

Knowledge Graphs for Multi-modal Learning: Survey and Perspective

Anonymous ACL submission

Abstract

Integrated with multi-modal learning, knowledge graphs (KGs) as structured knowledge repositories, enhance AI’s capability to process and understand complex, real-world data. This paper provides a comprehensive survey of cutting-edge research on KG-aware multi-modal learning, providing task definitions, evaluation benchmarks, and detailed insights into key breakthroughs. Furthermore, we also discuss current challenges, highlighting emerging trends and future research directions.

1 Introduction

The brain’s accumulation of memories over time is crucial for societal adaptation and survival, enabling meaningful actions and interactions. These memories can be categorized into two types. *Conditioned Reflexes*: Intuitive memory developed through repeated practice, enhancing analogical reasoning. Combined with sensory inputs like visual, auditory, and tactile data, it enables efficient performance of basic tasks, similar to many traditional multi-modal tasks. *Torso-to-tail Knowledge*: Less frequently encountered, requiring active memorization or deep contemplation. This knowledge highlights the importance of Knowledge Graphs (KGs) in capturing and structuring long-tail knowledge. Our study focuses on leveraging symbolic, structured knowledge within KGs to enhance Multi-Modal Learning (MML). Given their vital role in organizing long-tail knowledge and proven effectiveness in AI systems, integrating KGs with MML offers a promising approach to addressing the challenges inherent in multi-modal data integration.

As illustrated in Fig 1, individuals continuously process multi-modal information from the environment while absorbing and utilizing external knowledge. Despite extensive research within NLP communities, a systematic review of these approaches remains absent. Our survey fill this gap by synthesizing existing KG4MML methods, detailing key

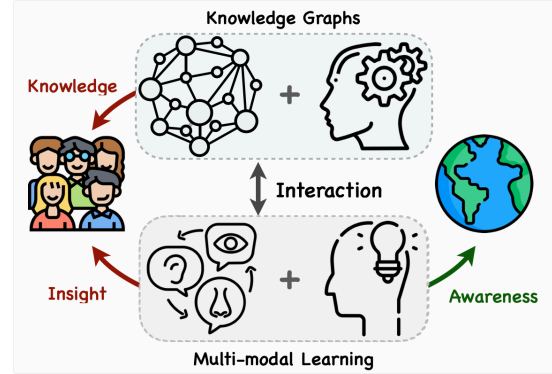


Figure 1: Knowledge Graphs for Multi-modal Learning. resources and breakthroughs. We balance these details to address content overlaps and focus on core challenges, ultimately emphasizing KG’s pivotal role in shaping the past, present, and future developments of multi-modal learning.

2 Preliminaries

Knowledge Graphs. KGs represent entities and their relationships in a graph structure, where nodes symbolize real-world entities or atomic values (attributes), and edges denote relations. Knowledge in KG is often captured in triples, with an ontology-based schema defining basic entity classes and their relations in a taxonomic structure. A KG is defined as $\mathcal{G} = \mathcal{E}, \mathcal{R}, \mathcal{T}$, with entities $\mathcal{E}$, relations $\mathcal{R}$, and statements $\mathcal{T}$. Statements include relational fact triples (h, r, t) , where h is the head entity, r is the relation, and t is the tail entity, or attribute triples (e, a, v) , where e is an entity, a is an attribute, and v is the attribute’s value. Attribute values can be literals such as strings or dates and may include metadata like labels and textual definitions.

Multi-modal Learning. We focus on visiolinguistic (VL) tasks involving text and image data, aiming to provide in-depth analysis and research continuity. Other modalities like video or biochemistry are less emphasized as VL methods can often

¹For a focused discussion, most method references & detailed descriptions are organized in the Appendix for readers interested in tracing the original sources.

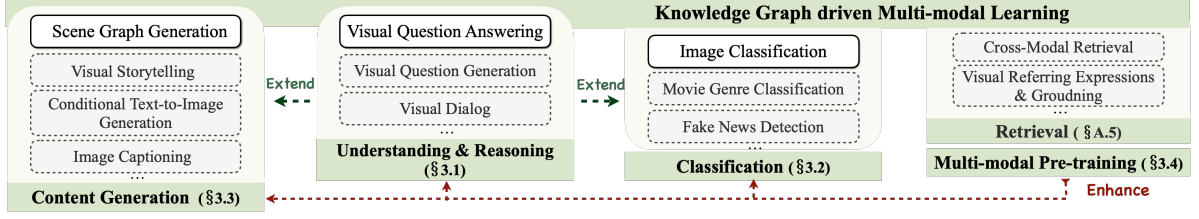


Figure 2: Comprehensive Overview of KG-driven Multi-modal Learning. Due to space constraints and task overlaps, we focus on the most representative sub-tasks in each category (Generation, Understanding & Reasoning, Classification) to maximize relevant content coverage. Additional content is analyzed in the Appendix¹.

be adapted to them. Thus, the input domain is $\mathcal{X} = \mathcal{X}^l \times \mathcal{X}^v$, with inputs $\hat{x} = (x^l, x^v)$, where x^l and x^v are language and visual data, respectively.

KGs Acquisition for Multi-modal Learning.

1) *Task-agnostic Sub-KG Extraction*: Practical applications often require localized knowledge to address specific tasks. Sub-KG extraction isolates minimal knowledge units or triplets from a large KG like WordNet (Miller, 1995) directly, reducing noise from irrelevant information using retrieval, routing, or semantic parsing algorithms.

2) *Task-Oriented KG Construction*: Sometimes, constructing task-specific KGs from scratch is necessary to meet unique requirements, either from datasets or by combining multiple KGs: (i) *Static Domain KGs Construction*: Creating stable, domain-specific KGs with predefined entities and relations to encapsulate crucial background knowledge, especially when general KGs are inadequate for a specific task. E.g., Zero-shot Image Classification tasks necessitate building KGs that capture visual attributes or taxonomy associations (Geng et al., 2021a). These KGs are designed to cover all relevant classification knowledge with textual data like class labels utilized to delineate class relationships, aiding in the formation of KG edges. (ii) *Dynamic Temporary KGs Construction*: Building dynamic, temporary KGs during task execution and leveraging KG reasoning algorithms for task support. E.g., establishing co-occurrence relations between classes (e.g., food ingredients) by analyzing their frequency in training datasets.

3 KG-driven Multi-modal Learning

3.1 Multi-modal Understanding & Reasoning

Visual Question Answering (VQA) (Antol et al., 2015) is a fundamental task in multi-modal learning, frequently used to evaluate multi-modal foundation models (Alayrac et al., 2022). Knowledge-based VQA (Wu et al., 2016) incorporates an external Knowledge Base (KB) for complex question analysis and deeper reasoning assistance (Wang

et al., 2018a). When using KGs as the KB, given an image-question pair (I, Q) , the goal of KG-based VQA is to derive an answer y via:

$$p(y|Q, I, \mathcal{G}, \Theta) = \underbrace{p(\mathcal{G}_{ret}|Q, I, \mathcal{G}; \Phi)}_{\text{Retriever (if have)}} \cdot \underbrace{p(A|Q, I, \mathcal{G}_{ret}; \Theta)}_{\text{Reader}},$$

where $\mathcal{G}$ is the background KG, $\mathcal{G}_{ret}$ is the retrieved sub-KG, while Φ refers to the model parameters used in the knowledge retrieval step.

Current KG-aware VQA research typically has four key stages for incorporating knowledge:

- **Knowledge Retrieval** focuses on extracting pertinent knowledge from various sources, including KGs and document collections like Wikipedia (Denoyer and Gallinari, 2006). The methods have evolved from early matching-based and dense embedding similarity approaches to learnable retrieval and Pre-trained Language Model (PLM) generation techniques, broadening the scope and efficiency of knowledge integration. Fig. 4 (a) illustrates the basic form of KG retrieval in VQA.

- (i) **Matching-based Retrieval** generally employs entity-level methods to identify key concepts within images and questions, linking these to relevant data to external large-scale KGs.

The extraction process from images involve identifying spatial positions (Zhu et al., 2020), visual object sizes and names (Narasimhan et al., 2018), high-level attributes like scene names, object parts, and human activities (Khademi et al., 2023). Additionally, image captions and OCR text strings are generated for supplement (Lin and Byrne, 2022). Questions and captions can be parsed for syntax analysis using various NLP tools like Dependency Parsers and Named Entity Recognizers (Wu and Mooney, 2022). Additional techniques include regular expressions (Wang et al., 2017), SpanSelector or query template selector (Wang et al., 2018a). During this stage, unimportant visual objects not present in the question or caption might be filtered out (Gardner et al., 2018). To associates these concepts with relevant entries in KGs, methods like greedy longest-string matching (Su et al., 2018),

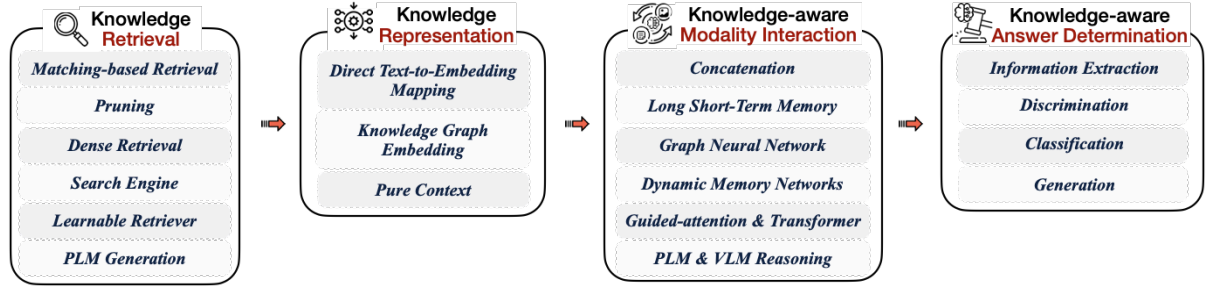


Figure 3: Current KG-aware Understanding & Reasoning research pipeline, which typically involves four key stages for incorporating knowledge. Note that studies often employ one or more of these stages in model design.

template matching (Wang et al., 2018a), and Multi-modal Entity Linking (Jain et al., 2021) are used. Then, fact triples can be collected by involving the first-order sub-KG from these identified concept nodes, which can sometimes extend to three-hop connections in character KGs (Shah et al., 2019). Alternatively, brief knowledge paths can be identified among the entities from I and Q via sub-KG construction (Wang et al., 2019). Generating RDF queries (e.g., SPARQL) often involves filling predefined templates with parsed question data, making it suitable for datasets with consistent question patterns (Wang et al., 2017). Term-based retrievers (e.g., TF-IDF and BM25) are another effective option, with their scoring reflecting the direct correlation between Q and fact triples (Luo et al., 2021).

(ii) Pruning. The pruning stage refines the coarse-grained sub-KG obtained from initial retrieval. This involves re-ranking candidate facts and may include assigning weights to nodes based on corresponding visual object sizes (Wang et al., 2019), ensuring each knowledge triple contains key elements from Q or auto-generated captions (Su et al., 2018; Wu and Mooney, 2022), or aligns with the relation type implied in Q (Yin et al., 2023a). Additionally, a learnable score function can be used to assess the compatibility between a fact and the Q - I representation (Ravi et al., 2023).

(iii) Dense Retrieval methods typically retrieve the most relevant top- k facts for a given Q - I pair. This technique utilizes embedding similarities to match questions and visual concepts with pre-flattened concise fact sentences (Narasimhan et al., 2018; Ziaeeafard and Lécué, 2020), simplifying the retrieval process without complex rules. Sometimes, dense retrieval can also serve as a mechanism for KG pruning, selectively excluding information that is likely irrelevant. Retrieval efficiency is frequently enhanced by employing open-source indexing engines like FAISS (Johnson et al., 2021), which facilitates the organization and indexing of

large-scale dense embeddings.

(iv) Search Engine serves as a valuable tool in VQA for accessing open knowledge and can complement other knowledge sources in ensemble methods. For instance, Marino et al. (2019) gather Wikipedia articles for each Q - I pair, selecting sentences that closely match the query by keyword frequency and then predicting answer presence and positioning within these articles. Given that this is not a primary focus for KG-based VQA, a more detailed discussion is provided in the appendix.

(v) Learnable Retrievers adapt to specific contexts, providing biased recall that highlights interactions between Q - I and knowledge components. These models require rigorous training, using either labeled data or autonomously generated labels. For example, UnifER (Guo et al., 2022b) compares reader loss from Q - I inputs to loss with additional retrieved knowledge, defining the difference as the *loss gap*. A negative *loss gap* indicates counterproductive knowledge, and the model uses this metric to iteratively refine both the retriever and the reader. **Cold Start** issues with asymmetric or randomly initialized retrievers can lead to irrelevant items and inadequate feedback. Hu et al. (2023) address this by creating an initial dataset with pseudo ground-truth knowledge from a large-scale image-caption dataset for pre-training. Chen et al. (2022c) use a BM25-based KG retriever initially to distill preference knowledge to the differentiable retriever.

(vi) PLM Generation. Recent research shows that PLMs can function as KBs when appropriately prompted (Petroni et al., 2019). Specifically, PROOFREAD (Zhou et al., 2023) uses ChatGPT to generate relevant Q and knowledge entries for each Q - I pair, storing them in a *Demo Bank* for demonstration reuse while ensuring case diversity; Additionally, several studies directly harness the knowledge embedded in PLMs for reasoning, skipping a separate knowledge retrieval step (Yang et al., 2022). E.g., CodeVQA (Subramanian et al., 2023)

involves a knowledge query module that utilizes a PLM to answer questions based on world knowledge embedded within its parameters.

- **Knowledge Representation** involves selecting the appropriate format for symbolic KGs to integrate with multi-modal models.

(i) **Direct Text-to-Embedding Mapping.** Some research treats e and r in KGs as words, embedding them into continuous vectors using methods like Glove. This enables the compression of knowledge components (e.g., triples) into fixed-size vectors using RNNs (Wang et al., 2019), (V)PLMs (Chevalier et al., 2023), or mean pooling (Chen et al., 2021d). Strategies like stop-word removal can further refine these embeddings by reducing noise from irrelevant words (Chen et al., 2022c). Additionally, some works convert fact collections into natural language sentences (Hu et al., 2023), allowing PLMs to encode them directly into fixed-length vectors.

(ii) **Knowledge Graph Embedding (KGE)** methods embeds facts and reveals semantic relationships among triples in an abstract space, facilitating initial fact embeddings (Zheng et al., 2021b; Han et al., 2023). During self-supervised training, signals from neighboring entities are embedded into each central entity’s unique representation. This approach facilitates the identification and integration of key entities into downstream tasks, effectively simulating the retrieval of specific sub-KGs without explicit retrieval steps.

(iii) **Pure Context.** In many cases, KG triples are maintained in their original textual format for direct participation in multi-modal reasoning. This includes using sub-KGs for RDF query-based answer retrieval (Wang et al., 2017) and serializing triples for joint reasoning with (V)PLMs (Gao et al., 2022; Dong et al., 2024). To handle lengthy input sequences predominantly composed of facts, which might shift the model’s focus away from other crucial cues, Ravi et al. (2023) summarizes information from each inference sentence into a single token representation using SBERT (Reimers and Gurevych, 2019); Wang et al. (2023g) select only factual summaries with higher contribution scores than the original Q as caption supplements.

- **Knowledge-aware Modality Interaction** is the core of KG-based multi-modal reasoning. To some extent, it mirrors human knowledge application in understanding the world (Fig. 4).

(i) **Concatenation** of multi-modal vectors provides a straightforward and effective approach for

modality fusion (Ramnath and and, 2020), combining different modality features into a single representation. This unified feature is typically refined with a MLP to enhance modality interaction and integration.

(ii) **LSTM** networks typically employ an LSTM encoder to process semantic inputs from I and Q , and an LSTM decoder for generating answers, initializing the hidden state with embeddings from attributes, captions, or external knowledge. As shown in Fig. 4 (b), Q is tokenized and fed into the system sequentially (Wu et al., 2016).

(iii) **GNNs** enhance concept connections in VQA by integrating representations from I , Q , and entities into cohesive networks, with each node (entity e) represented by a multi-modal concatenated embedding (Narasimhan et al., 2018). GNNs iteratively process these e embeddings, and the final learned representations are fed into an MLP, which assigns a binary label to each e indicating its relevance as an answer, as shown in Fig. 4 (c). Mucko (Zhu et al., 2020) independently processing distinct modality KGs, separately analyzing the visual scene KG, the semantic KG from image captions, and the common sense KG, which supports precise answer determination through Q -guided attention and cross-KG GNNs.

(iv) **Dynamic Memory Networks (DMNs)** (Kumar et al., 2016) filter critical information from localized small-scale knowledge triple embeddings (Fig. 4 (d)) to enable interactions across multiple data channels (Yin et al., 2023a). This process, akin to building a cache and performing secondary retrieval, is typically achieved through an attention-based mechanism. In particular, VKMN (Su et al., 2018) deconstructs each knowledge triple into three Key-Value pairs, e.g., (h, r) as the key and t as the value, improving reasoning performance by reducing interference from using only head and tail entities as keys for memory grounding.

(v) **Guided-Attention & Transformer.** The Transformer architecture, featuring multi-head attention, layered stacking, and residual connections, is widely used in multi-modal fusion (Vaswani et al., 2017). It enables knowledge embeddings to integrate seamlessly with other modalities. The guided-attention mechanism (Heo et al., 2022) enhances this by using distinct feature sets to direct attention unidirectionally, unlike self-attention’s symmetric interactions. This is intuitive, reflecting the unequal roles of different modalities and knowledge in fusion. Examples include knowledge-

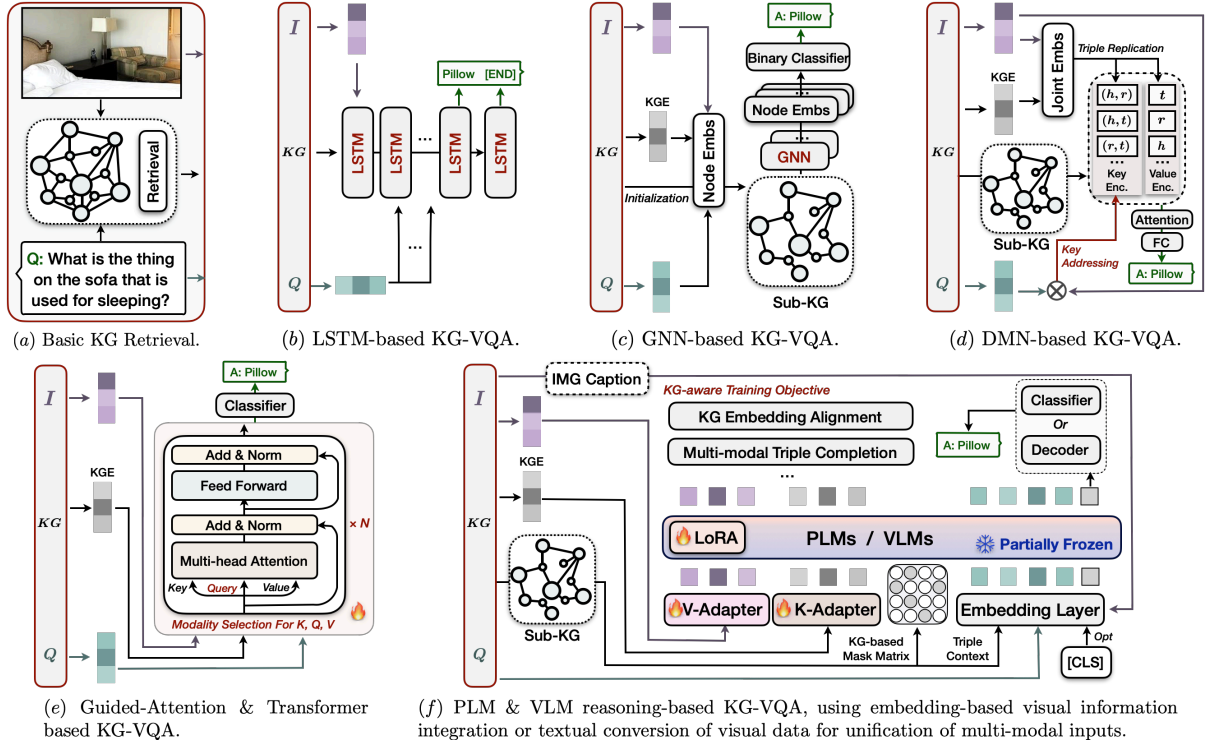


Figure 4: Current knowledge-aware modality interaction paradigms utilizing KG as the knowledge repositories.

guided visual/textual embeddings and Q -guided visual/knowledge embeddings (Fig. 4 (e)).

(vi) **PLM & VLM Reasoning** allow researchers to focus on data organization and training objective design without significantly altering the backbone structure, effectively utilizing the inherent parameterized knowledge in original models (Fig. 4 (f)):

(a) **Embedding-Based Visual Information Integration** involves converting visual data into embeddings compatible with existing (V)PLMs, like compressing patch or object features into fixed-length embeddings (Jaegle et al., 2021) or using adapters/projection heads for cross-modal alignment (Yin et al., 2023b). Specifically, RVL (Shevchenko et al., 2021) and KVQAmeta (García-Olano et al., 2022) inject the knowledge into VLMs via aligning the KG embeddings of e with corresponding textual phrase representations derived from the PLM’s embedding layers. MuKEA (Ding et al., 2022) uses VLM’s visual output embeddings as the head, question as the relation, and the answer as the tail to design a multi-modal triple completion training objective, leveraging implicit knowledge for reasoning.

(b) **Textual Conversion of Visual Data** involves converting all visual information into a textual format, like captions, allowing PLM reasoning on a uniform textual data collection that includes background knowledge, Q , and I (Hu et al., 2022).

• Knowledge-aware Answer Determination.

(i) **Information Extraction** methods typically focus on locating specific entities within the retrieved knowledge or existing I - Q pair, emphasizing content grounding. Among them, query-based methods (Wang et al., 2017) obtain the final answer through sub-KG inference, yielding benefits such as improved matching accuracy, relevance, and interpretability. Their effectiveness, however, depends on the model’s capability to parse queries and the KG’s completeness. Challenges arise with non-unique or difficult-to-find answers. To further rank the potential answers, heuristic rules can be employed, like matching score calculation (Wang et al., 2018a; Narasimhan and Schwing, 2018) and answer frequency assessment (Wang et al., 2018a).

(ii) **Discrimination**. Particularly suited in multi-choice VQA tasks, these methods use a discriminator for final selection among candidate answers. They are effective in narrowing down potential answers within a certain range (or in a sub-KG), often using GNN-alike backbones (Hussain et al., 2022) (Fig. 4 (c)). Furthermore, discriminators can be either MLP-based (Liu et al., 2022) or rule-based (Narasimhan and Schwing, 2018), with a notable limitation being time consumption, especially with extensive answer vocabularies.

(iii) **Classification**. In many VQA datasets, the range of possible answers is pre-determined, typically constrained by their frequency range or a minimum occurrence threshold. Consequently, many

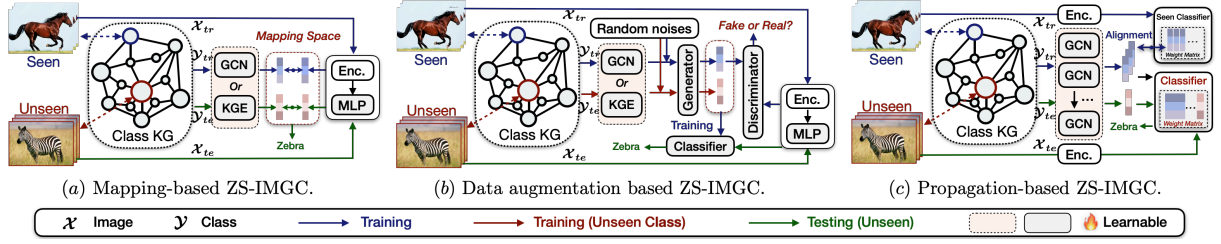


Figure 5: A comparison of previous paradigms for KG-based Zero-shot Image Classification (ZS-IMGC) methods.

studies reformulate VQA as a classification problem (Gardères et al., 2020; Song et al., 2023b), with the output dimension corresponding to the number of answer candidates. For (V)PLM-based methods, a classification (or projection) head is typically appended to the output [CLS] embedding, as shown in Fig. 4 (f) (He and Wang, 2023).

(iv) **Generation.** With the expansion of parameters and pre-training data in LMs (Schaeffer et al., 2023), generative models’ accuracy has significantly improved, effectively mitigating the disadvantages posed by the exact match criteria that constrained early LSTM-based methods. As a result, generative (V)PLM-based methods are now increasingly supplanting traditional classification-based approaches (Ghosal et al., 2023) (Tab. 1).

3.2 Multi-modal Classification

This section explores multi-modal classification tasks, particularly focusing on Zero-Shot Image Classification (ZS-IMGC), which classifies images from novel, unseen classes without prior specific training, as denoted by $\mathcal{Y}_{tr} \cap \mathcal{Y}_{te} = \emptyset$. In contrast, traditional image classification (x as an image and y as its class) assumes a closed world where $\mathcal{Y}_{tr} = \mathcal{Y}_{te} = \mathcal{Y}$, demanding extensive labeled images for both training and testing within known classes.

Early ZS-IMGC works use textual class characteristics (Zhu et al., 2018) and define shared descriptive attributes (Xian et al., 2018) to model inter-class relationships (see Fig. 10, left). Recently, KGs have become integral to ZS-IMGC by unifying various forms of above knowledge into a single graph, $\mathcal{G} = \{\mathcal{E}, \mathcal{R}, \mathcal{T}\}$, where $\mathcal{Y} \subset \mathcal{E}$. This integration not only encapsulates diverse and explicit class semantics but also enhances compatibility and interpretability (see Fig. 10, right). Specifically, studies like (Kampffmeyer et al., 2019) integrate hierarchical relationships from WordNet, while Roy et al. (2022) explore class knowledge from commonsense KGs (e.g., ConceptNet). Pahuja et al. (2024) enhance species classification by structuring it as a link prediction task within a Multi-Modal Knowledge Graph (MMKG), leverag-

ing visual cues and GPS coordinates to efficiently identify unseen classes, such as deducing that an image of a feline captured in Africa is likely a lion. Furthermore, ontologies can be utilized to define complex class relationships (Chen et al., 2020a) (e.g., *disjointness*), providing a channel to explicitly introduce rules for refinement.

Existing KG-driven ZS-IMGC approaches, which guide feature transfer from seen to unseen classes, can be categorized into three types:

Mapping-based methods develop mapping functions that align image inputs with KG-based class semantics into a shared vector space, simplifying the identification of the nearest class to a test image based on similarity metrics, as shown in Fig. 5 (a). For instance, HierSE (Li et al., 2015) employs a linear projection to map image features into a class embedding space derived from word embeddings of the class and its superclasses, using cosine similarity for comparison. Similarly, Chen et al. (2020a) use an OWL-based ontology to encode animal classes, and Akata et al. (2013) represent classes using multi-hot vectors of ancestors, reflecting class hierarchies in KGs.

Data Augmentation methods alleviate the sample shortage in ZS-IMGC by synthesizing images or features for unseen classes, transforming it into a supervised learning problem and reducing bias. These methods primarily use generative models like GANs and VAEs, leveraging feature-related KGs to simulate characteristics of unseen images. For example, OntoZSL (Geng et al., 2021a) synthesizes image features by blending class embeddings from an attribute-and-species KG with random noise vectors. This process, supervised by real features from annotated images, employs an adversarial discriminator to differentiate between real and generated features as shown in Fig. 5 (b).

Propagation-based methods leverage KG-structured inter-class relationships for knowledge transfer, using GNNs to propagate features from seen to unseen class nodes (Wang et al., 2018b; Geng et al., 2021b). Concretely, GNN models train to produce class-specific parameter vectors as

classifiers. These classifiers for unseen classes are estimated by aggregating those from neighboring seen classes within the graph (see Fig. 5 (c)).

3.3 Multi-modal Generation

Given visual images (x^v) or textual descriptions (x^l), the objective of KG-aware generation tasks is to generate, in a cross-modal manner, either a textual target y^l (e.g., caption and scene graph), a visual target y^v (e.g., image), leveraging the background KG $\mathcal{G}$ for foundational support².

Scene Graphs (SGs) are vital for scene understanding, cataloging objects and their interrelationships within a scene. These instances, ranging from people to places and objects, are described through attributes like shape, color, and pose (Chang et al., 2023). The relationships between these instances, often action-based or spatial, are expressed as (*subject, predicate, object*) triplets, paralleling the (*h, r, t*) and (*e, a, v*) triplets in KGs. Scene Graph Generation (SGG) serves as an intermediary task, unlike other multi-modal tasks with specific end goals, providing enhanced understanding and reasoning to support downstream tasks (Fig. 6).

Several works adopt the **KG representation learning** techniques into SGG scenario. For example, Yu et al. (2022) improve zero-shot performance in SGG by constructing a KG from training set SG triples, distinguishing existing (non-zero-shot) and missing (zero-shot) edges. They train a KG Embedding model to complete the graph and fills these missing edges, thereby integrating zero-shot triples similarly to their non-zero-shot counterparts. Others employ KGs for **triple prediction** to generate rich and expressive SGs. Specifically, Gu et al. (2019) utilize a knowledge-based module to identify relevant ConceptNet entities and retrieve commonsense facts, each assigned a weight indicating its real-world prevalence to filter candidate triples. Khan et al. (2022b) enrich SGs using CSKG (Ilievski et al., 2021), a substantial commonsense KG repository as shown in Fig. 6. This upgrades SGG with additional information on objects’ spatial proximity and potential interactions derived from external knowledge, improving higher-level reasoning and mitigating some missed or incorrect predictions made during SGG.

²We defer detailed discussion to Appendix A.5, noting significant overlaps with reasoning tasks like Image Captioning and Visual Storytelling, similar to VQA as discussed in § 3.1, and because KG-aware cross-modal image generation, where only limited work exists, is still in its early stages.

³Off-scene entities refer to those not part of the VG (Krishna et al., 2017) classes, as opposed to on-scene entities.

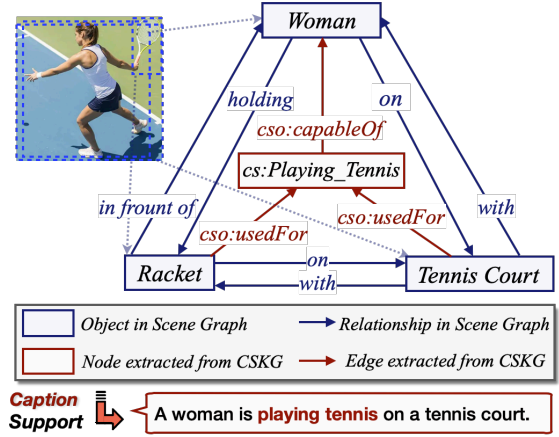


Figure 6: Scene Graph (SG) of an image equipped with CSKG (Ilievski et al., 2021). The SG (in blue) outlines objects and their relationships in the scene. Additional background knowledge from CSKG triples (in red) such as (*Woman, capableOf, Playing_Tennis*), enriches the SG with off-scene³ knowledge (Khan et al., 2022b). This facilitates higher-level inferences, enabling more accurate caption deductions, e.g., “A woman is playing tennis on a tennis court” (Hou et al., 2020).

3.4 KG-aware Multi-modal Pre-training

Currently, multi-modal LLMs (MLLMs) are widely recognized for their task-agnostic capabilities, but their lack of long-tail knowledge often leads to cross-modal hallucinations, causing inconsistencies and errors, as evidenced in Fig. 7. To mitigate this issue, KG-aware Multi-modal Pre-training offers a viable strategy for knowledge infusion.

Specifically, Med-VLP (Chen et al., 2022b) employs structured medical knowledge entities from UMLS KG (Bodenreider, 2004) as mediators to align image and text features (Li et al., 2021), using a whole-entity mask strategy (Sun et al., 2019a) instead of sub-word masking. It focuses the model’s attention on crucial medical information across modalities, enabling medical VLMs to gain domain-specific knowledge for knowledge-aware representations in downstream tasks. Ye et al. (2023) introduce a DANCE dataset where commonsense KG triples are turned into natural language riddles paired with images, aimed at embedding knowledge relations and linking KG entries (*h, r, t*) with images that depict related entities, where entities in the images are referred to as “this item”. KGTransformer (Zhang et al., 2023b) is a BERT-like KG pre-training model with objectives like triple-based masked relation/entity prediction and entity pair prediction (i.e., assessing whether two entities from a one-hop subgraph were previously connected by the same relation in $\mathcal{G}$), supporting downstream tasks like QA and ZS-IMGC.



Figure 7: Examples of limited cross-modal knowledge alignment ability in current multi-modal LLMs (Zha et al., 2023), as demonstrated by (a) BLIP-2 (Li et al., 2023b) and (b) MiniGPT-4 (Zhu et al., 2023), which lead to multi-modal hallucinations.

4 Challenges and Future Directions

Multi-modal Knowledge Graphs. Focusing solely on the benefits of traditional KGs for multi-modal tasks can be limiting due to the narrow scope of knowledge in single-modality KGs. Scenarios like detailing badge designs or architectural photos, which are challenging to describe with text alone, highlight the necessity for MMKGs which integrate knowledge symbols across modalities like text, images, and sound (typically as entities). However, most MMKG research remains focused on tasks within the In-MMKG framework like Multi-modal Knowledge Graph Completion (Zhao et al., 2022), with limited exploration into MMKG-driven tasks. We identify three main challenges hindering broader application: *non-uniform organization and ontology of MMKGs*, *substantial storage and processing overheads*, and issues of *data timeliness and completeness in MMKGs*, which are discussed in the Appendix. Currently, the few studies on MMKG-driven tasks primarily emphasize retrieval-related activities, exploiting MMKGs’ natural database-like functionalities, yet they do not fully exploit MMKGs’ structured multi-modal capabilities. For instance, Zha et al. (2023) enhance knowledge-based VQA by employing multi-modal concept descriptions and using MMKGs as “key:value” based retrieval KBs to support reasoning in MLLMs. Looking forward, to fully unlock the potential of large-scale MMKGs for multi-modal tasks, we must resolve key issues includ-

ing the effective construction and utilization of MMKGs that are well-suited to multi-modal tasks.

Large Language Models. The challenge of extracting sufficient visual knowledge, as identified by Chen et al. (2023a), alongside Zhou et al.’s (2023) finding that 43% of BLIP2 (Li et al., 2023b) errors on the A-OKVQA dataset (Schwenk et al., 2022) could be addressed with proper knowledge integration. Fine-tuning MLLMs with MMKGs can be realized via two main strategies: active KG routing for creating specific instructions (Wan et al., 2023), and the use of self-instructing techniques to autonomously generate multi-grained, multi-modal instructional data (Du et al., 2023). Besides, the structured multi-modal relational data inherent in MMKGs provides an essential foundation for investigating the visual extrapolation abilities of Large Vision Models (LVMs) (Bai et al., 2023) as well as MLLMs (Sun et al., 2023d).

The increasing risk of generating seemingly authentic but factually inaccurate web content in MLLMs, known as *hallucination* (Agrawal et al., 2023), is another concern. Future efforts could focus on combining MMKGs with hallucination detection or correction methods to enhance multi-modal task precision and leveraging KGs for rewriting input statements in a knowledge-aware manner to reduce factual hallucinations in MLLM reasoning (Wang et al., 2023a). This process could be aligned with CoT approaches like Think-on-Graph (ToG) (Sun et al., 2023a), improving MLLM’ abilities to interpret and interact with (MM)KGs, thus advancing towards human-like multi-modal proficiency and greater machine intelligence. Moreover, MMKGs can enhance multi-modal Retrieval Augmented Generation (RAG) by using various modalities as anchors (Song et al., 2023a), yielding more relevant and insightful results than traditional vector-based searches (Wu and Xie, 2023).

5 Conclusion and Vision

This paper provides an overview of KG-driven multi-modal learning, tracing the field’s evolution from past achievements through current trends to future developments. Our goal is to construct a systematic blueprint of the domain, providing a valuable reference for researchers currently involved in or planning to explore this area. Looking ahead, we envision a stronger synergy between MMKGs and MLLMs, aiming to create a robust interactive system powered by this dual-drive approach.

6 Limitations

In this study, we provide a survey of multi-modal learning with knowledge graphs. We discuss related surveys in Appendix A.1 and will continue adding more related approaches with more detailed analysis. Despite our best efforts, there may be still some limitations that remain in this paper.

References & Methods. Due to the page limit, we may have omitted some important references and cannot afford all the technical details. Our Literature Collection Methodology is shared in Appendix A.1. We primarily review cutting-edge methods from the past three years (mostly in 2023), sourced from major conferences and journals like ACL, EMNLP, NAACL, CVPR, NeurIPS, ICLR, and arXiv, etc., and we will continue to update our review with the latest research.

Benchmarks. Most of the benchmarks mentioned (e.g., Tab. 1 and Tab. 2) are gathered and categorized from the experimental part of mainstream works. In order to help readers quickly understand the tasks’ goals and formats from a unified perspective, the definition and boundary of each task may not be accurate enough. Additionally, considering the similarity and coupling between different tasks and the uneven number of related works, we may have overlooked some multi-modal tasks such as multi-modal summarization (Jangra et al., 2023). We also have not discussed specialized domains such as Medicine (Du et al., 2016) and Science, but we aim to address these gaps in the future.

Empirical Conclusions. We provide detailed comparisons and discussions on KG-driven multi-modal learning in § 3, listing some promising future directions in § 4. All conclusions are based on empirical analysis of existing works, which may not capture a broader perspective. As the field evolves, we will update our findings to reflect the latest developments.

References

Somak Aditya, Yezhou Yang, Chitta Baral, Cornelia Fermüller, and Yiannis Aloimonos. 2015. From images to sentences through scene description graphs using commonsense reasoning and knowledge. *CoRR*, abs/1511.03292.

Omar Adjali, Paul Grimal, Olivier Ferret, Sahar Ghanay, and Hervé Le Borgne. 2023. Explicit knowledge

integration for knowledge-aware visual question answering about named entities. In *ICMR*, pages 29–38. ACM.

Garima Agrawal, Tharindu Kumarage, Zeyad Alghami, and Huan Liu. 2023. Can knowledge graphs reduce hallucinations in llms? : A survey. *CoRR*, abs/2311.07914.

Zeynep Akata, Florent Perronnin, Zaïd Harchaoui, and Cordelia Schmid. 2013. Label-embedding for attribute-based classification. In *CVPR*, pages 819–826. IEEE Computer Society.

Jean-Baptiste Alayrac, Jeff Donahue, Pauline Luc, Antoine Miech, Iain Barr, Yana Hasson, Karel Lenc, Arthur Mensch, Katherine Millican, Malcolm Reynolds, Roman Ring, Eliza Rutherford, Serkan Cabi, Tengda Han, Zhitao Gong, Sina Samangooei, Marianne Monteiro, Jacob L. Menick, Sebastian Borgeaud, Andy Brock, Aida Nematzadeh, Sahand Sharifzadeh, Mikolaj Binkowski, Ricardo Barreira, Oriol Vinyals, Andrew Zisserman, and Karén Simonyan. 2022. Flamingo: a visual language model for few-shot learning. In *NeurIPS*.

Peter Anderson, Basura Fernando, Mark Johnson, and Stephen Gould. 2016. SPICE: semantic propositional image caption evaluation. In *ECCV (5)*, volume 9909 of *Lecture Notes in Computer Science*, pages 382–398. Springer.

Stanislaw Antol, Aishwarya Agrawal, Jiasen Lu, Margaret Mitchell, Dhruv Batra, C. Lawrence Zitnick, and Devi Parikh. 2015. VQA: visual question answering. In *ICCV*, pages 2425–2433. IEEE Computer Society.

Sören Auer, Christian Bizer, Georgi Kobilarov, Jens Lehmann, Richard Cyganiak, and Zachary G. Ives. 2007. Dbpedia: A nucleus for a web of open data. In *ISWC/ASWC*, volume 4825 of *Lecture Notes in Computer Science*, pages 722–735. Springer.

Yutong Bai, Xinyang Geng, Karttikeya Mangalam, Amir Bar, Alan L. Yuille, Trevor Darrell, Jitendra Malik, and Alexei A. Efros. 2023. Sequential modeling enables scalable learning for large vision models. *CoRR*, abs/2312.00785.

Max Bain, Arsha Nagrai, Andrew Brown, and Andrew Zisserman. 2020. Condensed movies: Story based retrieval with contextual embeddings. In *ACCV (5)*, volume 12626 of *Lecture Notes in Computer Science*, pages 460–479. Springer.

Collin F. Baker, Charles J. Fillmore, and John B. Lowe. 1998. The berkeley framenet project. In *COLING-ACL*, pages 86–90. Morgan Kaufmann Publishers / ACL.

Scott Barnett, Stefanus Kurniawan, Srikanth Thudumu, Zach Brannelly, and Mohamed Abdelrazek. 2024. Seven failure points when engineering a retrieval augmented generation system. *CoRR*, abs/2401.05856.

741	Hédi Ben-Younes, Rémi Cadène, Matthieu Cord, and	Xiaojun Chang, Pengzhen Ren, Pengfei Xu, Zhihui	797
742	Nicolas Thome. 2017. MUTAN: multimodal tucker	Li, Xiaojiang Chen, and Alex Hauptmann. 2023. A	798
743	fusion for visual question answering. In <i>ICCV</i> , pages	comprehensive survey of scene graphs: Generation	799
744	2631–2639. IEEE Computer Society.	and application. <i>IEEE Trans. Pattern Anal. Mach.</i>	800
		<i>Intell.</i> , 45(1):1–26.	801
745	Sumithra Bhakthavatsalam, Kyle Richardson, Niket Tan-	Danqi Chen and Christopher D. Manning. 2014. A	802
746	don, and Peter Clark. 2020. Do dogs have whiskers?	fast and accurate dependency parser using neural	803
747	A new knowledge base of haspart relations. <i>CoRR</i> ,	networks. In <i>EMNLP</i> , pages 740–750. ACL.	804
748	abs/2006.07510.		
749	Steven Bird, Ewan Klein, and Edward Loper. 2009. <i>Nat-</i>	Gongwei Chen, Leyang Shen, Rui Shao, Xiang Deng,	805
750	<i>ural Language Processing with Python</i> . O’Reilly.	and Liqiang Nie. 2023a. LION : Empowering mul-	806
		timodal large language model with dual-level visual	807
751	Olivier Bodenreider. 2004. The unified medical lan-	knowledge. <i>CoRR</i> , abs/2311.11860.	808
752	guage system (UMLS): integrating biomedical termi-		
753	nology. <i>Nucleic Acids Res.</i> , 32(Database-Issue):267–	Hong Chen, Yifei Huang, Hiroya Takamura, and Hideki	809
754	270.	Nakayama. 2021a. Commonsense knowledge aware	810
		concept selection for diverse and informative visual	811
755	Kurt D. Bollacker, Colin Evans, Praveen K. Paritosh,	storytelling. In <i>AAAI</i> , pages 999–1008. AAAI Press.	812
756	Tim Sturge, and Jamie Taylor. 2008. Freebase: a		
757	collaboratively created graph database for structuring	Jiali Chen, Zhenjun Guo, Jiayuan Xie, Yi Cai, and Qing	813
758	human knowledge. In <i>SIGMOD Conference</i> , pages	Li. 2023b. Deconfounded visual question generation	814
759	1247–1250. ACM.	with causal inference. In <i>ACM Multimedia</i> , pages	815
		5132–5142. ACM.	816
760	Antoine Bordes, Nicolas Usunier, Alberto García-	Jiaoyan Chen, Yuxia Geng, Zhuo Chen, Ian Horrocks,	817
761	Durán, Jason Weston, and Oksana Yakhnenko.	Jeff Z. Pan, and Huajun Chen. 2021b. Knowledge-	818
762	2013. Translating embeddings for modeling multi-	aware zero-shot learning: Survey and perspective. In	819
763	relational data. In <i>NIPS</i> , pages 2787–2795.	<i>IJCAI</i> , pages 4366–4373. ijcai.org.	820
764	Antoine Bosselut, Hannah Rashkin, Maarten Sap, Chai-	Jiaoyan Chen, Yuxia Geng, Zhuo Chen, Jeff Z. Pan,	821
765	tanya Malaviya, Asli Celikyilmaz, and Yejin Choi.	Yuan He, Wen Zhang, Ian Horrocks, and Huajun	822
766	2019. Comet: Commonsense transformers for auto-	Chen. 2023c. Zero-shot and few-shot learning with	823
767	matic knowledge graph construction. In <i>Proceedings</i>	knowledge graphs: A comprehensive survey. <i>Proc.</i>	824
768	<i>of the 57th Annual Meeting of the Association for</i>	<i>IEEE</i> , 111(6):653–685.	825
769	<i>Computational Linguistics</i> , pages 4762–4779.		
		Jiaoyan Chen, Freddy Lécué, Yuxia Geng, Jeff Z. Pan,	826
770	Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie	and Huajun Chen. 2020a. Ontology-guided semantic	827
771	Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind	composition for zero-shot learning. In <i>KR</i> , pages	828
772	Neelakantan, Pranav Shyam, Girish Sastry, Amanda	850–854.	829
773	Askell, Sandhini Agarwal, Ariel Herbert-Voss,	Jingjing Chen, Liangming Pan, Zhipeng Wei, Xiang	830
774	Gretchen Krueger, Tom Henighan, Rewon Child,	Wang, Chong-Wah Ngo, and Tat-Seng Chua. 2020b.	831
775	Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu,	Zero-shot ingredient recognition by multi-relational	832
776	Clemens Winter, Christopher Hesse, Mark Chen, Eric	graph convolutional network. In <i>AAAI</i> , pages 10542–	833
777	Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess,	10550. AAAI Press.	834
778	Jack Clark, Christopher Berner, Sam McCandlish,		
779	Alec Radford, Ilya Sutskever, and Dario Amodei.	Kan Chen, Jiyang Gao, and Ram Nevatia. 2018. Knowl-	835
780	2020. Language models are few-shot learners. In	edge aided consistency for weakly supervised phrase	836
781	<i>NeurIPS</i> .	grounding. In <i>CVPR</i> , pages 4042–4050. Computer	837
		Vision Foundation / IEEE Computer Society.	838
782	Yuqi Bu, Xin Wu, Liuwu Li, Yi Cai, Qiong Liu, and	Liangyu Chen, Bo Li, Sheng Shen, Jingkan Yang,	839
783	Qingbao Huang. 2023. Segment-level and category-	Chunyu Li, Kurt Keutzer, Trevor Darrell, and Zi-	840
784	oriented network for knowledge-based referring ex-	wei Liu. 2023d. Large language models are visual	841
785	pression comprehension. In <i>ACL (Findings)</i> , pages	reasoning coordinators. <i>CoRR</i> , abs/2310.15166.	842
786	8745–8757. Association for Computational Linguis-		
787	tics.	Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan,	843
		Henrique Pondé de Oliveira Pinto, Jared Kaplan,	844
788	Min Cao, Shiping Li, Juntao Li, Liqiang Nie, and Min	Harrison Edwards, Yuri Burda, Nicholas Joseph,	845
789	Zhang. 2022a. Image-text retrieval: A survey on	Greg Brockman, Alex Ray, Raul Puri, Gretchen	846
790	recent research and development. In <i>IJCAI</i> , pages	Krueger, Michael Petrov, Heidy Khlaaf, Girish Sas-	847
791	5410–5417. ijcai.org.	try, Pamela Mishkin, Brooke Chan, Scott Gray,	848
		Nick Ryder, Mikhail Pavlov, Alethea Power, Lukasz	849
792	Qingxing Cao, Bailin Li, Xiaodan Liang, Keze Wang,	Kaiser, Mohammad Bavarian, Clemens Winter,	850
793	and Liang Lin. 2022b. Knowledge-routed visual		
794	question reasoning: Challenges for deep represen-		
795	tation embedding. <i>IEEE Trans. Neural Networks</i>		
796	<i>Learn. Syst.</i> , 33(7):2758–2767.		

851	Philippe Tillet, Felipe Petroski Such, Dave Cummings, Matthias Plappert, Fotios Chantzis, Elizabeth Barnes, Ariel Herbert-Voss, William Hebgen Guss, Alex Nichol, Alex Paino, Nikolas Tezak, Jie Tang, Igor Babuschkin, Suchir Balaji, Shantanu Jain, William Saunders, Christopher Hesse, Andrew N. Carr, Jan Leike, Joshua Achiam, Vedant Misra, Evan Morikawa, Alec Radford, Matthew Knight, Miles Brundage, Mira Murati, Katie Mayer, Peter Welinder, Bob McGrew, Dario Amodei, Sam McCandlish, Ilya Sutskever, and Wojciech Zaremba. 2021c. Evaluating large language models trained on code. <i>CoRR</i> , abs/2107.03374.	
864	Tianqi Chen and Carlos Guestrin. 2016. Xgboost: A scalable tree boosting system. In <i>KDD</i> , pages 785–794. ACM.	
867	Tianshui Chen, Weihao Yu, Riquan Chen, and Liang Lin. 2019. Knowledge-embedded routing network for scene graph generation. In <i>CVPR</i> , pages 6163–6171. Computer Vision Foundation / IEEE.	
871	Xiang Chen, Duanzheng Song, Honghao Gui, Chengxi Wang, Ningyu Zhang, Fei Huang, Chengfei Lv, Dan Zhang, and Huajun Chen. 2023e. Unveiling the siren’s song: Towards reliable fact-conflicting hallucination detection. <i>CoRR</i> , abs/2310.12086.	
876	Xiaolin Chen, Xuemeng Song, Liqiang Jing, Shuo Li, Linmei Hu, and Liqiang Nie. 2022a. Multi-modal dialog systems with dual knowledge-enhanced generative pretrained language model. <i>CoRR</i> , abs/2207.07934.	
881	Yen-Chun Chen, Linjie Li, Licheng Yu, Ahmed El Kholy, Faisal Ahmed, Zhe Gan, Yu Cheng, and Jingjing Liu. 2020c. UNITER: universal image-text representation learning. In <i>ECCV (30)</i> , volume 12375 of <i>Lecture Notes in Computer Science</i> , pages 104–120. Springer.	
887	Zhanwen Chen, Saed Rezayi, and Sheng Li. 2023f. More knowledge, less bias: Unbiasing scene graph generation with explicit ontological adjustment. In <i>WACV</i> , pages 4012–4021. IEEE.	
891	Zhihong Chen, Guanbin Li, and Xiang Wan. 2022b. Align, reason and learn: Enhancing medical vision-and-language pre-training with knowledge. In <i>ACM Multimedia</i> , pages 5152–5161. ACM.	
895	Zhihong Chen, Ruifei Zhang, Yibing Song, Xiang Wan, and Guanbin Li. 2023g. Advancing visual grounding with scene knowledge: Benchmark and method. In <i>CVPR</i> , pages 15039–15049. IEEE.	
899	Zhuo Chen, Jiaoyan Chen, Yuxia Geng, Jeff Z. Pan, Zonggang Yuan, and Huajun Chen. 2021d. Zero-shot visual question answering using knowledge graph. In <i>ISWC</i> , volume 12922 of <i>Lecture Notes in Computer Science</i> , pages 146–162. Springer.	
904	Zhuo Chen, Yufeng Huang, Jiaoyan Chen, Yuxia Geng, Yin Fang, Jeff Z. Pan, Ningyu Zhang, and Wen Zhang. 2022c. Lako: Knowledge-driven visual question	
	answering via late knowledge-to-text injection. In <i>IJCKG</i> , pages 20–29. ACM.	907 908
	Zhuo Chen, Yufeng Huang, Jiaoyan Chen, Yuxia Geng, Wen Zhang, Yin Fang, Jeff Z. Pan, and Huajun Chen. 2023h. DUET: cross-modal semantic grounding for contrastive zero-shot learning. In <i>AAAI</i> , pages 405–413. AAAI Press.	909 910 911 912 913
	Jun Cheng, Fuxiang Wu, Yanling Tian, Lei Wang, and Dapeng Tao. 2020. Rifegan: Rich feature generation for text-to-image synthesis from prior knowledge. In <i>CVPR</i> , pages 10908–10917. Computer Vision Foundation / IEEE.	914 915 916 917 918
	Alexis Chevalier, Alexander Wettig, Anirudh Ajith, and Danqi Chen. 2023. Adapting language models to compress contexts. In <i>EMNLP</i> , pages 3829–3846. Association for Computational Linguistics.	919 920 921 922
	Jaemin Cho, Jie Lei, Hao Tan, and Mohit Bansal. 2021. Unifying vision-and-language tasks via text generation. In <i>ICML</i> , volume 139 of <i>Proceedings of Machine Learning Research</i> , pages 1931–1942. PMLR.	923 924 925 926
	Tat-Seng Chua, Jinhui Tang, Richang Hong, Haojie Li, Zhiping Luo, and Yantao Zheng. 2009. NUS-WIDE: a real-world web image database from national university of singapore. In <i>CIVR</i> . ACM.	927 928 929 930
	Yuhao Cui, Zhou Yu, Chunqi Wang, Zhongzhou Zhao, Ji Zhang, Meng Wang, and Jun Yu. 2021. ROSITA: enhancing vision-and-language semantic alignments via cross- and intra-modal knowledge integration. In <i>ACM Multimedia</i> , pages 797–806. ACM.	931 932 933 934 935
	Jia Deng, Wei Dong, Richard Socher, Li-Jia Li, Kai Li, and Li Fei-Fei. 2009. Imagenet: A large-scale hierarchical image database. In <i>CVPR</i> , pages 248–255. IEEE Computer Society.	936 937 938 939
	Ludovic Denoyer and Patrick Gallinari. 2006. The wikipedia xml corpus. In <i>ACM SIGIR Forum</i> , volume 40, pages 64–69. ACM New York, NY, USA.	940 941 942
	Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. BERT: pre-training of deep bidirectional transformers for language understanding. In <i>NAACL-HLT (1)</i> , pages 4171–4186. Association for Computational Linguistics.	943 944 945 946 947
	Yang Ding, Jing Yu, Bang Liu, Yue Hu, Mingxin Cui, and Qi Wu. 2022. Mukea: Multimodal knowledge extraction and accumulation for knowledge-based visual question answering. In <i>CVPR</i> , pages 5079–5088. IEEE.	948 949 950 951 952
	Junnan Dong, Qinggang Zhang, Huachi Zhou, Daochen Zha, Pai Zheng, and Xiao Huang. 2024. Modality-aware integration with large language models for knowledge-based visual question answering . <i>Preprint</i> , arXiv:2402.12728.	953 954 955 956 957
	Xiao Dong, Xunlin Zhan, Yunchao Wei, Xiaoyong Wei, Yaowei Wang, Minlong Lu, Xiaochun Cao, and Xiaodan Liang. 2023. Entity-graph enhanced cross-modal	958 959 960

961	pretraining for instance-level product retrieval. <i>IEEE</i>	Junyu Gao, Tianzhu Zhang, and Changsheng Xu. 2019.	1015
962	<i>Trans. Pattern Anal. Mach. Intell.</i> , 45(11):13117–	I know the relationships: Zero-shot action recognition	1016
963	13133.	via two-stream graph convolutional networks	1017
964	Xin Luna Dong. 2023. Generations of knowledge	and knowledge graphs. In <i>AAAI</i> , pages 8303–8311.	1018
965	graphs: The crazy ideas and the business impact.	AAAI Press.	1019
966	<i>CoRR</i> , abs/2308.14217.		
967	Zi-Yi Dou, Yichong Xu, Zhe Gan, Jianfeng Wang,	Shanghai Gao, Ming-Ming Cheng, Kai Zhao, Xin-Yu	1020
968	Shuohang Wang, Lijuan Wang, Chenguang Zhu,	Zhang, Ming-Hsuan Yang, and Philip H. S. Torr.	1021
969	Pengchuan Zhang, Lu Yuan, Nanyun Peng, Zicheng	2021. Res2net: A new multi-scale backbone archi-	1022
970	Liu, and Michael Zeng. 2022. An empirical study of	tecture. <i>IEEE Trans. Pattern Anal. Mach. Intell.</i> ,	1023
971	training end-to-end vision-and-language transform-	43(2):652–662.	1024
972	ers. In <i>CVPR</i> , pages 18145–18155. IEEE.		
973	Jiao Du, Weisheng Li, Ke Lu, and Bin Xiao. 2016. An	Noa Garcia, Mayu Otani, Chenhui Chu, and Yuta	1025
974	overview of multi-modal medical image fusion. <i>Neu-</i>	Nakashima. 2020. Knowit VQA: answering	1026
975	<i>rocomputing</i> , 215:3–20.	knowledge-based questions about videos. In <i>AAAI</i> ,	1027
976	Yifan Du, Hangyu Guo, Kun Zhou, Wayne Xin Zhao,	pages 10826–10834. AAAI Press.	1028
977	Jinpeng Wang, Chuyuan Wang, Mingchen Cai, Rui-		
978	hua Song, and Ji-Rong Wen. 2023. What makes	Diego García-Olano, Yasumasa Onoe, and Joydeep	1029
979	for good visual instructions? synthesizing complex	Ghosh. 2022. Improving and diagnosing knowledge-	1030
980	visual reasoning instructions for visual instruction	based visual question answering via entity enhanced	1031
981	tuning. <i>CoRR</i> , abs/2311.01487.	knowledge injection. In <i>WWW (Companion Volume)</i> ,	1032
982	Marco Federici, Anjan Dutta, Patrick Forré, Nate Kush-	pages 705–715. ACM.	1033
983	man, and Zeynep Akata. 2020. Learning robust rep-		
984	resentations via multi-view information bottleneck.	François Gardères, Maryam Ziaeeffard, Baptiste Abe-	1034
985	In <i>ICLR</i> . OpenReview.net.	loos, and Freddy Lécué. 2020. Conceptbert:	1035
986	Duoduo Feng, Xiangteng He, and Yuxin Peng. 2023.	Concept-aware representation for visual question	1036
987	MKVSE: multimodal knowledge enhanced visual-	answering. In <i>EMNLP (Findings)</i> , volume EMNLP	1037
988	semantic embedding for image-text retrieval. <i>ACM</i>	2020 of <i>Findings of ACL</i> , pages 489–498. Associa-	1038
989	<i>Trans. Multim. Comput. Commun. Appl.</i> , 19(5):162:1–	tion for Computational Linguistics.	1039
990	162:21.		
991	Jenny Rose Finkel, Trond Grenager, and Christopher D.	Matt Gardner, Joel Grus, Mark Neumann, Oyvind	1040
992	Manning. 2005. Incorporating non-local information	Tafjord, Pradeep Dasigi, Nelson F. Liu, Matthew E.	1041
993	into information extraction systems by gibbs sam-	Peters, Michael Schmitz, and Luke Zettlemoyer.	1042
994	pling. In <i>ACL</i> , pages 363–370. The Association for	2018. Allennlp: A deep semantic natural language	1043
995	Computer Linguistics.	processing platform. <i>CoRR</i> , abs/1803.07640.	1044
996	Andrea Frome, Gregory S. Corrado, Jonathon Shlens,		
997	Samy Bengio, Jeffrey Dean, Marc’Aurelio Ranzato,	Ning Ge, Yonghua Zhu, Xiaoyu Xiong, Binghui Zheng,	1045
998	and Tomás Mikolov. 2013. Devise: A deep visual-	and Jieyu Huang. 2021. Knhigan: Knowledge-	1046
999	semantic embedding model. In <i>NIPS</i> , pages 2121–	enhanced hierarchical generative adversarial network	1047
1000	2129.	for fine-grained text-to-image synthesis. In <i>ISCID</i> ,	1048
1001	Jingru Gan, Xinzhe Han, Shuhui Wang, and Qingming	pages 357–360. IEEE.	1049
1002	Huang. 2023. Open-set knowledge-based visual		
1003	question answering with inference paths. <i>CoRR</i> ,	Yuxia Geng, Jiaoyan Chen, Zhuo Chen, Jeff Z. Pan,	1050
1004	abs/2310.08148.	Zhiqian Ye, Zonggang Yuan, Yantao Jia, and Huajun	1051
1005	Feng Gao, Qing Ping, Govind Thattai, Aishwarya N.	Chen. 2021a. Ontozsl: Ontology-enhanced zero-	1052
1006	Reganti, Ying Nian Wu, and Prem Natarajan. 2022.	shot learning. In <i>WWW</i> , pages 3325–3336. ACM /	1053
1007	Transform-retrieve-generate: Natural language-	IW3C2.	1054
1008	centric outside-knowledge visual question answering.		
1009	In <i>CVPR</i> , pages 5057–5067. IEEE.	Yuxia Geng, Jiaoyan Chen, Zhiqian Ye, Zonggang	1055
1010	Jingying Gao, Qi Wu, Alan Blair, and Maurice Pag-	Yuan, Wei Zhang, and Huajun Chen. 2021b. Ex-	1056
1011	nuccio. 2023. Lora: A logical reasoning augmented	plainable zero-shot learning via attentive graph con-	1057
1012	dataset for visual question answering. In <i>Thirty-</i>	volutional network and knowledge graphs. <i>Semantic</i>	1058
1013	<i>seventh Conference on Neural Information Process-</i>	<i>Web</i> , 12(5):741–765.	1059
1014	<i>ing Systems Datasets and Benchmarks Track</i> .		
		Yuxia Geng, Jiaoyan Chen, Wen Zhang, Yajing Xu,	1060
		Zhuo Chen, Jeff Z. Pan, Yufeng Huang, Feiyu Xiong,	1061
		and Huajun Chen. 2022. Disentangled ontology em-	1062
		bedding for zero-shot learning. In <i>KDD</i> , pages 443–	1063
		453. ACM.	1064
		Yuxia Geng, Jiaoyan Chen, Xiang Zhuang, Zhuo Chen,	1065
		Jeff Z. Pan, Juan Li, Zonggang Yuan, and Huajun	1066
		Chen. 2023. Benchmarking knowledge-driven zero-	1067
		shot learning. <i>J. Web Semant.</i> , 75:100757.	1068

1069	Deepanway Ghosal, Navonil Majumder, Roy Ka-Wei	Yu-Jung Heo, Eun-Sol Kim, Woo Suk Choi, and	1124
1070	Lee, Rada Mihalcea, and Soujanya Poria. 2023. Lan-	Byoung-Tak Zhang. 2022. Hypergraph trans-	1125
1071	guage guided visual question answering: Elevate	former: Weakly-supervised multi-hop reasoning for	1126
1072	your multimodal language model using knowledge-	knowledge-based visual question answering. In <i>ACL</i>	1127
1073	enriched prompts. <i>CoRR</i> , abs/2310.20159.	(1), pages 373–390. Association for Computational	1128
		Linguistics.	1129
1074	José Manuél Gómez-Pérez and Raúl Ortega. 2019.	Ian Horrocks. 2008. Ontologies and the semantic web.	1130
1075	Look, read and enrich - learning from scientific fig-	<i>Communications of the ACM</i> , 51(12):58–67.	1131
1076	ures and their captions. In <i>K-CAP</i> , pages 101–108.		
1077	ACM.	Md. Zakir Hossain, Ferdous Sohel, Mohd Fairuz Shi-	1132
		ratuddin, and Hamid Laga. 2019. A comprehensive	1133
1078	Biao Gong, Xiaoying Xie, Yutong Feng, Yiliang Lv,	survey of deep learning for image captioning. <i>ACM</i>	1134
1079	Yujun Shen, and Deli Zhao. 2023. Uknow: A	<i>Comput. Surv.</i> , 51(6):118:1–118:36.	1135
1080	unified knowledge protocol for common-sense rea-		
1081	soning and vision-language pre-training. <i>CoRR</i> ,	Jingyi Hou, Xinxiao Wu, Yayun Qi, Wentian Zhao,	1136
1082	abs/2302.06891.	Jiebo Luo, and Yunde Jia. 2019. Relational reasoning	1137
		using prior knowledge for visual captioning. <i>CoRR</i> ,	1138
1083	Jiuxiang Gu, Handong Zhao, Zhe Lin, Sheng Li, Jian-	abs/1906.01290.	1139
1084	fei Cai, and Mingyang Ling. 2019. Scene graph		
1085	generation with external knowledge and image re-	Jingyi Hou, Xinxiao Wu, Xiaoxun Zhang, Yayun Qi,	1140
1086	construction. In <i>CVPR</i> , pages 1969–1978. Computer	Yunde Jia, and Jiebo Luo. 2020. Joint commonsense	1141
1087	Vision Foundation / IEEE.	and relation reasoning for image and video caption-	1142
		ing. In <i>AAAI</i> , pages 10973–10980. AAAI Press.	1143
1088	Xinyan Guan, Yanjiang Liu, Hongyu Lin, Yaojie Lu,	Chao-Chun Hsu, Zi-Yuan Chen, Chi-Yang Hsu, Chih-	1144
1089	Ben He, Xianpei Han, and Le Sun. 2023. Mitig-	Chia Li, Tzu-Yuan Lin, Ting-Hao Kenneth Huang,	1145
1090	ating large language model hallucinations via au-	and Lun-Wei Ku. 2020. Knowledge-enriched vi-	1146
1091	tonomous knowledge graph-based retrofitting. <i>CoRR</i> ,	sual storytelling. In <i>AAAI</i> , pages 7952–7960. AAAI	1147
1092	abs/2311.13314.	Press.	1148
1093	Liangke Gui, Borui Wang, Qiuyuan Huang, Alexan-	Chi-Yang Hsu, Yun-Wei Chu, Ting-Hao (Kenneth)	1149
1094	der Hauptmann, Yonatan Bisk, and Jianfeng Gao.	Huang, and Lun-Wei Ku. 2021. Plot and rework:	1150
1095	2022. KAT: A knowledge augmented transformer	Modeling storylines for visual storytelling. In <i>ACL/I-</i>	1151
1096	for vision-and-language. In <i>NAACL-HLT</i> , pages 956–	<i>JCNLP (Findings)</i> , volume ACL/IJCNLP 2021 of	1152
1097	968. Association for Computational Linguistics.	<i>Findings of ACL</i> , pages 4443–4453. Association for	1153
		Computational Linguistics.	1154
1098	Dan Guo, Hui Wang, and Meng Wang. 2022a. Context-	Yushi Hu, Hang Hua, Zhengyuan Yang, Weijia Shi,	1155
1099	aware graph inference with knowledge distillation	Noah A. Smith, and Jiebo Luo. 2022. Promptcap:	1156
1100	for visual dialog. <i>IEEE Trans. Pattern Anal. Mach.</i>	Prompt-guided task-aware image captioning. <i>CoRR</i> ,	1157
1101	<i>Intell.</i> , 44(10):6056–6073.	abs/2211.09699.	1158
1102	Dan Guo, Hui Wang, Hanwang Zhang, Zheng-Jun Zha,	Ziniu Hu, Ahmet Iscen, Chen Sun, Zirui Wang, Kai-	1159
1103	and Meng Wang. 2020. Iterative context-aware graph	Wei Chang, Yizhou Sun, Cordelia Schmid, David A.	1160
1104	inference for visual dialog. In <i>CVPR</i> , pages 10052–	Ross, and Alireza Fathi. 2023. Reveal: Retrieval-	1161
1105	10061. Computer Vision Foundation / IEEE.	augmented visual-language pre-training with multi-	1162
		source multimodal knowledge memory. In <i>CVPR</i> ,	1163
1106	Yangyang Guo, Liqiang Nie, Yongkang Wong, Yib-	pages 23369–23379. IEEE.	1164
1107	ing Liu, Zhiyong Cheng, and Mohan S. Kankanhalli.		
1108	2022b. A unified end-to-end retriever-reader frame-	Feicheng Huang, Zhixin Li, Shengjia Chen, Canlong	1165
1109	work for knowledge-based VQA. In <i>ACM Multime-</i>	Zhang, and Huifang Ma. 2020a. Image captioning	1166
1110	<i>dia</i> , pages 2061–2069. ACM.	with internal and external knowledge. In <i>CIKM</i> ,	1167
		pages 535–544. ACM.	1168
1111	Yudong Han, Jianhua Yin, Jianlong Wu, Yinwei Wei,	Qidong Huang, Xiaoyi Dong, Pan Zhang, Bin Wang,	1169
1112	and Liqiang Nie. 2023. Semantic-aware modular	Conghui He, Jiaqi Wang, Dahua Lin, Weiming	1170
1113	capsule routing for visual question answering. <i>IEEE</i>	Zhang, and Nenghai Yu. 2023a. OPERA: alleviating	1171
1114	<i>Trans. Image Process.</i> , 32:5537–5549.	hallucination in multi-modal large language models	1172
		via over-trust penalty and retrospection-allocation.	1173
1115	Tao He, Lianli Gao, Jingkuan Song, Jianfei Cai, and	<i>CoRR</i> , abs/2311.17911.	1174
1116	Yuan-Fang Li. 2020. Learning from the scene and		
1117	borrowing from the rich: Tackling the long tail in	Qingqiu Huang, Yu Xiong, Anyi Rao, Jiaze Wang, and	1175
1118	scene graph generation. In <i>IJCAI</i> , pages 587–593.	Dahua Lin. 2020b. Movienet: A holistic dataset for	1176
1119	ijcai.org.	movie understanding. In <i>ECCV (4)</i> , volume 12349 of	1177
1120	Xuehai He and Xin Wang. 2023. Multimodal graph	<i>Lecture Notes in Computer Science</i> , pages 709–727.	1178
1121	transformer for multimodal question answering. In	Springer.	1179
1122	<i>EACL</i> , pages 189–200. Association for Computa-		
1123	tional Linguistics.		

1180	Yan Huang, Yuming Wang, Yunan Zeng, and Liang Wang. 2022. MACK: multimodal aligned conceptual knowledge for unpaired image-text matching. In <i>NeurIPS</i> .	1235
1181		1236
1182		1237
1183		
1184	Yu Huang, Chenzhuang Du, Zihui Xue, Xuanyao Chen, Hang Zhao, and Longbo Huang. 2021. What makes multi-modal learning better than single (provably). In <i>NeurIPS</i> , pages 10944–10956.	1238
1185		1239
1186		1240
1187		1241
1188	Yufeng Huang, Jiji Tang, Zhuo Chen, Rongsheng Zhang, Xinfeng Zhang, Weijie Chen, Zeng Zhao, Tangjie Lv, Zhipeng Hu, and Wen Zhang. 2023b. Structure-clip: Enhance multi-modal language representations with structure knowledge. <i>CoRR</i> , abs/2305.06152.	1242
1189		1243
1190		1244
1191		1245
1192		1246
1193	Afzaal Hussain, Ifrah Maqsood, Muhammad Shahzad, and Muhammad Moazam Fraz. 2022. Multimodal knowledge reasoning for enhanced visual question answering. In <i>SITIS</i> , pages 224–230. IEEE.	1247
1194		1248
1195		1249
1196		1250
1197	Jena D. Hwang, Chandra Bhagavatula, Ronan Le Bras, Jeff Da, Keisuke Sakaguchi, Antoine Bosselut, and Yejin Choi. 2021. (comet-) atomic 2020: On symbolic and neural commonsense knowledge graphs. In <i>AAAI</i> , pages 6384–6392. AAAI Press.	1251
1198		1252
1199		1253
1200		1254
1201		1255
1202	Filip Ilievski, Pedro A. Szekely, and Bin Zhang. 2021. CSKG: the commonsense knowledge graph. In <i>ESWC</i> , volume 12731 of <i>Lecture Notes in Computer Science</i> , pages 680–696. Springer.	1256
1203		1257
1204		1258
1205		1259
1206	Gautier Izacard and Edouard Grave. 2021. Leveraging passage retrieval with generative models for open domain question answering. In <i>EACL</i> , pages 874–880. Association for Computational Linguistics.	1260
1207		1261
1208		
1209		
1210	Andrew Jaegle, Felix Gimeno, Andy Brock, Oriol Vinyals, Andrew Zisserman, and João Carreira. 2021. Perceiver: General perception with iterative attention. In <i>ICML</i> , volume 139 of <i>Proceedings of Machine Learning Research</i> , pages 4651–4664. PMLR.	1262
1211		1263
1212		1264
1213		1265
1214		1266
1215	Aman Jain, Mayank Kothiyari, Vishwajeet Kumar, Preethi Jyothi, Ganesh Ramakrishnan, and Soumen Chakrabarti. 2021. Select, substitute, search: A new benchmark for knowledge-augmented visual question answering. In <i>SIGIR</i> , pages 2491–2498. ACM.	1267
1216		
1217		
1218		
1219		
1220	Anubhav Jangra, Sourajit Mukherjee, Adam Jatowt, Sriparna Saha, and Mohammad Hasanuzzaman. 2023. A survey on multi-modal summarization. <i>ACM Comput. Surv.</i> , 55(13s):296:1–296:36.	1268
1221		1269
1222		1270
1223		1271
1224	Xiaoze Jiang, Siyi Du, Zengchang Qin, Yajing Sun, and Jing Yu. 2020. KBGN: knowledge-bridge graph network for adaptive vision-text reasoning in visual dialogue. In <i>ACM Multimedia</i> , pages 1265–1273. ACM.	1272
1225		1273
1226		
1227		
1228		
1229	Xueyao Jiang, Ailisi Li, Jiaqing Liang, Bang Liu, Rui Xie, Wei Wu, Zhixu Li, and Yanghua Xiao. 2022. Visualizable or non-visualizable? exploring the visualizability of concepts in multi-modal knowledge graph. In <i>DASFAA (I)</i> , volume 13245 of <i>Lecture Notes in Computer Science</i> , pages 180–187. Springer.	1274
1230		1275
1231		1276
1232		1277
1233		1278
1234		
	Jeff Johnson, Matthijs Douze, and Hervé Jégou. 2021. Billion-scale similarity search with gpus. <i>IEEE Trans. Big Data</i> , 7(3):535–547.	1279
		1280
	Justin Johnson, Ranjay Krishna, Michael Stark, Li-Jia Li, David A. Shamma, Michael S. Bernstein, and Li Fei-Fei. 2015. Image retrieval using scene graphs. In <i>CVPR</i> , pages 3668–3678. IEEE Computer Society.	1281
		1282
	Jun Cheng and Fuxiang Wu and Yanling Tian and Lei Wang and Dapeng Tao. 2022. Rifegan2: Rich feature generation for text-to-image synthesis from constrained prior knowledge. <i>IEEE Trans. Circuits Syst. Video Technol.</i> , 32(8):5187–5200.	1283
		1284
	Ehsan Kamalloo, Nouha Dziri, Charles L. A. Clarke, and Davood Rafiei. 2023. Evaluating open-domain question answering in the era of large language models. In <i>ACL (I)</i> , pages 5591–5606. Association for Computational Linguistics.	1285
		1286
	Michael Kampffmeyer, Yinbo Chen, Xiaodan Liang, Hao Wang, Yujia Zhang, and Eric P. Xing. 2019. Rethinking knowledge graph propagation for zero-shot learning. In <i>CVPR</i> , pages 11487–11496. Computer Vision Foundation / IEEE.	1287
		1288
	Gi-Cheon Kang, Junseok Park, Hwaran Lee, Byoung-Tak Zhang, and Jin-Hwa Kim. 2021. Reasoning visual dialog with sparse graph learning and knowledge transfer. In <i>EMNLP (Findings)</i> , pages 327–339. Association for Computational Linguistics.	1289
		1290
	Vladimir Karpukhin, Barlas Oguz, Sewon Min, Patrick S. H. Lewis, Ledell Wu, Sergey Edunov, Danqi Chen, and Wen-tau Yih. 2020. Dense passage retrieval for open-domain question answering. In <i>EMNLP (I)</i> , pages 6769–6781. Association for Computational Linguistics.	
	Mahmoud Khademi, Ziyi Yang, Felipe Frujeri, and Chenguang Zhu. 2023. Mm-reasoner: A multi-modal knowledge-aware framework for knowledge-based visual question answering. In <i>EMNLP (Findings)</i> , pages 6571–6581. Association for Computational Linguistics.	
	Muhammad Jaleed Khan, John G. Breslin, and Edward Curry. 2022a. Common sense knowledge infusion for visual understanding and reasoning: Approaches, challenges, and applications. <i>IEEE Internet Comput.</i> , 26(4):21–27.	
	Muhammad Jaleed Khan, John G. Breslin, and Edward Curry. 2022b. Expressive scene graph generation using commonsense knowledge infusion for visual understanding and reasoning. In <i>ESWC</i> , volume 13261 of <i>Lecture Notes in Computer Science</i> , pages 93–112. Springer.	
	Apoorv Khandelwal, Ellie Pavlick, and Chen Sun. 2023. Analyzing modular approaches for visual question decomposition. <i>CoRR</i> , abs/2311.06411.	
	Jin-Hwa Kim, Jaehyun Jun, and Byoung-Tak Zhang. 2018. Bilinear attention networks. In <i>NeurIPS</i> , pages 1571–1581.	

1291	Barbara Ann Kipfer. 1992. Roget’s 21st century the-	Mike Lewis, Yinhan Liu, Naman Goyal, Marjan	1347
1292	saurus in dictionary form: the essential reference for	Ghazvininejad, Abdelrahman Mohamed, Omer Levy,	1348
1293	home, school, or office. (<i>No Title</i>).	Veselin Stoyanov, and Luke Zettlemoyer. 2020.	1349
1294	Rajat Koner, Hang Li, Marcel Hildebrandt, Deepan	BART: denoising sequence-to-sequence pre-training	1350
1295	Das, Volker Tresp, and Stephan Günnemann. 2021.	for natural language generation, translation, and com-	1351
1296	Graphhopper: Multi-hop scene graph reasoning for	prehension. In <i>ACL</i> , pages 7871–7880. Association	1352
1297	visual question answering. In <i>ISWC</i> , volume 12922	for Computational Linguistics.	1353
1298	of <i>Lecture Notes in Computer Science</i> , pages 111–		
1299	127. Springer.	Guohao Li, Hang Su, and Wenwu Zhu. 2017. Incorpor-	1354
1300	Ranjay Krishna, Yuke Zhu, Oliver Groth, Justin John-	ating external knowledge to answer open-domain	1355
1301	son, Kenji Hata, Joshua Kravitz, Stephanie Chen,	visual questions with dynamic memory networks.	1356
1302	Yannis Kalantidis, Li-Jia Li, David A. Shamma,	<i>CoRR</i> , abs/1712.00733.	1357
1303	Michael S. Bernstein, and Li Fei-Fei. 2017. Vi-		
1304	sual genome: Connecting language and vision using	Guohao Li, Xin Wang, and Wenwu Zhu. 2020. Boost-	1358
1305	crowdsourced dense image annotations. <i>Int. J.</i>	ing visual question answering with context-aware	1359
1306	<i>Comput. Vis.</i> , 123(1):32–73.	knowledge aggregation. In <i>ACM Multimedia</i> , pages	1360
1307	Ankit Kumar, Ozan Irsoy, Peter Ondruska, Mohit Iyyer,	1227–1235. ACM.	1361
1308	James Bradbury, Ishaan Gulrajani, Victor Zhong, Ro-		
1309	main Paulus, and Richard Socher. 2016. Ask me	Jiaqi Li, Guilin Qi, Chuanyi Zhang, Yongrui Chen, Yim-	1362
1310	anything: Dynamic memory networks for natural lan-	ing Tan, Chenlong Xia, and Ye Tian. 2023a. Incorpor-	1363
1311	guage processing. In <i>ICML</i> , volume 48 of <i>JMLR</i>	ating domain knowledge graph into multimodal	1364
1312	<i>Workshop and Conference Proceedings</i> , pages 1378–	movie genre classification with self-supervised atten-	1365
1313	1387. JMLR.org.	tion and contrastive learning. In <i>ACM Multimedia</i> ,	1366
1314	Christoph H. Lampert, Hannes Nickisch, and Stefan	pages 3337–3345. ACM.	1367
1315	Harmeling. 2014. Attribute-based classification for		
1316	zero-shot visual object categorization. <i>IEEE Trans.</i>	Junnan Li, Dongxu Li, Silvio Savarese, and Steven C. H.	1368
1317	<i>Pattern Anal. Mach. Intell.</i> , 36(3):453–465.	Hoi. 2023b. BLIP-2: bootstrapping language-image	1369
1318	Zhenzhong Lan, Mingda Chen, Sebastian Goodman,	pre-training with frozen image encoders and large lan-	1370
1319	Kevin Gimpel, Piyush Sharma, and Radu Soricut.	guage models. In <i>ICML</i> , volume 202 of <i>Proceedings</i>	1371
1320	2020. ALBERT: A lite BERT for self-supervised	of <i>Machine Learning Research</i> , pages 19730–19742.	1372
1321	learning of language representations. In <i>ICLR</i> . Open-	PMLR.	1373
1322	Review.net.		
1323	Quoc V. Le and Tomáš Mikolov. 2014. Distributed rep-	Junnan Li, Dongxu Li, Caiming Xiong, and Steven	1374
1324	resentations of sentences and documents. In <i>ICML</i> ,	C. H. Hoi. 2022a. BLIP: bootstrapping language-	1375
1325	volume 32 of <i>JMLR Workshop and Conference Pro-</i>	image pre-training for unified vision-language under-	1376
1326	<i>ceedings</i> , pages 1188–1196. JMLR.org.	standing and generation. In <i>ICML</i> , volume 162 of	1377
1327	Chung-Wei Lee, Wei Fang, Chih-Kuan Yeh, and Yu-	<i>Proceedings of Machine Learning Research</i> , pages	1378
1328	Chiang Frank Wang. 2018. Multi-label zero-shot	12888–12900. PMLR.	1379
1329	learning with structured knowledge graphs. In <i>CVPR</i> ,		
1330	pages 1576–1585. Computer Vision Foundation /	Junnan Li, Ramprasaath R. Selvaraju, Akhilesh	1380
1331	IEEE Computer Society.	Gotmare, Shafiq R. Joty, Caiming Xiong, and	1381
1332	Jaejun Lee, Chanyoung Chung, Hochang Lee, Sungho	Steven Chu-Hong Hoi. 2021. Align before fuse: Vi-	1382
1333	Jo, and Joyce Jiyoung Whang. 2023. VISTA: visual-	sion and language representation learning with mo-	1383
1334	textual knowledge graph representation learning. In	mentum distillation. In <i>NeurIPS</i> , pages 9694–9705.	1384
1335	<i>EMNLP (Findings)</i> , pages 7314–7328. Association		
1336	for Computational Linguistics.	Liunian Harold Li, Mark Yatskar, Da Yin, Cho-Jui	1385
1337	Adam Lerer, Ledell Wu, Jiajun Shen, Timothée	Hsieh, and Kai-Wei Chang. 2019. Visualbert: A sim-	1386
1338	Lacroix, Luca Wehrstedt, Abhijit Bose, and Alex	ple and performant baseline for vision and language.	1387
1339	Peysakhovich. 2019. Pytorch-biggraph: A large	<i>CoRR</i> , abs/1908.03557.	1388
1340	scale graph embedding system. In <i>MLSys</i> . mlsys.org.		
1341	Paul Lerner, Olivier Ferret, Camille Guinaudeau,	Mingxiao Li and Marie-Francine Moens. 2022. Dy-	1389
1342	Hervé Le Borgne, Romaric Besançon, José G.	namic key-value memory enhanced multi-step graph	1390
1343	Moreno, and Jesús Lovón-Melgarejo. 2022. Viquae,	reasoning for knowledge-based visual question an-	1391
1344	a dataset for knowledge-based visual question an-	swering. In <i>AAAI</i> , pages 10983–10992. AAAI Press.	1392
1345	swering about named entities. In <i>SIGIR</i> , pages 3108–		
1346	3120. ACM.	Rengang Li, Cong Xu, Zhenhua Guo, Baoyu Fan, Runze	1393
		Zhang, Wei Liu, Yaqian Zhao, Weifeng Gong, and	1394
		Endong Wang. 2022b. AI-VQA: visual question an-	1395
		swering based on agent interaction with interpretabil-	1396
		ity. In <i>ACM Multimedia</i> , pages 5274–5282. ACM.	1397
		Tengpeng Li, Hanli Wang, Bin He, and Chang Wen	1398
		Chen. 2023c. Knowledge-enriched attention net-	1399
		work with group-wise semantic for visual storytelling.	1400
		<i>IEEE Trans. Pattern Anal. Mach. Intell.</i> , 45(7):8634–	1401
		8645.	1402

1403	Wenhui Li, Song Yang, Qiang Li, Xuanya Li, and	Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae	1456
1404	An-An Liu. 2023d. Commonsense-guided semantic	Lee. 2023c. Visual instruction tuning. <i>CoRR</i> ,	1457
1405	and relational consistencies for image-text retrieval.	abs/2304.08485.	1458
1406	<i>IEEE Transactions on Multimedia</i> .		
1407	Xiangyang Li and Shuqiang Jiang. 2019. Know more	Luping Liu, Meiling Wang, Xiaohai He, Linbo Qing,	1459
1408	say less: Image captioning based on scene graphs.	and Honggang Chen. 2022. Fact-based visual ques-	1460
1409	<i>IEEE Trans. Multim.</i> , 21(8):2117–2130.	tion answering via dual-process system. <i>Knowl.</i>	1461
		<i>Based Syst.</i> , 237:107650.	1462
1410	Xin Li, Dongze Lian, Zhihe Lu, Jiawang Bai, Zhibo	Zijun Long, George Killick, Richard McCreddie,	1463
1411	Chen, and Xinchao Wang. 2023e. Graphadapter:	and Gerardo Aragon-Camarasa. 2023. Multiway-	1464
1412	Tuning vision-language models with dual knowledge	adapater: Adapting large-scale multi-modal mod-	1465
1413	graph. <i>CoRR</i> , abs/2309.13625.	els for scalable image-text retrieval. <i>CoRR</i> ,	1466
		abs/2309.01516.	1467
1414	Xirong Li, Shuai Liao, Weiyu Lan, Xiaoyong Du, and	Di Lu, Spencer Whitehead, Lifu Huang, Heng Ji, and	1468
1415	Gang Yang. 2015. Zero-shot image tagging by hi-	Shih-Fu Chang. 2018a. Entity-aware image caption	1469
1416	erarchical semantic embedding. In <i>SIGIR</i> , pages	generation. In <i>EMNLP</i> , pages 4013–4023. Associa-	1470
1417	879–882. ACM.	tion for Computational Linguistics.	1471
1418	Yunxin Li, Longyue Wang, Baotian Hu, Xinyu Chen,	Jiale Lu, Lianggangxu Chen, Youqi Song, Shaohui	1472
1419	Wanqi Zhong, Chenyang Lyu, and Min Zhang.	Lin, Changbo Wang, and Gaoqi He. 2023. Prior	1473
1420	2023f. A comprehensive evaluation of GPT-4V	knowledge-driven dynamic scene graph generation	1474
1421	on knowledge-intensive visual question answering.	with causal inference. In <i>ACM Multimedia</i> , pages	1475
1422	<i>CoRR</i> , abs/2311.07536.	4877–4885. ACM.	1476
1423	Tsung-Yi Lin, Michael Maire, Serge J. Belongie, James	Jiasen Lu, Vedanuj Goswami, Marcus Rohrbach, Devi	1477
1424	Hays, Pietro Perona, Deva Ramanan, Piotr Dollár,	Parikh, and Stefan Lee. 2020. 12-in-1: Multi-task vi-	1478
1425	and C. Lawrence Zitnick. 2014. Microsoft COCO:	sion and language representation learning. In <i>CVPR</i> ,	1479
1426	common objects in context. In <i>ECCV (5)</i> , volume	pages 10434–10443. Computer Vision Foundation /	1480
1427	8693 of <i>Lecture Notes in Computer Science</i> , pages	IEEE.	1481
1428	740–755. Springer.		
1429	Weizhe Lin and Bill Byrne. 2022. Retrieval augmented	Keming Lu, Bowen Yu, Chang Zhou, and Jingren Zhou.	1482
1430	visual question answering with outside knowledge.	2024. Large language models are superpositions of	1483
1431	In <i>EMNLP</i> , pages 11238–11254. Association for	all characters: Attaining arbitrary role-play via self-	1484
1432	Computational Linguistics.	alignment.	1485
1433	Weizhe Lin, Zhilin Wang, and Bill Byrne. 2023. FVQA	Pan Lu, Lei Ji, Wei Zhang, Nan Duan, Ming Zhou, and	1486
1434	2.0: Introducing adversarial samples into fact-based	Jianyong Wang. 2018b. R-VQA: learning visual rela-	1487
1435	visual question answering. In <i>EACL (Findings)</i> ,	tion facts with semantic attention for visual question	1488
1436	pages 149–157. Association for Computational Lin-	answering. In <i>KDD</i> , pages 1880–1889. ACM.	1489
1437	guistics.		
1438	Yuanze Lin, Yujia Xie, Dongdong Chen, Yichong Xu,	Pan Lu, Swaroop Mishra, Tanglin Xia, Liang Qiu, Kai-	1490
1439	Chenguang Zhu, and Lu Yuan. 2022. REVIVE: re-	Wei Chang, Song-Chun Zhu, Oyvind Tafjord, Peter	1491
1440	gional visual representation matters in knowledge-	Clark, and Ashwin Kalyan. 2022. Learn to explain:	1492
1441	based visual question answering. In <i>NeurIPS</i> .	Multimodal reasoning via thought chains for science	1493
		question answering. In <i>NeurIPS</i> .	1494
1442	An-An Liu, Chenxi Huang, Ning Xu, Hongshuo Tian,	Man Luo, Yankai Zeng, Pratyay Banerjee, and Chitta	1495
1443	Jing Liu, and Yongdong Zhang. 2023a. Counterfac-	Baral. 2021. Weakly-supervised visual-retriever-	1496
1444	tual visual dialog: Robust commonsense knowledge	reader for knowledge-based question answering. In	1497
1445	learning from unbiased training. <i>IEEE Transactions</i>	<i>EMNLP (1)</i> , pages 6417–6431. Association for Com-	1498
1446	<i>on Multimedia</i> .	putational Linguistics.	1499
1447	An-An Liu, Zefang Sun, Ning Xu, Rongbao Kang, Jinbo	Ang Lv, Kaiyi Zhang, Shufang Xie, Quan Tu, Yuhan	1500
1448	Cao, Fan Yang, Weijun Qin, Shenyuan Zhang, Jiaqi	Chen, Ji-Rong Wen, and Rui Yan. 2023. Are we	1501
1449	Zhang, and Xuanya Li. 2023b. Prior knowledge	falling in a middle-intelligence trap? an analy-	1502
1450	guided text to image generation. <i>Pattern Recognition</i>	sis and mitigation of the reversal curse. <i>CoRR</i> ,	1503
1451	<i>Letters</i> .	abs/2311.07468.	1504
1452	Hanchao Liu, Wenyuan Xue, Yifei Chen, Dapeng Chen,	Maria Lymperaïou and Giorgos Stamou. 2022. A survey	1505
1453	Xiutian Zhao, Ke Wang, Liping Hou, Rongjun Li, and	on knowledge-enhanced multimodal learning. <i>CoRR</i> ,	1506
1454	Wei Peng. 2024. A survey on hallucination in large	abs/2211.12328.	1507
1455	vision-language models. <i>CoRR</i> , abs/2402.00253.		

1508	Chaitanya Malaviya, Chandra Bhagavatula, Antoine Bosselut, and Yejin Choi. 2020. Commonsense knowledge base completion with structural and semantic context. In <i>AAAI</i> , pages 2925–2933. AAAI Press.	Nihal V. Nayak and Stephen H. Bach. 2022. Zero-shot learning with common sense knowledge graphs. <i>Trans. Mach. Learn. Res.</i> , 2022.	1562
1509			1563
1510			1564
1511			
1512		Tuan-Phong Nguyen, Simon Razniewski, and Gerhard Weikum. 2021. Advanced semantics for commonsense knowledge extraction. In <i>WWW</i> , pages 2636–2647. ACM / IW3C2.	1565
1513	Shengyu Mao, Ningyu Zhang, Xiaohan Wang, Mengru Wang, Yunzhi Yao, Yong Jiang, Pengjun Xie, Fei Huang, and Huajun Chen. 2023. Editing personality for llms. <i>CoRR</i> , abs/2310.02168.		1566
1514			1567
1515		Weizhi Nie, Ruidong Chen, Weijie Wang, Bruno Lepri, and Nicu Sebe. 2023. T2TD: text-3d generation model based on prior knowledge guidance. <i>CoRR</i> , abs/2305.15753.	1568
1516			1569
1517	Kenneth Marino, Xinlei Chen, Devi Parikh, Abhinav Gupta, and Marcus Rohrbach. 2021. KRISP: integrating implicit and symbolic knowledge for open-domain knowledge-based VQA. In <i>CVPR</i> , pages 14111–14121. Computer Vision Foundation / IEEE.		1570
1518			1571
1519		Sofia Nikiforova, Tejaswini Deoskar, Denis Paperno, and Yoad Winter. 2022. Generating image captions with external encyclopedic knowledge. <i>CoRR</i> , abs/2210.04806.	1572
1520			1573
1521			1574
1522	Kenneth Marino, Mohammad Rastegari, Ali Farhadi, and Roozbeh Mottaghi. 2019. OK-VQA: A visual question answering benchmark requiring external knowledge. In <i>CVPR</i> , pages 3195–3204. Computer Vision Foundation / IEEE.		1575
1523			1576
1524		Timothy Ossowski and Junjie Hu. 2023. Multimodal prompt retrieval for generative visual question answering. <i>CoRR</i> , abs/2306.17675.	1577
1525			1578
1526			1579
1527	George A Miller. 1995. WordNet: A lexical database for english. <i>Communications of the ACM</i> , 38(11):39–41.	Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll L. Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, John Schulman, Jacob Hilton, Fraser Kelton, Luke Miller, Maddie Simens, Amanda Askell, Peter Welinder, Paul F. Christiano, Jan Leike, and Ryan Lowe. 2022. Training language models to follow instructions with human feedback. In <i>NeurIPS</i> .	1580
1528			1581
1529	Abhika Mishra, Akari Asai, Vidhisha Balachandran, Yizhong Wang, Graham Neubig, Yulia Tsvetkov, and Hannaneh Hajishirzi. 2024. Fine-grained hallucination detection and editing for language models.		1582
1530			1583
1531		Oded Ovadia, Menachem Brief, Moshik Mishaeli, and Oren Elisha. 2023. Fine-tuning or retrieval? comparing knowledge injection in llms. <i>CoRR</i> , abs/2312.05934.	1584
1532			1585
1533	Aditya Mogadala, Xiaoyu Shen, and Dietrich Klakow. 2020. Integrating image captioning with rule-based entity masking. <i>CoRR</i> , abs/2007.11690.		1586
1534			1587
1535		Vardaan Pahuja, Weidi Luo, Yu Gu, Cheng-Hao Tu, Hong-You Chen, Tanya Y. Berger-Wolf, Charles V. Stewart, Song Gao, Wei-Lun Chao, and Yu Su. 2024. Bringing back the context: Camera trap species identification as link prediction on multimodal knowledge graphs.	1588
1536	Debjyoti Mondal, Suraj Modi, Subhadarshi Panda, Rituraj Singh, and Godawari Sudhakar Rao. 2024. Kamcot: Knowledge augmented multimodal chain-of-thoughts reasoning.		1589
1537			1590
1538			1591
1539		Jeff Z. Pan, Simon Razniewski, Jan-Christoph Kalo, Sneha Singhania, Jiaoyan Chen, Stefan Dietze, Hajira Jabeen, Janna Omeljanenko, Wen Zhang, Matteo Lissandrini, Russa Biswas, Gerard de Melo, Angela Bonifati, Edlira Vakaj, Mauro Dragoni, and Damien Graux. 2023a. Large language models and knowledge graphs: Opportunities and challenges. <i>CoRR</i> , abs/2308.06374.	1592
1540	Sebastian Monka, Lavdim Halilaj, and Achim Rettinger. 2022. A survey on visual transfer learning using knowledge graphs. <i>Semantic Web</i> , 13(3):477–510.		1593
1541			1594
1542			1595
1543	Federico Monti, Davide Boscaini, Jonathan Masci, Emanuele Rodolà, Jan Svoboda, and Michael M. Bronstein. 2017. Geometric deep learning on graphs and manifolds using mixture model cnns. In <i>CVPR</i> , pages 5425–5434. IEEE Computer Society.		1596
1544			1597
1545		Shirui Pan, Linhao Luo, Yufei Wang, Chen Chen, Jiaapu Wang, and Xindong Wu. 2023b. Unifying large language models and knowledge graphs: A roadmap. <i>CoRR</i> , abs/2306.08302.	1598
1546			1599
1547			1600
1548	Nasrin Mostafazadeh, Ishan Misra, Jacob Devlin, Margaret Mitchell, Xiaodong He, and Lucy Vanderwende. 2016. Generating natural questions about an image. In <i>ACL (1)</i> . The Association for Computer Linguistics.		1601
1549			1602
1550		Ziqi Pang, Ziyang Xie, Yunze Man, and Yu-Xiong Wang. 2023. Frozen transformers in language models are effective visual encoder layers. <i>CoRR</i> , abs/2310.12973.	1603
1551			1604
1552			1605
1553	Medhini Narasimhan, Svetlana Lazebnik, and Alexander G. Schwing. 2018. Out of the box: Reasoning with graph convolution nets for factual visual question answering. In <i>NeurIPS</i> , pages 2659–2670.		1606
1554			1607
1555		Kishore Papineni, Salim Roukos, Todd Ward, and Wei-Jing Zhu. 2002. Bleu: a method for automatic evaluation of machine translation. In <i>ACL</i> , pages 311–318. ACL.	1608
1556			1609
1557	Medhini Narasimhan and Alexander G. Schwing. 2018. Straight to the facts: Learning knowledge base retrieval for factual visual question answering. In <i>ECCV (8)</i> , volume 11212 of <i>Lecture Notes in Computer Science</i> , pages 460–477. Springer.		1610
1558			1611
1559			1612
1560			1613
1561			1614
			1615
			1616
			1617

1618	Jinyoung Park, Ameen Patel, Omar Zia Khan, Hyunwoo J. Kim, and Joo-Kyung Kim. 2023. Graph-guided reasoning for multi-hop question answering in large language models. <i>CoRR</i> , abs/2311.09762.	1672
1619		1673
1620		1674
1621		1675
1622	Jeffrey Pennington, Richard Socher, and Christopher D. Manning. 2014. Glove: Global vectors for word representation. In <i>EMNLP</i> , pages 1532–1543. ACL.	1676
1623		1677
1624		1678
1625	Fabio Petroni, Tim Rocktäschel, Sebastian Riedel, Patrick S. H. Lewis, Anton Bakhtin, Yuxiang Wu, and Alexander H. Miller. 2019. Language models as knowledge bases? In <i>EMNLP/IJCNLP (1)</i> , pages 2463–2473. Association for Computational Linguistics.	1679
1626		1680
1627		1681
1628		1682
1629		1683
1630		1684
1631	Peng Qi, Yuhao Zhang, Yuhui Zhang, Jason Bolton, and Christopher D. Manning. 2020. Stanza: A python natural language processing toolkit for many human languages. In <i>ACL (demo)</i> , pages 101–108. Association for Computational Linguistics.	1685
1632		1686
1633		1687
1634		1688
1635		1689
1636	Shengsheng Qian, Jun Hu, Quan Fang, and Changsheng Xu. 2021. Knowledge-aware multi-modal adaptive graph convolutional networks for fake news detection. <i>ACM Trans. Multim. Comput. Commun. Appl.</i> , 17(3):98:1–98:23.	1690
1637		1691
1638		1692
1639		1693
1640		1694
1641	Shuofei Qiao, Ningyu Zhang, Runnan Fang, Yujie Luo, Wangchunshu Zhou, Yuchen Eleanor Jiang, Chengfei Lv, and Huajun Chen. 2024. AUTOACT: automatic agent learning from scratch via self-planning.	1695
1642		1696
1643		1697
1644		1698
1645	Tingting Qiao, Jing Zhang, Duanqing Xu, and Dacheng Tao. 2019. Learn, imagine and create: Text-to-image generation from prior knowledge. In <i>NeurIPS</i> , pages 885–895.	1699
1646		1700
1647		1701
1648		1702
1649	Yanyuan Qiao, Chaorui Deng, and Qi Wu. 2021. Referring expression comprehension: A survey of methods and datasets. <i>IEEE Trans. Multim.</i> , 23:4426–4440.	1703
1650		1704
1651		1705
1652	Chen Qu, Hamed Zamani, Liu Yang, W. Bruce Croft, and Erik G. Learned-Miller. 2021. Passage retrieval for outside-knowledge visual question answering. In <i>SIGIR</i> , pages 1753–1757. ACM.	1706
1653		1707
1654		1708
1655		1709
1656	Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J. Liu. 2020. Exploring the limits of transfer learning with a unified text-to-text transformer. <i>J. Mach. Learn. Res.</i> , 21:140:1–140:67.	1710
1657		1711
1658		1712
1659		1713
1660		1714
1661	Kiran Ramnath and Mark Hasegawa-Johnson and. 2020. Seeing is knowing! fact-based visual question answering using knowledge graph embeddings. <i>CoRR</i> , abs/2012.15484.	1715
1662		1716
1663		1717
1664		1718
1665	Sahithya Ravi, Aditya Chinchure, Leonid Sigal, Renjie Liao, and Vered Shwartz. 2023. VLC-BERT: visual question answering with contextualized commonsense knowledge. In <i>WACV</i> , pages 1155–1165. IEEE.	1719
1666		1720
1667		1721
1668		1722
1669		1723
1670	Joseph Redmon and Ali Farhadi. 2018. Yolov3: An incremental improvement. <i>CoRR</i> , abs/1804.02767.	1724
1671		1725
		1726
	Benjamin Z. Reichman, Anirudh Sundar, Christopher Richardson, Tamara Zubatiy, Prithwjit Chowdhury, Aaryan Shah, Jack Truxal, Micah Grimes, Dristi Shah, Woo Ju Chee, Saif Punjwani, Atishay Jain, and Larry Heck. 2023. Outside knowledge visual question answering version 2.0. In <i>ICASSP</i> , pages 1–5. IEEE.	
	Nils Reimers and Iryna Gurevych. 2019. Sentence-bert: Sentence embeddings using siamese bert-networks. In <i>EMNLP/IJCNLP (1)</i> , pages 3980–3990. Association for Computational Linguistics.	
	Abhinaba Roy, Deepanway Ghosal, Erik Cambria, Navonil Majumder, Rada Mihalcea, and Soujanya Poria. 2022. Improving zero-shot learning baselines with commonsense knowledge. <i>Cogn. Comput.</i> , 14(6):2212–2222.	
	Olga Russakovsky, Jia Deng, Hao Su, Jonathan Krause, Sanjeev Satheesh, Sean Ma, Zhiheng Huang, Andrej Karpathy, Aditya Khosla, Michael S. Bernstein, Alexander C. Berg, and Li Fei-Fei. 2015. Imagenet large scale visual recognition challenge. <i>Int. J. Comput. Vis.</i> , 115(3):211–252.	
	Tara Safavi and Danai Koutra. 2020. Codex: A comprehensive knowledge graph completion benchmark. In <i>EMNLP (1)</i> , pages 8328–8350. Association for Computational Linguistics.	
	Ander Salaberria, Gorka Azkune, Oier Lopez de Lacalle, Aitor Soroa, and Eneko Agirre. 2023. Image captioning for effective use of language models in knowledge-based visual question answering. <i>Expert Syst. Appl.</i> , 212:118669.	
	Alireza Salemi, Juan Altmayer Pizzorno, and Hamed Zamani. 2023a. A symmetric dual encoding dense retrieval framework for knowledge-intensive visual question answering. In <i>SIGIR</i> , pages 110–120. ACM.	
	Alireza Salemi, Mahta Rafiee, and Hamed Zamani. 2023b. Pre-training multi-modal dense retrievers for outside-knowledge visual question answering. In <i>ICTIR</i> , pages 169–176. ACM.	
	Tim Salimans, Ian J. Goodfellow, Wojciech Zaremba, Vicki Cheung, Alec Radford, and Xi Chen. 2016. Improved techniques for training gans. In <i>NIPS</i> , pages 2226–2234.	
	Maarten Sap, Ronan Le Bras, Emily Allaway, Chandra Bhagavatula, Nicholas Lourie, Hannah Rashkin, Brendan Roof, Noah A. Smith, and Yejin Choi. 2019. ATOMIC: an atlas of machine commonsense for if-then reasoning. In <i>AAAI</i> , pages 3027–3035. AAAI Press.	
	Raeid Saqr and Karthik Narasimhan. 2020. Multi-modal graph networks for compositional generalization in visual question answering. In <i>NeurIPS</i> .	
	Rylan Schaeffer, Brando Miranda, and Sanmi Koyejo. 2023. Are emergent abilities of large language models a mirage? <i>CoRR</i> , abs/2304.15004.	

1727	Michael Sejr Schlichtkrull, Thomas N. Kipf, Peter Bloem, Rianne van den Berg, Ivan Titov, and Max Welling. 2018. Modeling relational data with graph convolutional networks. In <i>ESWC</i> , volume 10843 of <i>Lecture Notes in Computer Science</i> , pages 593–607. Springer.	1779	Robyn Speer, Joshua Chin, and Catherine Havasi. 2017. Conceptnet 5.5: An open multilingual graph of general knowledge. In <i>Thirty-first AAAI conference on artificial intelligence</i> .	1780
1728		1781		1782
1729		1783	Krishna Srinivasan, Karthik Raman, Jiecao Chen, Michael Bendersky, and Marc Najork. 2021. WIT: wikipedia-based image text dataset for multimodal multilingual machine learning. In <i>SIGIR</i> , pages 2443–2449. ACM.	1784
1730		1785		1786
1731		1787		
1732		1788	Zhou Su, Chen Zhu, Yinpeng Dong, Dongqi Cai, Yurong Chen, and Jianguo Li. 2018. Learning visual knowledge memory networks for visual question answering. In <i>CVPR</i> , pages 7736–7745. Computer Vision Foundation / IEEE Computer Society.	1789
1733	Dustin Schwenk, Apoorv Khandelwal, Christopher Clark, Kenneth Marino, and Roozbeh Mottaghi. 2022. A-OKVQA: A benchmark for visual question answering using world knowledge. In <i>ECCV (8)</i> , volume 13668 of <i>Lecture Notes in Computer Science</i> , pages 146–162. Springer.	1790		1791
1734		1792		
1735		1793	Sanjay Subramanian, Medhini Narasimhan, Kushal Khangaonkar, Kevin Yang, Arsha Nagraani, Cordelia Schmid, Andy Zeng, Trevor Darrell, and Dan Klein. 2023. Modular visual question answering via code generation. In <i>ACL (2)</i> , pages 747–761. Association for Computational Linguistics.	1794
1736		1795		1796
1737		1797		1798
1738		1799	Fabian M. Suchanek, Gjergji Kasneci, and Gerhard Weikum. 2007. Yago: a core of semantic knowledge. In <i>WWW</i> , pages 697–706. ACM.	1800
1739	Sanket Shah, Anand Mishra, Naganand Yadati, and Partha Pratim Talukdar. 2019. KVQA: knowledge-aware visual question answering. In <i>AAAI</i> , pages 8876–8884. AAAI Press.	1801		
1740		1802	Jiashuo Sun, Chengjin Xu, Lumingyuan Tang, Saizhuo Wang, Chen Lin, Yeyun Gong, Heung-Yeung Shum, and Jian Guo. 2023a. Think-on-graph: Deep and responsible reasoning of large language model with knowledge graph. <i>CoRR</i> , abs/2307.07697.	1803
1741		1804		1805
1742		1806		
1743	Zhenwei Shao, Zhou Yu, Meng Wang, and Jun Yu. 2023. Prompting large language models with answer heuristics for knowledge-based visual question answering. In <i>CVPR</i> , pages 14974–14983. IEEE.	1807	Kai Sun, Yifan Ethan Xu, Hanwen Zha, Yue Liu, and Xin Luna Dong. 2023b. Head-to-tail: How knowledgeable are large language models (llm)? A.K.A. will llms replace knowledge graphs? <i>CoRR</i> , abs/2308.10168.	1808
1744		1809		1810
1745		1811		
1746		1812	Mengzhu Sun, Xi Zhang, Jianqiang Ma, Sihong Xie, Yazheng Liu, and Philip S. Yu. 2023c. Inconsistent matters: A knowledge-guided dual-consistency network for multi-modal rumor detection. <i>IEEE Trans. Knowl. Data Eng.</i> , 35(12):12736–12749.	1813
1747	Violetta Shevchenko, Damien Teney, Anthony R. Dick, and Anton van den Hengel. 2021. Reasoning over vision and language: Exploring the benefits of supplemental knowledge. <i>CoRR</i> , abs/2101.06013.	1814		1815
1748		1816		
1749		1817	Quan Sun, Yufeng Cui, Xiaosong Zhang, Fan Zhang, Qiying Yu, Zhengxiong Luo, Yueze Wang, Yongming Rao, Jingjing Liu, Tiejun Huang, and Xinlong Wang. 2023d. Generative multimodal models are in-context learners. <i>CoRR</i> , abs/2312.13286.	1818
1750		1819		1820
1751	Botian Shi, Lei Ji, Pan Lu, Zhendong Niu, and Nan Duan. 2019. Knowledge aware semantic concept expansion for image-text matching. In <i>IJCAI</i> , pages 5182–5189. ijcai.org.	1821		1822
1752		1823	Yu Sun, Shuohuan Wang, Yu-Kun Li, Shikun Feng, Xuyi Chen, Han Zhang, Xin Tian, Danxiang Zhu, Hao Tian, and Hua Wu. 2019a. ERNIE: enhanced representation through knowledge integration. <i>CoRR</i> , abs/1904.09223.	1824
1753		1825		1826
1754		1827	Zhiqing Sun, Zhi-Hong Deng, Jian-Yun Nie, and Jian Tang. 2019b. Rotate: Knowledge graph embedding by relational rotation in complex space. In <i>ICLR (Poster)</i> . OpenReview.net.	1828
1755	Zhan Shi, Yilin Shen, Hongxia Jin, and Xiaodan Zhu. 2022. Improving zero-shot phrase grounding via reasoning on external knowledge and spatial relations. In <i>AAAI</i> , pages 2253–2261. AAAI Press.	1829		1830
1756		1831	Hao Tan and Mohit Bansal. 2019. LXMERT: learning cross-modality encoder representations from transformers. In <i>EMNLP/IJCNLP (1)</i> , pages 5099–5110. Association for Computational Linguistics.	1832
1757		1833		1834
1758		1835		
1759	Qingyi Si, Yuchen Mo, Zheng Lin, Huishan Ji, and Weiping Wang. 2023. Combo of thinking and observing for outside-knowledge VQA. In <i>ACL (1)</i> , pages 10959–10975. Association for Computational Linguistics.			
1760				
1761				
1762				
1763				
1764	Sibei Yang and Guanbin Li and Yizhou Yu. 2021. Relationship-embedded representation learning for grounding referring expressions. <i>IEEE Trans. Pattern Anal. Mach. Intell.</i> , 43(8):2765–2779.			
1765				
1766				
1767				
1768	Linda B. Smith and Michael Gasser. 2005. The development of embodied cognition: Six lessons from babies. <i>Artif. Life</i> , 11(1-2):13–29.			
1769				
1770				
1771	Fangzhou Song, Bin Zhu, Yanbin Hao, Shuo Wang, and Xiangnan He. 2023a. CAR: consolidation, augmentation and regulation for recipe retrieval. <i>CoRR</i> , abs/2312.04763.			
1772				
1773				
1774				
1775	Lingyun Song, Jianao Li, Jun Liu, Yang Yang, Xuequn Shang, and Mingxuan Sun. 2023b. Answering knowledge-based visual questions via the exploration of question purpose. <i>Pattern Recognit.</i> , 133:109015.			
1776				
1777				
1778				

1835	Sinan Tan, Mengmeng Ge, Di Guo, Huaping Liu, and Fuchun Sun. 2023. Knowledge-based embodied question answering. <i>IEEE Trans. Pattern Anal. Mach. Intell.</i> , 45(10):11948–11960.	Peter Vickers, Nikolaos Aletras, Emilio Monti, and Loïc Barrault. 2021. In factuality: Efficient integration of relevant facts for visual question answering. In <i>ACL/IJCNLP (2)</i> , pages 468–475. Association for Computational Linguistics.	1889
1836			1890
1837			1891
1838			1892
1839	Niket Tandon, Gerard de Melo, Fabian M. Suchanek, and Gerhard Weikum. 2014. Webchild: harvesting and organizing commonsense knowledge from the web. In <i>WSDM</i> , pages 523–532. ACM.	Denny Vrandečić and Markus Krötzsch. 2014. Wiki-data: a free collaborative knowledgebase. <i>Commun. ACM</i> , 57(10):78–85.	1894
1840			1895
1841			1896
1842			
1843	Wei Tang, Liang Li, Xuejing Liu, Lu Jin, Jinhui Tang, and Zechao Li. 2023. Context disentangling and prototype inheriting for robust visual grounding. <i>IEEE Transactions on Pattern Analysis and Machine Intelligence</i> .	Fanqi Wan, Xinting Huang, Tao Yang, Xiaojun Quan, Wei Bi, and Shuming Shi. 2023. Explore-instruct: Enhancing domain-specific instruction coverage through active exploration. In <i>EMNLP</i> , pages 9435–9454. Association for Computational Linguistics.	1897
1844			1898
1845			1899
1846			1900
1847			1901
1848	Hongshuo Tian, Ning Xu, Yanhui Wang, Chenggang Yan, Bolun Zheng, Xuanya Li, and An-An Liu. 2023. Towards confidence-aware commonsense knowledge integration for scene graph generation. In <i>ICME</i> , pages 2255–2260. IEEE.	Hai Wan, Yonghao Luo, Bo Peng, and Wei-Shi Zheng. 2018. Representation learning for scene graph completion via jointly structural and visual embedding. In <i>IJCAI</i> , pages 949–956. ijcai.org.	1902
1849			1903
1850			1904
1851			1905
1852			
1853	Shengbang Tong, Zhuang Liu, Yuexiang Zhai, Yi Ma, Yann LeCun, and Saining Xie. 2024. Eyes wide shut? exploring the visual shortcomings of multimodal llms. <i>CoRR</i> , abs/2401.06209.	Haoran Wang, Dongliang He, Wenhao Wu, Boyang Xia, Min Yang, Fu Li, Yunlong Yu, Zhong Ji, Errui Ding, and Jingdong Wang. 2022a. CODER: coupled diversity-sensitive momentum contrastive learning for image-text retrieval. In <i>ECCV (36)</i> , volume 13696 of <i>Lecture Notes in Computer Science</i> , pages 700–716. Springer.	1906
1854			1907
1855			1908
1856			1909
1857	Kristina Toutanova and Danqi Chen. 2015. Observed versus latent features for knowledge base and text inference. In <i>CVSC</i> , pages 57–66. Association for Computational Linguistics.		1910
1858			1911
1859			1912
1860			
1861	Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, Aurélien Rodriguez, Armand Joulin, Edouard Grave, and Guillaume Lample. 2023. Llama: Open and efficient foundation language models. <i>CoRR</i> , abs/2302.13971.	Haoran Wang, Ying Zhang, Zhong Ji, Yanwei Pang, and Lin Ma. 2020a. Consensus-aware visual-semantic embedding for image-text matching. In <i>ECCV (24)</i> , volume 12369 of <i>Lecture Notes in Computer Science</i> , pages 18–34. Springer.	1913
1862			1914
1863			1915
1864			1916
1865			1917
1866			
1867			
1868	Quan Tu, Chuanqi Chen, Jinpeng Li, Yanran Li, Shuo Shang, Dongyan Zhao, Ran Wang, and Rui Yan. 2023. Characterchat: Learning towards conversational AI with personalized social support. <i>CoRR</i> , abs/2308.10278.	Huan Wang, Weiming Lu, and Zeyun Tang. 2019. Incorporating external knowledge to boost machine comprehension based question answering. In <i>ECIR (1)</i> , volume 11437 of <i>Lecture Notes in Computer Science</i> , pages 819–827. Springer.	1918
1869			1919
1870			1920
1871			1921
1872			1922
1873	Kohei Uehara and Tatsuya Harada. 2023. K-VQG: knowledge-aware visual question generation for common-sense acquisition. In <i>WACV</i> , pages 4390–4398. IEEE.	Jin Wang and Bo Jiang. 2021. Zero-shot learning via contrastive learning on dual knowledge graphs. In <i>ICCVW</i> , pages 885–892. IEEE.	1923
1874			1924
1875			1925
1876			
1877	Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, and Illia Polosukhin. 2017. Attention is all you need. In <i>NIPS</i> , pages 5998–6008.	Lei Wang, Jiabang He, Shenshen Li, Ning Liu, and Ee-Peng Lim. 2023a. Mitigating fine-grained hallucination by fine-tuning large vision-language models with caption rewrites. <i>CoRR</i> , abs/2312.01701.	1926
1878			1927
1879			1928
1880			1929
1881	Ramakrishna Vedantam, C. Lawrence Zitnick, and Devi Parikh. 2015. Cider: Consensus-based image description evaluation. In <i>CVPR</i> , pages 4566–4575. IEEE Computer Society.	Peng Wang, Dongyang Liu, Hui Li, and Qi Wu. 2020b. Give me something to eat: Referring expression comprehension with commonsense knowledge. In <i>ACM Multimedia</i> , pages 28–36. ACM.	1930
1882			1931
1883			1932
1884			1933
1885	Petar Velickovic, Guillem Cucurull, Arantxa Casanova, Adriana Romero, Pietro Liò, and Yoshua Bengio. 2018. Graph attention networks. In <i>ICLR (Poster)</i> . OpenReview.net.	Peng Wang, Qi Wu, Chunhua Shen, Anthony R. Dick, and Anton van den Hengel. 2017. Explicit knowledge-based reasoning for visual question answering. In <i>IJCAI</i> , pages 1290–1296. ijcai.org.	1934
1886			1935
1887			1936
1888			1937
		Peng Wang, Qi Wu, Chunhua Shen, Anthony R. Dick, and Anton van den Hengel. 2018a. FVQA: fact-based visual question answering. <i>IEEE Trans. Pattern Anal. Mach. Intell.</i> , 40(10):2413–2427.	1938
			1939
			1940
			1941

1942	Xiaodan Wang, Lei Li, Zhixu Li, Xuwu Wang, Xiangru	Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten	1995
1943	Zhu, Chengyu Wang, Jun Huang, and Yanghua Xiao.	Bosma, Brian Ichter, Fei Xia, Ed H. Chi, Quoc V. Le,	1996
1944	2023b. AGREE: aligning cross-modal entities for	and Denny Zhou. 2022. Chain-of-thought prompt-	1997
1945	image-text retrieval upon vision-language pre-trained	ing elicits reasoning in large language models. In	1998
1946	models. In <i>WSDM</i> , pages 456–464. ACM.	<i>NeurIPS</i> .	1999
1947	Xiaolong Wang, Yufei Ye, and Abhinav Gupta. 2018b.	Jialin Wu, Jiasen Lu, Ashish Sabharwal, and Roozbeh	2000
1948	Zero-shot recognition via semantic embeddings and	Mottaghi. 2022. Multi-modal answer validation for	2001
1949	knowledge graphs. In <i>CVPR</i> , pages 6857–6866.	knowledge-based VQA. In <i>AAAI</i> , pages 2712–2721.	2002
1950	Computer Vision Foundation / IEEE Computer Soci-	AAAI Press.	2003
1951	ety.	Jialin Wu and Raymond J. Mooney. 2022. Entity-	2004
1952	Xiaozhi Wang, Tianyu Gao, Zhaocheng Zhu, Zhengyan	focused dense passage retrieval for outside-	2005
1953	Zhang, Zhiyuan Liu, Juanzi Li, and Jian Tang. 2021.	knowledge visual question answering. In <i>EMNLP</i> ,	2006
1954	KEPLER: A unified model for knowledge embed-	pages 8061–8072. Association for Computational	2007
1955	ding and pre-trained language representation. <i>Trans.</i>	Linguistics.	2008
1956	<i>Assoc. Comput. Linguistics</i> , 9:176–194.	Likang Wu, Zhi Li, Hongke Zhao, Zhefeng Wang,	2009
1957	Xin Wang, Benyuan Meng, Hong Chen, Yuan Meng,	Qi Liu, Baoxing Huai, Nicholas Jing Yuan, and	2010
1958	Ke Lv, and Wenwu Zhu. 2023c. TIVA-KG: A multi-	Enhong Chen. 2023a. Recognizing unseen objects	2011
1959	modal knowledge graph with text, image, video and	via multimodal intensive knowledge graph propa-	2012
1960	audio. In <i>ACM Multimedia</i> , pages 2391–2399. ACM.	gation. In <i>Proceedings of the 29th ACM SIGKDD</i>	2013
1961	Yanan Wang, Michihiro Yasunaga, Hongyu Ren, Shinya	<i>Conference on Knowledge Discovery and Data Min-</i>	2014
1962	Wada, and Jure Leskovec. 2022b. VQA-GNN: rea-	<i>ing, KDD 2023, Long Beach, CA, USA, August 6-10,</i>	2015
1963	soning with multimodal semantic graph for visual	2023, pages 2618–2628. ACM.	2016
1964	question answering. <i>CoRR</i> , abs/2205.11501.	Penghao Wu and Saining Xie. 2023. V*: Guided vi-	2017
1965	Yikai Wang, Wenbing Huang, Fuchun Sun, Tingyang	sual search as a core mechanism in multimodal llms.	2018
1966	Xu, Yu Rong, and Junzhou Huang. 2020c. Deep mul-	<i>CoRR</i> , abs/2312.14135.	2019
1967	timodal fusion by channel exchanging. In <i>NeurIPS</i> .	Qi Wu, Peng Wang, Chunhua Shen, Anthony R. Dick,	2020
1968	Yizhong Wang, Yeganeh Kordi, Swaroop Mishra, Alisa	and Anton van den Hengel. 2016. Ask me anything:	2021
1969	Liu, Noah A. Smith, Daniel Khashabi, and Hannaneh	Free-form visual question answering based on knowl-	2022
1970	Hajishirzi. 2023d. Self-instruct: Aligning language	edge from external sources. In <i>CVPR</i> , pages 4622–	2023
1971	models with self-generated instructions. In <i>ACL (1)</i> ,	4630. IEEE Computer Society.	2024
1972	pages 13484–13508. Association for Computational	Sen Wu, Guoshuai Zhao, and Xueming Qian. 2023b.	2025
1973	Linguistics.	Resolving zero-shot and fact-based visual question	2026
1974	Youze Wang, Shengsheng Qian, Jun Hu, Quan Fang,	answering via enhanced fact retrieval. <i>IEEE Trans-</i>	2027
1975	and Changsheng Xu. 2020d. Fake news detection via	<i>actions on Multimedia</i> .	2028
1976	knowledge-driven multimodal graph convolutional	Wentao Wu, Hongsong Li, Haixun Wang, and	2029
1977	networks. In <i>ICMR</i> , pages 540–547. ACM.	Kenny Qili Zhu. 2012. Probbase: a probabilistic taxon-	2030
1978	Zeqing Wang, Wentao Wan, Runmeng Chen, Qiqing	omy for text understanding. In <i>SIGMOD Conference</i> ,	2031
1979	Lao, Minjie Lang, and Keze Wang. 2023e. To-	pages 481–492. ACM.	2032
1980	wards top-down reasoning: An explainable multi-	Yinan Wu, Xiaowei Wu, Junwen Li, Yue Zhang, Haofen	2033
1981	agent approach for visual question answering. <i>CoRR</i> ,	Wang, Wen Du, Zhidong He, Jingping Liu, and Tong	2034
1982	abs/2311.17331.	Ruan. 2023c. Mmpedia: A large-scale multi-modal	2035
1983	Zhiruo Wang, Jun Araki, Zhengbao Jiang, Md. Rizwan	knowledge graph. In <i>ISWC</i> , pages 18–37. Springer.	2036
1984	Parvez, and Graham Neubig. 2023f. Learning to filter	Zhibiao Wu and Martha Palmer. 1994. Verb seman-	2037
1985	context for retrieval-augmented generation. <i>CoRR</i> ,	tics and lexical selection. <i>arXiv preprint cmp-</i>	2038
1986	abs/2311.08377.	<i>lg/9406033</i> .	2039
1987	Zihao Wang, Junli Wang, and Changjun Jiang. 2022c.	Alexandros Xenos, Themis Stafylakis, Ioannis Patras,	2040
1988	Unified multimodal model with unlikelihood training	and Georgios Tzimiropoulos. 2023. A simple base-	2041
1989	for visual dialog. In <i>ACM Multimedia</i> , pages 4625–	line for knowledge-based visual question answering.	2042
1990	4634. ACM.	<i>CoRR</i> , abs/2310.13570.	2043
1991	Ziyue Wang, Chi Chen, Peng Li, and Yang Liu. 2023g.	Yongqin Xian, Christoph H. Lampert, Bernt Schiele,	2044
1992	Filling the image information gap for VQA: prompt-	and Zeynep Akata. 2019. Zero-shot learning - A	2045
1993	ing large language models to proactively ask ques-	comprehensive evaluation of the good, the bad and	2046
1994	tions. <i>CoRR</i> , abs/2311.11598.	the ugly. <i>IEEE Trans. Pattern Anal. Mach. Intell.</i> ,	2047
		41(9):2251–2265.	2048

2049	Yongqin Xian, Tobias Lorenz, Bernt Schiele, and Zeynep Akata. 2018. Feature generating networks for zero-shot learning. In <i>CVPR</i> , pages 5542–5551. Computer Vision Foundation / IEEE Computer Society.	model for visual storytelling. In <i>IJCAI</i> , pages 5356–5362. ijcai.org.	2104
2050			2105
2051		Sibei Yang, Guanbin Li, and Yizhou Yu. 2019b. Cross-modal relationship inference for grounding referring expressions. In <i>CVPR</i> , pages 4145–4154. Computer Vision Foundation / IEEE.	2106
2052			2107
2053			2108
2054	Yang Xiao, Yi Cheng, Jinlan Fu, Jiashuo Wang, Wenjie Li, and Pengfei Liu. 2023. How far are we from believable AI agents? A framework for evaluating the believability of human behavior simulation. <i>CoRR</i> , abs/2312.17115.		2109
2055		Song Yang, Qiang Li, Wenhui Li, Min Liu, Xuanya Li, and Anan Liu. 2023a. External knowledge dynamic modeling for image-text retrieval. In <i>ACM Multimedia</i> , pages 5330–5338. ACM.	2110
2056			2111
2057			2112
2058			2113
2059	Jiayuan Xie, Wenhao Fang, Yi Cai, Qingbao Huang, and Qing Li. 2022. Knowledge-based visual question generation. <i>IEEE Trans. Circuits Syst. Video Technol.</i> , 32(11):7547–7558.	Xiao Yang, Xiaojun Wu, and Tianyang Xu. 2021. DRSGN: dual revised semantic graph structured network for image-text matching. In <i>CCIS</i> , pages 230–242. IEEE.	2114
2060			2115
2061			2116
2062			2117
2063	Yiran Xing, Zai Shi, Zhao Meng, Gerhard Lakemeyer, Yunpu Ma, and Roger Wattenhofer. 2021. KM-BART: knowledge enhanced multimodal BART for visual commonsense generation. In <i>ACL/IJCNLP (1)</i> , pages 525–535. Association for Computational Linguistics.	Zhengyuan Yang, Zhe Gan, Jianfeng Wang, Xiaowei Hu, Yumao Lu, Zicheng Liu, and Lijuan Wang. 2022. An empirical study of GPT-3 for few-shot knowledge-based VQA. In <i>AAAI</i> , pages 3081–3089. AAAI Press.	2118
2064			2119
2065			2120
2066			2121
2067			2122
2068			
2069	Can Xu, Qingfeng Sun, Kai Zheng, Xiubo Geng, Pu Zhao, Jiazhan Feng, Chongyang Tao, and Daxin Jiang. 2023a. Wizardlm: Empowering large language models to follow complex instructions. <i>CoRR</i> , abs/2304.12244.	Zichao Yang, Xiaodong He, Jianfeng Gao, Li Deng, and Alexander J. Smola. 2016. Stacked attention networks for image question answering. In <i>CVPR</i> , pages 21–29. IEEE Computer Society.	2123
2070			2124
2071			2125
2072			2126
2073			
2074	Chang Xu, Dacheng Tao, and Chao Xu. 2013. A survey on multi-view learning. <i>CoRR</i> , abs/1304.5634.	Zonglin Yang, Xinya Du, Erik Cambria, and Claire Cardie. 2023b. End-to-end case-based reasoning for commonsense knowledge base completion. In <i>EACL</i> , pages 3491–3504. Association for Computational Linguistics.	2127
2075			2128
2076	Chunpu Xu, Min Yang, Chengming Li, Ying Shen, Xiang Ao, and Ruifeng Xu. 2021. Imagine, reason and write: Visual storytelling with graph knowledge and relational reasoning. In <i>AAAI</i> , pages 3022–3029. AAAI Press.		2129
2077			2130
2078			2131
2079		Liang Yao, Chengsheng Mao, and Yuan Luo. 2019. KG-BERT: BERT for knowledge graph completion. <i>CoRR</i> , abs/1909.03193.	2132
2080			2133
2081	Da Xu, Chuanwei Ruan, Evren Körpeoglu, Sushant Kumar, and Kannan Achan. 2020. Product knowledge graph embedding for e-commerce. In <i>WSDM</i> , pages 672–680. ACM.		2134
2082		Shuquan Ye, Yujia Xie, Dongdong Chen, Yichong Xu, Lu Yuan, Chenguang Zhu, and Jing Liao. 2023. Improving commonsense in vision-language models via knowledge graph riddles. In <i>CVPR</i> , pages 2634–2645. IEEE.	2135
2083			2136
2084			2137
2085	Fengli Xu, Jun Zhang, Chen Gao, Jie Feng, and Yong Li. 2023b. Urban generative intelligence (UGI): A foundational platform for agents in embodied city environment. <i>CoRR</i> , abs/2312.11813.		2138
2086			2139
2087		Chengxiang Yin, Zhengping Che, Kun Wu, Zhiyuan Xu, and Jian Tang. 2023a. Multi-clue reasoning with memory augmentation for knowledge-based visual question answering. <i>CoRR</i> , abs/2312.12723.	2140
2088			2141
2089	Silei Xu, Shicheng Liu, Theo Culhane, Elizaveta Pertseva, Meng-Hsi Wu, Sina J. Semnani, and Monica S. Lam. 2023c. Fine-tuned llms know more, hallucinate less with few-shot sequence-to-sequence semantic parsing over wikidata. In <i>EMNLP</i> , pages 5778–5791. Association for Computational Linguistics.		2142
2090			2143
2091		Zhenfei Yin, Jiong Wang, Jianjian Cao, Zhelun Shi, Dingning Liu, Mukai Li, Lu Sheng, Lei Bai, Xiaoshui Huang, Zhiyong Wang, Jing Shao, and Wanli Ouyang. 2023b. LAMM: language-assisted multimodal instruction-tuning dataset, framework, and benchmark. <i>CoRR</i> , abs/2306.06687.	2144
2092			2145
2093			2146
2094			2147
2095	Tao Xu, Pengchuan Zhang, Qiuyuan Huang, Han Zhang, Zhe Gan, Xiaolei Huang, and Xiaodong He. 2018. Attngan: Fine-grained text to image generation with attentional generative adversarial networks. In <i>CVPR</i> , pages 1316–1324. Computer Vision Foundation / IEEE Computer Society.		2148
2096			2149
2097		Gal Yona, Roei Aharoni, and Mor Geva. 2024. Narrowing the knowledge evaluation gap: Open-domain question answering with multi-granularity answers. <i>CoRR</i> , abs/2401.04695.	2150
2098			2151
2099			2152
2100			2153
2101	Pengcheng Yang, Fuli Luo, Peng Chen, Lei Li, Zhiyi Yin, Xiaodong He, and Xu Sun. 2019a. Knowledgeable storyteller: A commonsense-driven generative	Minji Yoon, Jing Yu Koh, Bryan Hooi, and Ruslan Salakhutdinov. 2023. Multimodal graph learning for generative tasks. <i>CoRR</i> , abs/2310.07478.	2154
2102			2155
2103			2156

2262	Yuchen Zhang, Xing Su, Jia Wu, Jian Yang, Hao Fan, and Xiaochuan Zheng. 2023d. Emoknow: Emotion- and knowledge-oriented model for COVID-19 fake news detection. In <i>ADMA (1)</i> , volume 14176 of <i>Lecture Notes in Computer Science</i> , pages 352–367. Springer.	2317
2263		2318
2264		2319
2265		2320
2266		
2267		
2268	Yuhong Zhang, Haitao Shu, Chenyang Bu, and Xuegang Hu. 2022d. A zero-shot learning method with a multi-modal knowledge graph. In <i>ICTAI</i> , pages 391–395. IEEE.	
2269		
2270		
2271		
2272	Zefan Zhang, Yi Ji, and Chunping Liu. 2023e. Knowledge-aware causal inference network for visual dialog. In <i>ICMR</i> , pages 253–261. ACM.	
2273		
2274		
2275	Zhi Zhang, Helen Yannakoudakis, Xiantong Zhen, and Ekaterina Shutova. 2023f. Ck-transformer: Commonsense knowledge enhanced transformers for referring expression comprehension. In <i>EACL (Findings)</i> , pages 2541–2551. Association for Computational Linguistics.	
2276		
2277		
2278		
2279		
2280		
2281	Lei Zhao, Lianli Gao, Yuyu Guo, Jingkuan Song, and Heng Tao Shen. 2021a. Skanet: Structured knowledge-aware network for visual dialog. In <i>ICME</i> , pages 1–6. IEEE.	
2282		
2283		
2284		
2285	Lei Zhao, Junlin Li, Lianli Gao, Yunbo Rao, Jingkuan Song, and Heng Tao Shen. 2023. Heterogeneous knowledge network for visual dialog. <i>IEEE Trans. Circuits Syst. Video Technol.</i> , 33(2):861–871.	
2286		
2287		
2288		
2289	Wentian Zhao, Yao Hu, Heda Wang, Xinxiao Wu, and Jiebo Luo. 2021b. Boosting entity-aware image captioning with multi-modal knowledge graph. <i>CoRR</i> , abs/2107.11970.	
2290		
2291		
2292		
2293	Yu Zhao, Xiangrui Cai, Yike Wu, Haiwei Zhang, Ying Zhang, Guoqing Zhao, and Ning Jiang. 2022. Mose: Modality split and ensemble for multimodal knowledge graph completion. In <i>EMNLP</i> , pages 10527–10536. Association for Computational Linguistics.	
2294		
2295		
2296		
2297		
2298	Wenbo Zheng, Lan Yan, Chao Gou, and Fei-Yue Wang. 2021a. Knowledge is power: Hierarchical-knowledge embedded meta-learning for visual reasoning in artistic domains. In <i>KDD</i> , pages 2360–2368. ACM.	
2299		
2300		
2301		
2302		
2303	Wenfeng Zheng, Lirong Yin, Xiaobing Chen, Zhiyang Ma, Shan Liu, and Bo Yang. 2021b. Knowledge base graph embedding module design for visual question answering model. <i>Pattern Recognit.</i> , 120:108153.	
2304		
2305		
2306		
2307	Ziqiang Zheng, Yiwei Chen, Jipeng Zhang, Tuan-Anh Vu, Huimin Zeng, Yue Him Wong Tim, and Sai-Kit Yeung. 2024. Exploring boundary of GPT-4V on marine analysis: A preliminary case study. <i>CoRR</i> , abs/2401.02147.	
2308		
2309		
2310		
2311		
2312	Yucheng Zhou, Pengfei Cao, Yubo Chen, Kang Liu, and Jun Zhao. 2023. Prompting vision language model with knowledge from large language model for knowledge-based VQA. <i>CoRR</i> , abs/2308.15851.	2321
2313		2322
2314		2323
2315		
2316		
	Yimin Zhou, Yiwei Sun, and Vasant G. Honavar. 2019. Improving image captioning by leveraging knowledge graphs. In <i>WACV</i> , pages 283–293. IEEE.	
	Deyao Zhu, Jun Chen, Xiaoqian Shen, Xiang Li, and Mohamed Elhoseiny. 2023. Minigpt-4: Enhancing vision-language understanding with advanced large language models. <i>CoRR</i> , abs/2304.10592.	2324
		2325
		2326
		2327
	Xiangru Zhu, Zhixu Li, Xiaodan Wang, Xueyao Jiang, Penglei Sun, Xuwu Wang, Yanghua Xiao, and Nicholas Jing Yuan. 2022. Multi-modal knowledge graph construction and application: A survey. <i>CoRR</i> , abs/2202.05786.	2328
		2329
		2330
		2331
		2332
	Yizhe Zhu, Mohamed Elhoseiny, Bingchen Liu, Xi Peng, and Ahmed Elgammal. 2018. A generative adversarial approach for zero-shot learning from noisy texts. In <i>CVPR</i> , pages 1004–1013. Computer Vision Foundation / IEEE Computer Society.	2333
		2334
		2335
		2336
		2337
	Yonghua Zhu, Ning Ge, Jieyu Huang, Yunwen Zhu, Binghui Zheng, and Wenjun Zhang. 2021. Enriching attributes from knowledge graph for fine-grained text-to-image synthesis. In <i>CSAE</i> , pages 80:1–80:6. ACM.	2338
		2339
		2340
		2341
		2342
	Yuke Zhu, Oliver Groth, Michael S. Bernstein, and Li Fei-Fei. 2016. Visual7w: Grounded question answering in images. In <i>CVPR</i> , pages 4995–5004. IEEE Computer Society.	2343
		2344
		2345
		2346
	Zihao Zhu, Jing Yu, Yujing Wang, Yajing Sun, Yue Hu, and Qi Wu. 2020. Mucko: Multi-layer cross-modal knowledge reasoning for fact-based visual question answering. In <i>IJCAI</i> , pages 1097–1103. ijcai.org.	2347
		2348
		2349
		2350
	Maryam Ziaeeefard and Freddy Lécué. 2020. Towards knowledge-augmented visual question answering. In <i>COLING</i> , pages 1863–1873. International Committee on Computational Linguistics.	2351
		2352
		2353
		2354

A Appendix

A.1 Literature Collection Methodology

For our paper, we source literature primarily from Google Scholar and arXiv. Google Scholar provides broad access to leading computer science conferences and journals, while arXiv serves as a key platform for preprints across various disciplines, including a significant repository recognized by the computer science community. We employ a systematic search strategy on these platforms, using relevant keyword combinations to assemble our references. We rigorously curate this collection, manually filtering out irrelevant papers and incorporating initially overlooked studies mentioned in their main texts. By exploiting Google Scholar’s citation tracking, we thoroughly augment our list through iterative depth and breadth traversal.

Organization. We begin by introducing the preliminaries, defining key concepts in KGs and multi-modal learning, and providing an overview of KG4MML settings (§ 2). Then we delve into various KG4MM tasks, detailing resources and key breakthroughs in recent years (§ 3), while balancing details to address content overlaps across tasks and focusing on core challenges (Fig. 2). Finally, we explore the future integration of multi-modal methods with (MM)KGs, proposing potential enhancements for previously discussed tasks (§ 4), especially considering the rapid development of multi-modal Large Language Models (LLMs).

Related work. Several surveys are closely related to our work. Monka et al. (2022) overview Knowledge Graph Embedding (KGE) methods and their integration with high-dimensional visual embeddings, emphasizing KGs’ role in visual information transfer. Lymperaïou and Stamou (2022) discuss enhancing multi-modal learning with knowledge but lack a systematic analysis of KG-driven multi-modal learning, including benchmarks, method comparisons, and task paradigms. Moreover, these studies all focus on developments up to 2022, missing the latest insights.

In response to the rapid advancements in AGI from 2022 to 2023, our survey places a strong emphasis on emerging areas like LLMs, aiming to fill critical knowledge gaps. Our goal is to provide a clear roadmap for future research, highlight challenges and opportunities, and systematically compare methodologies to inspire new ideas.

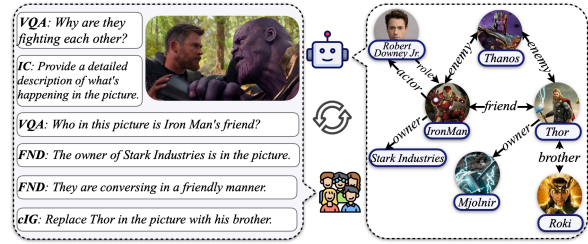


Figure 8: Applications of KGs in downstream multi-modal tasks within the context of the *Marvel Universe*.

A.2 Supplement for Preliminaries

Aiming to align with established literature, we begin with a widely-accepted definition of KG and its foundational operations, explore KGs enriched with ontologies from the semantic web perspective, and conclude with diverse interpretations and uses of KGs beyond the semantic web.

A.2.1 Knowledge Graph

Since their inception around 2007, Knowledge Graphs (KGs) have become pivotal in various academic domains, marked by foundational projects such as Yago (Suchanek et al., 2007), DBPedia (Auer et al., 2007), and Freebase (Bollacker et al., 2008). The integration of Google’s Knowledge Panels into web search in 2012 highlighted a significant milestone in the adoption of KGs. Today, KGs enhance search engines like Google and Bing and are integral to the functionality of voice assistants like Amazon Alexa and Apple Siri, reflecting their widespread business importance and increasing prevalence.

Structural Composition. KGs represent entities and relations using a graph structure, where nodes symbolize real-world entities or atomic values (attributes), and edges denote relations. Knowledge is often captured in triples, such as (*Hangzhou*, *locatedAt*, *China*). They utilize an ontology-based schema to define basic entity classes and their relations, usually in a taxonomic structure. This semi-structured nature merges structured data’s clear semantics (from ontologies) with the flexibility of unstructured data, allowing easy expansion through new classes and relations.

Accessibility and Advantages. KGs support a wide array of downstream applications, accessible primarily via *Lookup* and *Querying* methods. *Lookup* in KGs, also known as KG retrieval, identifies relevant entities or properties based on input strings, leveraging lexical indices (surface) from entity and relation labels. An example of this is

the DBpedia online lookup service ⁴. Alternatively, *Querying* returns results from input queries crafted in the RDF query language SPARQL⁵. These queries typically involve sub-graph patterns with variables, yielding matched entities, properties, literals, or complete sub-graphs.

Note that KGs, especially those with OWL ontologies, support symbolic reasoning, including consistency checks to identify logical conflicts and entailment reasoning to infer hidden knowledge via Description Logics. KGs also facilitate inter-domain connections. An example is the linkage between the *Movie* and *Music* domains through common entities like individuals who are both *actors* and *singers*. This interconnectivity not only enhances machine comprehension but also improves human understanding, benefiting applications like search, question answering, and recommendations. Furthermore, recent developments in LLMs highlight the crucial role of KGs, particularly in managing long-tailed knowledge, as evident in several studies (Dong, 2023; Sun et al., 2023b; Pan et al., 2023a,b).

Ontology. Within the semantic web, ontologies serve as KG schemas, utilizing languages like RDFS⁶ and OWL⁷ to ensure richer semantics and superior quality (Horrocks, 2008). Key features of ontologies include: (i) Hierarchical classes, often termed as concepts⁸; (ii) Properties that specify the terms used in relations; (iii) Hierarchies involving both concepts and relations; (iv) Constraints, including the domain and range of relations, as well as class disjointness; (v) Logical expressions that encompass relation composition.

Languages like RDF, RDFS, and OWL introduce built-in vocabularies to capture these knowledge elements, with predicates like *rdfs:subClassOf* denoting concept subsumption, and *rdf:type* indicating instance-concept associations. RDFS also provides annotation properties like *rdfs:label* and *rdfs:comment* for resource meta-information.

KG Scope Extension. Widely accepted KGs include WordNet (Miller, 1995), a lexical database defining word interrelations, and Concept-

Net (Speer et al., 2017), which archives common-sense knowledge interlinked by different terms.

In this paper, we extend the conventional view of KGs beyond standard-format entities and relations. This paper extends the conventional view of KGs beyond standard-format entities and relations. Besides, the ontology alone, often utilized to define domain knowledge including conceptualization and vocabularies like terms and taxonomies, is also considered a form of KG. Further elaborating on this expanded perspective, as outlined by Chen et al. (2023c), our scope includes simpler graph structures, such as basic taxonomies with hierarchical classes and graphs with weighted edges denoting quantitative relationships like similarity and distance between entities. Additionally, we categorize any structured data organized in a graph format with nodes that have explicit semantic interpretations as part of this broader KG definition. A prominent example is the Semantic Network, which connects various concepts with labeled edges to represent different relationships.

We also consider ontologies, basic taxonomies, and graphs with weighted edges as forms of KGs. Any structured data organized in a graph format with nodes having explicit semantic interpretations falls under this broader KG definition. An example is the Semantic Network, which connects various concepts with labeled edges to represent different relationships.

A.2.2 Multi-modal Learning.

Our world is perceived through diverse modalities, including sight, sound, movement, touch, and smell (Smith and Gasser, 2005). A “modality” typically refers to a specific type of data or information channel, characterized by sensory input or representation format. Each modality encapsulates unique features from specific sensory sources. It is intuitive that models, which integrate data from various modalities, generally surpass uni-modal models by accumulating more information. Multi-modal learning aims to develop a unified representation or mapping from multiple modalities to an output space, leveraging the complementarity and redundancy across modalities to improve prediction. The challenge lies in effectively aligning, fusing, and integrating information from various modalities to exploit their collective power.

Difference to Multi-view Learning. Unlike multi-view analysis, which suggests that each view

⁴<https://lookup.dbpedia.org/>

⁵<https://www.w3.org/TR/rdf-sparql-query/>

⁶RDF Schema, <https://www.w3.org/TR/rdf-schema/>

⁷Web Ontology Language, <https://www.w3.org/TR/owl2-overview/>

⁸To distinguish between *class* in machine learning tasks and *class* in KGs, we refer to the latter as *concept*.

(e.g., different perspectives of a flower) can independently yield accurate predictions (Xu et al., 2013; Federici et al., 2020), multi-modal learning contends with scenarios where the absence of one modality could impede task completion (Huang et al., 2021) (e.g., an image-lacking Visual Question Answering scenario). Additionally, multi-view learning typically involves varying perspectives of the same data type, originating from a single source, such as different features of image data. In contrast, multi-modal learning deals with disparate data types, like text and images, derived from multiple sources. In this paper, our exploration of multi-modal tasks and the application of multi-modal learning on KGs are grounded in this broader understanding of multi-modal learning.

Definition 1 Multi-modal Learning. Assume given data $\hat{x} = (x^{(1)}, \dots, x^{(K)})$ consists of K modalities, with $x^{(k)} \in \mathcal{X}^{(k)}$ representing the domain set of the k -th modality and $\mathcal{X} = \mathcal{X}^{(1)} \times \dots \times \mathcal{X}^{(K)}$. Let $\mathcal{Y}$ denote the target domain and $\mathcal{Z}$ represent a latent space. Denote $g : \mathcal{X} \mapsto \mathcal{Z}$ as the true mapping from the input space (utilizing all K modalities) to the latent space, and $q : \mathcal{Z} \mapsto \mathcal{Y}$ as the true task mapping. For example, in aggregation-based multi-modal fusion, g serves as an aggregation function built upon K separate sub-networks, and q is a multi-layer neural network (Wang et al., 2020c). In a learning task, a data pair $(\hat{x}, y) \in \mathcal{X} \times \mathcal{Y}$ is generated from an unknown distribution $\mathcal{D}$, such that

$$\mathbb{P}_{\mathcal{D}}(\hat{x}, y) = \mathbb{P}_{y|\hat{x}}(y | q \circ g(\hat{x})) \mathbb{P}_{\hat{x}}(\hat{x}), \quad (1)$$

where $q \circ g(\hat{x}) = q(g(\hat{x}))$ represents the composite function of q and g .

A.3 Supplement for Understanding & Reasoning Tasks

Multi-modal reasoning tasks, like knowledge-based VQA, demand knowledge that goes beyond regular daily experiences (Khan et al., 2022a). These tasks often delve into less common, long-tail knowledge domains that typically require intentional reflection, with KGs providing a crucial structured repository for this extensive, specialized knowledge.

A.3.1 Supplementary Information for VQA

Current KG-aware VQA research typically involves four key stages for incorporating knowledge. These stages, integral to the workflow of

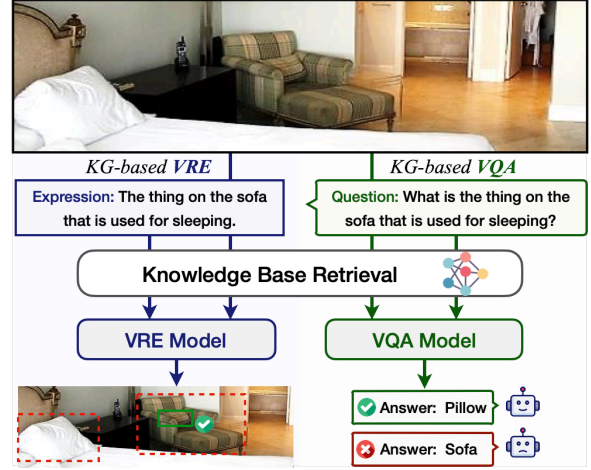


Figure 9: Illustration of KG-based Visual Question Answering (VQA) (§ 3.1) and Visual Referring Expressions (VRE) (§ A.6). To some extent, KG-based VRE can be viewed as an extension of KG-based VQA, incorporating an additional step of grounding answers.

KG-aware understanding and reasoning tasks, may be adopted individually or in combination across different studies to form a comprehensive approach. Those method references are provided and organized here for tracing the original sources:

Knowledge Retrieval. Implicit knowledge encoded in model parameters, typically pre-trained on large-scale datasets via self-supervised tasks, makes the use of a retriever **optional but still beneficial** for VQA.

1. Matching-based Retrieval:

- Identifying spatial positions (Shah et al., 2019; Zhu et al., 2020; Yu et al., 2020; Gardères et al., 2020; Marino et al., 2021; Lin et al., 2022).
- Identifying visual object sizes and names (Wang et al., 2017, 2018a; Narasimhan and Schwing, 2018; Wang et al., 2019; Narasimhan et al., 2018; Zhu et al., 2020; Ziaefard and Lécué, 2020; Yu et al., 2020; Li et al., 2020; Ramnath and and, 2020; Zhang et al., 2021a; Gardères et al., 2020; Zheng et al., 2021b; Vickers et al., 2021; Li and Moens, 2022; Zhang et al., 2022a; Ding et al., 2022; Hussain et al., 2022; Han et al., 2023; Song et al., 2023b; Ravi et al., 2023; You et al., 2023; Khademi et al., 2023; Yin et al., 2023a; Dong et al., 2024).
- Identifying high-level attributes like scene names, object parts, and human activities, using various pre-trained classifiers or

APIs⁹ (Wang et al., 2017; Wu et al., 2016; Narasimhan and Schwing, 2018; Wang et al., 2018a; Narasimhan et al., 2018; Yu et al., 2020; Ramnath and and, 2020; Marino et al., 2021; Hussain et al., 2022; Zhang et al., 2022a; Lin and Byrne, 2022; He and Wang, 2023; Song et al., 2023b; You et al., 2023; Khademi et al., 2023).

- Image captions and OCR text strings can be generated for information supplement (Wu et al., 2016; Su et al., 2018; Zhu et al., 2020; Yu et al., 2020; Salaberria et al., 2023; Chen et al., 2022c; Hussain et al., 2022; Gui et al., 2022; Lin and Byrne, 2022; Wu and Mooney, 2022; Lin et al., 2022; You et al., 2023; Zhou et al., 2023; Si et al., 2023; Khademi et al., 2023; Dong et al., 2024).
- Those question and captions can be parsed by NLP tools (e.g., NLTK (Bird et al., 2009), AllenNLP constituency parser (Gardner et al., 2018), Stanza (Qi et al., 2020), NLP Dependency Parser (Chen and Manning, 2014), Named Entity Recognizer (Finkel et al., 2005), LLMs (Dong et al., 2024)) for syntax analysis (Wang et al., 2019; Li et al., 2020; Saqur and Narasimhan, 2020; Cao et al., 2022b; Han et al., 2023; Wu and Mooney, 2022; Ravi et al., 2023), along with techniques like regular expressions (regex) (Wang et al., 2017), semantic graph parsing model (Zhu et al., 2020; Yu et al., 2020; Wu et al., 2022; He and Wang, 2023; Hussain et al., 2022), SpanSelector (Jain et al., 2021), or query template selector (Wang et al., 2018a; Narasimhan and Schwing, 2018; Narasimhan et al., 2018).
- Unimportant visual objects not present in the question or caption might be filtered out (Zhang et al., 2021a; Gardner et al., 2018).
- After extracting initial concepts from Q and I , two key mappings are established: the first links parsed objects in Q to their visual counterparts in I , and the second associates these concepts with relevant entries in KBs. This is achieved using methods like greedy longest-string matching (Wang et al., 2017; Su et al., 2018; Shevchenko

et al., 2021), template matching (Wang et al., 2018a), and Multi-modal Entity Linking methods (Zheng et al., 2021b; Jain et al., 2021; Wu and Mooney, 2022; You et al., 2023; Adjali et al., 2023). Techniques like face identification algorithms (Shah et al., 2019; Vickers et al., 2021; Heo et al., 2022; Lerner et al., 2022) and ViLBERT-multi-task (Lu et al., 2020) serve as effective tools for linking objects.

- Fact triples can be collected by involving the first-order sub-KG from these identified concept nodes (sometimes will be three-hops in character KG (Shah et al., 2019)) or by identifying brief knowledge paths among the entities from I and Q . This process requires constructing a temporary (local) sub-KG specific to the current Q - I pair (Wang et al., 2019; Su et al., 2018; Li et al., 2020). In addition, KG-Aug (Li et al., 2020) constructs a global sub-KG that links Q , I , and candidate answers in a unified knowledge-based semantic space. KAN (Zhang et al., 2021a) presents a weighting system for each fact to indicate the reliability of the corresponding knowledge piece. Heo et al. (2022) develop a hypergraph from the KG, using a triplet as the basic unit to preserve the higher-order semantics inherent in the KG.
- RDF query (e.g., SPARQL) generation often involves filling pre-defined templates with parsed question data, suitable for datasets with consistent question patterns (Wang et al., 2017). Those queries typically include both “ASK” and “SELECT” types, with “ASK” checking for a solution to the query pattern and “SELECT” returning variables from all matched solutions (Wang et al., 2017).
- Term-based (e.g., TF-IDF and BM25) retrievers is another good choice, with their scoring reflecting the direct correlation between the query and fact triplets. Luo et al. (2021) use image captions generated by a model, concatenating them with Q as a query for BM25-based document retrieval. LaKo (Chen et al., 2022c) presents a Stem-based BM25 approach, using word stems as the smallest semantic units to maximize knowledge extraction from limited

⁹<https://azure.microsoft.com/en-us/products/cognitive-services/vision-services>

VQA and KG resources. EnFoRe (Wu and Mooney, 2022) utilizes entity-augmented queries to recall passages via BM25, measuring an entity’s importance to the query by the relevance of these passages to the answer.

2. Retrieval Pruning:

- Re-ranking candidate facts and may include assigning weights to nodes based on corresponding visual object sizes (Wang et al., 2019), ensuring each knowledge triple contains key elements from Q or auto-generated captions (Su et al., 2018; Wu and Mooney, 2022), or aligns with the relation type implied in Q (Narasimhan et al., 2018; Zhu et al., 2020; Yu et al., 2020; Ramnath and and, 2020; Hussain et al., 2022; Yin et al., 2023a).
- A learnable score function can assess the compatibility between a fact representation and the Q - I representation (Narasimhan and Schwing, 2018; Hussain et al., 2022; Ravi et al., 2023).
- Global KG-level pruning is also practiced. For example, KRISP (Marino et al., 2021) gathers all symbolic entities from the VQA dataset, including questions, answers, and visual concepts recognized by visual systems, and incorporates only triples related to these concepts to the model training. LaKo (Chen et al., 2022c) streamlines the KG by creating a stem corpus specific to the VQA field, ensuring all KG triples contain at least one stem from this corpus. KAT (Gui et al., 2022) extracts a subset from Wikidata (Vrandecic and Krötzsch, 2014) covering common real-world objects, and RR-VEL (You et al., 2023) only retains triples in the KG that include candidate answers and visually detected entities in the training set images.

3. Search Engine:

- Marino et al. (2019) gather Wikipedia articles for each Q - I pair and select sentences closely matching the query based on key word frequency. Their ArticleNet predicts the presence and positioning of correct answers in these articles. Jain et al. (2021) utilize Google’s search engine to retrieve the top-10 relevant snippets for a Machine Reading Comprehension (MRC) module

based on a reformulated Q .

- MAVEx (Wu et al., 2022) enriches knowledge retrieval through Google APIs for category labels, OCR readings, and logo information, collecting sentences from Wikipedia articles that contain candidate answers. It also uses Google Image Search with candidate Q - A pairs to provide additional visual information. Luo et al. (2021) notice that snippet-level knowledge outperforms sentence-level, and select ten snippets for each Q - A query.

4. Dense Retrieval (Karpukhin et al., 2020):

- This technique utilizes embedding similarities to match questions and visual concepts with pre-flattened concise fact sentences (Narasimhan et al., 2018; Zhu et al., 2020; Yu et al., 2020; Ziaefard and Lécué, 2020; Wu et al., 2022; Li and Moens, 2022; Gao et al., 2022; Ossowski and Hu, 2023; Liu et al., 2022; You et al., 2023; Si et al., 2023).
- Retrieval efficiency is frequently enhanced by employing open-source indexing engines like FAISS (Johnson et al., 2021), which facilitates the organization and indexing of large-scale dense embeddings. The architectures involved are generally symmetrical or siamese to support shared embedding spaces, while asymmetrical designs are adopted for Cross-Modal Retrieval scenarios (e.g., CLIP-based retrieval).
- DMMGR (Li and Moens, 2022) ranks triplets based on the average cosine similarity between each word in the triplet and both the nouns in Q and the objects detected in I , excluding pairs with a zero average similarity.
- RR-VEL (You et al., 2023) assesses the similarity between Q and key entities across various knowledge triples, using combined similarity scores to rank the candidate triples. KAT (Gui et al., 2022) uses the CLIP model to encode patch-level image regions and knowledge entries for retrieval purposes.
- MAVEx (Wu et al., 2022) creates a concept pool for each Q - A pair, selecting facts containing potential answers identified by other VQA models. These facts are en-

coded using a pre-trained BERT model for re-ranking. Its subsequent work EnFoRe (Wu and Mooney, 2022) also prioritizes key entities in the Q - I pair, enhancing the knowledge retrieval process by focusing on entities crucial for answering the question.

- HKEML (Zheng et al., 2021a) applies 2D convolutional operations (Gao et al., 2021) to align the head and relation patterns in knowledge queries with Q , effectively mining implicit connections within the KG pertinent to Q - A pairs.

5. Learnable Retriever:

- Learnable Retriever refers to a trainable retrieval model that enhances the adaptability and compatibility in KG-based VQA settings (Chen et al., 2021d; Luo et al., 2021; Lin and Byrne, 2022; Ravi et al., 2023; Wu et al., 2023b; Adjali et al., 2023).
- Chen et al. (2021d) and Li et al. (2022b) separately aligning the joint embedding of the Q - I pair with the targets like relations in separate feature spaces. The prediction for relation type in the sub-KG pruning operation mentioned before is similar.
- VLC-BERT (Ravi et al., 2023) assigns similarity scores to inference facts for each Q based on their overlap with human-annotated answers. These scores act as weak signals, indicating the relevance of each fact to Q , thus guiding the training of the retriever.
- Luo et al. (2021) utilize DPR (Karpukhin et al., 2020) as a neural retriever, leveraging two BERT models for encoding the query and context. They further adapt DPR for visual domains with two variants: Image-DPR based on LXMERT (Tan and Bansal, 2019), and Caption-DPR, which modifies the DPR approach to suit visual content.
- LaKo (Chen et al., 2022c) explores a differentiable KG retriever, leveraging cross-attention scores between the token of the prediction output and input facts for iterative reader-retriever training.
- Addressing the challenge of slow convergence and sub-optimal performance in learnable retrievers, DEDR (Salemi et al., 2023a) employs a dual multi-modal encoder architecture with shared parameters

for both Q - I queries and knowledge content, starting from the same shared embedding space. It further explores both multi-modal and text-only retrievers, combining their results via an ensemble method. The training of these retrievers utilizes a multi-modal retrieval dataset provided by Qu et al. (2021) as a supervised corpus. Training for these retrievers is based on a supervised multi-modal retrieval dataset from Qu et al. (2021).

- REVEAL (Hu et al., 2023) integrates three data sources: WikiData KB (Vrandečić and Krötzsch, 2014), Wikipedia-Image-Text (WIT) (Srinivasan et al., 2021), and the VQA2.0 dataset (Antol et al., 2015). It utilizes a gating mechanism for optimal knowledge source selection and employs the perceiver architecture (Jaegle et al., 2021) to encode and compress knowledge items, enabling cascading multi-modal retrievers and joint reasoning.
- RAVQA (Lin and Byrne, 2022) treats retrieval content as negative if it does not aid in answer generation, using the rest as positive samples. This approach helps in training the retriever by defining relevant and irrelevant content. Additionally, it combines the retrieval probability with the reader’s answer prediction to determine the final result.
- **Cold Start** issues: REVEAL (Hu et al., 2023) creates an initial retrieval dataset with pseudo ground-truth knowledge, using a large-scale image-caption dataset (Srinivasan et al., 2021). For pre-training, REVEAL pairs passages with query images as pseudo ground-truth knowledge and, to align with VQA task formats, randomly truncates captions to predict the truncated content using the image and the remaining text. LaKo (Chen et al., 2022c) initially employs a BM25-based retriever for knowledge retrieval in the first training phase, allowing for preliminary distillation to the differentiable retriever to mitigate the cold start problem.

6. PLM Generation as the Retrieval:

- Implicit knowledge encoded in model parameters, typically pre-trained on large-scale datasets via self-supervised tasks,

makes the use of an explicit retriever optional for VQA completion. Several studies directly and implicitly harness the knowledge embedded in PLMs for reasoning, often skipping a separate knowledge retrieval step (Salaberria et al., 2023; Yang et al., 2022; Zhang et al., 2022a; Shao et al., 2023; Subramanian et al., 2023).

- KAT (Gui et al., 2022) and TwO (Si et al., 2023) take GPT-3 to retrieve implicit textual knowledge with supporting evidence; VLC-BERT (Ravi et al., 2023) uses COMET (Hwang et al., 2021), a LM trained on commonsense KGs, to generate contextual expansions instead of direct knowledge retrieval from KBs. Wang et al. (2023g) utilize ChatGPT to decompose Q , alleviating the issue of unfocused and lacking detailed image features in image captioning; MM-Reasoner (Khademi et al., 2023) employs LLMs to create rationales from multi-aspect visual descriptions (e.g., commonsense knowledge facts and external information). These rationales, alongside I and Q , are processed by a specifically fine-tuned Visual Language Model (VLM) designed to handle such enriched input.

Knowledge Representation involves selecting the appropriate format for symbolic KGs to integrate with multi-modal models. This decision is crucial for effectively infusing knowledge into multi-modal reasoning tasks.

1. Direct Text-to-Embedding Mapping:

- Some research treats entities and relations in KGs as words, using embedding methods like Glove (Pennington et al., 2014) to translate them into continuous vectors. This transformation enables the further compression of knowledge components (e.g., triples) into fixed-size vectors using Recurrent Neural Networks (RNNs) (Wang et al., 2019), (V)PLMs (You et al., 2023; Hu et al., 2023; Chevalier et al., 2023), or mean pooling (Narasimhan and Schwing, 2018; Narasimhan et al., 2018; Zhu et al., 2020; Yu et al., 2020; Chen et al., 2021d; Marino et al., 2021; Li and Moens, 2022; Wu et al., 2022; Hussain et al., 2022).
- When handling lengthy text from SPARQL queries, Wu et al. (2016) use Doc2Vec (Le

and Mikolov, 2014) to learn feature representations for variable-length texts.

- Techniques such as stop-word removal in Word2Vec can further refine knowledge representation, reducing the noise from irrelevant words in mean pooling (Narasimhan et al., 2018; Liu et al., 2022; Chen et al., 2022c).
- Some methods convert fact collections into natural language sentences via concatenating the relation and object entities (Ziaee-fard and Lécué, 2020; Zhang et al., 2021a; You et al., 2023; Hu et al., 2023), allowing direct encoding into fixed-length vectors by PLMs.

2. Knowledge Graph Embedding (KGE):

- KGE offers a practical approach to embed facts and reveal semantic relationships among triples in an abstract space. This technology is valuable for setting up initial (Su et al., 2018; Ramnath and and, 2020; Zheng et al., 2021b; Han et al., 2023) fact embeddings and gathering multi-modal knowledge (Ding et al., 2022).
- Cao et al. (2022b) train the RotatE model on the entire KG to get entity and relation features, modifying a guided-attention block to fuse those knowledge embeddings with I and Q features.
- Chen et al. (2021d) evaluate various embeddings including TransE-based KG embeddings, BERT-based node representations of ConceptNet (Yang et al., 2023b; Malaviya et al., 2020), and GloVe embeddings, finding that Word2Vec representations excel at mapping answers in smaller datasets.
- RVL (Shevchenko et al., 2021) utilizes the PyTorchBigGraph method (Lerer et al., 2019) for embedding the Wikidata KG, while KVQAmeta (García-Olano et al., 2022) employs Wikipedia2Vec for representing entities from Wikipedia, emphasizing KGE’s versatility in representing different knowledge sources.

3. Pure Context:

- In many cases, KG triples are maintained in their original textual format for direct participation in multi-modal reasoning. They serialize triples for joint reasoning with (V)PLMs (Vickers et al., 2021; Chen et al., 2022c; Gao et al., 2022; Yang et al., 2022;

Gui et al., 2022; Lin and Byrne, 2022; Wu and Mooney, 2022; Lin et al., 2022; Si et al., 2023; Ravi et al., 2023; You et al., 2023; Shao et al., 2023; Zhou et al., 2023; Wang et al., 2023g; Hu et al., 2022; Xenos et al., 2023; Khademi et al., 2023; Dong et al., 2024).

Knowledge-aware Modality Interaction.

1. Concatenation:

- This unified feature is typically refined with a Multi-Layer Perceptron (MLP) to enhance modality interaction and integration. In multi-modal fusion models like MUTAN (Ben-Younes et al., 2017), BAN (Kim et al., 2018), SAN (Yang et al., 2016) and ERMLP (Ramnath and and, 2020), feature concatenation is a preliminary step before being input to the MLP layer, crucial for sophisticated multi-modal analysis.

2. Long Short-Term Memory (LSTM) Network:

- LSTM Network is a foundational framework for integrating knowledge with multi-modal data.
- Sometimes LSTMs also act as standalone encoders for textual data (Narasimhan and Schwing, 2018; Narasimhan et al., 2018; Yu et al., 2020; Zhu et al., 2020; Li et al., 2020; Ramnath and and, 2020; Zhang et al., 2021a; Li and Moens, 2022), employing Glove (Pennington et al., 2014; Han et al., 2023) or PLMs (Devlin et al., 2019; Lan et al., 2020) for token embedding initialization. The output embeddings aid in subsequent stages of modality fusion, giving LSTM a pivotal role similar to those methods in *Direct Text-to-Embedding Mapping* paradigm.

3. Graph Neural Networks (GNNs):

- GNNs emphasize the connection of concepts in VQA by integrating representations from I , Q , and entities into cohesive networks, where each node (entity) is represented by an embedding that is a concatenation of different modalities (Narasimhan et al., 2018).
- Mucko (Zhu et al., 2020) diverges from traditional modality embedding concatenation by independently processing distinct modalities' KGs. This involves isolating and separately analyzing the visual scene

KG, the semantic KG from image captions, and the common sense KG, supporting precise answer determination through Q -guided attention and cross-KG convolution. The method of Q -guided KG node weighting has seen similar implementations in other studies (Yu et al., 2020; Li et al., 2020; Saqur and Narasimhan, 2020; Ziaee-fard and Lécué, 2020; Li and Moens, 2022; Liu et al., 2022; Hussain et al., 2022; Wang et al., 2022b).

- KG-Aug (Li et al., 2020) uses GCN to generate entity representations, which are then used to embed knowledge into the features of both Q and I .
- KRISP (Marino et al., 2021) applies a RGCN (Schlichtkrull et al., 2018) for symbolic knowledge reasoning, enhancing each entity with four inputs: *a*) A binary indicator for concept presence in Q ; *b*) Classifier probabilities for the concept's node, or zero if not detected in I , using various classifiers and detectors; *c*) A GloVe pooling representation of the concept; *d*) An implicit knowledge representation derived from a multi-modal pre-trained model (Li et al., 2019).
- VQA-GNN (Wang et al., 2022b) employs a multi-modal GNN with bidirectional fusion to update concept and scene graph nodes for answer prediction through inter-modal message passing.

4. Dynamic Memory Networks (DMNs):

- DMNs (Kumar et al., 2016) utilize an attention-based mechanism for filtering critical information from localized small-scale knowledge triple embeddings (Fig. 4 (d)), achieved by modeling interactions across multiple data channels (Wang et al., 2019; Shah et al., 2019; Han et al., 2023; Yin et al., 2023a).
- Through *triple replication*, VKMN (Su et al., 2018) deconstructs each knowledge triple into three Key-Value pairs, for instance, (h, r) as the key and t as the value, reducing interference caused by using only head and tail entities as keys for retrieval thereby improving reasoning performance.
- DMMGR (Li and Moens, 2022) follows this setting and further refines knowledge triple composition by using the average em-

bedding of a triple as a key and its individual elements as values for enhanced relevance assessment. These networks apply a multi-scale attention mechanism that initially evaluates the overall relevance of a triplet’s embedding, then assess the importance of each element, leading to more accurately recalled dynamic memories.

- GRUC (Yu et al., 2020) uses visual and semantic scene graphs as knowledge sources for external memory, iteratively updating multi-modal memories and employing a GRU module to refresh factual entity representations, incorporating inputs from previous entities and memory from the last time step.
- SUPER (Han et al., 2023) integrates a memory augmented component to retain and adjust key clues for answering questions, a method named *memory reactivation*. REVEAL (Hu et al., 2023) unifies multi-modal data by compressing each entry into a set number of value embeddings and a single key embedding for memory storage, achieving synchronous and stable updates between the memory encoder and main framework by re-encoding a portion (10%) of the retrieved knowledge items in each training iteration.

5. **Guided-Attention & Transformer:**

- Many studies (Ramnath and and, 2020; Gardères et al., 2020; Zhang et al., 2021a; Cao et al., 2022b; Wu et al., 2022; Heo et al., 2022) have adopted a guided-attention mechanism to merge knowledge embeddings with visual and textual features.

6. **PLM & VLM Reasoning:**

- **Embedding-Based Visual Information Integration:** This category includes methods that convert visual data into embeddings compatible with the input specifications of (V)PLMs (Dou et al., 2022). It involves techniques that restructure visual inputs into embeddings which seamlessly integrate with the model’s existing architecture, such as compressing patch or local object features into fixed-length embedding sets (Jaegle et al., 2021; Hu et al., 2023) or applying adapters or projection heads for cross-modal feature space align-

ment (Lin et al., 2022; Yin et al., 2023b). These visual embeddings, combined with textual inputs, are processed in the embedding layers of (V)PLMs (Jaegle et al., 2021; Ossowski and Hu, 2023) as shown in Fig. 4 (f). Some studies (Vickers et al., 2021; Luo et al., 2021; Guo et al., 2022b; Ravi et al., 2023; Salemi et al., 2023a) integrate retrieved knowledge content and questions with image regions of interest, subsequently fine-tuning VLMs end-to-end on the VQA dataset using ground truth answers for optimization. RVL (Shevchenko et al., 2021) and KVQAmeta (García-Olano et al., 2022) inject the knowledge into the VLMs via aligning the KG embedding with the corresponding textual phrase representations derived from the output summations of PLM’s embedding layers. MuKEA (Ding et al., 2022) uses the visual and language output sides of the LXMERT as the head and relation of a triple, respectively, pairing these with the ground truth answer as the tail entity. This association, aroused through the KGE method (e.g., TransE), leverages implicit knowledge within VLMs for reasoning. VLC-BERT (Ravi et al., 2023) uses a Q -guided multi-head attention block to fuse multiple knowledge representation vectors before feeding them into the VLM. He and Wang (2023) propose a graph-involved Q -attention mechanism, where V - Q guided graphs are built to direct VLM training by integrating a graph-aware mask matrix into the Transformers’ attention matrix. Pang et al. (2023) enhance a VLM’s ability for parametric knowledge injection by integrating the frozen Transformer layer of the LLM (LLaMA (Touvron et al., 2023)) between its cross-modal fusion and decoder modules.

- **Textual Conversion of Visual Data:** This category involves converting all visual information into a textual format, like captions shown in Fig. 4 (f), enabling the application of PLM reasoning to a uniform textual dataset that includes background knowledge, questions, and images (Salaberria et al., 2023; Chen et al., 2022c; Luo et al., 2021; Zhang et al., 2022a; Yang et al.,

3228
3229
3230
3231
3232
3233
3234
3235
3236
3237
3238
3239
3240
3241
3242
3243
3244
3245
3246
3247
3248
3249
3250
3251
3252
3253
3254
3255
3256
3257
3258
3259
3260
3261
3262
3263
3264
3265
3266
3267
3268
3269
3270
3271
3272
3273
3274
3275
3276
3277
3278

2022; Gao et al., 2022; Si et al., 2023; You et al., 2023; Zhou et al., 2023; Hu et al., 2022; Xenos et al., 2023). These works usually hold that text-only PLMs can effectively infer answers, even compensating for the loss of fine-grained visual features in image captions. Chen et al. (2022c) illustrate how encoder-decoder PLMs address the long-tail problem in answers and discrepancies between training and testing sets, while avoiding span prediction and directly generating free-form answers. Jain et al. (2021) reframe VQA as a MRC task, integrating search engines for additional context. TRiG (Gao et al., 2022) and TwO (Si et al., 2023) expand this approach to include object-level (e.g., object, attribute, and OCR labels) information alongside captions. Utilizing LLMs like GPT-3 with image captions, PICa (Yang et al., 2022) reveals that pure PLMs can achieve impressive performance in zero-shot and few-shot learning scenarios. KAT (Gui et al., 2022) further queries GPT-3 for providing reasoning evidence, aiming to extract deeper insights and implicit knowledge from GPT-3’s outputs to bolster the reasoning process. REVIVE (Lin et al., 2022) employs a Transformer encoder as an adapter to utilize fine-grained regional visual information. PROOFREAD (Zhou et al., 2023) utilizes XGBoost (Chen and Guestrin, 2016), a gradient-boosted decision tree model, as a knowledge perceiver to classify knowledge entries based on their contribution scores across various dimensions. The Fusion-in-Decoder (FiD) approach (Izacad and Grave, 2021), where knowledge is individually compressed in the encoder and then jointly utilized in the decoder for reasoning, is adopted by various studies (Gui et al., 2022; Gao et al., 2022; Chen et al., 2022c; Wu and Mooney, 2022; Lin et al., 2022; Salemi et al., 2023a; Si et al., 2023). This allows for the simultaneous input of a large corpus of uni-modal or multi-modal background knowledge into the (V)PLMs.

- To mitigate the loss of fine-grained visual details in caption-based conversion, Wang et al. (2023g) leverage the LLM’s reason-

ing capabilities to spotlight critical image details that that might be overlooked in captions. By decomposing the main Q into sub-questions and obtaining answers via a pre-trained VQA model, they identify and select those factual summaries with higher contribution scores than the original Q , supplementing the initial captions with these key details. This is similar to KAT (Gui et al., 2022) and TwO (Si et al., 2023), which apply In-Context Learning (ICL) in GPT-3 which employs a combination of Q , caption, and object labels as the prompt to generate implicit textual knowledge; PromptCap (Hu et al., 2022) introduces Q -guided caption generation to cover the visual details required by Q ; ASB (Xenos et al., 2023) identifies the image patches most relevant to Q and generates informative captions from these patches only; ColaFT (Chen et al., 2023d) prompts VLMs to generate captions and plausible answers separately, which are then concatenated with the instruction prompt, Q , and choices, forming a holistic context for LLMs to logically deduce the answer.

Knowledge-aware Answer Determination plays a crucial role in generating and predicting answers, often overlapping with *Knowledge-aware Modality Interaction*. Certain methods uniquely address both these aspects simultaneously, highlighting their intertwined nature.

1. **Information Extraction:**
 - To further rank the potential answers, some approaches implement heuristic rules, like matching score calculation (Wang et al., 2018a; Narasimhan and Schwing, 2018) and answer frequency assessment (Wang et al., 2018a).
2. **Discrimination:**
 - Such methods are effective when narrowing down potential answers within a certain range, often using GNN-alike models (Narasimhan et al., 2018; Liu et al., 2022; Hussain et al., 2022) as the backbones (Fig. 4 (c)).
 - Furthermore, discriminators can be either MLP-based (Wang et al., 2019; Narasimhan et al., 2018; Zhu et al., 2020; Yu et al., 2020; Liu et al., 2022) or rule-based (Narasimhan and Schwing, 2018). A

3279
3280
3281
3282
3283
3284
3285
3286
3287
3288
3289
3290
3291
3292
3293
3294
3295
3296
3297
3298
3299
3300
3301
3302
3303
3304
3305
3306
3307
3308
3309
3310
3311
3312
3313
3314
3315
3316
3317
3318
3319
3320
3321
3322
3323
3324
3325
3326
3327
3328
3329

notable limitation of this approach is time consumption, especially when dealing with extensive answer vocabularies.

3. Classification:

- Many studies reformulate the question-answering process as a classification problem, often employing a Fully Connected (FC) or MLP layer for answer prediction (Su et al., 2018; Marino et al., 2019; Shah et al., 2019; Li et al., 2020; Ziaeeefard and Lécué, 2020; Saqur and Narasimhan, 2020; Zhang et al., 2021a; Gardères et al., 2020; Zheng et al., 2021b; Li and Moens, 2022; Guo et al., 2022b; Han et al., 2023; Heo et al., 2022; Li et al., 2022b; Wang et al., 2022b; Song et al., 2023b), where the output dimension corresponds to the pre-defined number of answer candidates.
- Chen et al. (2021d) introduce an answer masking strategy that imposes direct knowledge-based constraints on the classifier’s predicted answer probabilities, thereby limiting the range of potential answers. This method parallels KRISP (Marino et al., 2021), which employs late fusion to integrate the implicit and symbolic components of the model, selecting the highest-scoring answer from the combined answer vectors. MAVEx (Wu et al., 2022) introduces an answer validation module that leverages knowledge features from retrieved *I*, ConceptNet, and Wikipedia for answer candidate validation.
- For (V)PLM-based methods, a classification (or projection) head is typically appended to the output [CLS] embedding (Fig. 4 (f)) (Shevchenko et al., 2021; Luo et al., 2021; Salaberria et al., 2023; García-Olano et al., 2022; Guo et al., 2022b; Ding et al., 2022; He and Wang, 2023; Ravi et al., 2023), often utilizing encoder-based backbones like LXMERT (Tan and Bansal, 2019) and BERT (Devlin et al., 2019). However, a significant trade-off common to classification-based approaches still exists, as noted by Chen et al. (2022c): the necessity to balance answer coverage and error rate, which hinges on pre-defining the answer candidate set according to its occurrence frequency.

4. Generation:

- Textual Generative Models have become increasingly important in VQA tasks, particularly for addressing questions with answers outside pre-defined vocabularies. Generative (V)PLM-based methods are now increasingly supplanting traditional classification-based approaches (Chen et al., 2022c; Yang et al., 2022; Gao et al., 2022; Gui et al., 2022; Lin and Byrne, 2022; Wu and Mooney, 2022; Zhang et al., 2022a; Lin et al., 2022; Salemi et al., 2023a; You et al., 2023; Si et al., 2023; Hu et al., 2023; Shao et al., 2023; Zhou et al., 2023; Ossowski and Hu, 2023; Wang et al., 2023g; Hu et al., 2022; Xenos et al., 2023; Ghosal et al., 2023; Khademi et al., 2023).
- These methods, using decoder-based or encoder-decoder based models like GPT-3 (Brown et al., 2020), T5 (Raffel et al., 2020), VL-T5 (Cho et al., 2021), and BLIP-2 (Li et al., 2023b), feed constructed prompts to implicitly retrieve knowledge and perform analytical reasoning. Answer generation often relies on greedy decoding or beam search strategies (Gao et al., 2022; Ravi et al., 2023; Khandelwal et al., 2023), with the former selecting the most probable token at each step and the latter maintaining a fixed-size beam to produce a list of ranked answer candidates. To improve few-shot learning performance in models with large parameters, such as GPT-3, strategies like incorporating high-quality ICL examples (Shao et al., 2023; Wang et al., 2023g; Hu et al., 2022; Xenos et al., 2023) and employing multi-query ensembles (Yang et al., 2022; Xenos et al., 2023) are effective. Prophet (Shao et al., 2023) enhances this process by first generating candidate answers using a standard VQA model, subsequently refined through GPT-3. Meanwhile, Cola-FT (Chen et al., 2023d) prompts VLMs to generate captions and plausible answers separately, then integrating them with the instructional prompt, question, and candidate options for LLM-based reasoning.
- As shown in Table 1, a noticeable increase in text-generation-based VQA methods is observed in the last two years. This trend

can also be attributed to the limitations of Exact Match answer evaluation manner in VQA benchmarks, which historically do not provide an advantage in evaluating the performance of open-ended generative models.

Recent advancements include CodeVQA (Subramanian et al., 2023), a training-free method prompting Codex (Chen et al., 2021c) with in-context examples to break down Q into Python code. This method leverages pre-defined visual modules in pre-trained VLMs, utilizing conditional logic and arithmetic. In line with NLP findings that LLMs improve performance at reasoning tasks when solving problems step-by-step (Wei et al., 2022; Yin et al., 2023b), VQA performance improves by decomposing Q and answering sub-questions sequentially. Khandelwal et al. (2023) propose *Successive Prompting*, where a LLM generates and resolves follow-up questions one at a time using a VLM, culminating in an answer to the original Q .

Metrics. In VQA performance evaluation, *Accuracy* (Acc), defined as the proportion of correctly answered test questions, is a predominant metric. The Georgia Tech Visual Intelligence Lab’s VQA Python API¹⁰ employs a standard technique for computing Acc is recommended in the VQA challenge (Antol et al., 2015):

$$\text{Acc}(ans) = \min(1, \#\{\text{human that said that } ans\}/3). \quad (2)$$

This metric assigns a soft score (ranging from 0 to 1) to each answer, based on a voting mechanism among multiple annotators. In contrast, the Exact Match (EM) metric treats all annotated answers as ground truth (GT), offering a less stringent evaluation criterion (Gao et al., 2022). Additionally, the WuPalmer similarity (WUPS) (Wu and Palmer, 1994) calculates the similarity between words based on their common sub-sequences in a taxonomy tree. A candidate answer is considered correct if its similarity to a reference word exceeds a specified threshold. Chen et al. (2022c) introduce Inclusion-based and Stem-based Acc metrics. The former considers an answer A correct if it includes or is included by a GT answer after normalization. The latter assesses correctness based on the intersection of stems between A and the GT A (e.g., the stem of “happy” and “happiness” is “happi”). Not that other NLP automatic evaluation metrics,

beyond assessing answer correctness, can also evaluate the model’s explanation quality. For example, generative metrics such as BLEU (Papineni et al., 2002), CIDEr (Vedantam et al., 2015), and METEOR measure the linguistic quality and relevance of rationale statements against a reference set. Originally developed for machine translation, these metrics provide insights into the generated explanations’ coherence and fluency, complementing the evaluation of answer correctness.

Question: What is he doing?
Answer: horseback riding
Candidate: riding a horse
Is the candidate correct? [yes/no]

Given the limitations of lexical matching metrics in evaluating open-domain VQA predictions from generative models, where entirely different words may convey the same meaning, (Khandelwal et al., 2023), Kamaloo et al. (2023) further propose an evaluation metric that leverages Instruct-GPT (Ouyang et al., 2022), prompting it with Q and the candidate answers to determine their correctness. An example of this process is illustrated in the adjacent code snippet.

Knowledge Base. The background KBs for knowledge-aware multi-modal reasoning frequently involves multiple KGs, each bringing its unique and complementary insights to the reasoning process. **Trivia** knowledge, such as DBpedia, provides facts about famous people, places, and events. **Commonsense** knowledge, offers insights into basic concepts like the composition of houses or parts of a wheel. **Scientific** knowledge, found in databases like hasPart KB, details classifications and properties, such as the genus of dogs or types of nutrients. Lastly, **situational** knowledge from resources like Visual Genome offers contextual data, e.g., typical locations of cars or common contents found in bowls.

- **ConceptNet** (Speer et al., 2017) encapsulates human commonsense knowledge, containing various relations including *usedFor*, *createdBy*, and *isA*, primarily generated from the Open Mind Common Sense (OMCS) project;
- **DBpedia** (Auer et al., 2007), constructed from Wikipedia, spans multiple fields relevant to daily life. In this KG, concepts are connected through categories and super-categories in accordance with the SKOS¹¹ Vocabulary;

¹⁰<https://github.com/GT-Vision-Lab/VQA>

¹¹<http://www.w3.org/2004/02/skos/>

- **WebChild** (Tandon et al., 2014) connects nouns with adjectives through fine-grained relations, such as *hasShape*, *faster*, *bigger*. This information is automatically extracted from the Web;
- **Wikidata** (Vrandečić and Krötzsch, 2014) offers extensive factual knowledge, including a broad range of topics about the world;
- **hasPart KB** (Bhakthavatsalam et al., 2020) documents relationships between objects, both common and scientific, such as (*Dog*, *hasPart*, *Whiskers*) and (*Molecules*, *hasPart*, *Atoms*);
- **Visual Genome (VG)** (Krishna et al., 2017) gathers scene graphs from real-life situations, focusing on spatial relationships, e.g., (*Boat*, *isOn*, *Water*), and common affordances, e.g., (*Person*, *sitsOn*, *Couch*);
- **ATOMIC** (Hwang et al., 2021) consists of over 1M knowledge triplets covering a range of topics, including physical-entity relations, event-centered relations, and social interactions.
- **CSKG** (Ilievski et al., 2021) is a large consolidated source that integrates common-sense knowledge from seven diverse and disjoint sources, including ConceptNet, Wikidata, ATOMIC, VG, Wordnet (Miller, 1995), Roger (Kipfer, 1992) and FrameNet (Baker et al., 1998).

BENCHMARKS: We select FVQA (Wang et al., 2018a) and OKVQA (Marino et al., 2019) as our primary datasets due to their critical contributions to advancing knowledge-aware VQA and their significant impact on the development of subsequent datasets. Table 1 presents a chronological analysis of relevant methods, detailing their performance, model paradigms, and design principles.

Discussion 1 *VQA datasets vary in their answer formats, ranging from multiple-choice, where models select from provided options, to open-ended formats that test a model’s understanding, reasoning, and independent answer generation or retrieval capabilities. Beyond answer formats, an important consideration in these datasets is the use of a Ground Truth (GT) set of facts for answering questions. Datasets like those in the FVQA series come with their own GT facts, while those in the OKVQA series do not. These facts should ideally be employed not for training purposes (such as pre-training a relation classifier) but for assessing the model’s proficiency in KG fact retrieval. Besides, selecting appropriate knowledge sources and*

methods for knowledge filtering is also crucial for model performance.

Furthermore, as indicated in Table 1, the comparison of VQA works can be influenced due to varying background KG sources and backbone models. Ensuring consistency in these aspects is essential for fair comparative analysis. It’s important for researchers to distinguish whether improvements are due to the quality of the KB, the method of KG integration, or the backbone’s inherent capabilities. These distinctions, often overlooked, are crucial to understanding genuine progress in the field. Relying solely on sophisticated visual, language, or multi-modal backbones to claim SOTA results, without addressing the uniformity of model parameters and ensuring fair comparisons, may compromise the credibility of the findings. Given VQA’s practical applications, additional factors such as time, space complexity, real-time consumption, and GPU requirements are also significant for a comprehensive evaluation of these models.

Resource: In analyzing the evolution of KG-aware VQA datasets, we categorize the developments into three main groups: FVQA-type, OKVQA-type, and others.

(i) **FVQA (Wang et al., 2018a):** The **KB-VQA** dataset (Wang et al., 2017) first evaluates VQA algorithms’ ability to leverage external knowledge for answering complex image-based questions. It consists of multiple *Q-A* pairs per image, crafted by five questioners using predefined templates. These pairs aim to probe knowledge levels that surpass mere visual observation by leveraging DBpedia as the knowledge source. Expanding on KB-VQA, FVQA (Wang et al., 2018a) includes more questions, images and integrates additional KGs such as ConceptNet and Webchild. Notably, FVQA is the first VQA dataset to provide supporting facts for each question (i.e., external knowledge facts, rather than visual relation facts in R-VQA (Lu et al., 2018b)), paving the way for developing more knowledgeable VQA systems. **Variants:** **ZS-F-VQA** (Chen et al., 2021d) targets Zero-shot VQA, designed to prevent overlap between training and testing answers, paying attention on answer bias and Out Of Vocabulary (OOV) issues. **KRVQA** (Cao et al., 2022b) imposes constraints to promote image context engagement over mere knowledge fact memorization; **FVQA 2.0** (Lin et al., 2023) increases dataset size and introduces adversarial question variants to balance the answer

Table 1: Comparison of Knowledge-based VQA accuracy results on OKVQA (Marino et al., 2019) and FVQA (Wang et al., 2018a). The icon  represents KG-based VQA methods;  indicates methods without KG utilization. The † symbol signifies methods pre-trained on VQA2.0 or similar datasets. * indicates results reported on dataset version 1.1, which differs from version 1.0 in answer stemming methods. Abbreviations used: Q (Question), V (Visual), w/ (with), KG (Knowledge Graph), CN (ConceptNet), WP (Wikipedia), WC (WebChild), WD (Wikidata), DBP (DBpedia), VG (VisualGenome), YG (YAGO), HP (hasPart KB), AT (ATOMIC (Sap et al., 2019)), AS (Ascent (Nguyen et al., 2021)), VLM (Visual-Language Model), GNN (Graph Neural Network), GAT (Graph Attention Network), MRC (Machine Reading Comprehension), MHA (Multi-head Attention), DMN (Dynamic Memory Network), DPR (Dense Passage Retriever), FiD (Fusion-in-Decoder), In-context Learning (ICL), GI (Google Image), GS (Google Search), Enc.Dec.(Encoder-Decoder), DC (Discrimination), IE (Information Extraction), CLS (Classification), TG (Text Generation), WIT (Wikipedia-Image-Text (Srinivasan et al., 2021)). For methods employing both PLM intrinsic knowledge and external KB, only the KB is listed in the knowledge source.

	Methods	Approaches (Paradigms)	Key Idea	Knowledge Source	FVQA	OKVQA
2018~2021	 QQmapping (Wang et al., 2018a)	RDF Query (IE)	Question-Query Mapping	DBP / CN / WC	56.91	-
	 STTF (Narasimhan and Schwing, 2018)	Relation Query (IE)	Scoring the Facts	DBP / CN / WC	62.20	-
	 OB (Narasimhan et al., 2018)	Fact Retrieval + GNN (DC)	Entity Graph + GCN	DBP / CN / WC	69.35	-
	 Marino et al. (2019)	Retrieval + ArticleNet (IE)	Web Retrieval + Knowledge Span Prediction	WP	-	27.84
	 KG-Aug (Li et al., 2020)	Fact Retrieval + GNN (CLS)	Augment Q&V Features w/ Entity Embedding	WD / CN	38.58	26.71
	 Chen et al. (2021d)	Alignment + Re-rank (CLS)	KG-aware Answer Masking for Validation	DBP / CN / WC	58.27	-
	 ERMLP (Ramnath and and, 2020)	KGE + Attention (IE)	Knowledge-guided Co-attention	DBP / CN / WC	60.82	-
	 Liu et al. (2022)	Fact Retrieval + GNN (DC)	Q-V Guided Cross-modal GAT	DBP / CN / WC	63.56	29.43
	 KAN (Zhang et al., 2021a)	Retrieval + Attention (CLS)	Question-guided MHA	CN	66.39	-
	 Mucko (Zhu et al., 2020)	Fact Retrieval + GNN (DC)	Question-guided Attention + Cross-KG GAT	DBP / CN / WC	73.06	29.20
	 GRUC (Yu et al., 2020)	Fact Retrieval + DMN (DC)	Fact-centered DMN + GRU	DBP / CN / WC	79.63	29.87
	 ConceptBert (Gardères et al., 2020)	GCN + Attention (CLS)	Compact Trilinear Interaction + MHA	CN	-	33.66
	 KRISP† (Marino et al., 2021)	GCN + VLM (CLS)	Global KG + RGCN + VisualBERT	HP / DBP / CN / VG	-	38.90*
	 MAVEx† (Wu et al., 2022)	Multi-retrieval + Re-rank (CLS)	Web Retrieval + ViLBERT + Answer Validation	WP / CN / GI	-	38.70*
	 RVL† (Shevchenko et al., 2021)	KGE Alignment + VLM (CLS)	Aligning VLM Text Embedding w/ KGE	WD / CN / PLM	54.27	39.04
	 Luo et al. (2021)†	Dense Retriever + PLM (IE)	RoBERTa + DPR Learning + MRC	GS	-	39.20*
	 PGVQA (Song et al., 2023b)	Retrieval + Re-rank (CLS)	KG-aware Answer Refinement	DBP / CN / WC	-	41.07
2022	 SUPER (Han et al., 2023)	Multi-modules + DMN (CLS)	Q-V Guided Modular Routing + DMN	CN	48.90	30.46
	 MKRE (Hussain et al., 2022)	Fact Retrieval + GNN (DC)	Question Guided Attention + Cross-KG GAT	DBP / CN / WC	73.06	-
	 DMMGR (Li and Moens, 2022)	Retrieval + DMN + GNN (CLS)	Caption + Multi-scale DMN + GAT	DBP / CN / WC	81.20	-
	 CBM-BERT† (Salaberria et al., 2023)	Caption + PLM (CLS)	Caption + BERT + Ensemble	PLM	-	36.00*
	 CBM-T5† (Salaberria et al., 2023)	Caption + PLM (TG)	Caption + T5 + Ensemble	PLM	-	40.80*
	 UniFer† (Guo et al., 2022b)	Fact Retrieval + VLM (CLS)	Loss Gap Driven DPR Learning + ViLT	CN	-	42.13*
	 MuKEA† (Ding et al., 2022)	VLM + KGE	LXMERT + Multi-modal TransE	VLM	-	42.59*
	 PICa† (Yang et al., 2022)	Caption + Decoder (TG)	Caption + Tag + ICL + GPT3	GPT3	-	43.30*
	 KVQAmeta† (García-Olano et al., 2022)	KGE Alignment + VLM (CLS)	Aligning VLM Embedding w/ Wikipedia2Vec	WP	-	43.67*
	 LaKo† (Chen et al., 2022c)	Retrieval + Enc.Dec. (TG)	Caption + Fact DPR + T5 + FiD	HP / DBP / CN / WC	-	47.01*
	 TRiG† (Gao et al., 2022)	Retrieval + Enc.Dec. (TG)	Caption + Tag + DPR + FiD	WP	-	49.24*
	 KAT† (Gui et al., 2022)	Retrieval + Enc.Dec. (TG)	Caption + Tag + ICL + GPT3 + DPR + FiD	WD / GPT3	-	53.09 *
	 EnFoRe† (Wu and Mooney, 2022)	Retrieval + Enc.Dec. (TG)	KAT + Entity Focused DPR	WD / GPT3	-	54.35*
	 RAVQA† (Lin and Byrne, 2022)	Retrieval + Enc.Dec. (TG)	DPR Learning + Ensemble	GS	-	54.48*
	 REVIVE† (Lin et al., 2022)	Retrieval + Enc.Dec. (TG)	KAT + Regional Visual	WD / GPT3	-	56.604*
2023	 VLC-BERT† (Ravi et al., 2023)	Fact Generation + VLM (CLS)	COMET Generation + MHA + VL-BERT	CN / AT	-	43.14*
	 DEDR† (Salemi et al., 2023a)	Retrieval + Enc.Dec. (TG)	Mutual Retriever Distillation + VL-T5 + FiD	WP	61.80	44.57*
	 RR-VEL† (You et al., 2023)	EL + Retrieval + Enc.Dec. (TG)	Ground Truth Referent in Q + T5	HP / CN / Ascent	65.59	49.48*
	 MCR-MemNN (Yin et al., 2023a)	Fact Retrieval + DMN (CLS)	Multi-clue Reasoning + Fact-centered DMN	DBP / CN / WC	70.92	-
	 CodeVQA† (Subramanian et al., 2023)	Code Generation + PLM (TG)	ICL + Codex + Modular Combination	Codex / VLM	-	53.50*
	 TwO† (Si et al., 2023)	Retrieval + Enc.Dec. (TG)	KAT + OFA Multi-modal Knowledge	WP / GPT3	-	57.57*
	 PROOFREAD† (Zhou et al., 2023)	Decoder + Re-rank (TG)	Answer-aware Knowledge Generation & Filter	ChatGPT	-	57.60*
	 REVEAL† (Hu et al., 2023)	Retrieval + Enc.Dec. (TG)	Multi-modal Retrieval + Gate + DMN + T5	WIT / WD / VQA2.0	-	59.10*
	 MM-Reasoner† (Khademi et al., 2023)	Fact Generation + Enc.Dec. (TG)	Vision APIs + ICL Rationales Generation	GPT-4	61.10	59.20*
	 Wang et al. (2023g)†	Q-decomposition + Decoder (TG)	Q Decomposition + Fact Refinement + PICa	ChatGPT	-	59.34*
	 PromptCap† (Hu et al., 2022)	Caption + Decoder (TG)	Q-guided Caption Generation + PICa	GPT-3	-	60.40*
	 Prophet† (Shao et al., 2023)	Caption + Decoder (TG)	MCAN + Answer Pruning + Answer-aware ICL	GPT3	-	61.10*
	 ASB† (Xenos et al., 2023)	Caption + Decoder (TG)	Q-guided Patch Caption Selection + PICa	LLaMA-13B	-	61.20*
	 Cola-FT† (Chen et al., 2023d)	Decoder (TG)	OFA + BLIP + LLM Answer Decision	FLAN-T5	-	62.40*
	 GPT-4V† (Li et al., 2023f)	Decoder (TG)	Prompt + GPT-4V	GPT-4V	-	64.28*

distribution;

(ii) **OKVQA** (Marino et al., 2019): Different from FVQA, OKVQA dataset focuses on open-world VQA, involving questions that implicitly require external knowledge without specifying a direct KB link or providing explicit KG triplets. Its broad knowledge scope makes it a benchmark alongside the VQA2.0 dataset (Antol et al., 2015). **Variants:** **OKVQA_{S3}** and **S3VQA** (Jain et al., 2021) enhance the original OK-VQA by incorporating questions that require object detection within images, with subsequent substitution of the detected object in the query and employing web searches to find answers; **A-OKVQA** (Schwenk et al., 2022) introduces a greater diversity of world knowledge and more reasoning steps to extend OK-VQA, further providing rationales for each question to aid in training explainable VQA models. **OKVQA2.0** (Reichman et al., 2023) refines OK-VQA with corrections and attaching Wikipedia sources to *Q-I* pairs; **ConceptVQA** (Gan et al., 2023) enriches OK-VQA with entity-level annotations aligned with ConceptNet entities and presents a unique challenge by ensuring its testing split features non-overlapping answers with the training set, similar to ZS-F-VQA (Chen et al., 2021d).

(iii) **Others:** Li et al. (2017) develop **Visual7W+KB** from the Visual7W test split images (Zhu et al., 2016), automating question creation using predefined templates and ConceptNet (Speer et al., 2017) for guidance; **KVQA** (Shah et al., 2019) incorporates world knowledge about named entities like *Barack Obama* and the *White House* from Wikidata (Vrandečić and Krötzsch, 2014), also employing face identification technology in image analysis; **ViQuAE** (Lerner et al., 2022) extends KVQA’s scope to include a broader range of entity types beyond just persons; **VCR** (Zellers et al., 2019) targets understanding human intentions in movie scenes with questions such as “*why is [PERSON] doing this?*”; **AI-VQA** (Li et al., 2022b) utilizes Visual Genome scene graphs and ATOMIC KG (Hwang et al., 2021) event knowledge, enriched by including volunteer-annotated QA pairs and detailed scene/object descriptions; **DANCE** (Ye et al., 2023) re-formats knowledge triples as natural language riddles paired with images, aiming to infuse visual language models with commonsense knowledge; Gao et al. (2023) introduce **LoRA**, a dataset focusing on formal and complex description logic reasoning in VQA. Centered around a KB related

to food and kitchen scenarios, LoRA aims to enhance the logical reasoning capabilities of VQA models, which are not adequately assessed by existing VQA datasets; **ScienceQA** (Lu et al., 2022), sourced from elementary and high school science curricula, includes 21,208 items along with lectures and explanations. It challenges models to generate coherent explanations across a wide range of subjects, setting it apart from OKVQA. Despite not incorporating a KG in its design, ScienceQA is pivotal for advancing knowledge-intensive multimodal models, which marks a significant step in the evolution of future KG-aware VQA methods.

In addition, KG-aware VQA can also extend to various scenarios beyond traditional settings. For example, **KnowIT VQA** (Garcia et al., 2020) contains video clips from “*The Big Bang Theory*” with associated knowledge-based QA pairs, annotated by those dedicated fans well-versed in the show’s content. **K-EQA** (Tan et al., 2023) employs a KB and 3D scene graphs, enabling an AI agent to navigate environments and answer environment-aware natural language queries.

A.3.2 Visual Question Generation

VQG (Xie et al., 2022; Chen et al., 2023b; Salemi et al., 2023b) leverages visual cues to generate questions, diverging from traditional VQA by prioritizing question creation. This process is crucial in educational applications, such as engaging children with questions about images to support learning. Early VQG models (Mostafazadeh et al., 2016) utilize RNNs to generate questions based solely on images, leading to questions that often lack specific focus. In the KG-aware VQG domain, volunteers create the K-VQG dataset (Uehara and Harada, 2023) by integrating external knowledge from resources like ConceptNet and Atomic (Hwang et al., 2021) with image content, using partially masked commonsense triplets to enrich questions with knowledge. Xie et al. (2022) develop a pipeline comprising a visual concept feature extractor, knowledge representation extractor, target object extractor, and a decoder. This setup, aligned with the process outlined in Fig. 3, integrates non-visual knowledge into VQG and employs FVQA for its evaluation. KECVQG (Chen et al., 2023b) utilizes a causal graph to analyze and correct spurious correlations in VQG by linking unbiased features with external knowledge, thereby disentangling visual features to lessen the impact of these correlations. Unlike VQA, VQG methods prioritize evaluation on mean-

ingfulness, logical soundness, and consistency with target knowledge over strict correctness (Uehara and Harada, 2023), often using NLP-style metrics like BLEU and CIDEr for assessment.

Discussion 2 *Evolving intelligent dialogues with chatbots remains a critical objective in VQG, particularly in empowering robots to formulate precise, knowledge-enriched questions that boost problem-solving capabilities for future advancements. Equally important is the progression towards more interactive KG-aware VQG systems, which can dynamically adapt their questioning strategies based on user interactions and feedback, marking a significant direction for future research. Moreover, with the ongoing rapid developments in VQA, transferring and adapting common problems and methods from VQA to VQG can catalyze further innovative breakthroughs in question generation technologies.*

A.3.3 Visual Dialog

VD (Chen et al., 2022a) extends the VQA task by adopting a multi-round format where a continuous series of Q - A pairs revolves around a single image. This setting shifts from the single-question focus of VQA to a dynamic, conversational interaction about the image, posing a challenge for agents to adaptively interpret evolving relationships among visual elements based on the dialogue context. VD methods typically leverage historical dialogue information as background knowledge (Guo et al., 2020; Jiang et al., 2020; Kang et al., 2021; Wang et al., 2022c; Zhao et al., 2023), employing visual graph construction, query-guided relation selection and GNN propagation for dialog reasoning. Guo et al. (2020, 2022a) introduce Q -conditioned attention to aggregate textual context from dialogue history, constructing a context-aware object graphs for Q -guided message passing. Similarly, KBGN (Jiang et al., 2020) uses cross-modal GNNs to bridge modal gaps and capture inter-modal semantics, retrieving information relevant to the current question from both vision and text sources.

Addressing the limitations of relying solely on internal knowledge from images and dialog history, some approaches integrate commonsense knowledge for enhanced conversational depth. These methods all align with the *KG-aware Understanding and Reasoning* paradigms we have previously outlined (Fig. 3). For example, (i) **Knowledge Retrieval**: SKANet (Zhao et al., 2021a) integrates commonsense knowledge from Concept-

Net into VD by using concept recognition and n-gram matching techniques to build a sub-KG. (ii) **Knowledge Representation**: KACI-Net (Zhang et al., 2023e) selects triplets with at least two entities or relations mentioned in a question and transforms them into textual format for subsequent processing. (iii) **Knowledge-aware Modality Interaction**: RMK (Zhang et al., 2022b) utilizes caption-based dense retrieval to fetch relevant facts from ConceptNet, injecting knowledge into the dialogues through sentence-level and graph-level cross-modal attention and embedding concatenation. (iv) **Knowledge-aware Answer Determination**: Acknowledging the issue of spurious correlations from unobserved confounders in retrieved knowledge, Liu et al. (2023a) construct a counterfactual commonsense-aware VD causal graph. This graph applies counterfactual reasoning to mitigate commonsense bias, reducing the effect of misleading or inaccurate commonsense in answer derivation.

Discussion 3 *Currently, the emphasis in knowledge-based VD mainly lies in using external commonsense knowledge, while other knowledge types, such as scientific and situational, have been relatively underexplored. However, the rise of LLMs is diminishing the distinction between VD and VQA, with In-context Learning techniques in VQA starting to overshadow the traditional role of context in dialogues. This shift prompts a need to reassess VD’s unique contribution and its path forward. As the boundaries between VD and VQA continue to merge, identifying and articulating the distinct potential of VD becomes imperative.*

A.4 Supplement for KG-driven Multi-modal Classification Tasks

Discussions are also extended to related multi-modal tasks like Fake News Detection and Movie Genre Classification, highlighting the diversity and wide-ranging applications within the field.

A.4.1 Supplementary Information for IMGC

Image Classification (IMGC) aims to identify objects within images and, with deep learning advancements, has even surpassed human performance in challenges like ImageNet ILSVRC (Rusakovsky et al., 2015). Consider a set of labeled training samples $\mathcal{D}_{tr} = \{(x, y) | x \in \mathcal{X}, y \in \mathcal{Y}\}$, a classifier aims to approximate a function $f : x \rightarrow y$ from input x to output label y with the assist of

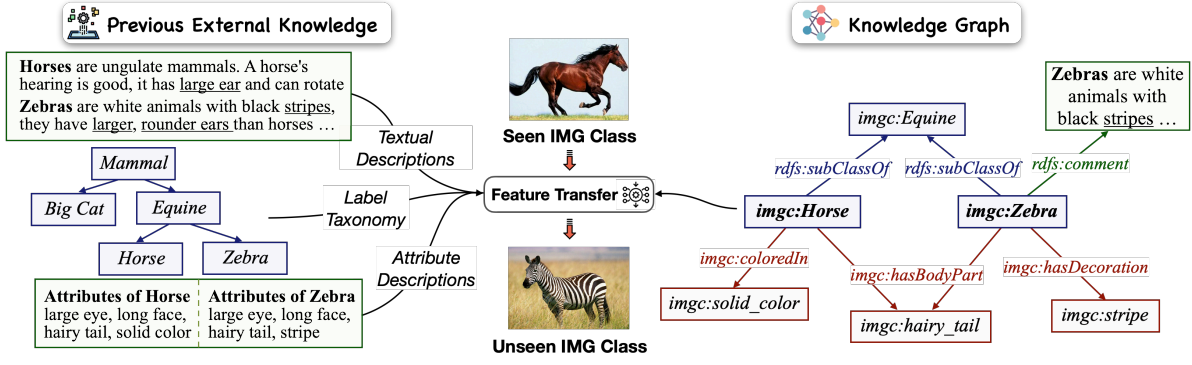


Figure 10: Comparison of previously used external knowledge (Left) and KG (Right) in Zero-shot Image Classification task (Geng et al., 2021a).

a background KG $\mathcal{G}$. This function should accurately predict the labels of samples in a testing set $\mathcal{D}_{te} = \{(x, y) | x \in \mathcal{X}', y \in \mathcal{Y}\}$, with $\mathcal{X} \cap \mathcal{X}' = \emptyset$.

In IMGC, x denotes an image, and y represents its class. Traditional IMGC follows a closed-world assumption, requiring extensive labeled images for both training and testing within known classes, i.e., $\mathcal{Y}_{tr} = \mathcal{Y}_{te} = \mathcal{Y}$. However, it is not feasible for newly emerging classes due to the impracticality of continuously annotating and retraining models with sufficient images for these classes. Consequently, there is a growing interest in Zero-Shot Image Classification (ZS-IMGC), which supports classifying images of novel, unseen classes without the need for specific training images, i.e., $\mathcal{Y}_{tr} \cap \mathcal{Y}_{te} = \emptyset$.

To handle these unseen classes, most existing ZS-IMGC methods adopt a knowledge transfer strategy (Chen et al., 2023c, 2021b): transferring labeled images, image features or model parameters from the seen classes in training set to unseen classes, guided by external knowledge that describes semantic relationships between classes. For instance, as illustrated in the left part of Fig. 10, consider the description of a “Zebra” as an animal with a horse-like body, tiger-like stripes, and black-and-white colors similar to a panda. Models can infer the appearance of a “Zebra” by combining features of these seen animals, even without direct exposure to its images. Briefly, ZS-IMGC relies on data from observed classes and class-specific semantic knowledge, with the external knowledge frequently embodying a modality distinct from the image data. This section reviews KG-based ZS-IMGC efforts to illustrate the typical practice of multi-modal learning in IMGC.

In ZS-IMGC literature, various forms of external knowledge are employed. Early ZS-IMGC works (Frome et al., 2013; Zhu et al., 2018) use textual class descriptions or names to model inter-class

relationships. Others (Xian et al., 2018; Lampert et al., 2014) utilize class attributes, annotating each class with descriptive characteristics, thereby defining semantic relationships through shared attributes (see Fig. 10, left). However, these approaches sometimes face limitations in capturing complete semantics (Geng et al., 2023).

In KG-aware ZS-IMGC, a KG is defined as $\mathcal{G} = \mathcal{E}, \mathcal{R}, \mathcal{T}$, with $\mathcal{Y} \subset \mathcal{E}$. This paradigm, representing semantic hierarchical relationships among classes, is instrumental in augmenting classification performance and interpretability. For example, studies like (Wang et al., 2018b; Kampffmeyer et al., 2019) integrate hierarchical relationships from WordNet, while (Roy et al., 2022; Nayak and Bach, 2022; Gao et al., 2019) explore class knowledge from commonsense KGs such as ConceptNet. KGs, due to their compatibility, can unify various knowledge forms, including textual description and discrete attributes, into a single graph (see Fig. 10, right).

Mapping-based Methods. Chen et al. (2020a) employ an OWL-based ontology for animal classes, encoding it via the OWL EL embedding method and learn a linear encoder to map image features to class embeddings, with a focus on reconstruction loss for reverse mapping. Kata et al. (Zeynep Akata and Florent Perronnin and Zaïd Harchaoui and Cordelia Schmid, 2016; Akata et al., 2013) map initial class encodings into the image feature space, representing classes as multi-hot vectors of ancestors based on class hierarchies. There are also some joint mapping methods that map both the class encoding and the image features. For example, DUET (Chen et al., 2023h), an end-to-end Transformer-based ZSL method, leverages cross-modal PLMs for fine-grained visual characteristic reorganization and discrimination with structured KGs serialized as input.

Although mapping-based methods, which often employ linear or nonlinear transformations, are straightforward to implement, they exhibit a bias towards seen classes when trained exclusively on their images. This bias is particularly problematic in generalized ZSL scenarios, where seen and unseen classes coexist, highlighting their inherent limitations.

Data augmentation Methods. DOZSL (Geng et al., 2022) employs a disentangled KG embedding module to enhance the quality of synthesized image features in OntoZSL (Geng et al., 2021a). TGG (Zhang et al., 2019) generates few-shot samples, where a GAN-based generation module is applied to generate an image-level graph. Zhang et al. (2022d) develop a MMKG by merging visual representations of classes with word embeddings, creating multi-modal class nodes and edges that explicitly model class correlations, thereby enhancing information transfer from seen to unseen class nodes.

Propagation-based Methods. Certain methods tackle KGs’ diverse relation types by using multi-relational GCNs (Chen et al., 2020b) or dividing multi-relation KGs into single-relation graphs with parameter-shared GCNs for feature propagation (Geng et al., 2022; Wang and Jiang, 2021; Wu et al., 2023a). For multi-label images, where models assign a probability score per class, GNNs utilize KG-implied correlations to propagate these scores from seen to unseen class nodes, as done by Lee et al. (Lee et al., 2018).

Discussion 4 *Incorporating diverse class semantics generally yield better ZS-IMGC results, even with basic methods. For instance, Table 2 illustrates that some mapping-based methods (Frome et al., 2013; Chen et al., 2020a), employing straightforward linear mapping functions, achieve better performance by integrating various class semantics such as attributes, hierarchy, and names, compared to GCNZ[†], which relies solely on class hierarchy. Significantly, enriching GCNZ[†] with more class semantics, markedly enhances its effectiveness (GCNZ[‡]). Furthermore, KG-based ZS-IMGC methods typically operate in a class transductive setting, where unseen **classes** are known during training (Fig. 5 (b)), in contrast to the conventional inductive approach that utilizes only seen class knowledge. These methods leverage a KG to bridge seen and unseen classes through semantic*

links. Additionally, although the generalized ZSL setting is well-recognized by many researchers, some studies adopt their own definitions, specifically testing only unseen class images while classifying them within a combined pool of both seen and unseen classes. This variation requires careful consideration in future work.

Table 2: Comparison of ZS-IMGC results across various datasets. We use *Acc* and *H* as the evaluation metrics of Standard ZSL and Generalized ZSL, respectively. DeVISE here is implemented with a KG which covers the semantics of class hierarchy, class attributes, attribute hierarchy and class names (see (Geng et al., 2023) for details). GCNZ[†] utilizes a KG with class hierarchy only, whereas GCNZ[‡] includes broader class semantics like attributes, as reported in (Geng et al., 2023). DOZSL (GAN) and DOZSL (GCN) are variants of DOZSL (Geng et al., 2022), representing generation-based and propagation-based ZSL learners, respectively. ImageNet results are tested on 2-hops unseen classes.

Dataset	Methods	Acc (%)	H (%)
ImageNet	GCNZ (Wang et al., 2018b)	19.8	–
	DGP (Kampffmeyer et al., 2019)	26.6	–
	FGP (Wu et al., 2023a)	26.4	–
ImNet-A	DeVISE (Frome et al., 2013)	33.62	26.01
	OntoZSL (Geng et al., 2021a)	39.00	32.15
	DOZSL (GAN) (Geng et al., 2022)	40.26	32.82
	DOZSL (GCN) (Geng et al., 2022)	38.69	32.12
	GCNZ [†] (Wang et al., 2018b; Geng et al., 2023)	33.95	26.68
	GCNZ [‡] (Wang et al., 2018b; Geng et al., 2023)	36.64	31.38
AwA2	DeVISE (Frome et al., 2013)	46.12	15.88
	OWL-based (Chen et al., 2020a)	58.90	–
	DUET (Chen et al., 2023b)	–	58.00
	OntoZSL (Geng et al., 2021a)	63.31	56.06
	DOZSL (GAN) (Geng et al., 2022)	66.36	57.62
	DOZSL (GCN) (Geng et al., 2022)	63.88	52.74
	GCNZ [†] (Wang et al., 2018b; Geng et al., 2023)	37.44	14.34
	GCNZ [‡] (Wang et al., 2018b; Geng et al., 2023)	62.98	31.98
	FGP (Wu et al., 2023a)	79.10	43.30

RESOURCES: Several open datasets and KG resources for KG-aware ZS-IMGC have been proposed:

(i) **ImageNet** (Deng et al., 2009): A large-scale database with 14M images across 21K classes each aligned with a WordNet (Miller, 1995) entity. It leverages class hierarchies as KG-based knowledge, where the graph only contains one type of relation, i.e., *subClassOf*. From (Xian et al., 2019), a subset of 1K classes serves as seen classes, with unseen classes determined by their distance in the WordNet graph. ImageNet is widely used for ZS-IMGC benchmarking, albeit with a single-relational KG limitation.

(ii) **ImNet-A** and **ImNet-O**: Subsets of ImageNet by Geng et al. (2021a, 2023). ImNet-A contains 80 animal classes, and ImNet-O has 35 general object classes. Each comes with a KG combining multiple knowledge types, including

class attribute, class name, commonsense knowledge from ConceptNet, class hierarchy (taxonomy) from WordNet, and logical relationships such as *disjointness*.

(iii) **AwA2** (Xian et al., 2019): A coarse-grained animal classification dataset with 50 animal classes and 37,322 images, plus 85 expert-annotated attributes. Its classes align with WordNet entities for taxonomy-based KGs. Geng et al. (2023) equip AwA2 with a KG similar to ImNet-A and ImNet-O. Chen et al. (2020a) utilize OWL2 to model complex class semantics.

(iv) **NUS-WIDE** (Chua et al., 2009): A multi-label image classification dataset, where each image contains multiple objects. In works like (Lee et al., 2018), NUS-WIDE is accompanied by a KG with three types of label relations, including a super-subordinate correlation from WordNet, positive and negative correlations computed by label similarities such as WUP similarity.

BENCHMARKS: Single-label image classification in ZS-IMGC includes two evaluation settings: (i) **Standard ZSL:** Focuses solely on unseen class samples, using *Macro Accuracy*, which is calculated as the average of individual class accuracies (correct predictions to total samples ratio), as the metric. (ii) **Generalized ZSL (GZSL):** Evaluates both seen and unseen class samples, hence more challenging. Here, two *Macro Accuracies*, Acc_s for seen classes and Acc_u for unseen classes, are measured. The key performance indicator is the harmonic mean $H = (2 \times Acc_s \times Acc_u) / (Acc_s + Acc_u)$, ensuring a balance between the two.

In Table 2, we compile results from key methods across the three groups for benchmarks ZS-IMGC task. Incorporating diverse class semantics generally yield better ZS-IMGC results, even with basic methods. For instance, Table 2 illustrates that some mapping-based methods (Frome et al., 2013; Chen et al., 2020a), employing straightforward linear mapping functions, achieve better performance by integrating various class semantics such as attributes, hierarchy, and names, compared to GCNZ[†], which relies solely on class hierarchy. Significantly, enriching GCNZ[†] with more class semantics, markedly enhances its effectiveness (GCNZ[‡]).

A.4.2 Fake News Detection

FND, also termed Rumor Detection, addresses the proliferation of misleading multimedia content on social media to ensure the dissemination of trust-

worthy information. Unlike standard text classification, FND challenges involve discerning falsehoods across diverse subjects. While traditional deep learning approaches in FND emphasize text, they frequently neglect the significance of visual content and background knowledge. Thus, a comprehensive integration of text, visuals, and knowledge is essential for precise FND.

KMGCN (Wang et al., 2020d) employs Entity Linking to map entities from social media posts to concepts from Probase (Wu et al., 2012) and YAGO KGs. It constructs a graph with post words as nodes and incorporates visual words from images (detected by a pre-trained object detector (Redmon and Farhadi, 2018)), weighting edges with Point-wise Mutual Information to emphasize word correlations. A GCN is then utilized to model semantic interactions, employing global mean pooling for the final binary classification of multimedia posts. Extending these insights, KMAGCN (Qian et al., 2021) integrates the visual modality through a late fusion paradigm, employing feature-level attention to more accurately delineate the interplay between visual and textual content. As a dual-consistency network, KDCN (Sun et al., 2023c) identifies inconsistencies in both cross-modal and content-knowledge aspects with Freebase as the reference KG. It finds that entities in rumor posts are more distantly connected within KGs compared to non-rumors, providing a clear marker for distinction. EmoKnow (Zhang et al., 2023d) advances COVID-19 FND by incorporating WiKi-Data5M (Wang et al., 2021) as an external knowledge source. It uses PLMs for text analysis, extracts emotion features, and identifies relevant linked entities, utilizing TransE (Bordes et al., 2013) for entity representation, with an MLP-based classifier to combine these multi-modal inputs.

Discussion 5 *The progression of LLMs is transforming many classification tasks into ones focused on inference and directive question-answering, emphasizing the crucial role of selecting knowledge sources. Moreover, given the frequent association of fake news with political content, their urgency of timely news highlights the need for conducting research on knowledge updating and lifelong learning in FND.*

A.4.3 Movie Genre Classification

MMGC models integrate visual, textual, and metadata information to predict movie genres, representing each genre as an element in a binary vector for

multi-genre classification of movies.

Traditional methods (Bain et al., 2020; Huang et al., 2020b) primarily rely on features extracted from images and texts alone. A recent work, IDKG (Li et al., 2023a), integrates a domain-specific MMKG created from metadata fields such as titles, genres, casts, and directors. Its motivation stems from recognizing the relational patterns in metadata, like the tendency for “Nolan to direct science fiction movies” or “Emma to not often feature in comedies”. It use a translation model to merge KG embeddings with other modality features, guided by attention mechanism.

Discussion 6 *The success of a domain-specific KG largely depends on the metadata’s quality and completeness, where addressing scalability is vital especially for large movie datasets. Additionally, there is scope for creating interactive and personalized Movie Genre Classification systems. By integrating user feedback and preferences, the system can be tailored to individual tastes, offering personalized genre suggestions. Techniques such as reinforcement learning and user modeling could be utilized to customize the genre classification process, thus further enhancing user experience and satisfaction.*

A.5 Supplement for KG-driven Multi-modal Generation Tasks

Some of the reasoning methods previously discussed, including those used in VQA tasks, are based on generative approaches. This section focuses on tasks where content generation is **strictly necessary** for task completion.

A.5.1 Scene Graph Generation

Introduced by Johnson et al. (2015), Scene Graphs (SGs) form a crucial data structure for scene understanding, cataloging object instances within a scene and delineating their interrelationships. These instances, ranging from people to places and objects, are described through attributes like shape, color, and pose (Chang et al., 2023). The relationships between these instances, often action-based or spatial, are expressed as (*subject*, *predicate*, *object*) triplets, paralleling the (*h*, *r*, *t*) and (*e*, *a*, *v*) triplets in KGs. Scene Graph Generation (SGG) serves as an intermediary task, unlike other multi-modal tasks with specific end goals, providing enhanced understanding and reasoning to support downstream tasks (Huang et al., 2023b; Koner et al., 2021; Yu et al., 2021).

Grasping all relationships in SGG training data is challenging, yet crucial, and leveraging prior knowledge significantly aids in effectively learning relationship representations from limited data, thereby enhancing detection, recognition, and overall accuracy of SGG. One effective approach is using **language priors**. By leveraging semantic word embeddings, these priors adjust relationship prediction probabilities, thus augmenting visual relationship identification. For example, even with infrequent occurrences in training data, like interactions between *people* and *elephants*, language priors can assist the inference of similar relationships, such as “a person riding an elephant”, by studying more common examples like “a person riding a horse” (Chang et al., 2023). This also helps mitigate the long tail effect in visual relationships (He et al., 2020). Another approach involves **statistical priors**, leveraging the structural regularity inherent in visual scenes, as highlighted in (Zellers et al., 2018). These priors capitalize on typical object-relation statistical correlations, such as “people wearing shoes” or “mountains being near water”.

Several works adopt the KG representation learning techniques into SGG scenario. For example, RLSV (Wan et al., 2018) uses existing SGs and images to predict new relationships between entities, targeting SG completion and blending KG embedding methods with SG characteristics in a structural-visual embedding model. Yu et al. (2022) improve zero-shot performance in SGG by constructing a KG from training set SG triples, distinguishing existing (non-zero-shot) and missing (zero-shot) edges. They train a KG Embedding model to complete the graph and fills these missing edges, thereby integrating zero-shot triples similarly to their non-zero-shot counterparts. GLAT (Zareian et al., 2020b) separates perception and commonsense into two models, training on annotated SGs with a BERT-like masking approach (akin to KG pre-training (Yao et al., 2019)) for element prediction. This method, when added to any SGG model, can rectify errors in SGs by harnessing the synergy of perception and commonsense.

Some SGG studies (Chen et al., 2019; Gu et al., 2019; Zareian et al., 2020a; Khan et al., 2022b; Chen et al., 2023f; Lu et al., 2023) also employ KGs for **triple prediction**, utilizing them to generate rich and expressive SGs. Specifically, KERN (Chen et al., 2019) leverages structured KGs to capture statistical correlations between object

pairs and relationships, boosting SGG by contextualizing and stabilizing relationship predictions to address distribution imbalances. Gu et al. (2019) utilize a knowledge-based module to identify relevant ConceptNet entities and retrieve commonsense facts, each assigned a weight indicating its real-world prevalence to filter candidate triples. Then a Dynamic Memory Network (Kumar et al., 2016) is applied for multi-hop reasoning on these facts, enabling inference of the most probable SG triples. GB-Net (Zareian et al., 2020a) views SGs as image-conditioned versions of commonsense KGs, shifting the focus from traditional entity and predicate classification to linking these two graph types. Utilizing a graph-based neural network, GB-Net iteratively propagates and refines information between and within both graphs, effectively bridging scene and commonsense knowledge. Khan et al. (2022b) enrich SGs using CSKG (Ilievski et al., 2021), a substantial commonsense KG repository. By employing graph embeddings to assess the similarity of object nodes, their approach enables graph refinement and enrichment as shown in Fig. 6. This upgrades SGG with additional information on objects’ spatial proximity and potential interactions derived from external knowledge, improving higher-level reasoning and mitigating some missed or incorrect predictions made during SGG. Explicit Ontological Adjustment framework (Chen et al., 2023f) mitigates predicate biases using knowledge priors from ConceptNet and Wikidata, refining relationship detection by integrating an edge matrix from the KG into a GNN model. Tian et al. (2023) add a branch for independent label confidence estimation in SGG network, which assesses the difficulty of visual recognition. This branch balances the need for commonsense knowledge in diverse scenes, especially for relations like “throwing” that require supplementary knowledge compared to more straightforward spatial relations like “sitting on”.

Discussion 7 *In the realm of SGG, KGs are instrumental in mitigating relationship bias and the long-tail phenomenon within training sets, serving as a form of refinement. However, existing SGG methods still face challenges in complex scenarios where the spatial distance between objects may be significant enough to disregard potential interactions. Enhancing scene graph integrity could be achieved by incorporating larger-scale images to recognize relationships between distantly located*

objects (Hossain et al., 2019). Moreover, extending SGG to identify human interactions, both in terms of object relations and social dynamics, would enrich scene comprehension and widen its practical uses, thereby aiding in the development of MMKG. Additionally, leveraging structured features from SGs in training LLMs represents a promising strategy to boost multi-modal learning, capitalizing on the combined strengths of SGs, KGs, and language priors.

A.5.2 Image Captioning

Image Captioning (IC) (Hossain et al., 2019) is a pivotal multi-modal learning task, aiming to describe images in natural language. In IC, KGs can provide essential prior knowledge, including commonsense semantic correlations and constraints among objects, guiding the construction of semantic graphs for meaningful caption generation, even when certain elements are not visually present (Fig. 6). Furthermore, since each image in the training data typically comes with only a few ground truth captions, models often lack the cues necessary to uncover implicit intentions. KGs can significantly bridge this gap by offering essential fact-checking support.

Rule-based methods (Aditya et al., 2015; Lu et al., 2018a; Huang et al., 2020a) primarily incorporate KG knowledge into caption models through Entity Linking and symbolic rules, often supplemented by inter-concept co-occurrence scores. Aditya et al. (2015) pioneer the application of KG into IC, identifying relevant events from a KG based on detected visual concepts, then constructing a Scene Description Graph (SDG) with predefined rules, from which captions are generated using NLG tools. Lu et al. (2018a) use a CNN-LSTM model to create a template caption from the input image, followed by employing a KG-based collective inference algorithm to populate the template with specific named entities, sourced from hashtags. Instead of directly integrating semantic knowledge into the neural network layers, Huang et al. (2020a) input retrieved triples from ConceptNet during the word generation stage for next-word prediction, augmenting the probabilities of potential words identified within the semantic knowledge corpus.

Embedding-based methods (Hou et al., 2019, 2020; Li and Jiang, 2019; Mogadala et al., 2020; Zhang et al., 2021c; Zhong and Wang, 2023) typ-

ically employ networks such as GNNs or RNNs to efficiently encode retrieved knowledge as vectors, subsequently incorporating these vectors into the caption generation process. Hou et al. (2019, 2020) utilize human commonsense knowledge to support object relationship reasoning in IC, avoiding the need for pre-trained detectors. Using Visual Genome as an external KG, they map densely sampled regions from images into low-dimensional vectors, and then, guided by the KG, form a temporary semantic graph. This graph enhances GNN-based relational reasoning for captioning and iteratively refines the KG itself. CNet-NIC (Zhou et al., 2019) connects ConceptNet entries with identified image objects to enrich descriptions and infer non-explicit visual information. This method enhances the semantic depth of object recognition module outputs, integrating the embeddings of knowledge terms and image features to initialize a RNN for IC generation. Interpret-IC (Mogadala et al., 2020) selects local objects in an image based on human-interpretable rules, ensuring captions reflect only those objects of human interest. During training, entities not present in standard captions are masked to align the model with human preferences. Zhang et al. (2020) employ a chest abnormality KG with prior **chest X-ray knowledge** to support radiology report generation. In this KG, entity features are initialized with CNN-extracted features of frontal and lateral chest X-ray images, where the application of GCN mean pooling yields graph-level features that contributes to generating radiology reports. Zhao et al. (2021b) utilizes an MMKG that associates visual objects with named entities for IC, incorporating external **multi-modal knowledge** sourced from Wikipedia and Google Images. This MMKG, once processed through a GAT (Velickovic et al., 2018), feeds its final layer’s output into a Transformer decoder which enables entity-aware caption generation. Nikiforova et al. (2022) propose a dataset from the Geograph project¹², including geographic coordinates of photo locations. Concentrating on **encyclopedic knowledge**, they extract facts from DBpedia and use a retriever to prioritize facts for possible caption inclusion. These knowledge triples, combined with the image and geographic context, are then utilized in an encoder-decoder IC pipeline.

¹²<http://www.geograph.org.uk/>

A.5.3 Visual Storytelling

Visual Storytelling (VST) transcends traditional Image Captioning by transforming a series of pictures into a cohesive narrative, demanding both the recognition of contexts within and across images and overcoming narrative monotony. KGs are crucial here, enhancing story diversity, rationality, and coherence.

KG-Story (Hsu et al., 2020) links concept terms from images across scenes using background KGs like FrameNet (Baker et al., 1998) and Visual Genome, refined by a PLM for sequential image storytelling. Yang et al. (2019a) develop a vision-aware directional encoding schema, integrating essential commonsense knowledge from ConceptNet for concept in each image. The enhanced snapshot representations, augmented with attentive knowledge, processed in a GRU-based framework for final VST. Building upon this, MCSM (Chen et al., 2021a) applies pruning rules and two concept selection modules to refine commonsense knowledge facts and facilitate sentence generation for each image using a visual-adaptor-equipped BART (Lewis et al., 2020). Further, PR-VIST (Hsu et al., 2021) represents image sequences as story graphs to identify the best storyline path and develop a discriminator model for outputting story quality scores, aligning the narratives with human preferences. IRW (Xu et al., 2021) utilizes imaginary key concepts derived from each image for entity mention detection, retrieving candidate fact triples from ConceptNet to form a sub-KG. This sub-KG, along with the constructed scene and event graph for each image, is integrated using separate GCNs, adaptively contributing to the VST process. KAGS (Li et al., 2023c) involves a knowledge-enriched attention network with a group-wise semantic model for globally consistent VST guidance.

Discussion 8 *The advent of Multi-modal LLMs (MLLMs) has enriched the knowledge embedded in pre-trained models, often diminishing the need for KGs to supply coarse-grained commonsense knowledge for those IC and VST tasks. This development highlights the need for KGs offering finer-grained or specific commonsense knowledge to address model hallucination issues. Moreover, for VST task, maintaining coherence between pictures and scenes is essential, where KGs are vital for linking disparate scenes and enriching scene transitions with background knowledge. Several methods have innovated with data-centric KG enhance-*

ments, such as deriving background KGs from story collections in training corpora (Hsu et al., 2021), or creating event graphs through image selection from the training set that resemble the query image, subsequently using Information Extraction tools to construct events for each sentence associated with an image (Xu et al., 2021). While these strategies are pioneering, they introduce challenges in ensuring equitable model comparisons due to varied dependencies on external knowledge sources, suggesting the need for separate evaluation of such data-centric methods.

A.5.4 Conditional Text-to-Image Generation

Conditional Text-to-Image Generation (cIG) aims to transform textual descriptions into visually realistic images, where KGs could supply detailed prior knowledge and commonsense elements not originally present in the datasets. LeicaGAN (Qiao et al., 2019) establishes a shared semantic space that enables text embeddings to convey visual information, by integrating a text-image encoder for semantic, texture, and color understanding, alongside a text-mask encoder for shaping layout through segmentation masks. During the image imagination phase, it merges the outputs of these encoders with added Gaussian noise to enhance diversity. Here, a cascaded attentive generator produces detailed and realistic images, ensuring semantic and visual coherence through adversarial learning. Many following works (Cheng et al., 2020; Jun Cheng and Fuxiang Wu and Yanling Tian and Lei Wang and Dapeng Tao, 2022; Liu et al., 2023b) treat image-caption pairs in training sets as KB entries, enriching captions by selecting and refining relevant items from this KB, thereby aiding in feature extraction and enabling more accurate cIG. Concretely, KnHiGAN (Ge et al., 2021) and AttRiGAN (Zhu et al., 2021) present a Knowledge-enhanced Hierarchical GAN, employing a KG to enrich text descriptions for detailed generative input. This task-specific KG is constructed from training sample attributes, formatted in RDF triples (Geng et al., 2021a, 2023). For 3D cIG, T2TD (Nie et al., 2023) involves a text-3D KG that correlates text with 3D shapes and textual attributes, utilizing these elements as prior knowledge. During 3D generation, it retrieves the knowledge based on text descriptions and employs a causal module to select shape information relevant to the text.

Discussion 9 While metrics like the Inception score (Salimans et al., 2016) and R-precision (Xu

et al., 2018) are commonly used for evaluating the diversity of generated images and the semantic consistency between input text and generated images, current evaluation methods for generated images still lack critical assessment at the knowledge and commonsense level (Huang et al., 2023b). Bridging this gap presents a critical direction for future research.

A.6 KG-driven Multi-modal Retrieval Tasks

Definition 2 *KG-aware Retrieval* aims to utilize textual descriptions (x^l) for ranking similar visual images (x^v), or vice versa, including the sorting and retrieval of all relevant images or region proposals within an image. Utilizing a background KG $\mathcal{G}$, this approach transcends mere appearance-based retrieval by incorporating non-visual attributes, striving for a human-level semantic understanding, especially in scenarios lacking precise targets.

A.6.1 Cross-Modal Retrieval

Cross-Modal Retrieval (CMR) focuses on fetching data across different modalities, such as images, text, audio, or video, in response to a query from another modality. Specifically, this section explores Image-Text Retrieval, aiming to identify semantically similar instances across visual and textual modalities.

Image-Text Matching (ITM) vs. Image-Text Retrieval (ITR): ITM and ITR are closely related yet differ mainly in their application: ITM evaluates relevance between an image and text, often used in image-caption correspondence (Huang et al., 2023b; Gómez-Pérez and Ortega, 2019), while ITR focuses on finding relevant matches in larger datasets based on textual or visual queries, crucial for visual search engines, digital asset management, and automated content generation (Cao et al., 2022a). Both ITM and ITR leverage similar underlying technologies, metrics, and datasets such as Flickr30k (Young et al., 2014) and MSCOCO (Lin et al., 2014), which feature extensive labeled images with captions. In cross-modal pre-training, ITM serves as a foundational task, honing the model’s ability to semantically correlate images and text, thereby improving its effectiveness in ITR (Zhang et al., 2021b; Li et al., 2022a, 2023b). This pre-training ranges from coarse-grained matching (assessing general semantic relatedness) to fine-grained matching (aligning specific image regions with text). Such granularity

enhances the nuanced understanding and retrieval capabilities for pre-trained models, bridging the gap between general-purpose models and the specific demands of ITR tasks.

Early CMR research often overlook long-tail and occluded semantic concepts in images (Yang et al., 2021; Cao et al., 2022a). Recent advancements (Shi et al., 2019; Wang et al., 2020a, 2022a; Li et al., 2023d; Yang et al., 2023a) tend to fix this by leveraging knowledge from frequently co-occurring concept pairs in Visual Genome’s scene graphs (Krishna et al., 2017) or image captioning corpus. They create Scene Concept Graphs (SCGs) using heuristic or rule-based tools such as language parsers (Anderson et al., 2016), aiming to capture fine-grained details. Shi et al. (2019) initially identify broad concepts, further refined into detailed ones via SCG’s co-occurrence relationships, followed by a concept prediction module for accurate labeling. EKDM (Yang et al., 2023a) employs an iterative concept filtering module that progressively incorporates candidate concepts into a static global representation in a dynamic manner, which uses the significance scores of these concepts to set the fusion order, integrating higher-scored concepts first. CVSE (Wang et al., 2020a, 2022a) utilizes a GNN for semantic correlation propagation in SCG, enriching concept representations with commonsense knowledge via weighted embedding summation. A confidence scaling function is introduced to mitigate long-tail distribution challenges. CSRC (Li et al., 2023d) further employs a multi-head self-attention mechanism to selectively focus on deeper conceptual emphasis, while MACK (Huang et al., 2022) eliminates the need for paired domain data during training.

Note that the background KGs in these works are typically derived from large-scale multi-modal datasets, rather than directly utilizing public KGs. However, **a reliance on mere word co-occurrences for entity similarities can be misleading**, like wrongly linking “man” and “dog” due to frequent co-occurrences. Utilizing WordNet’s noun hierarchies helps distinguish such entities. Additionally, MMKGs could address this by capturing inter-modal co-occurrence relations, like temporal, causal, and logical connections. For instance, “washing” with “tap” or “cutting” with “knife” in image-text pairs enhances semantic understanding across modalities. Building on this perspective, Fig. 11 illustrates MMKG-based approach MKVSE (Feng et al., 2023), which en-

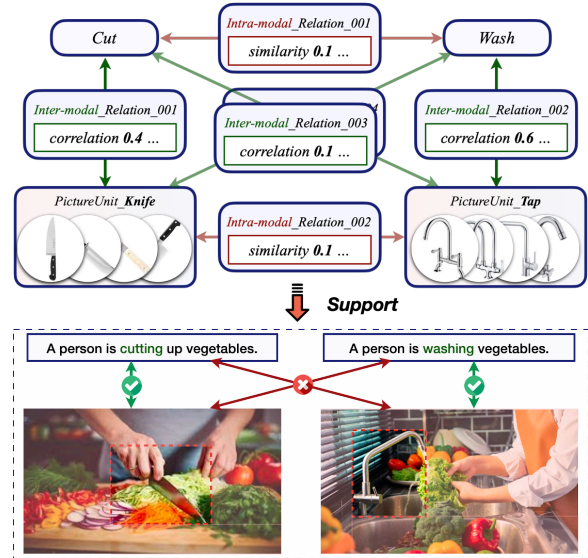


Figure 11: We illustrates the MMKG-supported Image-Text Retrieval process (Feng et al., 2023). For simplicity, all URI prefixes and certain relations (*sourceImg* and *targetImg*) from the *PictureRelation* (*Inter-modal_Relation* and *Intra-modal_Relation*) entity are omitted. This entity’s values indicate intra-modal path similarities or inter-modal co-occurrence correlations, essential for training a model (e.g., multi-modal GCN) to produce knowledgeable image or text representations. Note: In cases of multiple images within a picture unit, mean pooling is used for a unified feature representation.

hances image-text semantic connections, especially for images with indirect textual descriptions. It scores intra- and inter-modal relations in MMKGs using WordNet path similarity (calculated by NLTK (Bird et al., 2009)) and co-occurrence correlations, improving ITR through GNN-based embeddings. Moreover, Yang et al. (2023a) focus on a common limitation in visual concept modeling, where **varying spatial locations are often inaccurately linked by fixed relationships**, like “man-on-bike” for any proximity of “man” and “bike” in an image. They spatial information from a geometric graph (Monti et al., 2017) to discern spatial relations between image regions and employ a location CNN model to refine visual-semantic representations. EGE-CMP (Dong et al., 2023) is a entity-graph enhanced cross-modal pre-training framework that leverages entity knowledge extracted from captions instead of human labeling. It focuses on learning instance-level feature representations by infusing real semantic information into visual-text alignment, improving text-image cross-modal alignment.

Discussion 10 Current VLMs face challenges in fine-grained cross-modal semantic matching. Wang et al. (2023b) tackle this issue by using contrastive

learning for aligning entities from Visual Genome in ITR, enhancing cross-modal sensitivity with entity masking. We note that a shift towards knowledge-guided strategies rather than relying solely on co-occurrence in VLM training could significantly improve retrieval and matching of fine-grained, long-tail objects, potentially leading to advanced semantic grounding (Chen et al., 2023h) and wider applications. However, only limited studies (Feng et al., 2023) have considered the role of external knowledge like WordNet’s semantic structures. Besides, as discussed in §A.3.1, various types of KGs, including trivia, commonsense, scientific, and situational knowledge, offer unique and complementary insights for reasoning processes. But the prevalent focus on co-occurrence information captures a fraction of commonsense knowledge. Looking ahead, exploiting long-tail knowledge from diverse large-scale KBs holds significant potential for enhancing models’ generalization capabilities across various domains and real-world scenarios.

A.6.2 Visual Referring Expressions & Grounding

This section revisits KG-aware approaches in Visual Referring Expressions (also known as Phrase Grounding or Referring Expression Comprehension) and Visual Grounding. While CMR typically entails matching across diverse textual and visual contexts, VRE and VG focus on aligning fine-grained features within specific textual-visual pairs. From a certain point of view, these tasks are akin to adding an extra step of grounding answers in the conventional KG-based VQA, as illustrated in Fig. 9.

Visual Referring Expressions (VRE) vs. Visual Grounding (VG): VRE and VG (Qiao et al., 2021) integrate linguistic and visual information, differing in focus (Qiao et al., 2021): VRE identifies and localizes a specific image region that corresponds to a given textual expression, typically involving a detailed description of one object. Conversely, VG is about localizing various object regions linked to multiple noun phrases in a sentence, aiming to establish fine-grained alignment between vision and language. Despite these differences, both tasks require deep semantic language interpretation and manage ambiguities inherent in natural language and visual perception, relying on extensive annotated datasets. The line between VRE and VG often blurs in research, with some

approaches (Yang et al., 2019b; Sibe Yang and Guanbin Li and Yizhou Yu, 2021; Tang et al., 2023) merging their key aspects: VRE’s precise object localization and VG’s broad contextual analysis.

KAC Net (Chen et al., 2018) utilizes the knowledge from pre-trained fixed category detectors, essential for selecting relevant proposals and ensuring visual consistency, to filter out unrelated proposals in VG progress. Shi et al. (2022) tackle **zero-shot VRE**, where visual examples of queried object categories in the test set are not shown in the training set (i.e., open-vocabulary scene); they achieve this by dynamically building MMKGs using commonsense knowledge from WordNet and ConceptNet, combined with situational knowledge from Visual Genome. Query-derived entities, detected objects, and predefined relationships are integrated into these MMKGs, employing GCN for node representation and defining eight spatial relations to assist localization of noun phrases.

The **KB-Ref** dataset (Wang et al., 2020b) emphasizes commonsense knowledge, with its construction process inspired by the F-VQA dataset (Wang et al., 2018a), which involves creating a commonsense KG. Concretely, volunteers craft referring expressions for queried objects based on facts from this KG, deliberately avoiding the use of specific object names. Building upon the KB-Ref dataset, ECIFA (Wang et al., 2020b) introduces a multi-hop facts attention module from the KG and a matching module that utilizes expression-object scores for accurate grounding; CK-Transformer (Zhang et al., 2023f), leveraging the UNITER (Chen et al., 2020c) as its backbone, selects top-K retrieved facts from the KG for a given expression and visual region candidates, encoding these into multi-modal features to compute matching scores for each candidate. Bu et al. (2023) observe that knowledge-based Referring Expressions often consist of two segments: visual segments (e.g., “on the sofa” in Fig. 9), interpretable directly from visual content like color and shape, and knowledge segments (e.g., “used for sleeping” in Fig. 9), requiring additional information beyond visuals like function and non-visual attributes. To mitigate similarity bias, they introduce the SLCO network, which uses knowledge segments for category retrieval and visual segments for object grounding.

The **SK-VG** dataset (Chen et al., 2023g) targets at scene knowledge-guided VG, using movie scene images from the VCR dataset (Zellers et al., 2019). Designed to promote reasoning beyond mere image

content, SK-VG employs a detailed two-stage annotation process: firstly, generating story descriptions for each image, and secondly, crafting referring expressions tied to these stories and images, accompanied by object bounding box annotations. These annotations are crafted to ensure knowledge relevance to the scene context, uniqueness for accurate object identification, and diversity in both lexical use and objects. [Chen et al. \(2023g\)](#) further provide two benchmarking algorithms: a one-stage approach that embeds knowledge into image features prior to query interaction, and a two-stage method that extracts features from images and text, subsequently employing structured linguistic data for computing region-entity similarity.

Discussion 11 *A good VRE and VG system can benefit various downstream tasks such as VQA, CMR, and IMGC. [Chen et al. \(2023h\)](#) develop a cross-modal semantic grounding network for ZS-IMGC, aimed at disentangling semantic attributes from images via a self-supervised method. This technique bridges knowledge from PLMs to visual models without needing region-attribute supervision. By leveraging AWA2-KG ([Geng et al., 2023](#)) for fine-grained labeling, it connects species to their attributes (e.g., “zebra” to “striped”) and uses KG serialization to blend structured knowledge into cross-modal grounding. The network also incorporates attribute-level contrastive learning to tackle attribute imbalance and co-occurrence, thus refining the distinction of fine-grained visual features across images from both seen and unseen classes. This highlights the value of KGs in Visual Grounding tasks, serving as a natural knowledge organizer and a conduit for transferring VG principles to related tasks without specialized annotations.*

A.7 Supplement for KG-driven Multi-modal Pre-training

In this section, our primary focus is on pre-training definitions related to Transformer-based models, aligning with the current mainstream discourse in AI community. Other paradigms, such as Poincaré embedding pre-training ([Xu et al., 2020](#)), are not covered in this discussion.

We highlight that multi-modal reasoning and generation tasks often require an extensive range of specialized knowledge, typically involving long-tail information that goes beyond everyday experiences. KGs are crucial in these scenarios, serving

as structured repositories for such diverse knowledge. However, there exists a notable gap between KGs and multi-modal tasks, as current methods frequently depend on indirect approaches like modal transformation for knowledge representation, retrieval, and interaction in multi-modal contexts. A significant challenge arises in tasks requiring visual common sense, where models may falter due to limited cross-modal alignment capabilities, leading to multi-modal hallucinations as evidenced in Fig. 7. Recent works ([Zha et al., 2023](#)) demonstrate that MMKGs can effectively bridge this gap, enhancing the potential of multi-modal methods and offering a robust solution to multi-modal hallucinations in the era of LLMs. Specifically, [Zha et al. \(2023\)](#) introduce M²ConceptBase, a multi-modal conceptual MMKG. They develop a pipeline using M²ConceptBase to improve knowledge-based VQA performance by retrieving multi-modal concept descriptions and crafting instructions to refine answers with MLLMs.

Structure Knowledge aware Pre-training. The integration of structured knowledge into multi-modal content understanding has gradually gained momentum, drawing inspiration from advancements in the NLP field. KM-BART ([Xing et al., 2021](#)) adapts the BART ([Lewis et al., 2020](#)) model to multi-modal tasks by incorporating a pre-trained visual feature extractor. It tackles knowledge-based commonsense generation by using COMET ([Bosse-lut et al., 2019](#)) to augment image-caption datasets with commonsense context. The enriched datasets, combined with a next-token prediction target, empower KM-BART to deduce events and character intentions from image-text pairs. ERNIE-ViL ([Yu et al., 2021](#)) incorporates Scene Graph (SG) knowledge into a VLM, enhancing visual scene comprehension by adding SG completion and prediction tasks (covering objects, attributes, and relationships) during its multi-modal pre-training stage. ROSITA ([Cui et al., 2021](#)) strengthens semantic alignments across visual and language modalities by employing a unified SG shared between the input image and text. Existing VLMs often struggle with Image-Text Matching tasks that demand an understanding of reversed roles or actions, evident in scenarios like “An astronaut rides a horse” versus “A horse rides an astronaut” (refer to § A.6.1). To tackle this, Structure-CLIP ([Huang et al., 2023b](#)) improves structured multi-modal representation learning by leveraging SGs to generate semantic

negative examples.

Knowledge Graph aware Pre-training. KG-Transformer (Zhang et al., 2023b) is pre-trained on KGs including WN18RR (Sun et al., 2019b), FB15k-237 (Toutanova and Chen, 2015), and CoDEX (Safavi and Koutra, 2020), with pre-training objectives like masked relation/entity prediction and entity pair prediction. This model can be applied to ZS-IMGC, framing the task as determining the match score between input images and target classes. It undergoes fine-tuning using AwA-KG (Geng et al., 2023), with a pre-trained ResNet serving as the vision encoder and image representations further transformed through a trainable matrix.

Discussion 12 *Current KG-equipped VLMs primarily use triple contexts to enhance multi-modal data, with a few examples, like KGTransformer (Zhang et al., 2023b), incorporating KG’s structural information into pre-training. However, its application is limited to Zero-shot Image Classification, using a uni-modal approach during pre-training. Future research in this domain can focus on four key areas: Firstly, scaling up KG to exploit its rich knowledge and structural traits, rethinking the long-tail phenomenon in multi-modal pre-training data and expanding the knowledge scope to involve world knowledge. Secondly, the integration of MMKGs, which are further discussed in § 4. Third, exploring unique pre-training paradigms suited for (MM)KGs to fully harness the value of structured knowledge in multi-modal pre-training. Fourth, extending to more downstream tasks to align with the latest advancements in AGI, utilizing MLLMs (Liu et al., 2023c).*

A.8 Supplement for Future Directions

Focusing solely on the benefits that traditional KGs bring to multi-modal tasks can be inherently limiting due to the restricted scope of knowledge captured in single-modality KGs. In evaluating KG-aware multi-modal tasks, it’s crucial to discern the unique advantages of multi-modal knowledge, especially compared to large-scale textual or multi-modal corpora. Specifically for image modalities, MMKGs can be categorized into two types: A-MMKGs where images serve as entity attributes, and N-MMKGs where images are independent entities with their own relationships (Zhu et al., 2022). A pivotal question is whether structured (MM)KGs offer irreplaceable benefits that maximize their po-

tential. Additionally, we should consider whether non-LLM models augmented by (MM)KG can rival or outperform MLLMs in specific tasks, providing compelling reasons to support the future development.

A.8.1 KG4MML Tasks.

Multi-modal Content Generation. Current applications of MMKGs in multi-modal content generation are quite limited. Most existing efforts only integrate KGs to supplement additional context beyond datasets or to connect different visual scenes. Future developments should aim for larger, more detailed MMKGs to employ multi-modal structural data in training, fostering more controlled and logically coherent generation and mitigating hallucinations.

Multi-modal Task Integration. Different domains currently evolve independently with limited cross-interaction. In Cross-Modal Retrieval (CMR), (MM)KGs are widely employed for information enhancement, whereas in knowledge-based VQA, the focus is mainly on dense vector retrieval and modality conversion techniques. This highlights the potential for future advancements like integrating KG-based CMR methods into KG-based VQA. In a similar vein, generation tasks can enhance retrieval, reasoning, and discrimination, with knowledge-enhanced discrimination tasks playing a key role in refining answers for other tasks. As knowledge-intensive multi-modal tasks gain prominence, merging these distinct domains with (MM)KG at the core will become crucial.

Challenges in Scaling MMKG for Multi-modal Tasks. MMKG-driven tasks often emphasize retrieval-related activities, leveraging the natural database-like capabilities of MMKGs. However, the utilization of large-scale MMKGs in varied tasks, especially reasoning, is still nascent with limited exploratory studies. For example, Zha et al. (2023) enhance knowledge-based VQA by employing multi-modal concept descriptions and integrating MLLMs for refined answers. Nevertheless, these methods only use MMKGs as “key:value” based retrieval databases, not fully leveraging their multi-modal structured capabilities.

The constrained utilization of MMKGs in diverse tasks can be attributed to several factors:

- **Non-Uniform Organization and Ontology of MMKGs:** Current MMKGs, lacking a stan-

dardized format, vary significantly in their focal points and the knowledge domains they cover for each downstream task. Predominantly, MMKGs cater to encyclopedic or trivia knowledge (Gong et al., 2023; Zhang et al., 2023a; Wu et al., 2023c; Zha et al., 2023), with commonsense and scientific related MMKGs (Wang et al., 2023c; Lee et al., 2023) being notably scarce. Moreover, the “non-visualizable” nature of some abstract knowledge components restricts their practical application (Jiang et al., 2022; Wu et al., 2023c).

- **Storage and Processing Overheads:** The substantial storage space requirements and extended processing times for large-scale MMKGs hinder their extensive adoption. Conversely, small-scale MMKGs frequently offer limited value for cross-task generalization.
- **Data Timeliness and Completeness Issues in MMKGs** heightens the risk of multi-modal hallucinations.
- **Comparative Advantages of LLMs and MLLMs:** LLMs and MLLMs excel in generalizability and AGI potential across various domains (Zhang et al., 2024), complementing the interpretability and editing flexibility of MMKGs. While MMKGs bring unique value, their development, maintenance, and application also involve certain costs. The evolving feedback from downstream tasks will continue to shape the industry’s perspective on their respective roles and potentials.

Unlocking the Potential of Large-Scale MMKGs for Multi-Modal Tasks.

- **Integration with Non-text Modalities:** Future downstream tasks driven by large-scale MMKGs can integrate methods from current KG-driven VQA methods, placing equal emphasis on non-textual modalities. This may further involve using modality projection or adapters for cross-modal alignment (Li et al., 2023e; Long et al., 2023), along with multi-modal GNN methods (Yoon et al., 2023) and modal feature decoupling techniques to enrich the granularity and hierarchy of multi-modal information (Chen et al., 2023h).
- **Rich Semantic MMKG Construction:** MMKG data can transcend traditional specialized or general formats. By developing task-specific pipelines, multi-modal datasets can be converted into MMKGs with enhanced

semantics, using existing KGs as foundational references or bridges. This process can not only augments MLLM training with structured multi-modal input but also enriches the MMKG community with valuable, semantically rich datasets.

- **Reconstruction of Multi-Modal Tasks with LLM:** Combining LLM’s text understanding and generation capabilities, multi-modal tasks can be restructured. Transforming KG-driven multi-modal tasks into in-MMKG-tasks, such as Multi-modal Knowledge Graph Construction, Multi-modal Entity Alignment, can enhance domain integration. There are already some attempts in this direction (Pahuja et al., 2024).

A.8.2 Large Language Models.

The academic definition of LLMs, often associated with models possessing extensive parameters such as LLaMA-7B (Touvron et al., 2023), remains broad. These models’ emergent abilities and Zero-shot Learning capabilities edge them closer to achieving AGI, underscoring their importance in NLP and multi-modal domains.

1. Fine-Tuning:

- MMKGs provide a rich source of structured multi-modal data for Supervised Fine-Tuning (SFT) of Multi-modal LLMs (MLLMs), Training techniques effective for MMKGs in VLMs can also be applied to MLLMs, especially in domain-specific applications (Zheng et al., 2024; Zhang et al., 2023c).
- Leverage self-instructing techniques to autonomously evolve and generate multi-grained, multi-modal instructional data (Wang et al., 2023d; Xu et al., 2023a; Du et al., 2023; Yona et al., 2024)
- Furthermore, MMKG data can be utilized to further explore the concept of multi-modal reversal curse (Lv et al., 2023), where the ordering of knowledge entities in training data influences model comprehension, potentially limiting the model’s understanding.

2. Hallucination:

- As LLMs rapidly advance, the risk of generating seemingly authentic but factually inaccurate web content is increasing. This phenomenon, known as *hallucination* (Agrawal et al., 2023), often arises

from outdated or incorrect training encountered during the model training process, or from the frequent co-occurrence bindings of objects, affecting both LLMs and MLLMs (Liu et al., 2024).

- Hallucination in MLLMs: (Huang et al., 2023a; Tong et al., 2024; Liu et al., 2024)
- To combat this, LAMM (Yin et al., 2023b) incorporates 42K KG facts from Wikipedia and leveraged the Bamboo dataset (Zhang et al., 2022c) to refine commonsense knowledge in Q&A, underscoring the role of quality (MM)KGs in mitigating LLM hallucinations (Agrawal et al., 2023; Xu et al., 2023c). Developing robust hallucination detectors (Chen et al., 2023e; Mishra et al., 2024) is crucial for identifying and curbing errors in LLM outputs. Future efforts could focus on pairing MMKGs with detection methods to improve multi-modal task precision and leveraging (MM)KGs for knowledge-aware statement rewriting to diminish factual hallucinations in LLM reasoning (Guan et al., 2023; Wang et al., 2023a).

3. Agent:

- Multi-agent Collaboration (Xu et al., 2023b; Xiao et al., 2023; Lu et al., 2024), simulating human cognitive processes, can dissect VQA reasoning paths and engage multiple (M)LLMs in collective problem-solving (Wang et al., 2023e; Qiao et al., 2024). In this framework, KGs can initialize agent personalities (Mao et al., 2023; Tu et al., 2023), providing a structured basis for intuitively designing character brains, enriching the interaction between agents and enhancing their collective reasoning capabilities.
- Chain-of-thought (CoT) reasoning (Wei et al., 2022) significantly improves LLMs' complex reasoning abilities by incorporating intermediate reasoning steps. This progress has catalyzed the emergence of various KG-focused applications (Park et al., 2023; Sun et al., 2023a). For example, Sun et al. (2023a) demonstrate how LLMs can be used to interactively navigate KGs to extract knowledge for reasoning. Their Think-on-Graph (ToG) approach utilizes beam search to identify effective reasoning

paths within KGs. Merging these innovations with MMKGs promises to expand the scope of tasks, especially in improving the ability of models to interpret and interact with diverse data types, such as images and text (Mondal et al., 2024). This integration moves us closer to achieving human-like multi-modal proficiency and paves the way for advanced machine intelligence.

4. Retrieval Augmented Generation:

- Retrieval Augmented Generation (RAG) (Ovadia et al., 2023) systems enhance (M)LLMs by incorporating long-tail knowledge beyond their parameter limits. However, excessive document retrieval can lead to contextually inappropriate answers (Barnett et al., 2024), increasing hallucination risks unless carefully designed prompts are used (Wang et al., 2023f). The high information density and structured organization in KGs can mitigate this issue.
- Moreover, MMKGs can further aid multi-modal RAG by using various modalities as anchors (Song et al., 2023a), offering more relevant and explanatorily powerful results than vector-based searches (Wu and Xie, 2023; Yu et al., 2023).