv preprint arXiv:2401.01286*.

Zihan Zhang, Meng Fang, Ling Chen, Mohammad-Reza Namazi-Rad, and Jun Wang. 2023c. [How do large language models capture the ever-changing world knowledge? a review of recent advances](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 8289–8311, Singapore. Association for Computational Linguistics.

Wayne Xin Zhao, Jing Liu, Ruiyang Ren, and Ji-Rong Wen. 2024. [Dense text retrieval based on pretrained language models: A survey](#). *ACM Trans. Inf. Syst.*, 42(4).

Ce Zheng, Lei Li, Qingxiu Dong, Yuxuan Fan, Zhiyong Wu, Jingjing Xu, and Baobao Chang. 2023. [Can we edit factual knowledge by in-context learning?](#) In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 4862–4876, Singapore. Association for Computational Linguistics.

Wenzhen Zheng, Wenbo Pan, Xu Xu, Libo Qin, Li Yue, and Ming Zhou. 2024. [Breaking language barriers: Cross-lingual continual pre-training at scale](#). In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 7725–7738, Miami, Florida, USA. Association for Computational Linguistics.

Zexuan Zhong, Zhengxuan Wu, Christopher Manning, Christopher Potts, and Danqi Chen. 2023. [MQuAKE: Assessing knowledge editing in language models via multi-hop questions](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 15686–15702, Singapore. Association for Computational Linguistics.

Wenxuan Zhou, Sheng Zhang, Hoifung Poon, and Muhao Chen. 2023. [Context-faithful prompting for large language models](#). In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pages 14544–14556, Singapore. Association for Computational Linguistics.

Chen Zhu, Ankit Singh Rawat, Manzil Zaheer, Srinadh Bhojanapalli, Daliang Li, Felix Yu, and Sanjiv Kumar. 2020. [Modifying memories in transformer models](#). *arXiv preprint arXiv:2012.00363*.

Daniel M Ziegler, Nisan Stiennon, Jeffrey Wu, Tom B Brown, Alec Radford, Dario Amodei, Paul Christiano, and Geoffrey Irving. 2019. [Fine-tuning language models from human preferences](#). *arXiv preprint arXiv:1909.08593*.

Antonio Jimeno Yepes, Yao You, Jan Milczek, Sebastian Laverde, and Renyu Li. 2024. [Financial report chunking for effective retrieval augmented generation](#). *arXiv preprint arXiv:2402.05131*.

Pengfei Yu and Heng Ji. 2024. [Information association for language model updating by mitigating LM-logical discrepancy](#). In *Proceedings of the 28th Conference on Computational Natural Language Learning*, pages 117–129, Miami, FL, USA. Association for Computational Linguistics.

Zichun Yu, Chenyan Xiong, Shi Yu, and Zhiyuan Liu. 2023. [Augmentation-adapted retriever improves generalization of language models as generic plug-in](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 2421–2436, Toronto, Canada. Association for Computational Linguistics.

Daoguang Zan, Bei Chen, Dejian Yang, Zeqi Lin, Minsu Kim, Bei Guan, Yongji Wang, Weizhu Chen, and Jian-Guang Lou. 2022. [CERT: Continual pre-training on sketches for library-oriented code generation](#). In *Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence, IJCAI-22*, pages 2369–2375. International Joint Conferences on Artificial Intelligence Organization.

Boyu Zhang, Hongyang Yang, Tianyu Zhou, Muhammad Ali Babar, and Xiao-Yang Liu. 2023a. [Enhancing financial sentiment analysis via retrieval augmented large language models](#). In *Proceedings of the Fourth ACM International Conference on AI in Finance, ICAIF ’23*, pages 349–356. Association for Computing Machinery.

Han Zhang, Lin Gui, Yuanzhao Zhai, Hui Wang, Yu Lei, and Ruifeng Xu. 2023b. [Copf: Continual learning human preference through optimal policy fitting](#). *arXiv preprint arXiv:2310.15694*.

Han Zhang, Yu Lei, Lin Gui, Min Yang, Yulan He, Hui Wang, and Ruifeng Xu. 2024a. [Cppo: Continual learning for reinforcement learning with human feedback](#). In *The Twelfth International Conference on Learning Representations*.

Miaoran Zhang, Mingyang Wang, Jesujoba Alabi, and Dietrich Klakow. 2024b. [AAdAM at SemEval-2024 task 1: Augmentation and adaptation for multilingual semantic textual relatedness](#). In *Proceedings of the 18th International Workshop on Semantic Evaluation (SemEval-2024)*, pages 800–810, Mexico City, Mexico. Association for Computational Linguistics.

Michael Zhang and Eunsol Choi. 2023. [Mitigating temporal misalignment by discarding outdated facts](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 14213–14226, Singapore. Association for Computational Linguistics.

Ningyu Zhang, Yunzhi Yao, Bozhong Tian, Peng Wang, Shumin Deng, Mengru Wang, Zekun Xi, Shengyu Mao, Jintian Zhang, Yuansheng Ni, et al. 2024c. [A](#)

A Appendix

A.1 Comprehensive Taxonomy of Methods

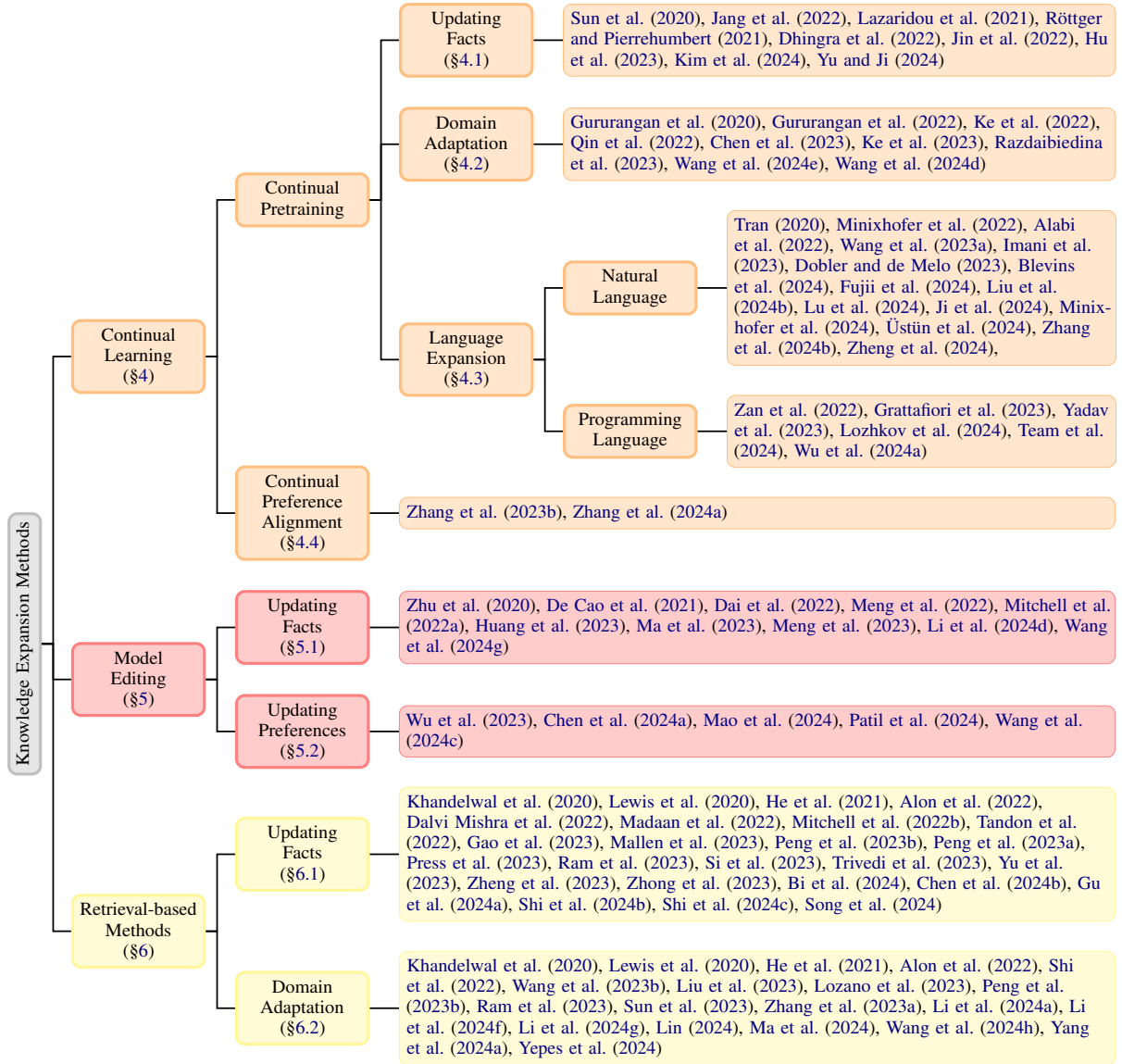


Figure 2: Taxonomy of methods for expanding LLM knowledge.