

Glider: Global and Local Instruction-Driven Expert Router

Anonymous ACL submission

Abstract

The development of performant pre-trained models has driven the advancement of routing-based expert models tailored to specific tasks. However, these methods often favor generalization over performance on held-in tasks. This limitation adversely impacts practical applicability, as real-world deployments require robust performance across both known and novel tasks. We observe that current token-level routing mechanisms neglect the global semantic context of the input task. To address this, we propose a novel method, **Global and Local Instruction Driven Expert Router (GLIDER)** that proposes a multi-scale routing mechanism, encompassing a semantic global router and a learned local router. The global router leverages recent LLMs’ semantic reasoning capabilities to generate task-specific instructions from the input query, guiding expert selection across all layers. This global guidance is complemented by a local router that facilitates token-level routing decisions within each module, enabling finer control and enhanced performance on unseen and challenging tasks. Our experiments using T5-based expert models for T0 and FLAN tasks demonstrate that GLIDER achieves substantially improved held-in performance while maintaining strong generalization on held-out tasks. Additionally, we perform ablations experiments to dive deeper into the components of GLIDER and plot routing distributions to show that GLIDER can effectively retrieve the correct expert for held-in tasks while also demonstrating compositional capabilities for held-out tasks. Our experiments highlight the importance of our multi-scale routing that leverages LLM-driven semantic reasoning for MoErging methods.

1 Introduction

The emergence of highly capable large language models (LLMs) has marked an increased attention in downstream task specialization. This spe-

cialization often leverages parameter-efficient fine-tuning (PEFT) techniques, such as LoRA (Hu et al., 2021), which introduce minimal trainable parameters (“adapters”) to adapt pre-trained LLMs for specific tasks. The compact size of these specialized PEFT modules enables easy sharing, which has led to the distribution of an evergrowing number of adapters on various platforms.

This proliferation of expert models, *i.e.* specialized adapters, has led to the development of methods for re-using such experts to improve performance or generalization (Muqeeth et al., 2024; Ostapenko et al., 2024; Huang et al., 2024a). Central to these approaches are routing mechanisms that adaptively select relevant experts for a particular task or query. These routing methods have been referred to as “Model MoErging” (Yadav et al., 2024) since they frequently share methodologies and ideas with mixture-of-experts (MoE) models (Shazeer et al., 2017; Fedus et al., 2022; Du et al., 2022) and model merging (Yadav et al., 2023b,a; Ilharco et al., 2022). However, MoE methods train experts jointly from scratch (Gupta et al., 2022) while MoErging utilizes a decentralized, community-sourced pool of pre-trained experts. Furthermore, it departs from traditional model merging techniques by dynamically and adaptively combining these experts, optimizing performance at the query or task level. MoErging methods offer three key advantages: (1) They support decentralized model development by reusing and routing among independently trained experts, reducing reliance on centralized resources. (2) They facilitate modular capability expansion and “transparency” in updates as they either add or modify specialized expert models. (3) They allow for compositional generalization by recombining fine-grained skills from various experts, extending the system’s abilities to new unseen tasks beyond the capabilities of the individual expert models.

Most MoErging (Chronopoulou et al., 2023;

Contributor: Training LoRA Experts & Task Vectors

Aggregator: Creating Router Function for Inference

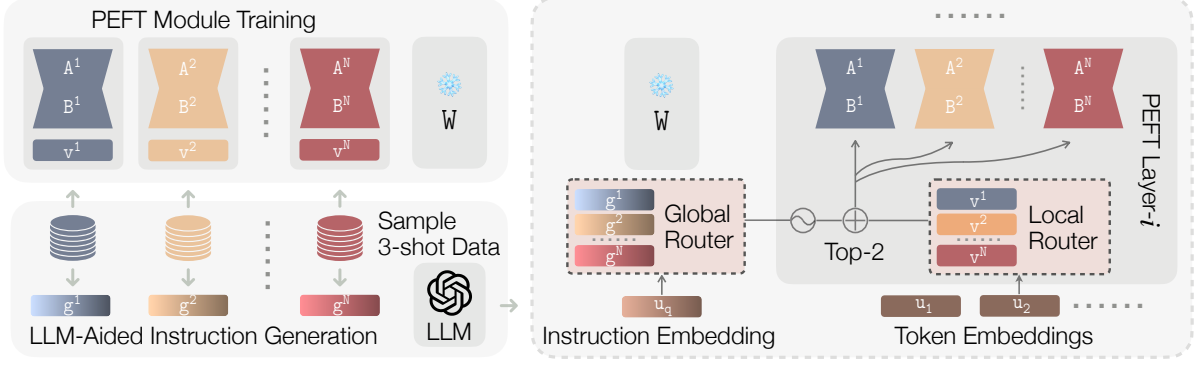


Figure 1: Overview of our method. Contributor (left): Each contributor utilizes local data to train several components: the PEFT module (comprising A_i and B_i), task vectors (v_i), and global routing vectors (g_i). For the latter, an LLM is employed to generate semantically-informed instructions based on 3 randomly selected examples, which are then embedded into g_i . Aggregator (right): The aggregator utilizes local and global task vectors to construct local routers [$\bar{v}^1; \dots; \bar{v}^N$] and a global router [$\bar{g}^1; \dots; \bar{g}^N$], respectively. For each query, the global router uses an LLM-generated instruction embedding to produce the global routing score. This score is then scaled and combined with the local routing score, enabling fine-grained control over expert selection.

Muqeeth et al., 2024; Zhao et al., 2024b) methods prioritize either known or unseen tasks, limiting real-world applicability where both are critical. Real-world queries often span domains and defy clean categorization into predefined task boundaries. For instance, translating and analyzing text requires collaboration between multiple experts rather than selecting a single specialized model. Current approaches struggle in such scenarios. Phatgoose demonstrates this tradeoff, excelling on unseen tasks but underperforming on known ones.

We hypothesize that this gap arises from the model’s token-level routing mechanism. We show that for the held-in tasks, the independent routing decisions at each layer, based solely on individual token embeddings, lack sufficient global context to retrieve the correct expert for all tokens at every module. This leads to suboptimal routing, which may propagate noise through the network, further hindering accurate expert utilization in deeper layers. This highlights a critical limitation of token-level approaches to handling held-in tasks, which hence falls short of the goal of building a routing system that seamlessly handles arbitrary queries. We believe that adding a global routing mechanism based on semantic task information can aid the token-level router for the correct retrieval of held-in tasks. Hence, we ask the question.

(Q) Can we leverage LLMs to generate semantics-aware task instructions to guide routing mechanism to facilitate both specialization and generalization?

This paper addresses the challenges by inves-

tigating the potential of leveraging the inherent reasoning and generalization capabilities of LLMs to guide the routing process in an MoE-like model composed of specialized LoRA modules. We introduce, **Global and Local Instruction Driven Expert Router (GLIDER)** that hinges on a multi-scale routing mechanism that contains both local and global routers to select top-2 expert models as shown in Figure 1. The global router leverages LLM-generated, semantics-aware instructions (see Appendix B.2) for each input query to score expert models. This high-level guidance is then complemented by a learned local router, which makes token-level routing decisions at each module, enabling fine-grained control and improving performance on the challenging held-out tasks. Through this framework, we highlight the crucial role of LLM reasoning in unlocking the compositional generalization capabilities of MoE models.

To test the effectiveness of our GLIDER method, we follow Phatgoose (Muqeeth et al., 2024) and use T5 models (Raffel et al., 2020) to create expert models for T0 held-in (Sanh et al., 2022) and FLAN tasks (Longpre et al., 2023) and test performance on T0 held-in & held-out (Sanh et al., 2022) and big-bench lite (BIG-bench authors, 2023) & hard tasks (Suzgun et al., 2022). Our key contributions and findings are:

- We introduce GLIDER, which employs LLM-guided multi-scale global and local attention. Our experiments show that GLIDER outperforms previous methods, significantly improving performance on held-in tasks (e.g. 6.6% over Phatgoose on T0 held-in) while also en-

hancing zero-shot held-out compositional generalization (*e.g.* 0.9% on T0 held-out).

- We find that without LLM assistance, MoE models underperform individual specialized models on held-in tasks by 8.2%. Incorporating semantic-aware instructions enables GLIDER to achieve comparable performance, demonstrating the LLM’s capacity to effectively infer task identity and guide module selection without explicit task labels.
- GLIDER also maintains strong performance on held-out tasks, showcasing its adaptability and generalization capabilities. Our work highlights the critical role of LLMs in enhancing MoE models’ compositional generalization, advancing the development of more robust and versatile AI systems capable of handling both familiar and novel tasks.

2 Related Works

The abundance of specialized expert models has spurred the development of techniques to leverage “experts” models for enhanced performance and generalization. Yadav et al. (2024) called such techniques as “MoErging”¹ methods which rely on adaptive routing mechanisms to select relevant experts for specific tasks or queries. These methods can be broadly classified into four categories based on the design of their routing mechanisms.

Embedding-Based Routing: This category encompasses methods that derive routing decisions from learned embeddings of expert training data. These methods typically compare a query embedding against the learned expert embeddings to determine the optimal routing path. Examples include AdapterSoup (Chronopoulou et al., 2023), Retrieval of Experts (Jang et al., 2023), LoraRetriever (Zhao et al., 2024b), Mo’LoRA (Maxine, 2023), the embedding-based approach of Airoboros (Durbin, 2024), and Dynamic Adapter Merging (Cheng et al., 2024).

Classifier-Based Routing: This category consists of methods that train a router to function as a classifier. This router is trained to predict the optimal routing path based on features extracted from expert datasets or unseen data. Representative methods in this category include Zooter (Lu et al., 2023), Branch-Train-Mix (Sukhbaatar et al., 2024),

Routing with Benchmark Datasets (Shnitzer et al., 2023), Routoo (Mohammadshahi et al., 2024), and RouteLLM (Ong et al., 2024). The key distinction between embedding-based and classifier-based routing lies in the router’s architecture and training methodology. While embedding-based routing often employs a nearest neighbor approach, classifier-based routing typically relies on logistic regression or analogous classification techniques.

Task-Specific Routing: This category focuses on methods tailored to enhance performance on specific target tasks. These methods learn a task-specific routing distribution over the target dataset to optimize performance for the given task. Methods include LoraHub (Huang et al., 2023), LoRA-Flow (Wang et al., 2024), AdapterFusion (Pfeiffer et al., 2021), π -Tuning (Wu et al., 2023), Co-LLM (Shen et al., 2024), Weight-Ensembling MoE (Tang et al., 2024), MoLE (Wu et al., 2024), MeteorA (Xu et al., 2024), PEMT (Lin et al., 2024), MixDA (Diao et al., 2023), and Twin-Merging (Lu et al., 2024).

Routerless Methods: This final category encompasses methods that do not rely on an explicitly trained router. Instead, these methods often employ alternative mechanisms, such as heuristics or rule-based systems, for routing decisions. Examples include Arrow (Ostapenko et al., 2024), PHAT-GOOSE (Muqeeth et al., 2024), the “ask an LLM” routing of Airoboros (Durbin, 2024) and LlamaIndex (Liu, 2024). Phatgoose and Arrow use only local routers, in contrast, GLIDER uses both local and global guidance for routing.

3 Problem Statement

In our work, we aim to build a routing mechanism capable of performing well on diverse queries from various tasks, including both seen and unseen tasks. For each query/token and module, this routing mechanism dynamically selects a model from a large pool of specialized expert models to achieve high performance. To facilitate modular development, we adopt a *contributor-aggregator* framework (Yadav et al., 2024) where individual contributors create specialized expert models from a generalist model for their respective tasks and distribute these models to others for public usage. The aggregator builds a routing mechanism over the expert models that shared by the contributor to direct queries to the most relevant experts. Follow-

¹See *e.g.* https://huggingface.co/spaces/open-llm-leaderboard/open_llm_leaderboard

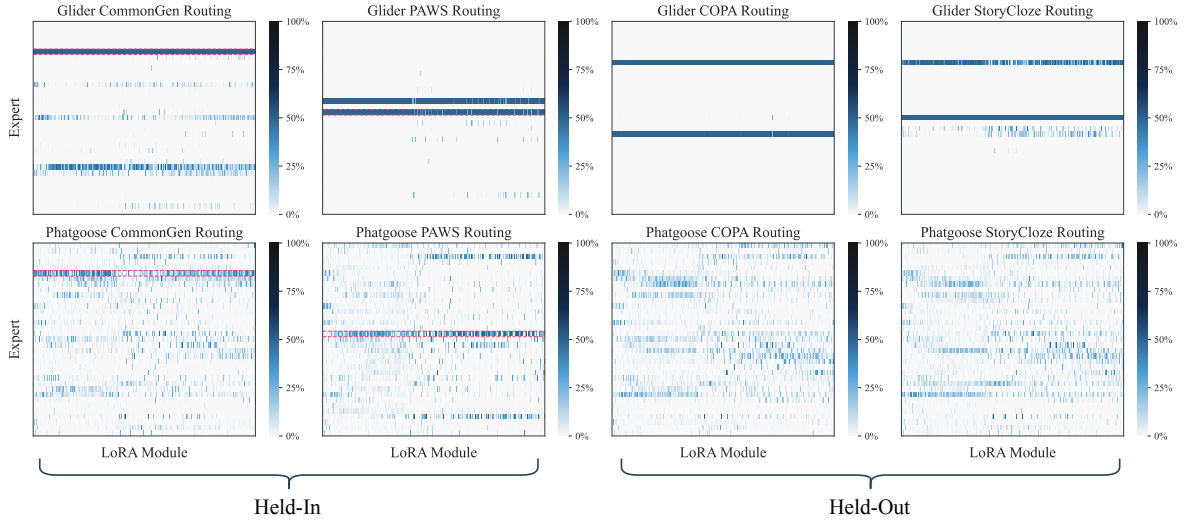


Figure 2: We present routing heatmaps for GLIDER and Phatgoose on two held-in and two held-out tasks. For held-in tasks, oracle experts are marked with red dashed lines. GLIDER selects oracle experts more frequently than Phatgoose for held-in tasks, leading to improvements of 3.3% on CommonGen and 6.5% on PAWS. For held-out tasks, GLIDER also tends to select the most relevant experts across most LoRA modules, resulting in improvements of 2.2% on COPA and 5.8% on StoryCloze.

ing recent works (Muqeeth et al., 2024; Ostapenko et al., 2024), we use parameter-efficient finetuning (PEFT) (Liu et al., 2022; Sung et al., 2022; Poth et al., 2023) methods like LoRA (Hu et al., 2022) to train the expert models. Since PEFT typically has lower computational and communication costs than full-model finetuning (Hu et al., 2022; Liu et al., 2022), the use of PEFT makes it easier to participate and contribute. PEFT methods introduce modules throughout the model – for example, LoRA (Hu et al., 2022) introduces a low-rank update at every linear layer in the model. We refer to each of these updates as a *module*. Subsequently, the trained expert models and additional information are shared with the aggregators. The aggregator’s job is to collect these expert models and the additional information and design the post-hoc routing mechanism. This mechanism will effectively direct incoming queries to the most appropriate expert model for each token and at each module to ensure optimal performance on both seen and unseen tasks. This approach allows for the seamless integration of new capabilities by adding expert models to the existing pool. Next, we formally define our contributor-aggregator framework.

Let us assume that there are N contributors, $\{c_1, c_2, \dots, c_N\}$, and each contributor c_i has access to a task-specific datasets $\mathcal{D}_i$. Each contributor, c_i , follows the predefined training protocol $\mathcal{T}$ provided by the aggregator. The training protocol ($\mathcal{T}$) takes in a base model (θ_{base}) and a dataset ($\mathcal{D}_i$). It returns the expert model parameters (ϕ_i) along with any additional information (Ψ_i) that needs to

be shared with the aggregators, for example, the gate vectors described in Section 4.1. Specifically, $\{\phi_i, \Psi_i\} \leftarrow \mathcal{T}(\theta_{\text{base}}, \mathcal{D}_i)$. All contributors share this information with the aggregator, which creates a pool of models containing $\{(\phi_i, \Psi_i)\}_{i=1}^N$. The aggregators ($\mathcal{A}$) then uses these expert models and the auxiliary information to create a routing mechanism $\mathcal{R}(\cdot)$ that takes the user query q as the input and return routing path describing how the information is routed through the given set of expert models. Formally, $\mathcal{R}(\cdot) \leftarrow \mathcal{A}(\{(\phi_i, \Psi_i)\}_{i=1}^N)$. The function $\mathcal{R}(\cdot)$ describe the full path of input query by making various choices about 1) expert input granularity, choosing to route per-token, per-query, or per-task, 2) expert depth granularity, opting for either per-module or model-level routing, and 3) selecting between sparse or dense routing. Finally, the aggregator uses the routing mechanism to answer incoming queries.

4 Methodology

To recap, our goal is to build a MoErging method that dynamically routes queries to a diverse pool of specialized expert models, addressing the challenge of effectively handling queries from various tasks and ensuring both held-in and held-out performance. Our proposed method, **Global and Local Instruction Driven Expert Router (GLIDER)**, leverages a combination of local and global routing vectors to achieve this goal. Specifically, contributors train task-specific routing vectors, while an LLM generates global semantic task instructions, which are then converted to global instruction routing

vectors. During inference, these local and global routing vectors are combined to perform top-k discrete routing, directing queries to the most suitable expert model. This process is visualized in Figure 1 and described in detail below.

4.1 Expert Training Protocol

Our expert training protocol $\mathcal{T}$ takes as input the base model parameters, θ_{base} , and a dataset $\mathbf{d}$ and performs three steps to obtain the required output. First, we train the LoRA experts (ϕ) and then the local routing vectors ($\mathbf{l}$) while keeping the LoRA experts fixed. Finally, we train the global routing vector ($\mathbf{g}$) by using an LLM and an embedding model. Formally, in our case, $\phi, \Psi = \{\mathbf{l}, \mathbf{g}\} \leftarrow \mathcal{T}(\theta_{\text{base}}, \mathbf{d})$ which are then shared with the aggregators to create the routing mechanism. We described these steps in detail below.

PEFT Training of Expert Model. GLIDER is compatible with expert models trained using parameter-efficient finetuning methods (e.g. LoRA (Hu et al., 2022), Adapters (Houlsby et al., 2019)) that introduce small trainable modules throughout the model. We focus on PEFT experts because they typically have lower computational and communication costs than full-model finetuning (Yadav et al., 2023a), making it easier to train and share expert models. Following Phatgoose (Muqeeth et al., 2024), this work specifically focuses on LoRA (Hu et al., 2022) due to its widespread use. LoRA introduces a *module* comprising the trainable matrices $\mathbf{B} \in \mathbb{R}^{d \times r}$ and $\mathbf{A} \in \mathbb{R}^{r \times n}$ in parallel to each linear layer with parameters $\mathbf{W} \in \mathbb{R}^{d \times n}$. Given the i^{th} input token activation $\mathbf{u}_i$, LoRA modifies the output of the linear layer from $\mathbf{W}\mathbf{u}_i$ to $\mathbf{W}\mathbf{u}_i + \frac{\alpha}{r} \cdot \mathbf{B}\mathbf{A}\mathbf{u}_i$ where α is a constant and usually is set to 1. During training, the matrices $\mathbf{A}$ and $\mathbf{B}$ are trainable, while the original linear layer $\mathbf{W}$ is kept frozen. We denote the final trained expert parameters with $\phi = \{(\mathbf{A}_1, \mathbf{B}_1), \dots, (\mathbf{A}_m, \mathbf{B}_m)\}$, where m is the number of modules in the model.

Training Local Routing Vectors. Following Phatgoose (Muqeeth et al., 2024), after training the PEFT modules on their dataset, a local router is introduced before each PEFT module. This router, employing a shared vector across all queries and tokens, dynamically determines the utilization of the PEFT module based on the input token activations. The router is trained for a small number of steps using the same dataset and objective as

the PEFT module while keeping the expert PEFT parameters fixed. This process effectively learns to associate the token activation patterns with the learned expert model. For LoRA, the local router, represented by a trainable vector $\mathbf{v} \in \mathbb{R}^d$, controls the contribution of the PEFT module to the final output. This results in a modified linear layer of the form $\mathbf{W}\mathbf{u}_i + \frac{\alpha}{r} \cdot \mathbf{B}\mathbf{A}\mathbf{u}_i \cdot \text{sigmoid}(\mathbf{v}^T \mathbf{u}_i)$, where $\alpha, \mathbf{W}, \mathbf{B}$, and $\mathbf{A}$ are frozen, and the local router $\mathbf{v}$ is learned. We denote the final local routing vectors as $\mathbf{l} = \{\mathbf{v}_1, \dots, \mathbf{v}_m\}$ where m is the number of modules in the model.

Creating LLM-Aided Global Routing Vector.

The local routing vectors capture the intricate relationships between token activations and expert models, enabling efficient query routing in cases where no dedicated expert is available. Conversely, for queries corresponding to held-in tasks, direct retrieval of the relevant expert model is preferred to process the full query. For this purpose, we create a global routing vector that utilizes an LLM to generate a semantically-informed instruction, termed as task description, which effectively captures the essence of the kind of queries the expert can handle. We prompt an LLM with three randomly selected in-context examples to generate this task description. We used the gpt-4-turbo model along with the prompt provided in Appendix B. The resulting task description is then embedded using an off-the-shelf embedding model, specifically the nomic-embed-text-v1.5 model, to produce a global routing vector for the task. We denote the global routing vector as $\mathbf{g} \in \mathbb{R}^{d_g}$.

4.2 GLIDER: Inference Expert Aggregation

Following training, all contributors share their expert models along with the auxiliary information comprising of the local and global routing vectors, $\{\phi^t, \mathbf{l}^t, \mathbf{g}^t\}_{t=1}^N$, where t indexes the input tokens with the aggregators. The GLIDER method subsequently leverages this information to perform inference on arbitrary queries.

Local Router. Before each input module m , a separate local router weight $\mathbf{L}_m \in \mathbb{R}^{N \times d}$ is inserted to make local per-token, per-module routing decisions. For a given module m and expert model c , we have $\bar{\mathbf{v}}_m^c = \frac{\mathbf{v}_m^c - \mu(\mathbf{v}_m^c)}{\sigma(\mathbf{v}_m^c)}$, where $\mu(\cdot)$ and $\sigma(\cdot)$ denote the mean and standard deviation respectively. Next, we obtain the local router for module m by stacking these standardised local routing vectors as $\mathbf{L}_m = [\bar{\mathbf{v}}_m^1, \dots, \bar{\mathbf{v}}_m^N] \in \mathbb{R}^{N \times d}$. Next, for each token i with

activation u_i coming into module m , we standardise it to obtain $\bar{u}_i = \frac{u_i - \mu(u_i)}{\sigma(u_i)}$. We then compute the local affinity scores, $s_m^{\text{loc}} \in \mathbb{R}^N$ between the local router L_m and u_i as $s_m^{\text{loc}} = \text{cos-sim}(L_m, u_i)$.

Global Router. The global router aims to capture task semantics to retrieve relevant experts for any given input query. We create the global router weight $G \in \mathbb{R}^{N \times d_g}$ by stacking the global routing vectors from all the expert models as $G = [g^1; \dots; g^N]$. This router is not a part of the base model and is added before the model to independently process the full query. Given an input query u along with three few-shot input-output pairs of similar queries, we prompt an LLM (gpt-4-turbo) using the template provided in Appendix B to obtain a task description for the query. We then embed this task description using the same embedding model (nomic-embed-text-v1.5) to obtain the vector $q_u \in \mathbb{R}^{d_g}$. We then compute the global affinity score, $s^{\text{glob}} \in \mathbb{R}^N$, by computing the cosine similarity as $s^{\text{glob}} = \text{cos-sim}(G, q_u)$.

Combining Global and Local Router. At each module m , we have the global and local affinity scores s^{glob} and s_m^{loc} respectively. Following Phatgoose (Muqeeth et al., 2024), we scale the local scores with a factor of $1/\sqrt{N}$. However, the global router’s main goal is to retrieve the correct expert for the held-in tasks. Therefore, we first check if the expert with the highest global affinity score ($\max(s^{\text{glob}})$) is above a threshold (p). If such experts exist, then we set a high α to enforce retrieval and vice versa. Hence, we propose to scale the global scores with α , where $\alpha = \gamma \cdot \mathbb{I}_{\{\max(s^{\text{glob}}) - p > 0\}} + \beta$, where p is the cosine similarity threshold, and γ and β are scaling hyperparameters. Using our ablation experiments in Section 5.4, we set $p = 0.8$, $\gamma = 100$ and $\beta = 3$. We then obtain the final affinity score $s \in \mathbb{R}^N = \alpha \cdot s^{\text{glob}} + s_m^{\text{loc}}/\sqrt{N}$. Then GLIDER selects the top- k experts after performing softmax over the final affinity score s as $\mathcal{E}_{\text{top}} = \text{top-k}(\text{softmax}(s))$. Finally, the output of the module for token activation u_i is computed as $wu_i + \sum_{k \in \mathcal{E}_{\text{top}}} s_k \cdot B_k A_k u_i$.

5 Experiments

5.1 Setting

Dataset. Our experiments utilize the multitask prompted training setup (**T0-HI**) introduced by Sanh et al. (2021), which has become a standard

benchmark for evaluating held-in performance as well as generalization to unseen tasks (Chung et al., 2022; Longpre et al., 2023; Jang et al., 2023; Zhou et al., 2022). Phatgoose (Muqeeth et al., 2024) shows how local routing can be used for generalization to unseen domain, hence, following them, we employ LM-adapted T5.1.1 XL (Lester et al., 2021) as our base model which is a 3B parameter variant of T5 (Raffel et al., 2020) further trained on the C4 dataset using a standard language modeling objective. For held-out evaluations, we follow Phatgoose (Muqeeth et al., 2024) and use three held-out benchmark collections. We use the T0 held-out (**T0-HO**) datasets used in Sanh et al. (2021) and the two subsets of BIG-bench (BIG-bench authors, 2023). Specifically, we use BIG-bench Hard (**BBH**) (Suzgun et al., 2022), consisting of 23 challenging datasets, and BIG-bench Lite (**BBL**) (BIG-bench authors, 2023), a lightweight 24-dataset proxy for the full benchmark. Similar to Muqeeth et al. (2024), we exclude certain BIG-bench datasets due to tokenization incompatibility with the T5 tokenizer.

Expert Creation. To create the pool of expert module for routing, we follow Muqeeth et al. (2024) and use two distinct dataset collections: ❶ T0 Held-In (Sanh et al., 2021) consisting of the 36 held-in prompted datasets for tasks from the T0 training procedure. ❷ The “FLAN Collection” (Longpre et al., 2023) which significantly expands the T0 tasks by incorporating prompted datasets from SuperGLUE (Wang et al., 2019a), Super Natural Instructions (Wang et al., 2022b), dialogue datasets, and Chain-of-Thought datasets (Wei et al., 2022b). Following Muqeeth et al. (2024), we create 166 specialized models from the FLAN Collection. For each dataset in these collections, we train Low-Rank Adapters (LoRAs) (Hu et al., 2021) modules resulting in pools of 36 and 166 expert models for T0 Held-In and FLAN, respectively. Similar to Phatgoose, we use a rank of $r = 16$ and train for 1000 steps using the AdamW optimizer (Loshchilov and Hutter, 2017) with a learning rate of 5×10^{-3} and a warmup ratio of 0.06. After training the LoRA module, we freeze it and train the local routing vectors for an additional 100 steps with the same hyperparameters. Finally, following prior work (Shazeer et al., 2016; Du et al., 2022; Lepikhin et al., 2020), GLIDER performs top- k routing with $k = 2$.

Table 1: Performance evaluated on the T0 set and FLAN set. We present the performance on both held-in tasks (*i.e.* T0-HI) and held-out tasks (*i.e.* T0-HO, BBH, and BBL). We compare the following methods: (1) performance upper bound, *i.e.* Oracle Expert; (2) zero-shot baselines, *i.e.* Multi-Task Fine-Tuning, Expert Merging, Arrow, and Phatgoose; (3) few-shot baselines, *i.e.* LoRA Hub and GLIDER. We mark the best performance besides the upper bound (*i.e.*, Oracle Expert) in **bold**.

Method	T0				FLAN	
	T0-HI	T0-HO	BBH	BBL	BBH	BBL
Oracle Expert	69.60	51.60	34.90	36.60	38.90	45.40
Multi-Task Fine-Tuning	55.90	51.60	34.90	36.60	38.90	45.40
Expert Merging	30.73	45.40	35.30	36.00	34.60	34.00
Arrow	39.84	55.10	33.60	34.50	30.60	29.60
Phatgoose	61.42	56.90	34.90	37.30	35.60	35.20
LoRA Hub	31.90	46.85	31.35	31.18	34.50	30.54
GLIDER	68.04	57.78	35.29	37.46	35.07	35.52

5.2 Baselines

Expert Merging. Model Merging (Yadav et al., 2023b; Choshen et al., 2022) involves averaging the parameters of multiple models or modules to create a single aggregate model. We merge by multiplying the LoRA matrices and then taking an unweighted average of all the experts within the pool. It is important to note that this merging strategy requires homogeneous expert module architectures; in contrast, GLIDER can accommodate heterogeneous expert modules.

Arrow. Following Ostapenko et al. (2024), we employ a routing mechanism where gating vectors are derived from LoRA expert modules. Specifically, the first right singular vector of the outer product of each module’s LoRA update (BA) serves as its gating vector. Input routing is determined by a probability distribution based on the absolute dot product between the input representation and each gating vector. We utilize top- k routing with $k = 2$.

Phatgoose. Phatgoose (Muqeeth et al., 2024) first learn the LoRA modules for each, followed by learning a sigmoid gating vector similar to our local router. During inference, they make routing decisions for each token independently for all modules. Specifically, they first standardize the input token activations and gating vectors from all experts and then perform similarity-based top-2 routing.

LoRA Hub. LoraHub (Huang et al., 2023) method performs gradient-free optimization using few-shot task samples to learn mixing coefficients for different expert models while keeping them fixed. Once the coefficients are learned, they merge the experts with the learned weight and route through the merged expert.

Multi-task Fine-Tuning. Multitask training is a

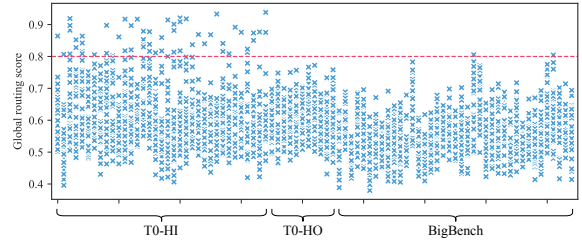


Figure 3: Global routing scores for tasks in the T0 set. The red horizontal line indicates our design threshold of 0.8. Each column represents an evaluated task from T0-HI, T0-HO, BigBench using T0 held-in experts. All global routing scores for each task are plotted, corresponding to the 35 experts in total.

proven method for enhancing zero-shot generalization (Sanh et al., 2021; Wei et al., 2022a) but is infeasible given our problem setting and data access limitations. We include it as a baseline using publicly available models. Specifically, we utilize the T0-3B model (Sanh et al., 2021) for the T0 Held-In datasets, given its training on a matching dataset collection. For FLAN, a directly comparable publicly available model is unavailable; therefore, we report FLAN-T5 XL results trained on a different, undisclosed dataset mixture, while acknowledging the limitations of this indirect comparison.

Oracle. Following (Jang et al., 2023) and (Muqeeth et al., 2024), we employ an Oracle routing scheme as a performance upper bound. This scheme selects the expert exhibiting optimal performance on a given evaluation dataset, thus representing a non-zero-shot approach.

5.3 Main Results

Table 1 presents the comparison results among our GLIDER and six baselines on both held-in and held-out settings. We report the average performance across all tasks for each setting, please see Appendix D for each tasks metric. To further illustrate the performance, we also include the results of Oracle Expert, which has extra access to the task identities of expert modules and evaluated datasets and can be regarded as an *upper bound*.

T0 Setting. In the T0 task set, the following observations can be drawn: ❶ For the held-in tasks, *i.e.* T0-HI, GLIDER significantly outperforms other baselines and almost matches the performance of Oracle Expert upper bound. ❷ For T0-HO and BBL tasks, GLIDER achieves the best performance among all the methods, including Oracle Expert upper bound. ❸ GLIDER has negligible lower performance, *i.e.* 0.01%, compared to the Expert Merging baseline in BBH but outperforms it by around

Table 2: Ablation on the instruction coefficient α . We mark the best performance in **bold** and the performance corresponding to the selected α by GLIDER in **blue**.

α	T0			
	T0-HI	T0-HO	BBH	BBL
0	61.42	56.90	34.90	37.30
1	62.20	57.04	35.05	37.79
3	63.40	57.78	35.29	37.46
10	65.52	57.98	34.80	37.04
100	68.04	53.22	31.73	34.97
1000	66.88	52.91	30.71	34.31
3000	66.69	52.37	30.03	33.24

12% on T0-HO and 1.5% on BBL. Besides Expert Merging, GLIDER outperforms all other methods on BBH, including the Oracle Expert upper bound.

5.4 Ablation Study and Further Investigation

Ablation on the global routing scale α . To illustrate how the specialization and generalization abilities change as we scale the coefficient α of the global routing score, we conduct the ablation study of α ranging $\{1, 3, 10, 100, 1000, 3000\}$. As shown in Table 2, we present experimental results of the T0 task set on both held-in and held-out tasks. For held-in tasks, *i.e.* T0-HI, GLIDER can select the optimal α to scale the global routing score. For held-out tasks, *i.e.* $\{T0-HO, BBH, BBL\}$, GLIDER produce either the optimal α (for BBH) or the sub-optimal α with slightly lower performance to the optimal ones (for T0-HO and BBL). Lastly, note that Phatgoose correspond to the setting where there is no global semantics used, *i.e.*, $\alpha = 0$.

Ablation on the routing strategy. There exists a trade-off between performance and efficiency when using different top-k routing strategies (Rachamandran and Le, 2019). To investigate the impact of routing strategy in GLIDER, we evaluate top-k routing of k in $\{1, 2, 3\}$. Moreover, we further evaluate the top-p routing (Huang et al., 2024c; Zeng et al., 2024) of p in $\{25\%, 50\%, 75\%\}$, where each token selects experts with higher routing probabilities until the cumulative probability exceeds threshold p. As shown in Table 3, we can draw the following conclusions: (1) For top-k routing, k = 2 shows comparable or better performance than k = 3, particularly for T0-HO and BBH, while offering improved efficiency. (2) For top-p routing, higher p values consistently yield better performance at the cost of efficiency. Therefore, we use top-2 routing in GLIDER by default.

Investigation on the threshold design of global scores. As in Section 4, we compute the scale

Table 3: Ablation on the routing strategy. GLIDER employs top-2 routing. We mark the best performance among top-k and top-p routing in **bold**, respectively.

Method	T0			
	T0-HI	T0-HO	BBH	BBL
Top-1	67.96	56.07	33.91	35.82
Top-2	68.04	57.78	35.39	37.46
Top-3	68.06	57.52	35.08	38.55
Top-25%	67.98	56.53	34.10	36.32
Top-50%	67.95	57.25	35.07	37.49
Top-75%	68.02	57.86	35.38	38.65

α for global scores using the formula $\alpha = \gamma * \mathbb{I}_{\{\max(\text{sg}^{\text{glob}}) - 0.8 > 0\}} + \beta$, where we establish a threshold of 0.8 to differentiate evaluated tasks. Figure 3 presents the global routing scores for each task in the T0 set to motivate the rationale behind this design. For all held-in tasks (*i.e.*, T0-HI), at least one expert (typically the oracle expert trained on the evaluated task) achieves global routing scores exceeding 0.8. Consequently, GLIDER applies a higher $\alpha = 100$, enabling effective identification of tasks corresponding to a specifically trained expert and enhancing retrieval of this oracle expert. For nearly all held-out tasks (*i.e.*, T0-HO and BigBench), no global routing score surpasses 0.8, prompting GLIDER to utilize a lower $\alpha = 3$. Two exceptions among the held-out tasks are `bbq_lite_json` and `strange_stories` in BigBench, where one score marginally exceeds 0.8 in each case. For these two, GLIDER employs the higher $\alpha = 100$, resulting in performance improvements of 1.3% and 2.9% respectively over $\alpha = 3$, thus showing the effectiveness of our design.

6 Conclusion

This paper introduces GLIDER, a novel multi-scale routing mechanism that incorporates both global semantic and local token-level routers. By leveraging the semantic reasoning capabilities of LLMs for global expert selection and refining these choices with a learned local router, GLIDER addresses the limitations of existing methods that often perform poorly on held-in tasks. Our empirical evaluation on T0 and FLAN benchmarks, using T5-based experts, demonstrates that GLIDER achieves substantial improvements in held-in task performance while maintaining competitive generalization on held-out tasks. These findings suggest that incorporating global semantic task context into routing mechanisms is crucial for building robust and practically useful routing-based systems.

7 Limitation

The main limitation of GLIDER lies in its heavy dependence on large language models (specifically GPT-4) for generating semantic task descriptions. This reliance introduces potential accessibility barriers due to API costs. Furthermore, investigating the application of GLIDER to other modalities beyond language tasks, such as vision or multi-modal expert models, could unlock new capabilities for specialized model routing.

References

Wikiquote, russian proverbs.

Abubakar Abid, Maheen Farooqi, and James Zou. 2021. [Persistent anti-Muslim bias in large language models](#). *arXiv preprint*.

Alessandro Achille, Michael Lam, Rahul Tewari, Avinash Ravichandran, Subhansu Maji, Charles C Fowlkes, Stefano Soatto, and Pietro Perona. 2019. Task2vec: Task embedding for meta-learning. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 6430–6439.

Joshua Ackerman and George Cybenko. 2020. [A survey of neural networks and formal languages](#). *arXiv preprint*.

Aishwarya Agrawal, Jiasen Lu, Stanislaw Antol, Margaret Mitchell, C. Lawrence Zitnick, Dhruv Batra, and Devi Parikh. 2015. [VQA: Visual question answering](#). *arXiv preprint*.

Akshay Agrawal, Shane Barratt, and Stephen Boyd. 2020. [Learning convex optimization models](#). *arXiv preprint*.

Scott Alexander. 2020. [A very unlikely chess game](#).

Mohammad Aliannejadi, Hamed Zamani, Fabio Crestani, and W. Bruce Croft. 2019. [Asking clarifying questions in open-domain information-seeking conversations](#). In *Proceedings of the 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval*, New York, NY, USA. Association for Computing Machinery.

Miltiadis Allamanis, Earl T. Barr, Premkumar Devanbu, and Charles Sutton. 2018. [A survey of machine learning for big code and naturalness](#). *ACM Comput. Surv.*, 51(4).

Miltiadis Allamanis, Pankajan Chanthirasegaran, Pushmeet Kohli, and Charles Sutton. 2016. [Learning continuous semantic representations of symbolic expressions](#). *arXiv preprint*.

Uri Alon, Shaked Brody, Omer Levy, and Eran Yahav. 2018. [code2seq: Generating sequences from structured representations of code](#). *arXiv preprint*.

Uri Alon, Roy Sadaka, Omer Levy, and Eran Yahav. 2020. [Structural language models of code](#). In *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pages 245–256. PMLR.

Miriam Amin and Manuel Burghardt. 2020. [A survey on approaches to computational humor generation](#). In *Proceedings of the 4th Joint SIGHUM Workshop on Computational Linguistics for Cultural Heritage, Social Sciences, Humanities and Literature*, pages 29–41, Online. International Committee on Computational Linguistics.

Aida Amini, Saadia Gabriel, Shanchuan Lin, Rik Koncel-Kedziorski, Yejin Choi, and Hannaneh Hajishirzi. 2019. [MathQA: Towards interpretable math word problem solving with operation-based formalisms](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 2357–2367, Minneapolis, Minnesota. Association for Computational Linguistics.

Prithviraj Ammanabrolu, William Broniec, Alex Mueller, Jeremy Paul, and Mark O. Riedl. 2019. [Toward automated quest generation in text-adventure games](#). *arXiv preprint*.

Prithviraj Ammanabrolu, Wesley Cheung, Dan Tu, William Broniec, and Mark O. Riedl. 2020. [Bringing stories alive: Generating interactive fiction worlds](#). *arXiv preprint*.

Dario Amodei, Chris Olah, Jacob Steinhardt, Paul Christiano, John Schulman, and Dan Mané. 2016. [Concrete problems in AI safety](#). *arXiv preprint*.

Brandon Amos and J. Zico Kolter. 2017. [Optnet: Differentiable optimization as a layer in neural networks](#). *arXiv preprint*.

Issa Annamoradnejad and Gohar Zoghi. 2020. [ColBERT: Using BERT sentence embedding for humor detection](#). *arXiv preprint*.

Mikel Artetxe, Sebastian Ruder, and Dani Yogatama. 2020. [On the cross-lingual transferability of monolingual representations](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 4623–4637, Online. Association for Computational Linguistics.

Shima Asaadi, Saif Mohammad, and Svetlana Kiritchenko. 2019. [Big BiRD: A large, fine-grained, bigram relatedness dataset for examining semantic composition](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 505–516, Minneapolis, Minnesota. Association for Computational Linguistics.

Pepa Atanasova, Jakob Grue Simonsen, Christina Lioma, and Isabelle Augenstein. 2020. [Generating fact](#)

766	checking explanations. In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 7352–7364, Online. Association for Computational Linguistics.	819
767		820
768		821
769		
770	Salvatore Attardo. 2017. Humor in language . In <i>Oxford Research Encyclopedia of Linguistics</i> . Oxford University Press.	822
771		823
772		824
773	Sören Auer, Christian Bizer, Georgi Kobilarov, Jens Lehmann, Richard Cyganiak, and Zachary Ives. 2007. Dbpedia: A nucleus for a web of open data . In <i>The Semantic Web</i> .	825
774		
775		
776		
777	R H. Baayen, R Piepenbrock, and L Gulikers. 1995. Celex2 ldc96l14 .	826
778		827
779	Yasaman Bahri, Ethan Dyer, Jared Kaplan, Jaehoon Lee, and Utkarsh Sharma. 2021. Explaining neural scaling laws .	828
780		829
781		830
782	Anton Bakhtin, Sam Gross, Myle Ott, Yuntian Deng, Marc’Aurelio Ranzato, and Arthur Szlam. 2019. Real or fake? Learning to discriminate machine from human generated text . <i>arXiv preprint</i> .	831
783		832
784		833
785		834
786	Matej Balog, Alexander L. Gaunt, Marc Brockschmidt, Sebastian Nowozin, and Daniel Tarlow. 2016. Deepcoder: Learning to write programs . <i>arXiv preprint</i> .	835
787		836
788		837
789	Satanjeev Banerjee and Ted Pedersen. 2003. Extended gloss overlaps as a measure of semantic relatedness . In <i>IJCAI’03: Proceedings of the 18th International Joint Conference on Artificial Intelligence</i> , page 805–810, San Francisco. Morgan Kaufmann.	838
790		839
791		840
792		841
793		842
794	Roy Bar-Haim, Ido Dagan, Bill Dolan, Lisa Ferro, Danilo Giampiccolo, Bernardo Magnini, and Idan Szpektor. 2006. The second pascal recognising textual entailment challenge. In <i>Proceedings of the second PASCAL challenges workshop on recognising textual entailment</i> , volume 6, pages 6–4. Venice.	843
795		844
796		845
797		846
798		847
799		848
800	Oren Barkan and Noam Koenigstein. 2016. ITEM2VEC: Neural item embedding for collaborative filtering . In <i>2016 IEEE 26th International Workshop on Machine Learning for Signal Processing (MLSP)</i> , pages 1–6, Piscataway, NJ. Institute of Electrical and Electronics Engineers.	849
801		850
802		
803		
804		
805		
806	Solon Barocas and Andrew D. Selbst. 2016. Big data’s disparate impact . <i>California Law Review</i> , 104(3):671–732.	851
807		852
808		853
809	Max Bartolo, Alastair Roberts, Johannes Welbl, Sebastian Riedel, and Pontus Stenetorp. 2020. Beat the ai: Investigating adversarial human annotation for reading comprehension. <i>Transactions of the Association for Computational Linguistics</i> , 8:662–678.	854
810		855
811		856
812		857
813		
814	Sumit Basu and Janara Christensen. 2013. Teaching classification boundaries to humans . In <i>Proceedings of the AAAI Conference on Artificial Intelligence</i> , volume 27, pages 109–115, Menlo Park, CA. Association for the Advancement of Artificial Intelligence.	858
815		859
816		860
817		
818		
	Jason Baumgartner, Savvas Zannettou, Brian Keegan, Megan Squire, and Jeremy Blackburn. 2020. The Pushshift Reddit dataset . <i>arXiv preprint</i> .	861
		862
		863
		864
		865
	Nihat Bayat and Gökhan Çetinkaya. 2020. The relationship between inference skills and reading comprehension . <i>TED EĞİTİM VE BİLİM (Education and Science)</i> , 45(203):177–190.	866
		867
		868
		869
		870
		871
		872
	Mayur J. Bency, Ahmed H. Qureshi, and Michael C. Yip. 2019. Neural path planning: Fixed time, near-optimal path generation via oracle imitation . In <i>2019 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)</i> , pages 3965–3972, Piscataway, NJ. Institute of Electrical and Electronics Engineers.	
	Emily M. Bender, Timnit Gebru, Angelina McMillan-Major, and Shmargaret Shmitchell. 2021. On the dangers of stochastic parrots: Can language models be too big? In <i>Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency, FAccT ’21</i> , page 610–623, New York, NY, USA. Association for Computing Machinery.	
	Emily M. Bender and Alexander Koller. 2020. Climbing towards NLU: On meaning, form, and understanding in the age of data . In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 5185–5198, Online. Association for Computational Linguistics.	
	Luisa Bentivogli, Peter Clark, Ido Dagan, and Danilo Giampiccolo. 2009. The fifth pascal recognizing textual entailment challenge. In <i>TAC</i> .	
	Jean Berko. 1958. The child’s learning of english morphology . <i><i>WORD</i></i> , 14(2-3):150–177.	
	Tarek R. Besold, Artur d’Avila Garcez, Sebastian Bader, Howard Bowman, Pedro Domingos, Pascal Hitzler, Kai-Uwe Kuehnberger, Luis C. Lamb, Daniel Lowd, Priscila Machado Vieira Lima, Leo de Penning, Gadi Pinkas, Hoifung Poon, and Gerson Zaverucha. 2017. Neural-symbolic learning and reasoning: A survey and interpretation . <i>arXiv preprint</i> .	
	Gregor Betz, Christian Voigt, and Kyle Richardson. 2020. Critical thinking for language models . <i>arXiv preprint</i> .	
	Chandra Bhagavatula, Ronan Le Bras, Chaitanya Malaviya, Keisuke Sakaguchi, Ari Holtzman, Han-nah Rashkin, Doug Downey, Scott Wen-tau Yih, and Yejin Choi. 2019. Abductive commonsense reasoning . <i>arXiv preprint</i> .	
	Sumithra Bhakthavatsalam, Daniel Khashabi, Tushar Khot, Bhavana Dalvi Mishra, Kyle Richardson, Ashish Sabharwal, Carissa Schoenick, Oyvind Tafjord, and Peter Clark. 2021. Think you have solved direct-answer question answering? Try ARC-DA, the direct-answer AI2 reasoning challenge . <i>arXiv preprint</i> .	

873	Satwik Bhattamishra, Kabir Ahuja, and Navin Goyal.	Laria Reynolds, Jonathan Tow, Ben Wang, and	929
874	2020a. On the ability and limitations of transformers	Samuel Weinbach. 2022. GPT-NeoX-20B: An open-	930
875	to recognize formal languages . In <i>Proceedings of the</i>	source autoregressive language model . In <i>Proceed-</i>	931
876	<i>2020 Conference on Empirical Methods in Natural</i>	<i>ings of the ACL Workshop on Challenges & Perspec-</i>	932
877	<i>Language Processing (EMNLP)</i> , pages 7096–7116,	<i>tives in Creating Large Language Models</i> .	933
878	Online. Association for Computational Linguistics.		
879	Satwik Bhattamishra, Kabir Ahuja, and Navin Goyal.	Vladislav Blinov, Valeria Bolotova-Baranova, and Pavel	934
880	2020b. On the practical ability of recurrent neural	Braslavski. 2019. Large dataset and language model	935
881	networks to recognize hierarchical languages . In	fun-tuning for humor recognition . In <i>Proceedings of</i>	936
882	<i>Proceedings of the 28th International Conference</i>	<i>the 57th Annual Meeting of the Association for Com-</i>	937
883	<i>on Computational Linguistics</i> , pages 1481–1494,	<i>putational Linguistics</i> , pages 4027–4032, Florence,	938
884	Barcelona, Spain (Online). International Committee	Italy. Association for Computational Linguistics.	939
885	on Computational Linguistics.		
886	Surya Bhupatiraju, Rishabh Singh, Abdel-rahman Mo-	John Blitzer, Mark Dredze, and Fernando Pereira. 2007.	940
887	hamed, and Pushmeet Kohli. 2017. Deep API pro-	Large scale sentiment analysis for news and blogs .	941
888	grammer: Learning to program with APIs . <i>arXiv</i>	In <i>Proceedings of the International Conference on</i>	942
889	<i>preprint</i> .	<i>Weblogs and Social Media (ICWSM)</i> .	943
890	Alan W. Biermann. 1978. The inference of regular	Su Lin Blodgett, Solon Barocas, Hal Daumé III, and	944
891	LISP programs from examples . <i>IEEE Transactions</i>	Hanna Wallach. 2020. Language (technology) is	945
892	<i>on Systems, Man, and Cybernetics</i> , 8(8):585–600.	power: A critical survey of “bias” in NLP . In <i>Pro-</i>	946
893	BIG-bench authors. 2023. Beyond the imitation game:	ceedings of the 58th Annual Meeting of the Asso-	947
894	Quantifying and extrapolating the capabilities of lan-	ciation for Computational Linguistics , pages 5454–	948
895	guage models . <i>Transactions on Machine Learning</i>	5476, Online. Association for Computational Lin-	949
896	<i>Research</i> .	guistics.	950
897	Joachim Bingel and Anders Søgaard. 2017. Identi-	Yulia V. Bodrova. 2007. <i>Russian Proverbs and Sayings</i>	951
898	fying beneficial task relations for multi-task learn-	<i>and Their English Equivalents</i> . AST, Moscow.	952
899	ing in deep neural networks. <i>arXiv preprint</i>	Nicholas Boillot. 2019. Vector forms as a foreign lan-	953
900	<i>arXiv:1702.08303</i> .	guage .	954
901	Julia Birke and Anoop Sarkar. 2006. A clustering ap-	Paul F. Boller, Jr. and John George. 1989. <i>They Never</i>	955
902	proach for nearly unsupervised recognition of nonlit-	<i>Said It: A Book of Fake Quotes, Misquotes, and</i>	956
903	eral language . In <i>11th Conference of the European</i>	<i>Misleading Attributions</i> . Oxford University Press,	957
904	<i>Chapter of the Association for Computational Lin-</i>	Oxford.	958
905	<i>guistics</i> , pages 329–336, Trento, Italy. Association	Tolga Bolukbasi, Kai-Wei Chang, James Y. Zou,	959
906	for Computational Linguistics.	Venkatesh Saligrama, and Adam T. Kalai. 2016. Man	960
907	Sebastian Bischoff, Niklas Deckers, Marcel Schliebs,	is to computer programmer as woman is to home-	961
908	Ben Thies, Matthias Hagen, Efsthios Stamatas,	maker? Debiasing word embeddings . In <i>Advances in</i>	962
909	Benno Stein, and Martin Potthast. 2020. The im-	<i>Neural Information Processing Systems</i> , volume 29.	963
910	portance of suppressing domain style in authorship	Curran Associates, Inc.	964
911	analysis . <i>arXiv preprint</i> .	Shikha Bordia and Samuel R. Bowman. 2019. Identify-	965
912	Yonatan Bisk, Rowan Zellers, Ronan Le Bras, Jian-	ing and reducing gender bias in word-level language	966
913	feng Gao, and Yejin Choi. 2020. PIQA: reasoning	models . <i>arXiv preprint</i> .	967
914	about physical commonsense in natural language . In	Antoine Bosselut, Hannah Rashkin, Maarten Sap, Chai-	968
915	<i>Proceedings of the AAAI Conference on Artificial In-</i>	tanya Malaviya, Asli Celikyilmaz, and Yejin Choi.	969
916	<i>telligence</i> , volume 34, pages 7423–7439, Menlo Park,	2019. COMET: Commonsense transformers for auto-	970
917	CA. Association for the Advancement of Artificial	matic knowledge graph construction . <i>arXiv preprint</i> .	971
918	Intelligence.	Michael Bostock, Vadim Ogievetsky, and Jeffrey Heer.	972
919	Yuri Bizzoni and Shalom Lappin. 2018. Predicting	2011. D³ data-driven documents . <i>IEEE Trans-</i>	973
920	human metaphor paraphrase judgments with deep	<i>actions on Visualization and Computer Graphics</i> ,	974
921	neural networks . In <i>Proceedings of the Workshop on</i>	17(12):2301–2309.	975
922	<i>Figurative Language Processing</i> , pages 45–55, New	Samuel R. Bowman, Gabor Angeli, Christopher Potts,	976
923	Orleans, Louisiana. Association for Computational	and Christopher D. Manning. 2015. A large anno-	977
924	Linguistics.	tated corpus for learning natural language inference .	978
925	Sid Black, Stella Biderman, Eric Hallahan, Quentin An-	In <i>Proceedings of the 2015 Conference on Empiri-</i>	979
926	thony, Leo Gao, Laurence Golding, Horace He, Con-	<i>cal Methods in Natural Language Processing</i> , pages	980
927	nor Leahy, Kyle McDonell, Jason Phang, Michael	632–642, Lisbon, Portugal. Association for Compu-	981
928	Pieler, USVSN Sai Prashanth, Shivanshu Purohit,	tational Linguistics.	982

983	Samuel R. Bowman and George Dahl. 2021. What will it take to fix benchmarking in natural language understanding? In <i>Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies</i> , pages 4843–4855, Online. Association for Computational Linguistics.	1040
984		1041
985		1042
986		1043
987		
988		1044
989		1045
		1046
990	Bozhidar Bozhanov and Ivan Derzhanski. 2013. Rosetta stone linguistic problems . In <i>Proceedings of the Fourth Workshop on Teaching NLP and CL</i> , pages 1–8, Sofia, Bulgaria. Association for Computational Linguistics.	1047
991		1048
992		1049
993		1050
994		1051
995	Matko Bošnjak, Tim Rocktäschel, Jason Naradowsky, and Sebastian Riedel. 2017. Programming with a differentiable Forth interpreter . In <i>Proceedings of the 34th International Conference on Machine Learning</i> , volume 70, pages 547–556.	1052
996		1053
997		1054
998		
999		
1000	Gwern Branwen. 2020. GPT-3 creative fiction . <i>Gwern.net</i> .	1055
1001		1056
		1057
1002	Luke Breitfeller, Emily Ahn, David Jurgens, and Yulia Tsvetkov. 2019. Finding microaggressions in the wild: A case for locating elusive phenomena in social media posts . In <i>Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)</i> , pages 1664–1674, Hong Kong, China. Association for Computational Linguistics.	1058
1003		1059
1004		1060
1005		1061
1006		
1007		1062
1008		1063
1009		1064
1010		
1011	Ralf Brown. 2014. Non-linear mapping for improved identification of 1300+ languages . In <i>Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)</i> , pages 627–632, Doha, Qatar. Association for Computational Linguistics.	1065
1012		1066
1013		1067
1014		1068
1015		1069
1016		1070
		1071
1017	Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D. Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Chris Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. 2020. Language models are few-shot learners . In <i>Advances in Neural Information Processing Systems</i> , volume 33, pages 1877–1901. Curran Associates, Inc.	1072
1018		1073
1019		1074
1020		1075
1021		1076
1022		1077
1023		1078
1024		
1025		1079
1026		1080
1027		1081
1028		
1029		1082
1030		1083
		1084
1031	Thomas Bugnyar, Stephan A. Reber, and Cameron Buckner. 2016. Ravens attribute visual access to unseen competitors . <i>Nature Communications</i> , 7:article 10506.	1085
1032		
1033		
1034		
1035	Franck Burlot, Yves Scherrer, Vinit Ravishankar, Ondřej Bojar, Stig-Arne Grönroos, Maarit Koponen, Tommi Nieminen, and François Yvon. 2018. The WMT’18 morphEval test suites for English-Czech, English-German, English-Finnish and Turkish-English . In <i>Proceedings of the Third Conference on Machine Translation: Shared Task Papers</i> , pages 546–560, Belgium, Brussels. Association for Computational Linguistics.	1086
1036		1087
1037		1088
1038		1089
1039		1090
		1091
		1092
	Corrado Böhm. 1964. On a family of Turing machines and the related programming language. <i>ICC Bulletin</i> , 3:187–194.	
	Lucas Caccia, Edoardo Ponti, Zhan Su, Matheus Pereira, Nicolas Le Roux, and Alessandro Sordoni. 2023. Multi-head adapter routing for cross-task generalization. In <i>Thirty-seventh Conference on Neural Information Processing Systems</i> .	
	Kate Cain and Jane V. Oakhill. 1999. Inference making ability and its relation to comprehension failure . <i>Reading and Writing</i> , 11(5–6):489–503.	
	Maya Cakmak and Andrea L. Thomaz. 2014. Eliciting good teaching from humans for machine learners . <i>Artificial Intelligence</i> , 217:198–215.	
	Aylin Caliskan, Joanna J. Bryson, and Arvind Narayanan. 2017. Semantics derived automatically from language corpora contain human-like biases . <i>Science</i> , 356(6334):183–186.	
	Josep Call and Michael Tomasello. 2008. Does the chimpanzee have a theory of mind? 30 years later . <i>Trends in Cognitive Sciences</i> , 12:187–192.	
	Tracy Canfield. 2010. Machine translation of Klingon .	
	Nathanael Chambers. 2012. Labeling documents with timestamps: Learning from their time expressions . In <i>Proceedings of the 50th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 98–106, Jeju Island, Korea. Association for Computational Linguistics.	
	Sharath Chandra Guntuku, Mingyang Li, Louis Tay, and Lyle H. Ungar. 2019. Studying cultural differences in emoji usage across the East and the West . In <i>Proceedings of the International AAAI Conference on Web and Social Media</i> , volume 13, pages 226–235, Menlo Park, CA. Association for the Advancement of Artificial Intelligence.	
	Nick Chater and Paul Vitányi. 2003. Simplicity: A unifying principle in cognitive science? <i>Trends in Cognitive Sciences</i> , 7:19–22.	
	Antonio Chella, Arianna Pipitone, Alain Morin, and Famira Racy. 2020. Developing self-awareness in robots via inner speech . <i>Frontiers in Robotics and AI</i> , 7.	
	Howard Chen, Alane Suhr, Dipendra Misra, Noah Snaveley, and Yoav Artzi. 2019a. Touchdown: Natural language navigation and spatial reasoning in visual street environments . In <i>2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)</i> , pages 12530–12539, Piscataway, NJ. Institute of Electrical and Electronics Engineers.	

1093	Mark Chen, Alec Radford, Rewon Child, Jeffrey Wu, Heewoo Jun, David Luan, and Ilya Sutskever. 2020. Generative pretraining from pixels . In <i>Proceedings of the 37th International Conference on Machine Learning</i> , volume 119 of <i>Proceedings of Machine Learning Research</i> , pages 1691–1703. PMLR.	1146
1094		1147
1095		1148
1096		1149
1097		1150
1098		
1099	Peng-Yu Chen and Von-Wun Soo. 2018. Humor recognition using deep learning . In <i>Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 2 (Short Papers)</i> , pages 113–117, New Orleans, Louisiana. Association for Computational Linguistics.	1151
1100		1152
1101		1153
1102		
1103		1154
1104		1155
1105		1156
1106	Ricson Chen. 2020. Transformers play chess .	1157
1107		1158
1108	Xinyun Chen, Chang Liu, and Dawn Song. 2019b. Execution-guided neural program synthesis. https://openreview.net/pdf?id=H1gf0iAqYm .	
1109		
1110	Feng Cheng, Ziyang Wang, Yi-Lin Sung, Yan-Bo Lin, Mohit Bansal, and Gedas Bertasius. 2024. DAM: Dynamic adapter merging for continual video qa learning. <i>arXiv preprint arXiv:2403.08755</i> .	1159
1111		1160
1112		1161
1113		1162
1114		1163
1115	Rewon Child, Scott Gray, Alec Radford, and Ilya Sutskever. 2019. Generating long sequences with sparse transformers . <i>arXiv preprint</i> .	
1116		
1117	Won Ik Cho, Ji Won Kim, Seok Min Kim, and Nam Soo Kim. 2019. On measuring gender bias in translation of gender-neutral pronouns . In <i>Proceedings of the First Workshop on Gender Bias in Natural Language Processing</i> , pages 173–181, Florence, Italy. Association for Computational Linguistics.	1164
1118		1165
1119		1166
1120		
1121		1167
1122		1168
1123	Eunsol Choi, He He, Mohit Iyyer, Mark Yatskar, Wentau Yih, Yejin Choi, Percy Liang, and Luke Zettlemoyer. 2018. QuAC: Question answering in context . <i>arXiv preprint</i> .	1169
1124		1170
1125		1171
1126		
1127	François Chollet. 2019. On the measure of intelligence .	1172
1128		1173
1129	François Chollet. 2020. Abstraction and reasoning challenge .	1174
1130		1175
1131	Noam Chomsky and Marcel P. Schützenberger. 1959. The algebraic theory of context-free languages . In P. Braffort and D. Hirschberg, editors, <i>Computer Programming and Formal Systems</i> , volume 26 of <i>Studies in Logic and the Foundations of Mathematics</i> , pages 118–161. Elsevier.	1176
1132		1177
1133		1178
1134		1179
1135		1180
1136	Leshem Choshen, Elad Venezian, Noam Slonim, and Yoav Katz. 2022. Fusing finetuned models for better pretraining. <i>arXiv preprint arXiv:2204.03044</i> .	1181
1137		1182
1138		1183
1139	Alexandra Chronopoulou, Matthew E Peters, Alexander Fraser, and Jesse Dodge. 2023. Adaptersoup: Weight averaging to improve generalization of pretrained language models. <i>arXiv preprint arXiv:2302.07027</i> .	1184
1140		1185
1141		1186
1142		
1143	Casey Chu, Andrey Zhmoginov, and Mark Sandler. 2017. CycleGAN, a master of steganography . <i>arXiv preprint</i> .	1187
1144		1188
1145		1189
		1190
		1191
		1192
		1193
		1194
		1195
		1196
		1197
		1198
		1199
		1200

1201	Andrew Cropper and Stephen H. Muggleton. 2016.	Marie-Catherine de Marneffe, Mandy Simons, and Ju-	1254
1202	Metagol system .	dith Tonhauser. 2019. The CommitmentBank: Inves-	1255
1203	Andrew Cropper, Alireza Tamaddoni-Nezhad, and	tigating projection in naturally occurring discourse .	1256
1204	Stephen H. Muggleton. 2016. Meta-interpretive	<i>Proceedings of Sinn und Bedeutung</i> , 23(2):107–124.	1257
1205	learning of data transformation programs . In <i>In-</i>	Scott Deerwester, Susan T. Dumais, George W. Furnas,	1258
1206	<i>ductive Logic Programming</i> , pages 46–59, Cham.	Thomas K. Landauer, and Richard Harshman. 1990.	1259
1207	Springer.	Indexing by latent semantic analysis . <i>Journal of the</i>	1260
1208	Joe Cruse. 2015. Emoji usage in TV conversation . <i>Twit-</i>	<i>American Society for Information Science</i> , 41(6):391–	1261
1209	<i>ter blog</i> .	407.	1262
1210	Marc-Alexandre Côté, Ákos Kádár, Xingdi Yuan, Ben	Judith Degen, Robert D. Hawkins, Caroline Graf, Elisa	1263
1211	Kybartas, Tavian Barnes, Emery Fine, James Moore,	Kreiss, and Noah D. Goodman. 2020. When redun-	1264
1212	Ruo Yu Tao, Matthew Hausknecht, Layla El Asri,	dancy is useful: A Bayesian approach to “overinfor-	1265
1213	Mahmoud Adada, Wendy Tay, and Adam Trischler.	mative” referring expressions . <i>Psychological Review</i> ,	1266
1214	2018. TextWorld: A learning environment for text-	127:591–621.	1267
1215	based games . <i>arXiv preprint</i> .	Sunipa Dev, Tao Li, Jeff M. Phillips, and Vivek Sriku-	1268
1216	Ido Dagan, Oren Glickman, and Bernardo Magnini.	mar. 2020. On measuring and mitigating biased in-	1269
1217	2005. The pascal recognising textual entailment chal-	ferences of word embeddings . In <i>Proceedings of</i>	1270
1218	lenge. In <i>Machine Learning Challenges Workshop</i> ,	<i>the AAAI Conference on Artificial Intelligence</i> , vol-	1271
1219	pages 177–190. Springer.	ume 34, pages 7659–7666, Menlo Park, CA. Associ-	1272
1220	Jim Daley. 2021. White Chicago cops use force more	ation for the Advancement of Artificial Intelligence.	1273
1221	often than Black officers . <i>Scientific American</i> .	Jacob Devlin, Ming-Wei Chang, Kenton Lee, and	1274
1222	Sahith Dambekodi, Spencer Frazier, Prithviraj Am-	Kristina Toutanova. 2018. BERT: Pre-training of	1275
1223	manabrolu, and Mark O. Riedl. 2020. Playing text-	deep bidirectional transformers for language under-	1276
1224	based games with common sense . <i>arXiv preprint</i> .	standing . <i>arXiv preprint</i> .	1277
1225	Pradeep Dasigi, Nelson F. Liu, Ana Marasovic, Noah A.	Jacob Devlin, Jonathan Uesato, Surya Bhupatiraju,	1278
1226	Smith, and Matt Gardner. 2019. Quoref: A read-	Rishabh Singh, Abdel-rahman Mohamed, and Push-	1279
1227	ing comprehension dataset with questions requir-	meet Kohli. 2017. RobustFill: Neural program learn-	1280
1228	ing coreferential reasoning . In <i>Proceedings of the</i>	ing under noisy I/O . In <i>Proceedings of the 34th</i>	1281
1229	<i>2019 Conference on Empirical Methods in Natu-</i>	<i>International Conference on Machine Learning</i> , vol-	1282
1230	<i>ral Language Processing and the 9th International</i>	ume 70, pages 990–998, New York, NY, USA. Asso-	1283
1231	<i>Joint Conference on Natural Language Processing</i>	ciation for Computing Machinery.	1284
1232	<i>(EMNLP-IJCNLP)</i> .	Bhuwan Dhingra, Kathryn Mazaitis, and William W.	1285
1233	Wayne Davis. 2019. Implicature. In Edward N. Zalta,	Cohen. 2017. Quasar: Datasets for question answer-	1286
1234	editor, <i>The Stanford Encyclopedia of Philosophy</i> , Fall	ing by search and reading . <i>arXiv preprint</i> .	1287
1235	2019 edition. Metaphysics Research Lab, Stanford	Kaustubh Dhole, Gurdeep Singh, Priyadarshini P. Pai,	1288
1236	University.	and Sukanta Mondal. 2014. Sequence-based predic-	1289
1237	Marie-Catherine de Marneffe, Christopher D. Man-	tion of protein–protein interaction sites with l1-logreg	1290
1238	ning, and Christopher Potts. 2012. Did it happen?	classifier . <i>Journal of Theoretical Biology</i> , 348:47–	1291
1239	The pragmatic complexity of veridicality assessment .	54.	1292
1240	<i>Computational Linguistics</i> , 38(2):301–333.	Shizhe Diao, Tianyang Xu, Ruijia Xu, Jiawei Wang,	1293
1241	Marie-Catherine de Marneffe, Anna N. Rafferty, and	and T. Zhang. 2023. Mixture-of-domain-adapters:	1294
1242	Christopher D. Manning. 2008. Finding contradic-	Decoupling and injecting domain knowledge to pre-	1295
1243	tions in text . In <i>Proceedings of ACL-08: HLT</i> , pages	trained language models’ memories . In <i>Annual Meet-</i>	1296
1244	1039–1047, Columbus, Ohio. Association for Com-	ing of the Association for Computational Linguistics .	1297
1245	putational Linguistics.	Emily Dinan, Angela Fan, Adina Williams, Jack Ur-	1298
1246	Marie-Catherine de Marneffe, Mandy Simons, and Ju-	banek, Douwe Kiela, and Jason Weston. 2019.	1299
1247	dith Tonhauser. 2017. The commitmentbank: Inves-	Queens are powerful too: Mitigating gender bias	1300
1248	tigating projection in naturally occurring discourse .	in dialogue generation . <i>arXiv preprint</i> .	1301
1249	In <i>Proceedings of the 15th Conference of the Euro-</i>	William B Dolan and Chris Brockett. 2005. Automati-	1302
1250	<i>pean Chapter of the Association for Computational</i>	cally constructing a corpus of sentential paraphrases.	1303
1251	<i>Linguistics: Volume 2, Short Papers</i> , pages 48–54,	In <i>Proceedings of the Third International Workshop</i>	1304
1252	Valencia, Spain. Association for Computational Lin-	<i>on Paraphrasing (IWP2005)</i> .	1305
1253	guistics.	Tiansi Dong, Chengjiang Li, Christian Bauckhage,	1306
		Juanzi Li, Stefan Wrobel, and Armin B. Cre-	1307
		mers. 2020. Learning syllogism with Euler neural-	1308
		networks . <i>arXiv preprint</i> .	1309

Nan Du, Yanping Huang, Andrew M Dai, Simon Tong, Dmitry Lepikhin, Yuanzhong Xu, Maxim Krikun, Yanqi Zhou, Adams Wei Yu, Orhan Firat, et al. 2022. Glam: Efficient scaling of language models with mixture-of-experts. In <i>International Conference on Machine Learning</i> , pages 5547–5569. PMLR.	1367
Dheeru Dua, Yizhong Wang, Pradeep Dasigi, Gabriel Stanovsky, Sameer Singh, and Matt Gardner. 2019. DROP: A reading comprehension benchmark requiring discrete reasoning over paragraphs. In <i>Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)</i> , pages 2368–2378, Minneapolis, Minnesota. Association for Computational Linguistics.	1373
Liam Dugan, Daphne Ippolito, Arun Kirubakaran, and Chris Callison-Burch. 2020. RoFT: A tool for evaluating human detection of machine-generated text. In <i>Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations</i> , pages 189–196, Online. Association for Computational Linguistics.	1382
Jesse Duniety, Greg Burnham, Akash Bharadwaj, Owen Rambow, Jennifer Chu-Carroll, and Dave Ferrucci. 2020. To test machine comprehension, start by defining comprehension. In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 7839–7859, Online. Association for Computational Linguistics.	1383
Jon Durbin. 2024. aioboros: Customizable implementation of the self-instruct paper. https://github.com/jondurbin/aioboros .	1384
Esin Durmus, He He, and Mona Diab. 2020. FEQA: A question answering evaluation framework for faithfulness assessment in abstractive summarization. In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 5055–5070, Online. Association for Computational Linguistics.	1385
Nouha Dziri, Andrea Madotto, Osmar Zaiane, and Avishek Joey Bose. 2021. Neural path hunter: Reducing hallucination in dialogue systems via path grounding. <i>arXiv preprint</i> .	1386
Javid Ebrahimi, Dhruv Gelda, and Wei Zhang. 2020. How can self-attention networks recognize Dyck-n languages? <i>arXiv preprint</i> .	1387
Bora Edizel, Aleksandra Piktus, Piotr Bojanowski, Rui Ferreira, Edouard Grave, and Fabrizio Silvestri. 2019. Misspelling oblivious word embeddings. <i>arXiv preprint</i> .	1388
Daniel Edmiston and Karl Stratos. 2018. Compositional morpheme embeddings with affixes as functions and stems as arguments. In <i>Proceedings of the Workshop on the Relevance of Linguistic Structure in Neural Architectures for NLP</i> , pages 1–5, Melbourne, Australia. Association for Computational Linguistics.	1389
Avia Efrat, Uri Shaham, Dan Kilman, and Omer Levy. 2021. Cryptonite: A cryptic crossword benchmark for extreme ambiguity in language. In <i>Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing</i> , pages 4186–4192, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.	1390
Liat Ein Dor, Alon Halfon, Yoav Kantor, Ran Levy, Yosi Mass, Ruty Rinott, Eyal Shnarch, and Noam Slonim. 2018. Semantic relatedness of Wikipedia concepts – benchmark data and a working solution. In <i>Proceedings of the Eleventh International Conference on Language Resources and Evaluation (LREC 2018)</i> , Miyazaki, Japan. European Language Resources Association.	1391
Ben Eisner, Tim Rocktäschel, Isabelle Augenstein, Matko Bošnjak, and Sebastian Riedel. 2016. emoji2vec: Learning emoji representations from their description. In <i>Proceedings of The Fourth International Workshop on Natural Language Processing for Social Media</i> , pages 48–54, Austin, TX, USA. Association for Computational Linguistics.	1392
Ahmed El-Kishky, Frank Xu, Aston Zhang, and Jiawei Han. 2019. Parsimonious morpheme segmentation with an application to enriching word embeddings. <i>arXiv preprint</i> .	1393
Ran El-Yaniv and David Yanay. 2013. Semantic sort: A supervised approach to personalized semantic relatedness.	1394
Yanai Elazar, Nora Kassner, Shauli Ravfogel, Abhilasha Ravichander, Eduard Hovy, Hinrich Schütze, and Yoav Goldberg. 2021. Measuring and improving consistency in pretrained language models. <i>arXiv preprint</i> .	1395
Kevin Ellis and Sumit Gulwani. 2017. Learning to learn programs from examples: Going beyond program structure. In <i>Proceedings of the Twenty-Sixth International Joint Conference on Artificial Intelligence, IJCAI-17</i> , pages 1638–1645.	1396
Kevin Ellis, Catherine Wong, Maxwell Nye, Mathias Sable-Meyer, Luc Cary, Lucas Morales, Luke Hewitt, Armando Solar-Lezama, and Joshua B. Tenenbaum. 2020. Dreamcoder: Growing generalizable, interpretable knowledge with wake-sleep Bayesian program learning. <i>arXiv preprint</i> .	1397
Richard Evans, Jose Hernandez-Orallo, Johannes Welbl, Pushmeet Kohli, and Marek Sergot. 2019. Making sense of sensory input. <i>arXiv preprint</i> .	1398
Richard Evans, David Saxton, David Amos, Pushmeet Kohli, and Edward Grefenstette. 2018. Can neural networks understand logical entailment? <i>arXiv preprint</i> .	1399
Matan Eyal, Tal Baumel, and Michael Elhadad. 2019. Question answering as an automatic evaluation metric for news article summarization. In <i>Proceedings</i>	1400

1422	of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers), pages 3938–3948, Minneapolis, Minnesota. Association for Computational Linguistics.	1477
1423		1478
1424		1479
1425		1480
1426		1481
1427	Alexander R. Fabbri, Irene Li, Tianwei She, Suyi Li, and Dragomir R. Radev. 2019. Multi-news: A large-scale multi-document summarization dataset and abstractive hierarchical model . In <i>Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics</i> .	1482
1428		1483
1429		1484
1430		
1431		1485
1432		1486
1433	Felix Faltings, Michel Galley, Gerold Hintz, Chris Brockett, Chris Quirk, Jianfeng Gao, and Bill Dolan. 2020. Text editing by command . <i>arXiv preprint</i> .	1487
1434		1488
1435		1489
1436	Angela Fan, Mike Lewis, and Yann Dauphin. 2018. Hierarchical neural story generation . In <i>Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 889–898, Melbourne, Australia. Association for Computational Linguistics.	1490
1437		1491
1438		1492
1439		1493
1440		
1441		
1442	Xiaochao Fan, Hongfei Lin, Liang Yang, Yufeng Diao, Chen Shen, Yonghe Chu, and Yanbo Zou. 2020. Humor detection via an internal and external neural network . <i>Neurocomputing</i> , 394:105–111.	1494
1443		1495
1444		1496
1445		1497
1446	William Fedus, Barret Zoph, and Noam Shazeer. 2022. Switch transformers: Scaling to trillion parameter models with simple and efficient sparsity. <i>Journal of Machine Learning Research</i> , 23(120).	1498
1447		1499
1448		1500
1449		1501
1450	Bjarke Felbo, Alan Mislove, Anders Søgaard, Iyad Rahwan, and Sune Lehmann. 2017. Using millions of emoji occurrences to learn any-domain representations for detecting sentiment, emotion and sarcasm . In <i>Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing</i> . Association for Computational Linguistics.	1502
1451		1503
1452		1504
1453		1505
1454		1506
1455		1507
1456		1508
1457	Christiane Fellbaum, editor. 1998. <i>WordNet: An Electronic Lexical Database</i> . MIT Press, Cambridge, MA.	1509
1458		1510
1459		
1460	John K. Feser, Swarat Chaudhuri, and Isil Dillig. 2015. Synthesizing data structure transformations from input-output examples . In <i>Proceedings of the 36th ACM SIGPLAN Conference on Programming Language Design and Implementation, PLDI ’15</i> , page 229–239, New York, NY, USA. Association for Computing Machinery.	1511
1461		1512
1462		1513
1463		1514
1464		1515
1465		1516
1466		
1467	Susan T. Fiske. 1993. Controlling other people: The impact of power on stereotyping . <i>American Psychologist</i> , 48:621–628.	1517
1468		1518
1469		1519
1470	Dawn P. Flanagan and Shauna G. Dixon. 2014. The Cattell-Horn-Carroll theory of cognitive abilities . In <i>Encyclopedia of Special Education</i> . John Wiley & Sons, Ltd.	1520
1471		1521
1472		1522
1473		
1474	Pierre Flener and Ute Schmid. 2008. An introduction to inductive programming . <i>Artificial Intelligence Review</i> , 29:45–62.	1523
1475		1524
1476		
	Jerry A. Fodor. 1975. <i>The Language of Thought</i> . Harvard University Press, Cambridge, MA.	1525
		1526
	Jerry A. Fodor and Zenon W. Pylyshyn. 1988. Connectionism and cognitive architecture: A critical analysis . <i>Cognition</i> , 28(1):3–71.	1527
		1528
	Mark Forsyth. 2014. <i>The Elements of Eloquence: Secrets of the Perfect Turn of Phrase</i> . Berkley, New York.	1529
		1530
	Jonathan Frankle, Gintare Karolina Dziugaite, Daniel Roy, and Michael Carbin. 2020. Linear mode connectivity and the lottery ticket hypothesis. In <i>International Conference on Machine Learning</i> , pages 3259–3269. PMLR.	
	Lea Frermann, Shay B. Cohen, and Mirella Lapata. 2018. Whodunnit? Crime drama as a case for natural language understanding . <i>Transactions of the Association for Computational Linguistics</i> , 6:1–15.	
	Kathryn J. Friedlander and Philip A. Fine. 2018. “The penny drops”: Investigating insight through the medium of cryptic crosswords . <i>Frontiers in Psychology</i> , 9.	
	Martins Frolovs. 2019. Teaching GPT-2 transformer a sense of humor: How to fine-tune large transformer models on a single GPU in PyTorch . <i>Towards Data Science, Medium</i> .	
	Saadia Gabriel, Asli Celikyilmaz, Rahul Jha, Yejin Choi, and Jianfeng Gao. 2020. Go figure: A meta evaluation of factuality in summarization . <i>arXiv preprint</i> .	
	Evgeniy Gabrilovich and Shaul Markovitch. 2007. Computing semantic relatedness using Wikipedia-based explicit semantic analysis . In <i>IJCAI’07: Proceedings of the 20th International Joint Conference on Artificial Intelligence</i> , page 1606–1611, San Francisco. Morgan Kaufmann.	
	Ge Gao, Eunsol Choi, Yejin Choi, and Luke Zettlemoyer. 2018. Neural metaphor detection in context . In <i>Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing</i> , pages 607–613, Brussels, Belgium. Association for Computational Linguistics.	
	Leo Gao, Stella Biderman, Sid Black, Laurence Golding, Travis Hoppe, Charles Foster, Jason Phang, Horace He, Anish Thite, Noa Nabeshima, Shawn Presser, and Connor Leahy. 2021. The pile: An 800GB dataset of diverse text for language modeling . <i>arXiv preprint</i> .	
	Artur d’Avila Garcez and Luis C. Lamb. 2020. Neurosymbolic AI: The 3rd wave . <i>arXiv preprint</i> .	
	Matt Gardner, Yoav Artzi, Victoria Basmova, Jonathan Berant, Ben Bogin, Sihao Chen, Pradeep Dasigi, Dheeru Dua, Yanai Elazar, Ananth Gottumukkala, Nitish Gupta, Hanna Hajishirzi, Gabriel Ilharco, Daniel Khashabi, Kevin Lin, Jiangming Liu, Nelson F. Liu, Phoebe Mulcaire, Qiang Ning, Sameer	

1643	Bogdan Gliwa, Iwona Mochol, Maciej Biesek, and Aleksander Wawer. 2019. Samsum corpus: A human-annotated dialogue dataset for abstractive summarization . In <i>Proceedings of the 2nd Workshop on New Frontiers in Summarization</i> .	1698
1644		1699
1645		1700
1646		
1647		
1648	Yoav Goldberg. 2019. Assessing BERT’s syntactic abilities . <i>arXiv preprint</i> .	
1649		
1650	Arthur S. Goldberger. 1972. Structural equation methods in the social sciences . <i>Econometrica</i> , 40(6):979–1001.	
1651		
1652		
1653	Hila Gonen and Yoav Goldberg. 2019. Lipstick on a pig: Debiasing methods cover up systematic gender biases in word embeddings but do not remove them . In <i>Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)</i> , pages 609–614, Minneapolis, Minnesota. Association for Computational Linguistics.	
1654		
1655		
1656		
1657		
1658		
1659		
1660		
1661		
1662	Roberto González-Ibáñez, Smaranda Muresan, and Nina Wacholder. 2011. Identifying sarcasm in Twitter: A closer look . In <i>Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies</i> , pages 581–586, Portland, Oregon, USA. Association for Computational Linguistics.	
1663		
1664		
1665		
1666		
1667		
1668		
1669	Noah D. Goodman and Michael C. Frank. 2016. Pragmatic language interpretation as probabilistic inference . <i>Trends in Cognitive Sciences</i> , 20:818–829.	
1670		
1671		
1672	Karthik Gopalakrishnan, Behnam Hedayatnia, Qinqiang Chen, Anna Gottardi, Sanjeev Kwatra, Anu Venkatesh, Raefer Gabriel, and Dilek Hakkani-Tür. 2019. Topical-chat: Towards knowledge-grounded open-domain conversations . In <i>Proc. Interspeech 2019</i> , pages 1891–1895.	
1673		
1674		
1675		
1676		
1677		
1678	Karthik Gopalakrishnan, Behnam Hedayatnia, Longshaokan Wang, Yang Liu, and Dilek Hakkani-Tür. 2020. Are neural open-domain dialog systems robust to speech recognition errors in the dialog history? An empirical study . In <i>Proc. Interspeech 2020</i> , pages 911–915.	
1679		
1680		
1681		
1682		
1683		
1684	Andrew S. Gordon. 2010. Choice of plausible alternatives (COPA) .	
1685		
1686	David Graff, Junbo Kong, Ke Chen, and Kazuaki Maeda. 2003. English gigaword . <i>Linguistic Data Consortium, Philadelphia</i> , 4(1):34.	
1687		
1688		
1689	Alex Graves, Greg Wayne, and Ivo Danihelka. 2014. Neural Turing machines . <i>arXiv preprint</i> .	
1690		
1691	Alex Graves, Greg Wayne, Malcom Reynolds, Tim Harley, Ivo Danihelka, Agnieszka Grabska-Barwińska, Sergio Gómez Colmenarejo, Edward Grefenstette, Tiago Ramalho, John Agapiou, Adrià Puigdomènech Badia, Karl Moritz Hermann, Yori Zwols, Georg Ostrovski, Adam Cain, Helen King, Christopher Summerfield, Phil Blunsom, Koray Kavukcuoglu, and Demis Hassabis. 2016. Hybrid computing using a neural network with dynamic external memory . <i>Nature</i> , 538:471–476.	1698
1692		1699
1693		1700
1694		
1695		
1696		
1697		
	C. Cordell Green, Richard J. Waldinger, David R. Barstow, Robert Elschlager, Douglas B. Lenat, Brian P. McCune, David E. Shaw, and Louis I. Steinberg. 1974. Progress report on program-understanding systems (AIM-240) .	1701
		1702
		1703
		1704
		1705
	Cordell Green. 1981. Application of theorem proving to problem solving . In Bonnie Lynn Webber and Nils J. Nilsson, editors, <i>Readings in Artificial Intelligence</i> , pages 202–222. Morgan Kaufmann.	1706
		1707
		1708
		1709
	H. Paul Grice and Peter F. Strawson. 1956. In defense of a dogma . <i>The Philosophical Review</i> , 65(2):141–158.	1710
		1711
	Aditya Grover, Eric Wang, Aaron Zweig, and Stefano Ermon. 2019. Stochastic optimization of sorting networks via continuous relaxations . <i>arXiv preprint</i> .	1712
		1713
		1714
	Jiatao Gu, Hany Hassan, Jacob Devlin, and Victor O. K. Li. 2018. Universal neural machine translation for extremely low resource languages . <i>arXiv preprint</i> .	1715
		1716
		1717
	Kristina Gulordava, Piotr Bojanowski, Edouard Grave, Tal Linzen, and Marco Baroni. 2018. Colorless green recurrent networks dream hierarchically . In <i>Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)</i> , pages 1195–1205, New Orleans, Louisiana. Association for Computational Linguistics.	1718
		1719
		1720
		1721
		1722
		1723
		1724
		1725
		1726
	Sumit Gulwani. 2011. Automating string processing in spreadsheets using input-output examples . <i>SIGPLAN Not.</i> , 46(1):317–330.	1727
		1728
		1729
	Sumit Gulwani, William R. Harris, and Rishabh Singh. 2012. Spreadsheet data manipulation using examples . <i>Commun. ACM</i> , 55(8):97–105.	1730
		1731
		1732
	Sumit Gulwani, José Hernández-Orallo, Emanuel Kitzelmann, Stephen H. Muggleton, Ute Schmid, and Benjamin Zorn. 2015. Inductive programming meets the real world . <i>Commun. ACM</i> , 58(11):90–99.	1733
		1734
		1735
		1736
	Sumit Gulwani, Oleksandr Polozov, and Rishabh Singh. 2017a. Program synthesis . <i>Foundations and Trends in Programming Languages</i> , 4(1–2):1–119.	1737
		1738
		1739
	Sumit Gulwani, Oleksandr Polozov, and Rishabh Singh. 2017b. Program Synthesis . NOW, Boston.	1740
		1741
	Aditya Gupta, Jiacheng Xu, Shyam Upadhyay, Diyi Yang, and Manaal Faruqui. 2021. Disfl-QA: A benchmark dataset for understanding disfluencies in question answering . In <i>Findings of the Association for Computational Linguistics: ACL-IJCNLP 2021</i> , pages 3309–3319, Online. Association for Computational Linguistics.	1742
		1743
		1744
		1745
		1746
		1747
		1748

1749	Shashank Gupta, Subhabrata Mukherjee, Krishan Subudhi, Eduardo Gonzalez, Damien Jose, Ahmed H Awadallah, and Jianfeng Gao. 2022. Sparsely activated mixture-of-experts are robust multi-task learners. <i>arXiv preprint arXiv:2204.07689</i> .	1804
1750		1805
1751		1806
1752		1807
1753		1808
1754	Isidor S. Gvarjalaže and Dzhuansher I. Mchedlishvili. 1971. <i>English Proverbs and Sayings</i> . Vysshaya shkola, Moscow.	1809
1755		1810
1756		1811
1757	Samuel Gyasi Obeng. 1996. The proverb as a mitigating and politeness strategy in Akan discourse . <i>Anthropological Linguistics</i> , 38(3):521–549.	1812
1758		1813
1759		1814
1760	Ivan Habernal, Henning Wachsmuth, Iryna Gurevych, and Benno Stein. 2018. The argument reasoning comprehension task: Identification and reconstruction of implicit warrants . In <i>Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)</i> , pages 1930–1940, New Orleans, Louisiana. Association for Computational Linguistics.	1815
1761		1816
1762		1817
1763		1818
1764		1819
1765		1820
1766		1821
1767		1822
1768		1823
1769	Michael Hahn. 2020. Theoretical limitations of self-attention in neural sequence models . <i>Transactions of the Association for Computational Linguistics</i> , 8:156–171.	1824
1770		1825
1771		1826
1772		1827
1773	Rowan Hall Maudslay, Hila Gonen, Ryan Cotterell, and Simone Teufel. 2019. It’s all in the name: Mitigating gender bias with name-based counterfactual data substitution . In <i>Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)</i> , pages 5267–5275, Hong Kong, China. Association for Computational Linguistics.	1828
1774		1829
1775		1830
1776		1831
1777		1832
1778		1833
1779		1834
1780		1835
1781		1836
1782	Joseph Y. Halpern. 2016. <i>Actual causality</i> . MIT Press, Cambridge, MA.	1837
1783		1838
1784	Rujun Han, Xiang Ren, and Nanyun Peng. 2020. ECONET: Effective continual pretraining of language models for event temporal reasoning . <i>arXiv preprint</i> .	1839
1785		1840
1786		1841
1787		1842
1788	Maria Hanzén. 2007. When in Rome, do as the Romans do: Proverbs as a part of EFL teaching . Master’s thesis, Jönköping University, School of Education and Communication, Jönköping.	1843
1789		1844
1790		1845
1791		1846
1792	Yiding Hao, William Merrill, Dana Angluin, Robert Frank, Noah Amsel, Andrew Benz, and Simon Mendelsohn. 2018. Context-free transductions with neural stacks . <i>arXiv preprint</i> .	1847
1793		1848
1794		1849
1795		1850
1796	Francesca G.E. Happé. 1994. An advanced test of theory of mind: Understanding of story characters thoughts and feelings by able autistic, mentally handicapped, and normal children and adults . <i>Journal of Autism and Developmental Disorders</i> , 24:129–154.	1851
1797		1852
1798		1853
1799		1854
1800		1855
1801	F. Maxwell Harper and Joseph A. Konstan. 2015. The MovieLens datasets: History and context . <i>ACM Trans. Interact. Intell. Syst.</i> , 5(4).	1856
1802		1857
1803		1858
	Behnam Hedayatnia, Karthik Gopalakrishnan, Seokhwan Kim, Yang Liu, Mihail Eric, and Dilek Hakkani-Tur. 2020. Policy-driven neural response generation for knowledge-grounded dialog systems . In <i>Proceedings of the 13th International Conference on Natural Language Generation</i> , pages 412–421, Dublin, Ireland. Association for Computational Linguistics.	
	Irene Heim. 1983. On the projection problem for pre-suppositions. In Paul Portner and Barbara H. Partee, editors, <i>Formal Semantics - The Essential Readings</i> , pages 249–260. Blackwell, Oxford.	
	Mikael Henaff, Jason Weston, Arthur Szlam, Antoine Bordes, and Yann LeCun. 2016. Tracking the world state with recurrent entity networks . <i>arXiv preprint</i> .	
	Lisa Anne Hendricks, Kaylee Burns, Kate Saenko, Trevor Darrell, and Anna Rohrbach. 2018. Women also snowboard: Overcoming bias in captioning models . In <i>Proceedings of the European Conference on Computer Vision (ECCV)</i> , pages 771–787, Cham. Springer.	
	Dan Hendrycks, Steven Basart, Saurav Kadavath, Mantas Mazeika, Akul Arora, Ethan Guo, Collin Burns, Samir Puranik, Horace He, Dawn Song, and Jacob Steinhardt. 2021a. Measuring coding challenge competence with APPS . <i>arXiv</i> , arXiv:2105.09938.	
	Dan Hendrycks, Collin Burns, Steven Basart, Andrew Critch, Jerry Li, Dawn Song, and Jacob Steinhardt. 2020. Aligning AI with shared human values . <i>arXiv preprint</i> .	
	Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. 2021b. Measuring massive multitask language understanding . In <i>International Conference on Learning Representations</i> .	
	Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. 2021c. Measuring mathematical problem solving with the MATH dataset . <i>arXiv preprint</i> .	
	Tom Henighan, Jared Kaplan, Mor Katz, Mark Chen, Christopher Hesse, Jacob Jackson, Heewoo Jun, Tom B. Brown, Prafulla Dhariwal, Scott Gray, Chris Hallacy, Benjamin Mann, Alec Radford, Aditya Ramesh, Nick Ryder, Daniel M. Ziegler, John Schulman, Dario Amodei, and Sam McCandlish. 2020. Scaling laws for autoregressive generative modeling . <i>arXiv preprint</i> .	
	Joseph Henrich, Steven J. Heine, and Ara Norenzayan. 2010. The weirdest people in the world? <i>Behavioral and Brain Sciences</i> , 33(2-3):61–83.	
	Karl Moritz Hermann, Tomas Kocisky, Edward Grefenstette, Lasse Espeholt, Will Kay, Mustafa Suleyman, and Phil Blunsom. 2015a. Teaching machines to read and comprehend . <i>arXiv preprint arXiv:1506.03340</i> .	

1859	Karl Moritz Hermann, Tomáš Kociský, Edward Grefen-	Yufang Hou, Katja Markert, and Michael Strube. 2013.	1913
1860	stette, Lasse Espeholt, Will Kay, Mustafa Suleyman,	Global inference for bridging anaphora resolution .	1914
1861	and Phil Blunsom. 2015b. Teaching machines to read	In <i>Proceedings of the 2013 Conference of the North</i>	1915
1862	and comprehend . In <i>Advances in Neural Information</i>	<i>American Chapter of the Association for Computa-</i>	1916
1863	<i>Processing Systems</i> , volume 28. Curran Associates,	<i>tional Linguistics: Human Language Technologies</i> ,	1917
1864	Inc.	pages 907–917, Atlanta, Georgia. Association for	1918
		Computational Linguistics.	1919
1865	Danny Hernandez, Jared Kaplan, Tom Henighan, and	Neil Houlsby, Andrei Giurgiu, Stanislaw Jastrzebski,	1920
1866	Sam McCandlish. 2021. Scaling laws for transfer .	Bruna Morrone, Quentin De Laroussilhe, Andrea	1921
1867	<i>arXiv preprint</i> .	Gesmundo, Mona Attariyan, and Sylvain Gelly. 2019.	1922
1868	Álvaro Hernandez, Suhan Woo, Héctor Corrales, Ig-	Parameter-efficient transfer learning for NLP . In <i>In-</i>	1923
1869	nacio Parra, Euntai Kim, D. Fernández Llorca, and	<i>ternational Conference on Machine Learning</i> , pages	1924
1870	Miguel A. Sotelo. 2020. 3D-DEEP: 3-dimensional	2790–2799.	1925
1871	deep-learning based on elevation patterns for road		
1872	scene interpretation . In <i>2020 IEEE Intelligent Ve-</i>	China Household Management Research Center,	1926
1873	<i>hicles Symposium (IV)</i> , Piscataway, NJ. Institute of	Ministry of Public Security. 2019. National name	1927
1874	Electrical and Electronics Engineers.	report 2018. http://news.cpd.com.cn/n18151/	1928
		(Accessed 3	1929
1875	Jonathan Herzig, Pawel Krzysztof Nowak, Thomas	March 2021).	1930
1876	Müller, Francesco Piccinno, and Julian Eisenschlos.		
1877	2020. TaPas: Weakly supervised table parsing via	China Household Management Research Center, Min-	1931
1878	pre-training . In <i>Proceedings of the 58th Annual Meet-</i>	istry of Public Security. 2020. National name	1932
1879	<i>ing of the Association for Computational Linguistics</i> ,	report 2019. https://www.mps.gov.cn/n2254314/	1933
1880	pages 4320–4333, Online. Association for Computa-	n6409334/c6874817/content.html (Accessed 3	1934
1881	tional Linguistics.	March 2021).	1935
1882	John Hewitt, Michael Hahn, Surya Ganguli, Percy	China Household Management Research Center, Min-	1936
1883	Liang, and Christopher D. Manning. 2020. RNNs	istry of Public Security. 2021. National name	1937
1884	can generate bounded hierarchical languages with	report 2020. https://www.mps.gov.cn/n2253534/	1938
1885	optimal memory . <i>arXiv preprint</i> .	n2253535/c7725981/content.html (Accessed 3	1939
		March 2021).	1940
1886	Mireille Hildebrandt. 2018. Algorithmic regulation and	Hrisztalina Hrisztova-Gotthardt and Melita	1941
1887	the rule of law . <i>Philosophical Transactions of the</i>	Aleksa Varga. 2015. Introduction to Paremiol-	1942
1888	<i>Royal Society A: Mathematical, Physical and Engi-</i>	ogy: A Comprehensive Guide to Proverb Studies . De	1943
1889	<i>neering Sciences</i> , 376(2128):20170355.	Gruyter Open, Warsaw.	1944
1890	Keith J. Holyoak. 2012. Analogy and relational reason-	Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan	1945
1891	ing . In Keith J. Holyoak and Robert G. Morrison,	Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and	1946
1892	editors, <i>The Oxford Handbook of Thinking and Rea-</i>	Weizhu Chen. 2021. Lora: Low-rank adaptation of	1947
1893	<i>soning</i> . Oxford University Press, Oxford.	large language models. In <i>International Conference</i>	1948
1894	Richard P. Honeck. 1997. <i>A Proverb in Mind: The</i>	<i>on Learning Representations</i> .	1949
1895	<i>Cognitive Science of Proverbial Wit and Wisdom</i> .		
1896	Lawrence Erlbaum Associates, Mahwah, NJ.	Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan	1950
1897	Alexandra Horowitz. 2017. Smelling themselves: Dogs	Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and	1951
1898	investigate their own odours longer when modified	Weizhu Chen. 2022. LoRA: Low-rank adaptation of	1952
1899	in an “olfactory mirror” test . <i>Behavioural Processes</i> ,	large language models . In <i>International Conference</i>	1953
1900	143:17–24.	<i>on Learning Representations</i> .	1954
1901	Mohammad Javad Hosseini, Hannaneh Hajishirzi, Oren	Chengsong Huang, Qian Liu, Bill Yuchen Lin, Tianyu	1955
1902	Etzioni, and Nate Kushman. 2014. Learning to solve	Pang, Chao Du, and Min Lin. 2023. Lorahub: Effi-	1956
1903	arithmetic word problems with verb categorization .	cient cross-task generalization via dynamic lora	1957
1904	In <i>Proceedings of the 2014 Conference on Empirical</i>	composition. <i>arXiv preprint arXiv:2307.13269</i> .	1958
1905	<i>Methods in Natural Language Processing (EMNLP)</i> ,		
1906	pages 523–533, Doha, Qatar. Association for Com-	Chengsong Huang, Qian Liu, Bill Yuchen Lin, Tianyu	1959
1907	putational Linguistics.	Pang, Chao Du, and Min Lin. 2024a. Lorahub: Effi-	1960
1908	Yufang Hou. 2020. Bridging anaphora resolution as	cient cross-task generalization via dynamic lora	1961
1909	question answering . In <i>Proceedings of the 58th An-</i>	composition . <i>Preprint</i> , arXiv:2307.13269.	1962
1910	<i>nnual Meeting of the Association for Computational</i>		
1911	<i>Linguistics</i> , pages 1428–1438, Online. Association	Daniel Huang, Prafulla Dhariwal, Dawn Song, and Ilya	1963
1912	for Computational Linguistics.	Sutskever. 2018. Gamepad: A learning environment	1964
		for theorem proving . <i>arXiv preprint</i> .	1965

1966	Haoxu Huang, Fanqi Lin, Yingdong Hu, Shengjie Wang, and Yang Gao. 2024b. Copa: General robotic manipulation through spatial constraints of parts with foundation models . <i>Preprint</i> , arXiv:2403.08248.	2022
1967		2023
1968		
1969		
1970	Lifu Huang, Ronan Le Bras, Chandra Bhagavatula, and Yejin Choi. 2019. Cosmos QA: Machine reading comprehension with contextual commonsense reasoning . In <i>Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)</i> , pages 2391–2401, Hong Kong, China. Association for Computational Linguistics.	2024
1971		2025
1972		2026
1973		2027
1974		2028
1975		2029
1976		2030
1977		
1978		
1979	Quzhe Huang, Zhenwei An, Nan Zhuang, Mingxu Tao, Chen Zhang, Yang Jin, Kun Xu, Kun Xu, Liwei Chen, Songfang Huang, and Yansong Feng. 2024c. Harder tasks need more experts: Dynamic routing in moe models . <i>Preprint</i> , arXiv:2403.07652.	2031
1980		2032
1981		2033
1982		2034
1983		2035
1984	Drew A. Hudson and Christopher D. Manning. 2019. GQA: A new dataset for real-world visual reasoning and compositional question answering . <i>arXiv preprint</i> .	2036
1985		2037
1986		2038
1987		
1988	Thad Hughes and Daniel Ramage. 2007. Lexical semantic relatedness with random graph walks . In <i>Proceedings of the 2007 Joint Conference on Empirical Methods in Natural Language Processing and Computational Natural Language Learning (EMNLP-CoNLL)</i> , pages 581–589, Prague, Czech Republic. Association for Computational Linguistics.	2039
1989		2040
1990		2041
1991		2042
1992		2043
1993		2044
1994		
1995	David Hume. 1739–1740. <i>A Treatise of Human Nature</i> . John Noon, London.	2045
1996		2046
1997	Dieuwke Hupkes, Verna Dankers, Mathijs Mul, and Elia Bruni. 2020. Compositionality decomposed: How do neural networks generalise? <i>Journal of Artificial Intelligence Research</i> , 67:757–795.	2047
1998		2048
1999		2049
2000		2050
2001	Annamarie W. Huttunen, Geoffrey K. Adams, and Michael L. Platt. 2017. Can self-awareness be taught? Monkeys pass the mirror test – again . <i>Proceedings of the National Academy of Sciences</i> , 114(13):3281–3283.	2051
2002		
2003		
2004		
2005		
2006	David Huynh and Stefano Mazzocchi. 2012. OpenRefine. https://openrefine.org/ .	2052
2007		2053
2008	Gabriel Ilharco, Marco Tulio Ribeiro, Mitchell Wortsman, Suchin Gururangan, Ludwig Schmidt, Hananeh Hajishirzi, and Ali Farhadi. 2022. Editing models with task arithmetic. <i>arXiv preprint arXiv:2212.04089</i> .	2054
2009		2055
2010		2056
2011		2057
2012		2058
2013	Instagram Engineering. 2015. Emojiengineering part 1: Machine learning for emoji trends . <i>Medium</i> .	2059
2014		2060
2015	Daphne Ippolito, Daniel Duckworth, Chris Callison-Burch, and Douglas Eck. 2020. Automatic detection of generated text is easiest when humans are fooled . In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 1808–1822, Online. Association for Computational Linguistics.	2061
2016		2062
2017		2063
2018		2064
2019		2065
2020		
2021		
	Geoffrey Irving, Paul Christiano, and Dario Amodei. 2018. AI safety via debate . <i>arXiv preprint</i> .	2066
		2067
		2068
		2069
	Gautier Izacard and Edouard Grave. 2020. Leveraging passage retrieval with generative models for open domain question answering . <i>arXiv preprint</i> .	2070
		2071
		2072
	Kushal Jain, Adwait Deshpande, Kumar Shridhar, Felix Laumann, and Ayushman Dash. 2020. Indic-transformers: An analysis of transformer language models for indian languages . <i>arXiv preprint</i> .	2073
		2074
		2075
		2076
		2077
	Joel Jang, Seungone Kim, Seonghyeon Ye, Doyoung Kim, Lajanugen Logeswaran, Moontae Lee, Kyungjae Lee, and Minjoon Seo. 2023. Exploring the benefits of training expert language models over instruction tuning. <i>arXiv preprint arXiv:2302.03202</i> .	
	Mario Jarmasz. 2012. Roget’s Thesaurus as a lexical resource for natural language processing . Master’s thesis, University of Ottawa, Ottawa.	
	Prashant Jayannavar, Anjali Narayan-Chen, and Julia Hockenmaier. 2020. Learning to execute instructions in a Minecraft dialogue . In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 2589–2602, Online. Association for Computational Linguistics.	
	Paloma Jeretic, Alex Warstadt, Suvarat Bhooshan, and Adina Williams. 2020. Are natural language inference models IMPPRESSive? Learning IMPLicature and PRESupposition . In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 8690–8705, Online. Association for Computational Linguistics.	
	Jay J. Jiang and David W. Conrath. 1997. Semantic similarity based on corpus statistics and lexical taxonomy . In <i>Proceedings of the 10th Research on Computational Linguistics International Conference</i> , pages 19–33, Taipei, Taiwan. The Association for Computational Linguistics and Chinese Language Processing (ACLCLP).	
	Nanjiang Jiang and Marie-Catherine de Marneffe. 2019. Do you know that Florence is packed with visitors? Evaluating state-of-the-art models of speaker commitment . In <i>Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics</i> , pages 4208–4213, Florence, Italy. Association for Computational Linguistics.	
	Xisen Jin, Xiang Ren, Daniel Preotiuc-Pietro, and Pengxiang Cheng. 2022. Dataless knowledge fusion by merging weights of language models. <i>arXiv preprint arXiv:2212.09849</i> .	
	Matt Gardner Johannes Welbl, Nelson F. Liu. 2017. Crowdsourcing multiple choice science questions. <i>arXiv:1707.06209v1</i> .	
	Justin Johnson, Bharath Hariharan, Laurens van der Maaten, Li Fei-Fei, C. Lawrence Zitnick, and Ross Girshick. 2016. CLEVR: A diagnostic dataset for compositional language and elementary visual reasoning . <i>arXiv preprint</i> .	

2189	Andrew Kim, Maxim Ruzmaykin, Aaron Truong, and Adam Summerville. 2019. Cooperation and code-names: Understanding natural language processing via codenames . In <i>Proceedings of the AAAI Conference on Artificial Intelligence and Interactive Digital Entertainment</i> , volume 15, pages 160–166, Menlo Park, CA. Association for the Advancement of Artificial Intelligence.	2244
2190		2245
2191		2246
2192		2247
2193		2248
2194		
2195		
2196		
2197	Yoon Kim, Yacine Jernite, David Sontag, and Alexander M. Rush. 2015. Character-aware neural language models . <i>arXiv preprint</i> .	
2198		
2199		
2200	Milton King and Paul Cook. 2020. Evaluating approaches to personalizing language models . In <i>Proceedings of the 12th Language Resources and Evaluation Conference</i> , pages 2461–2469, Marseille, France. European Language Resources Association.	
2201		
2202		
2203		
2204		
2205	Christo Kirov and Ryan Cotterell. 2018. Recurrent Neural Networks in Linguistic Theory: Revisiting Pinker and Prince (1988) and the Past Tense Debate . <i>Transactions of the Association for Computational Linguistics</i> , 6:651–665.	
2206		
2207		
2208		
2209		
2210	Emanuel Kitzelmann. 2010. Inductive programming: A survey of program synthesis techniques . In <i>Approaches and Applications of Inductive Programming</i> , pages 50–73, Berlin. Springer.	
2211		
2212		
2213		
2214	Joshua Knobe. 2003. Intentional action and side effects in ordinary language . <i>Analysis</i> , 63.	
2215		
2216	Vid Kocijan, Ana-Maria Cretu, Oana-Maria Camburu, Yordan Yordanov, and Thomas Lukasiewicz. 2019. A surprisingly robust trick for the Winograd schema challenge . In <i>Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics</i> , pages 4837–4842, Florence, Italy. Association for Computational Linguistics.	
2217		
2218		
2219		
2220		
2221		
2222		
2223	Tomáš Kočický, Jonathan Schwarz, Phil Blunsom, Chris Dyer, Karl Moritz Hermann, Gábor Melis, and Edward Grefenstette. 2018. The NarrativeQA reading comprehension challenge . <i>Transactions of the Association for Computational Linguistics</i> , 6:317–328.	
2224		
2225		
2226		
2227		
2228	Jan Kocoń, Piotr Miłkowski, and Kamil Kanclerz. 2021. MultiEmo: Multilingual, multilevel, multidomain sentiment analysis corpus of consumer reviews . In <i>Computational Science – ICCS 2021</i> , pages 297–312, Cham. Springer.	
2229		
2230		
2231		
2232		
2233	Alexander W. Kocurek and Ethan Jerzak. 2021. Counterlogicals as counterconventionals . <i>Journal of Philosophical Logic</i> , 50:673–704.	
2234		
2235		
2236	Alexander W. Kocurek, Ethan Jerzak, and Rachel Etta Rudolph. 2020. Against conventional wisdom . <i>Philosophers’ Imprint</i> , 20(22):1–27.	
2237		
2238		
2239	Moshe Koppel and Jonathan Schler. 2004. Authorship verification as a one-class classification problem . In <i>Proceedings of the Twenty-First International Conference on Machine Learning</i> , page 62, New York, NY, USA. Association for Computing Machinery.	
2240		
2241		
2242		
2243		
	Jarmo Korhonen. 2009. Sprichwörter und zweisprachige lexikographie: Deutsch-schwedische und deutsch-finnische wörterbücher im vergleich. In C. Földes, editor, <i>Phraseologie disziplinär und interdisziplinär</i> , pages 537–549. Gunter Narr Verlag.	2249
		2250
		2251
		2252
		2253
		2254
		2255
		2256
	Dimitrios Kotsakos, Theodoros Lappas, Dimitrios Kotzias, Dimitrios Gunopulos, Nattiya Kanhabua, and Kjetil Nørvåg. 2014. A burstiness-aware approach for document dating . In <i>Proceedings of the 37th International ACM SIGIR Conference on Research & Development in Information Retrieval, SIGIR ’14</i> , page 1003–1006, New York, NY, USA. Association for Computing Machinery.	2257
		2258
		2259
	Samuel Kounev, Jeffrey O. Kephart, Aleksandar Milenkoski, and Xiaoyun Zhu, editors. 2017. <i>Self-Aware Computing Systems</i> . Springer, Cham.	2260
		2261
		2262
		2263
	Sarah E. Kreps, Miles McCain, and Miles Brundage. 2020. All the news that’s fit to fabricate: AI-generated text as a tool of media misinformation . <i>SSRN</i> .	2264
		2265
		2266
	Kalpesh Krishna, Aurko Roy, and Mohit Iyer. 2021. Hurdles to progress in long-form question answering . <i>arXiv preprint</i> .	2267
		2268
		2269
		2270
		2271
		2272
		2273
	Wojciech Kryscinski, Bryan McCann, Caiming Xiong, and Richard Socher. 2020. Evaluating the factual consistency of abstractive text summarization . In <i>Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)</i> , pages 9332–9346, Online. Association for Computational Linguistics.	2274
		2275
		2276
		2277
		2278
		2279
		2280
		2281
		2282
	Tom Kwiatkowski, Jennimaria Palomaki, Olivia Redfield, Michael Collins, Ankur Parikh, Chris Alberti, Danielle Epstein, Illia Polosukhin, Jacob Devlin, Kenton Lee, Kristina N. Toutanova, Llion Jones, Matthew Kelcey, Ming-Wei Chang, Andrew M. Dai, Jakob Uszkoreit, Quoc V. Le, and Slav Petrov. 2019. Natural questions: A benchmark for question answering research . <i>Transactions of the Association for Computational Linguistics</i> , 7:453–466.	2283
		2284
		2285
		2286
	Niklas Kühl, Marc Goutier, Lucas Baier, Clemens Wolff, and Dominik Martin. 2020. Human vs. supervised machine learning: Who learns patterns faster? <i>arXiv preprint</i> .	2287
		2288
		2289
		2290
	Heinrich Küttler, Nantas Nardelli, Alexander H. Miller, Roberta Raileanu, Marco Selvatici, Edward Grefenstette, and Tim Rocktäschel. 2020. The NetHack learning environment . <i>arXiv preprint</i> .	2291
		2292
	Kevin Lacker. 2020. Giving GPT-3 a Turing test . <i>Kevin Lacker’s blog</i> .	2293
		2294
		2295
		2296
		2297
		2298
		2299
	Guokun Lai, Qizhe Xie, Hanxiao Liu, Yiming Yang, and Eduard Hovy. 2017. RACE: Large-scale Reading comprehension dataset from examinations . In <i>Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing</i> , pages 785–794, Copenhagen, Denmark. Association for Computational Linguistics.	

2300	Brenden M. Lake and Marco Baroni. 2017. Generalization without systematicity: On the compositional skills of sequence-to-sequence recurrent networks . <i>arXiv preprint</i> .	2354
2301		2355
2302		2356
2303		2357
2304	Brenden M. Lake and Gregory L. Murphy. 2020. Word meaning in minds and machines . <i>arXiv preprint</i> .	2358
2305		2359
2306		2360
2307	Brenden M. Lake, Tomer D. Ullman, Joshua B. Tenenbaum, and Samuel J. Gershman. 2017. Building machines that learn and think like people . <i>Behavioral and Brain Sciences</i> , 40:e253.	2361
2308		2362
2309		2363
2310	George Lakoff and Mark Johnson. 2008. <i>Metaphors We Live By</i> . University of Chicago Press, Chicago.	2364
2311		2365
2312	Yair Lakretz, Théo Desbordes, Jean-Rémi King, Benoît Crabbé, Maxime Oquab, and Stanislas Dehaene. 2021a. Can RNNs learn recursive nested subject-verb agreements? <i>arXiv preprint</i> .	2366
2313		2367
2314		2368
2315		2369
2316	Yair Lakretz, Dieuwke Hupkes, Alessandra Vergallito, Marco Marelli, Marco Baroni, and Stanislas Dehaene. 2021b. Mechanisms for handling nested dependencies in neural-network language models and humans . <i>Cognition</i> , 213:104699. Special Issue in Honour of Jacques Mehler, Cognition’s founding editor.	2370
2317		2371
2318		2372
2319		2373
2320		2374
2321		2375
2322	Yair Lakretz, German Kruszewski, Theo Desbordes, Dieuwke Hupkes, Stanislas Dehaene, and Marco Baroni. 2019. The emergence of number and syntax units in LSTM language models . In <i>Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)</i> , pages 11–20, Minneapolis, Minnesota. Association for Computational Linguistics.	2376
2323		2377
2324		2378
2325		2379
2326		2380
2327		2381
2328		2382
2329		2383
2330		2384
2331	Guillaume Lample and François Charton. 2019. Deep learning for symbolic mathematics . <i>arXiv preprint</i> .	2385
2332		2386
2333	Zhenzhong Lan, Mingda Chen, Sebastian Goodman, Kevin Gimpel, Piyush Sharma, and Radu Soricut. 2019. ALBERT: A lite BERT for self-supervised learning of language representations . <i>arXiv preprint</i> .	2387
2334		2388
2335		2389
2336		2390
2337	Matthew Le, Y-Lan Boureau, and Maximilian Nickel. 2019. Revisiting the evaluation of theory of mind through question answering . In <i>Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)</i> , pages 5872–5877, Hong Kong, China. Association for Computational Linguistics.	2391
2338		2392
2339		2393
2340		2394
2341		2395
2342		2396
2343		2397
2344		2398
2345	Remi Lebret, David Grangier, and Michael Auli. 2016a. Neural text generation from structured data with application to the biography domain . <i>Preprint</i> , arXiv:1603.07771.	2399
2346		2400
2347		2401
2348		2402
2349	Rémi Lebret, David Grangier, and Michael Auli. 2016b. The wikibio corpus: A corpus of biographical texts for natural language generation . In <i>Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing</i> .	2403
2350		2404
2351		2405
2352		2406
2353		2407
		2408
		2409
	Nayeon Lee, Yejin Bang, Andrea Madotto, and Pascale Fung. 2021. Towards few-shot fact-checking via perplexity . In <i>Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies</i> , pages 1971–1981, Online. Association for Computational Linguistics.	
	Nayeon Lee, Belinda Z. Li, Sinong Wang, Wen-tau Yih, Hao Ma, and Madian Khabsa. 2020. Language models as fact checkers? In <i>Proceedings of the Third Workshop on Fact Extraction and VERification (FEVER)</i> , pages 36–41, Online. Association for Computational Linguistics.	
	Jan Leike, David Krueger, Tom Everitt, Miljan Martić, Vishal Maini, and Shane Legg. 2018. Scalable agent alignment via reward modeling: A research direction . <i>arXiv preprint</i> .	
	Dmitry Lepikhin, HyoukJoong Lee, Yuanzhong Xu, Dehao Chen, Orhan Firat, Yanping Huang, Maxim Krikun, Noam Shazeer, and Zhifeng Chen. 2020. Gshard: Scaling giant models with conditional computation and automatic sharding. <i>arXiv preprint arXiv:2006.16668</i> .	
	Brian Lester, Rami Al-Rfou, and Noah Constant. 2021. The power of scale for parameter-efficient prompt tuning . <i>Preprint</i> , arXiv:2104.08691.	
	Iddo Lev, Bill MacCartney, Christopher Manning, and Roger Levy. 2004. Solving logic puzzles: From robust processing to precise semantics . In <i>Proceedings of the 2nd Workshop on Text Meaning and Interpretation</i> , pages 9–16, Barcelona, Spain. Association for Computational Linguistics.	
	Hector Levesque, Ernest Davis, and Leora Morgenstern. 2012. The winograd schema challenge. In <i>Thirteenth International Conference on the Principles of Knowledge Representation and Reasoning</i> .	
	Omer Levy, Minjoon Seo, Eunsol Choi, and Luke Zettlemoyer. 2017. Zero-shot relation extraction via reading comprehension . In <i>Proceedings of the 21st Conference on Computational Natural Language Learning (CoNLL 2017)</i> , pages 333–342, Vancouver, Canada. Association for Computational Linguistics.	
	Ran Levy, Liat Ein-Dor, Shay Hummel, Ruty Rinott, and Noam Slonim. 2015. TR9856: A multi-word term relatedness benchmark . In <i>Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference on Natural Language Processing (Volume 2: Short Papers)</i> , pages 419–424, Beijing, China. Association for Computational Linguistics.	
	Sharon Levy, Michael Saxon, and William Yang Wang. 2021. Investigating memorization of conspiracy theories in text generation . <i>arXiv preprint</i> .	
	Patrick Lewis, Barlas Oguz, Ruty Rinott, Sebastian Riedel, and Holger Schwenk. 2020a. MLQA: Evaluating cross-lingual extractive question answering . In	

2410	<i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 7315–7330, Online. Association for Computational Linguistics.	2466
2411		2467
2412		2468
2413		2469
2414	Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, Sebastian Riedel, and Douwe Kiela. 2020b. Retrieval-augmented generation for knowledge-intensive NLP tasks . <i>arXiv preprint</i> .	2470
2415		2471
2416		2472
2417		
2418		
2419		
2420	Patrick Lewis, Pontus Stenetorp, and Sebastian Riedel. 2020c. Question and answer test-train overlap in open-domain question answering datasets . <i>arXiv preprint</i> .	
2421		
2422		
2423		
2424	Pingzhi Li, Zhenyu Zhang, Prateek Yadav, Yi-Lin Sung, Yu Cheng, Mohit Bansal, and Tianlong Chen. 2023. Merge, then compress: Demystify efficient smoe with hints from its routing policy. <i>arXiv preprint arXiv:2310.01334</i> .	
2425		
2426		
2427		
2428		
2429	Tao Li, Daniel Khashabi, Tushar Khot, Ashish Sabharwal, and Vivek Srikumar. 2020a. UNQOVERing stereotyping biases via underspecified questions . In <i>Findings of the Association for Computational Linguistics: EMNLP 2020</i> , pages 3475–3489, Online. Association for Computational Linguistics.	
2430		
2431		
2432		
2433		
2434		
2435	Yanran Li, Hui Su, Xiaoyu Shen, Wenjie Li, Ziqiang Cao, and Shuzi Niu. 2017. DailyDialog: A manually labelled multi-turn dialogue dataset . In <i>Proceedings of the Eighth International Joint Conference on Natural Language Processing (Volume 1: Long Papers)</i> , pages 986–995, Taipei, Taiwan. Asian Federation of Natural Language Processing.	
2436		
2437		
2438		
2439		
2440		
2441		
2442	Yiwei Li, G Brian Golding, and Lucian Ilie. 2020b. DELPHI: Accurate deep ensemble model for protein interaction sites prediction . <i>Bioinformatics</i> , 37(7):896–904.	
2443		
2444		
2445		
2446	Yuhua Li, Zuhair A. Bandar, and David Mclean. 2003. An approach for measuring semantic similarity between words using multiple information sources . <i>IEEE Transactions on Knowledge and Data Engineering</i> , 15(4):871–882.	
2447		
2448		
2449		
2450		
2451	Chao-Chun Liang, Yu-Shiang Wong, Yi-Chung Lin, and Keh-Yih Su. 2018. A meaning-based statistical English math word problem solver . In <i>Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)</i> , pages 652–662, New Orleans, Louisiana. Association for Computational Linguistics.	
2452		
2453		
2454		
2455		
2456		
2457		
2458		
2459	Paul Pu Liang, Irene Mengze Li, Emily Zheng, Yao Chong Lim, Ruslan Salakhutdinov, and Louis-Philippe Morency. 2020a. Towards debiasing sentence representations . In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 5502–5515, Online. Association for Computational Linguistics.	
2460		
2461		
2462		
2463		
2464		
2465		
	Zujie Liang, Weitao Jiang, Haifeng Hu, and Jiaying Zhu. 2020b. Learning to contrast the counterfactual samples for robust visual question answering . In <i>Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)</i> , pages 3285–3292, Online. Association for Computational Linguistics.	2466
		2467
		2468
		2469
		2470
		2471
		2472
	Bill Yuchen Lin, Seyeon Lee, Rahul Khanna, and Xiang Ren. 2020a. Birds have four legs?! NumerSense: probing numerical commonsense knowledge of pre-trained language models . In <i>Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)</i> , pages 6862–6868, Online. Association for Computational Linguistics.	2473
		2474
		2475
		2476
		2477
		2478
		2479
	Bill Yuchen Lin, Ziyi Wu, Yichi Yang, Dong-Ho Lee, and Xiang Ren. 2021a. Riddlesense: Reasoning about riddle questions featuring linguistic creativity and commonsense knowledge . <i>arXiv preprint</i> .	2480
		2481
		2482
		2483
	Bill Yuchen Lin, Wangchunshu Zhou, Ming Shen, Pei Zhou, Chandra Bhagavatula, Yejin Choi, and Xiang Ren. 2020b. CommonGen: A constrained text generation challenge for generative commonsense reasoning . <i>Preprint</i> , arXiv:1911.03705.	2484
		2485
		2486
		2487
		2488
	Bill Yuchen Lin, Wangchunshu Zhou, Ming Shen, Pei Zhou, Chandra Bhagavatula, Yejin Choi, and Xiang Ren. 2020c. CommonGen: A constrained text generation challenge for generative commonsense reasoning . In <i>Findings of the Association for Computational Linguistics: EMNLP 2020</i> , pages 1823–1840, Online. Association for Computational Linguistics.	2489
		2490
		2491
		2492
		2493
		2494
		2495
	Kevin Lin, Oyvind Tafjord, Peter Clark, and Matt Gardner. 2019a. Reasoning over paragraph effects in situations . In <i>Proceedings of the 2nd Workshop on Machine Reading for Question Answering</i> , pages 58–62, Hong Kong, China. Association for Computational Linguistics.	2496
		2497
		2498
		2499
		2500
		2501
	Kevin Lin, Ben Tan, Swabha Swayamdipta, Noah A. Smith, and Yejin Choi. 2019b. Ropes: Reading comprehension over paragraphs . In <i>Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)</i> .	2502
		2503
		2504
		2505
		2506
		2507
		2508
	Stephanie Lin, Jacob Hilton, and Owain Evans. 2021b. TruthfulQA: Measuring how models mimic human falsehoods . <i>arXiv preprint</i> .	2509
		2510
		2511
	Zhisheng Lin, Han Fu, Chenghao Liu, Zhuo Li, and Jianling Sun. 2024. Pemt: Multi-task correlation guided mixture-of-experts enables parameter-efficient transfer learning. <i>arXiv preprint arXiv:2402.15082</i> .	2512
		2513
		2514
		2515
	Wang Ling, Dani Yogatama, Chris Dyer, and Phil Blunsom. 2017. Program induction by rationale generation: Learning to solve and explain algebraic word problems . <i>arXiv preprint</i> .	2516
		2517
		2518
		2519

2520	Tal Linzen, Emmanuel Dupoux, and Yoav Goldberg.	Shayne Longpre, Le Hou, Tu Vu, Albert Webson,	2575
2521	2016. Assessing the ability of LSTMs to learn syntax-sensitive dependencies . <i>Transactions of the Association for Computational Linguistics</i> , 4:521–535.	Hyung Won Chung, Yi Tay, Denny Zhou, Quoc V Le, Barret Zoph, Jason Wei, et al. 2023. The flan collection: Designing data and methods for effective instruction tuning. <i>arXiv preprint arXiv:2301.13688</i> .	2576
2522			2577
2523			2578
2524	Haokun Liu, Derek Tam, Mohammed Muqeeth, Jay Mohita, Tenghao Huang, Mohit Bansal, and Colin A Raffel. 2022. Few-shot parameter-efficient fine-tuning is better and cheaper than in-context learning. <i>Advances in Neural Information Processing Systems</i> , 35:1950–1965.	Ilya Loshchilov and Frank Hutter. 2017. Decoupled weight decay regularization . In <i>International Conference on Learning Representations</i> .	2579
2525			2580
2526			2581
2527			2582
2528			
2529		Nicholas Lourie, Ronan Le Bras, Chandra Bhagavatula, and Yejin Choi. 2021. Unicorn on rainbow: A universal commonsense reasoning model on a new multitask benchmark . <i>arXiv preprint</i> .	2583
2530	Jerry Liu. 2024. LlamaIndex, a data framework for your LLM applications. https://github.com/run-llama/llama_index .		2584
2531			2585
2532			2586
2533	Jiachang Liu, Dinghan Shen, Yizhe Zhang, Bill Dolan, Lawrence Carin, and Weizhu Chen. 2021a. What makes good in-context examples for GPT-3? <i>arXiv preprint</i> .	Nicholas Lourie, Ronan Le Bras, and Yejin Choi. 2020. Scruples: A corpus of community ethical judgments on 32,000 real-life anecdotes . <i>arXiv preprint</i> .	2587
2534			2588
2535			2589
2536			
2537	Jian Liu, Leyang Cui, Hanmeng Liu, Dandan Huang, Yile Wang, and Yue Zhang. 2020a. LogiQA: A challenge dataset for machine reading comprehension with logical reasoning . <i>arXiv preprint</i> .	Kaiji Lu, Piotr Mardziel, Fangjing Wu, Preetam Amancharla, and Anupam Datta. 2020. Gender bias in neural natural language processing . In Vivek Nigam, Tajana Ban Kirigin, Carolyn Talcott, Joshua Guttman, Stepan Kuznetsov, Boon Thau Loo, and Mitsuhiro Okada, editors, <i>Logic, Language, and Security</i> . Springer, Cham.	2590
2538			2591
2539			2592
2540			2593
2541	Nelson F. Liu, Tony Lee, Robin Jia, and Percy Liang. 2021b. Can small and synthetic benchmarks drive modeling innovation? A retrospective study of question answering modeling approaches . <i>arXiv preprint</i> .		2594
2542			2595
2543			2596
2544			
2545	Tianyu Liu, Yizhe Zhang, Chris Brockett, Yi Mao, Zhifang Sui, Weizhu Chen, and Bill Dolan. 2021c. A token-level reference-free hallucination detection benchmark for free-form text generation . <i>arXiv preprint</i> .	Keming Lu, Hongyi Yuan, Runji Lin, Junyang Lin, Zheng Yuan, Chang Zhou, and Jingren Zhou. 2023. Routing to the expert: Efficient reward-guided ensemble of large language models. <i>arXiv preprint arXiv:2311.08692</i> .	2597
2546			2598
2547			2599
2548			2600
2549			2601
2550	Ye Liu, Shaika Chowdhury, Chenwei Zhang, Cornelia Caragea, and Philip S. Yu. 2020b. Interpretable multi-step reasoning with knowledge extraction on complex healthcare question answering . <i>arXiv preprint</i> .	Zhenyi Lu, Chenghao Fan, Wei Wei, Xiaoye Qu, Danyang Chen, and Yu Cheng. 2024. Twin-merging: Dynamic integration of modular expertise in model merging. <i>arXiv preprint arXiv:2406.15479</i> .	2602
2551			2603
2552			2604
2553			2605
2554	Yinhan Liu, Jiatao Gu, Naman Goyal, Xian Li, Sergey Edunov, Marjan Ghazvininejad, Mike Lewis, and Luke Zettlemoyer. 2020c. Multilingual denoising pre-training for neural machine translation . <i>arXiv preprint</i> .	Jelena Luketina, Nantas Nardelli, Gregory Farquhar, Jakob N. Foerster, Jacob Andreas, Edward Grefenstette, Shimon Whiteson, and Tim Rocktäschel. 2019. A survey of reinforcement learning informed by natural language . In <i>Proceedings of the Twenty-Eighth International Joint Conference on Artificial Intelligence, IJCAI-19</i> , page 6309–6317.	2606
2555			2607
2556			2608
2557			2609
2558			2610
2559	Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. 2019. RoBERTa: A robustly optimized BERT pretraining approach . <i>arXiv preprint</i> .		2611
2560		Mingyu Derek Ma, Jiao Sun, Mu Yang, Kung-Hsiang Huang, Nuan Wen, Shikhar Singh, Rujun Han, and Nanyun Peng. 2021. EventPlus: A temporal event understanding pipeline . <i>arXiv preprint</i> .	2612
2561			2613
2562			2614
2563			2615
2564	Hector Llorens, Nathanael Chambers, Naushad UzZaman, Nasrin Mostafazadeh, James Allen, and James Pustejovsky. 2015. SemEval-2015 task 5: QA TempEval - evaluating temporal information understanding with question answering . In <i>Proceedings of the 9th International Workshop on Semantic Evaluation (SemEval 2015)</i> , pages 792–800, Denver, Colorado. Association for Computational Linguistics.	Andrew L. Maas, Raymond E. Daly, Peter T. Pham, Dan Huang, Andrew Y. Ng, and Christopher Potts. 2011a. Learning word vectors for sentiment analysis . In <i>Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies</i> .	2616
2565			2617
2566			2618
2567			2619
2568			2620
2569			2621
2570			2622
2571		Andrew L. Maas, Raymond E. Daly, Peter T. Pham, Dan Huang, Andrew Y. Ng, and Christopher Potts. 2011b. Learning word vectors for sentiment analysis . In <i>Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies</i> .	2623
2572	Lajanugen Logeswaran, Honglak Lee, and Samy Bengio. 2018. Content preserving text generation with attribute controls . <i>arXiv preprint</i> .		2624
2573			2625
2574			2626
			2627
			2628

2629	Andrea Madotto, Zhaojiang Lin, Genta Indra Winata, and Pascale Fung. 2021. Few-shot bot: Prompt-based learning for dialogue systems .	2683	Philip Massey, Patrick Xia, David Bamman, and Noah A. Smith. 2015. Annotating character relationships in literary texts . <i>arXiv preprint</i> .	2684
2630		2685		
2631				
2632	Andrea Madotto, Zihan Liu, Zhaojiang Lin, and Pascale Fung. 2020. Language models as few-shot learner for task-oriented dialogue systems . <i>arXiv preprint</i> .	2686	Michael S Matena and Colin A Raffel. 2022. Merging models with fisher-weighted averaging. <i>Advances in Neural Information Processing Systems</i> , 35:17703–17716.	2687
2633		2688		
2634		2689		
2635	Eric Malmi, Sebastian Krause, Sascha Rothe, Daniil Mirylenka, and Aliaksei Severyn. 2019. Encode, tag, realize: High-precision text editing . <i>arXiv preprint</i> .	2690	Matthew Matero, Akash Idnani, Youngseo Son, Salvatore Giorgi, Huy Vu, Mohammad Zamani, Parth Limbachiya, Sharath Chandra Guntuku, and H. Andrew Schwartz. 2019. Suicide risk assessment with multi-level dual-context language and BERT . In <i>Proceedings of the Sixth Workshop on Computational Linguistics and Clinical Psychology</i> , pages 39–44, Minneapolis, Minnesota. Association for Computational Linguistics.	2691
2636		2692		
2637		2693		
2638	Eric Malmi, Daniele Pighin, Sebastian Krause, and Mikhail Kozhevnikov. 2018. Automatic prediction of discourse connectives . In <i>Proceedings of the Eleventh International Conference on Language Resources and Evaluation (LREC 2018)</i> , Miyazaki, Japan. European Language Resources Association.	2694		
2639		2695		
2640		2696		
2641		2697		
2642		2698		
2643				
2644	Jihang Mao and Wanli Liu. 2019. A BERT-based approach for automatic humor detection and scoring . In <i>Proceedings of the Iberian Languages Evaluation Forum (IberLEF 2019)</i> , pages 197–202.	2699	Maxine. 2023. Llama-2, mo’ lora. https://crumbly.medium.com/llama-2-molara-f5f909434711 .	2700
2645				
2646				
2647				
2648	Gary Marcus. 2020. The next decade in AI: Four steps towards robust artificial intelligence . <i>arXiv preprint</i> .	2701	Chandler May, Alex Wang, Shikha Bordia, Samuel R. Bowman, and Rachel Rudinger. 2019. On measuring social biases in sentence encoders . In <i>Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)</i> , pages 622–628, Minneapolis, Minnesota. Association for Computational Linguistics.	2702
2649		2703		
2650	Gary Marcus and Ernest Davis. 2020. GPT-3, bloviator: OpenAI’s language generator has no idea what it’s talking about . <i>MIT Technology Review</i> .	2704		
2651		2705		
2652		2706		
2653	Mitchell Marcus, Grace Kim, Mary Ann Marcinkiewicz, Robert MacIntyre, Ann Bies, Mark Ferguson, Karen Katz, and Britta Schasberger. 1994. The Penn Treebank: Annotating predicate argument structure . In <i>Human Language Technology: Proceedings of a Workshop held at Plainsboro, New Jersey, March 8-11, 1994</i> .	2707		
2654		2708		
2655		2709	Andrew Mayne. 2020. OpenAI API alchemy: Emoji storytelling . <i>Andrew Mayne blog</i> .	2710
2656		2711	Joshua Maynez, Shashi Narayan, Bernd Bohnet, and Ryan McDonald. 2020. On faithfulness and factuality in abstractive summarization . <i>arXiv preprint</i> .	2712
2657		2713		
2658		2714	Eric Mays, Fred J. Damerau, and Robert L. Mercer. 1991. Context based spelling correction . <i>Information Processing & Management</i> , 27(5):517–522.	2715
2659		2716		
2660	Katja Markert, Yufang Hou, and Michael Strube. 2012. Collective classification for fine-grained information status . In <i>Proceedings of the 50th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 795–804, Jeju Island, Korea. Association for Computational Linguistics.	2717	Momoh Karmah Mbogba, Zeeshan Haider, S.M. Chapal Hossain, Daobin Huang, Kashan Memon, Fazil Panhwar, Zeling Lei, and Gang Zhao. 2018. The application of convolution neural network based cell segmentation during cryopreservation . <i>Cryobiology</i> , 85:95–104.	2718
2661		2719		
2662		2720		
2663		2721		
2664		2722		
2665				
2666	Kim Marriott, Bongshin Lee, Matthew Butler, Ed Cutrell, Kirsten Ellis, Cagatay Goncu, Marti Hearst, Kathleen McCoy, and Danielle Albers Szafir. 2021. Inclusive data visualization for people with disabilities: A call to action . <i>Interactions</i> , 28(3):47–51.	2723	Julian McAuley, Christopher Targett, Qinfeng Shi, and Anton van den Hengel. 2015. Image-based recommendations on styles and substitutes . <i>arXiv preprint</i> .	2724
2667		2725		
2668		2726	Bryan McCann, Nitish Shirish Keskar, Caiming Xiong, and Richard Socher. 2018. The natural language decathlon: Multitask learning as question answering . <i>arXiv preprint</i> .	2727
2669		2728		
2670		2729		
2671	Rebecca Marvin and Tal Linzen. 2018. Targeted syntactic evaluation of language models . In <i>Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing</i> , pages 1192–1202, Brussels, Belgium. Association for Computational Linguistics.	2730	James L. McClelland, Felix Hill, Maja Rudolph, Jason Baldridge, and Hinrich Schütze. 2019. Extending machine language models toward human-level language understanding . <i>arXiv preprint</i> .	2731
2672		2732		
2673		2733		
2674				
2675				
2676				
2677	Javier Marín, Aritro Biswas, Ferda Ofli, Nicholas Hynes, Amaia Salvador, Yusuf Aytar, Ingmar Weber, and Antonio Torralba. 2021. Recipe1M+: A dataset for learning cross-modal embeddings for cooking recipes and food images . <i>IEEE Transactions on Pattern Analysis and Machine Intelligence</i> , 43(1):187–203.	2734	R Thomas McCoy, Robert Frank, and Tal Linzen. 2018. Revisiting the poverty of the stimulus: Hierarchical generalization without a hierarchical bias in recurrent neural networks. <i>arXiv preprint arXiv:1802.09091</i> .	2735
2678		2736		
2679		2737		
2680				
2681				
2682				

- R. Thomas McCoy, Robert Frank, and Tal Linzen. 2020. [Does Syntax Need to Grow on Trees? Sources of Hierarchical Inductive Bias in Sequence-to-Sequence Networks](#). *Transactions of the Association for Computational Linguistics*, 8:125–140.
- R. Thomas McCoy, Ellie Pavlick, and Tal Linzen. 2019. [Right for the wrong reasons: Diagnosing syntactic heuristics in natural language inference](#). *arXiv preprint*.
- Brendan McMahan, Eider Moore, Daniel Ramage, Seth Hampson, and Blaise Agüera y Arcas. 2017. Communication-efficient learning of deep networks from decentralized data. In *Artificial intelligence and statistics*.
- Shikib Mehri and Maxine Eskenazi. 2020. [USR: An unsupervised and reference free evaluation metric for dialog generation](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 681–707, Online. Association for Computational Linguistics.
- Christine Palm Meister. 2007. [Phraseologie des schwedischen](#). In H. Burger et al., editor, *Phraseologie/Phrasology*, volume 2, pages 673–681. De Gruyter Mouton.
- Francisco S. Melo, Carla Guerra, and Manuel Lopes. 2018. [Interactive optimal teaching with unknown learners](#). In *Proceedings of the Twenty-Seventh International Joint Conference on Artificial Intelligence, IJCAI-18*, pages 2567–2573.
- Julia Mendelsohn, Yulia Tsvetkov, and Dan Jurafsky. 2020. [A framework for the computational linguistic analysis of dehumanization](#). *Frontiers in Artificial Intelligence*, 3.
- Yuanliang Meng, Anna Rumshisky, and Alexey Romanov. 2017. [Temporal information extraction for question answering using syntactic dependencies in an LSTM-based architecture](#). In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, pages 887–896, Copenhagen, Denmark. Association for Computational Linguistics.
- Stephen Merity, Caiming Xiong, James Bradbury, and Richard Socher. 2016. [Pointer sentinel mixture models](#). *arXiv preprint*.
- William Merrill. 2020. [On the linguistic capacity of real-time counter automata](#). *arXiv preprint*.
- Elliot Meyerson and Risto Miikkulainen. 2017. [Beyond shared hierarchies: Deep multitask learning through soft layer ordering](#). *ArXiv*, abs/1711.00108.
- Wolfgang Mieder. 2019. ["Andere zeiten, andere lehren": Sprach-und kulturgeschichtliche betrachtungen zum sprichwort](#). In K. Steyer, editor, *Wortverbindungen - mehr oder weniger fest*, pages 415–438. De Gruyter, Berlin.
- Rada Mihalcea and Carlo Strapparava. 2005. [Making computers laugh: Investigations in automatic humor recognition](#). In *Proceedings of Human Language Technology Conference and Conference on Empirical Methods in Natural Language Processing*, pages 531–538, Vancouver, British Columbia, Canada. Association for Computational Linguistics.
- John Miller, Karl Krauth, Benjamin Recht, and Ludwig Schmidt. 2020. [The effect of natural distribution shift on question answering models](#). In *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pages 6905–6916. PMLR.
- Tristan Miller and Iryna Gurevych. 2015. [Automatic disambiguation of English puns](#). In *Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 719–729, Beijing, China. Association for Computational Linguistics.
- Tristan Miller, Christian Hempelmann, and Iryna Gurevych. 2017. [SemEval-2017 task 7: Detection and interpretation of English puns](#). In *Proceedings of the 11th International Workshop on Semantic Evaluation (SemEval-2017)*, pages 58–68, Vancouver, Canada. Association for Computational Linguistics.
- David Milne and Ian H. Witten. 2008. [An effective, low-cost measure of semantic relatedness obtained from Wikipedia links](#). In *Proceeding of AAAI Workshop on Wikipedia and Artificial Intelligence: An Evolving Synergy*, page 25–30, Menlo Park. Association for the Advancement of Artificial Intelligence.
- Republic of China Ministry of the Interior. 2018. National name statistical analysis. <https://www.ris.gov.tw/documents/data/5/2/107namestat.pdf> (Accessed 3 March 2021).
- Swaroop Mishra, Daniel Khashabi, Chitta Baral, and Hannaneh Hajishirzi. 2021. [Cross-task generalization via natural language crowdsourcing instructions](#). *arXiv preprint arXiv:2104.08773*, arXiv:2104.08773.
- Ishan Misra, Abhinav Shrivastava, Abhinav Kumar Gupta, and Martial Hebert. 2016. [Cross-stitch networks for multi-task learning](#). *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 3994–4003.
- Margaret Mitchell, Kees van Deemter, and Ehud Reiter. 2010. [Natural reference to objects in a visual domain](#). In *Proceedings of the 6th International Natural Language Generation Conference*. Association for Computational Linguistics.
- Margaret Mitchell, Kees van Deemter, and Ehud Reiter. 2013. [Generating expressions that refer to visible objects](#). In *Proceedings of the 2013 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 1174–1184, Atlanta, Georgia. Association for Computational Linguistics.

2849	Volodymyr Mnih, Koray Kavukcuoglu, David Silver,	Mohammed Muqeeth, Haokun Liu, and Colin Raffel.	2906
2850	Alex Graves, Ioannis Antonoglou, Daan Wierstra,	2023. Soft merging of experts with adaptive routing.	2907
2851	and Martin Riedmiller. 2013. Playing Atari with	<i>arXiv preprint arXiv:2306.03745</i> .	2908
2852	deep reinforcement learning . <i>arXiv preprint</i> .		
2853	Elham Mohammadi, Hessam Amini, and Leila Kosseim.	Yoichi Murakami and Kenji Mizuguchi. 2010. Apply-	2909
2854	2019. CLaC at CLPsych 2019: Fusion of neural fea-	ing the naïve Bayes classifier with kernel density	2910
2855	tures and predicted class probabilities for suicide risk	estimation to the prediction of protein–protein inter-	2911
2856	assessment based on online posts. In <i>Proceedings</i>	action sites . <i>Bioinformatics</i> , 26(15):1841–1848.	2912
2857	<i>of the Sixth Workshop on Computational Linguistics</i>		
2858	<i>and Clinical Psychology</i> , pages 34–38, Minneapolis,	Gregory L. Murphy. 1988. Comprehending complex	2913
2859	Minnesota. Association for Computational Linguis-	concepts . <i>Cognitive Science</i> , 12(4):529–562.	2914
2860	tics.		
2861	Alireza Mohammadshahi, Ali Shaikh, and Majid	Moin Nadeem, Anna Bethke, and Siva Reddy. 2020.	2915
2862	Yazdani. 2024. Routoo: Learning to route	StereoSet: Measuring stereotypical bias in pretrained	2916
2863	to large language models effectively . <i>Preprint</i> ,	language models . <i>arXiv preprint</i> .	2917
2864	<i>arXiv:2401.13979</i> .		
2865	Michael Mohler, Mary Brunson, Bryan Rink, and Marc	Aakanksha Naik, Abhilasha Ravichander, Norman	2918
2866	Tomlinson. 2016. Introducing the LCC metaphor	Sadeh, Carolyn Rose, and Graham Neubig. 2018.	2919
2867	datasets . In <i>Proceedings of the Tenth International</i>	Stress test evaluation for natural language inference .	2920
2868	<i>Conference on Language Resources and Evaluation</i>	<i>arXiv preprint</i> .	2921
2869	<i>(LREC’16)</i> , pages 4221–4227, Portorož, Slovenia.		
2870	European Language Resources Association.	Ramanujapuram Narasimhachar. 1988. <i>History of Kan-</i>	2922
2871	Jane Morris and Graeme Hirst. 1991. Lexical cohe-	<i>nada Literature: Readership Lectures</i> . Asian Educa-	2923
2872	sion computed by thesaural relations as an indicator	tional Services, New Dehli.	2924
2873	of the structure of text . <i>Computational Linguistics</i> ,		
2874	17(1):21–48.	Shashi Narayan, Shay B. Cohen, and Mirella Lapata.	2925
2875	Nasrin Mostafazadeh, Nathanael Chambers, Xiaodong	2018. Don’t give me the details, just the summary!	2926
2876	He, Devi Parikh, Dhruv Batra, Lucy Vanderwende,	Topic-aware convolutional neural networks for ex-	2927
2877	Pushmeet Kohli, and James Allen. 2016a. A corpus	treme summarization . In <i>Proceedings of the 2018</i>	2928
2878	and evaluation framework for deeper understanding	<i>Conference on Empirical Methods in Natural Lan-</i>	2929
2879	of commonsense stories . <i>arXiv preprint</i> .	<i>guage Processing</i> , pages 1797–1807, Brussels, Bel-	2930
2880	Nasrin Mostafazadeh, Michael Roth, Annie Louis,	gium. Association for Computational Linguistics.	2931
2881	Nathanael Chambers, and James Allen. 2016b. A	Ziad S. Nasreddine, Natalie A. Phillips, Valérie	2932
2882	corpus and cloze evaluation for deeper understanding	Bédirian, Simon Charbonneau, Victor Whitehead,	2933
2883	of commonsense stories . In <i>Proceedings of the 2016</i>	Isabelle Collin, Jeffrey L. Cummings, and Howard	2934
2884	<i>Conference of the North American Chapter of the</i>	Chertkow. 2005. The Montreal Cognitive Assess-	2935
2885	<i>Association for Computational Linguistics: Human</i>	ment, MoCA: A brief screening tool for mild cogni-	2936
2886	<i>Language Technologies</i> , pages 839–849, San Diego,	tive impairment . <i>Journal of the American Geriatrics</i>	2937
2887	California. Association for Computational Linguis-	<i>Society</i> , 53(4):695–699.	2938
2888	tics.		
2889	Lili Mou, Zhao Meng, Rui Yan, Ge Li, Yan Xu,	Aida Nematzadeh, Kaylee Burns, Erin Grant, Alison	2939
2890	Lu Zhang, and Zhi Jin. 2016. How transferable are	Gopnik, and Tom Griffiths. 2018. Evaluating theory	2940
2891	neural networks in nlp applications? In <i>Conference</i>	of mind in question answering . In <i>Proceedings of the</i>	2941
2892	<i>on Empirical Methods in Natural Language Process-</i>	<i>2018 Conference on Empirical Methods in Natural</i>	2942
2893	<i>ing</i> .	<i>Language Processing</i> , pages 2392–2400, Brussels,	2943
2894	Karl Mulligan, Robert Frank, and Tal Linzen. 2021.	Belgium. Association for Computational Linguistics.	2944
2895	Structure here, bias there: Hierarchical generaliza-	Nam Nguyen and Yunsong Guo. 2007. Comparisons	2945
2896	tion by jointly learning syntactic transformations . In	of sequence labeling algorithms and extensions . In	2946
2897	<i>Proceedings of the Society for Computation in Lin-</i>	<i>Proceedings of the 24th International Conference on</i>	2947
2898	<i>guistics 2021</i> , pages 125–135, Online. Association	<i>Machine Learning</i> , page 681–688, New York, NY,	2948
2899	for Computational Linguistics.	USA. Association for Computing Machinery.	2949
2900	Mohammed Muqeeth, Haokun Liu, Yufan Liu, and	Allen Nie, Erin Bennett, and Noah Goodman. 2019.	2950
2901	Colin Raffel. 2024. Learning to route among special-	DisSent: Learning sentence representations from ex-	2951
2902	ized experts for zero-shot generalization . In <i>Proceed-</i>	plicit discourse relations . In <i>Proceedings of the 57th</i>	2952
2903	<i>ings of the 41st International Conference on Machine</i>	<i>Annual Meeting of the Association for Computational</i>	2953
2904	<i>Learning</i> , volume 235 of <i>Proceedings of Machine</i>	<i>Linguistics</i> , pages 4497–4510, Florence, Italy. Asso-	2954
2905	<i>Learning Research</i> , pages 36829–36846. PMLR.	ciation for Computational Linguistics.	2955
		Yixin Nie, Adina Williams, Emily Dinan, Mohit Bansal,	2956
		Jason Weston, and Douwe Kiela. 2020. Adversarial	2957
		NLI: A new benchmark for natural language under-	2958
		standing . In <i>Proceedings of the 58th Annual Meet-</i>	2959
		<i>ing of the Association for Computational Linguistics</i> ,	2960

2961	pages 4885–4901, Online. Association for Computational Linguistics.	
2962		
2963	Marilyn Nippold, Melissa Allen, and Dixon Kirsch.	
2964	2001. Proverb comprehension as a function of reading proficiency in preadolescents . <i>Language Speech and Hearing Services in Schools</i> , 32:90.	
2965		
2966		
2967	Masaaki Nishino, Sho Takase, Tsutomu Hirao, and Masaaki Nagata. 2019. Generating natural anagrams: Towards language generation under hard combinatorial constraints . In <i>Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)</i> , pages 6408–6412, Hong Kong, China. Association for Computational Linguistics.	
2968		
2969		
2970		
2971		
2972		
2973		
2974		
2975		
2976	David Noever, Matt Ciolino, and Josh Kalin. 2020. The chess transformer: Mastering play using generative language models .	
2977		
2978		
2979	David Nolan and Akina Mikami. 2013. "The things that we have to do": Ethics and instrumentality in humanitarian communication . <i>Global Media and Communication</i> , 9(1):53–70.	
2980		
2981		
2982		
2983	Klaus Oberauer, Robin Hörnig, Andrea Weidenfeld, and Oliver Wilhelm. 2005. Effects of directionality in deductive reasoning, II. Premise integration and conclusion evaluation . <i>The Quarterly Journal of Experimental Psychology Section A</i> , 58(7):1225–1247.	
2984		
2985		
2986		
2987		
2988	Klaus Oberauer and Oliver Wilhelm. 2000. Effects of directionality in deductive reasoning, I. The comprehension of single relational premises . <i>Journal of Experimental Psychology: Learning, Memory, and Cognition</i> , 26(6):1702–1712.	
2989		
2990		
2991		
2992		
2993	The Working Committee on the Revision of the National Standard Occupational Classification. 2015. <i>Standard Occupational Classification of the People's Republic of China</i> . China Labour and Social Security Publishing House. http://www.jiangmen.gov.cn/bmpd/jmsr1zyhshbzj/zwfw/bmjd/jdks/content/post_2334804.html (Accessed 4 June 2022).	
2994		
2995		
2996		
2997		
2998		
2999		
3000		
3001	Isaac Ong, Amjad Almahairi, Vincent Wu, Wei-Lin Chiang, Tianhao Wu, Joseph E. Gonzalez, M Waleed Kadous, and Ion Stoica. 2024. Routellm: Learning to route llms with preference data . <i>Preprint</i> , arXiv:2406.18665.	
3002		
3003		
3004		
3005		
3006	Silviu Oprea and Walid Magdy. 2020. iSarcasm: A dataset of intended sarcasm . In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 1279–1289, Online. Association for Computational Linguistics.	
3007		
3008		
3009		
3010		
3011	Peter-Michael Osera and Steve Zdancewic. 2015. Type-and-example-directed program synthesis . In <i>Proceedings of the 36th ACM SIGPLAN Conference on Programming Language Design and Implementation, PLDI '15</i> , page 619–630, New York, NY, USA. Association for Computing Machinery.	
3012		
3013		
3014		
3015		
3016		
	Oleksiy Ostapenko, Zhan Su, Edoardo Maria Ponti, Laurent Charlin, Nicolas Le Roux, Matheus Pereira, Lucas Caccia, and Alessandro Sordoni. 2024. Towards modular llms by building and reusing a library of loras. <i>arXiv preprint arXiv:2405.11157</i> .	3017
		3018
		3019
		3020
		3021
	Jahna C. Otterbacher, Dragomir R. Radev, and Airong Luo. 2002. Revisions that improve cohesion in multi-document summaries: A preliminary study . In <i>Proceedings of the ACL-02 Workshop on Automatic Summarization</i> , pages 27–44, Philadelphia, Pennsylvania, USA. Association for Computational Linguistics.	3022
		3023
		3024
		3025
		3026
		3027
	Artidoro Pagnoni, Vidhisha Balachandran, and Yulia Tsvetkov. 2021. Understanding factuality in abstractive summarization with FRANK: A benchmark for factuality metrics . In <i>Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies</i> , pages 4812–4829, Online. Association for Computational Linguistics.	3028
		3029
		3030
		3031
		3032
		3033
		3034
		3035
	Kartikey Pant and Tanvi Dadu. 2020. Sarcasm detection using context separators in online discourse . <i>arXiv preprint</i> .	3036
		3037
		3038
	Tae Jin Park, Naoyuki Kanda, Dimitrios Dimitriadis, Kyu J. Han, Shinji Watanabe, and Shrikanth Narayanan. 2021. A review of speaker diarization: Recent advances with deep learning . <i>arXiv preprint</i> .	3039
		3040
		3041
		3042
	Arkil Patel, Satwik Bhattamishra, and Navin Goyal. 2021. Are NLP models really able to solve simple math word problems? In <i>Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies</i> , pages 2080–2094, Online. Association for Computational Linguistics.	3043
		3044
		3045
		3046
		3047
		3048
		3049
	Anthony M. Paul. 1970. Figurative language . <i>Philosophy & Rhetoric</i> , 3(4):225–248.	3050
		3051
	Ali Payani and Faramarz Fekri. 2019. Learning algorithms via neural logic networks . <i>arXiv preprint</i> .	3052
		3053
	Judea Pearl. 1988. <i>Probabilistic Reasoning in Intelligent Systems: Networks of Plausible Inference</i> . Morgan Kaufmann, San Francisco.	3054
		3055
		3056
	Judea Pearl. 2000. <i>Causality: Models, Reasoning, and Inference</i> . Cambridge University Press, Cambridge.	3057
		3058
	Devin Pelser and Hugh Murrell. 2019. Deep and dense sarcasm detection . <i>arXiv preprint</i> .	3059
		3060
	Ethan Perez, Douwe Kiela, and Kyunghyun Cho. 2021. True few-shot learning with language models . <i>arXiv preprint</i> .	3061
		3062
		3063
	Carla Perez Almendros, Luis Espinosa Anke, and Steven Schockaert. 2020. Don't patronize me! An annotated dataset with patronizing and condescending language towards vulnerable communities . In <i>Proceedings of the 28th International Conference on Computational Linguistics</i> , pages 5891–5902, Barcelona, Spain (Online). International Committee on Computational Linguistics.	3064
		3065
		3066
		3067
		3068
		3069
		3070
		3071

3072	Fabio Petroni, Tim Rocktäschel, Patrick Lewis, An-	Simone Paolo Ponzetto and Michael Strube. 2007.	3127
3073	ton Bakhtin, Yuxiang Wu, Alexander H. Miller, and	Knowledge derived from Wikipedia for computing	3128
3074	Sebastian Riedel. 2019. Language models as knowl-	semantic relatedness . <i>Journal of Artificial Intelli-</i>	3129
3075	edge bases? <i>arXiv preprint</i> .	<i>gence Research</i> , 30:181–212.	3130
3076	Dessislava Petrova-Antonova and Rumyana Tancheva.	Octavian Popescu and Carlo Strapparava. 2015. Se-	3131
3077	2020. Data cleaning: A case study with OpenRefine	mEval 2015, task 7: Diachronic text evaluation . In	3132
3078	and Trifecta Wrangler . In <i>Quality of Information and</i>	<i>Proceedings of the 9th International Workshop on Se-</i>	3133
3079	<i>Communications Technology</i> , pages 32–40, Cham.	<i>mantic Evaluation (SemEval 2015)</i> , pages 870–878,	3134
3080	Springer.	Denver, Colorado. Association for Computational	3135
3081	Jonas Pfeiffer, Aishwarya Kamath, Andreas Rücklé,	Linguistics.	3136
3082	Kyunghyun Cho, and Iryna Gurevych. 2021.	Soujanya Poria, Devamanyu Hazarika, Navonil Ma-	3137
3083	AdapterFusion: Non-destructive task composition	jumder, Gautam Naik, Erik Cambria, and Rada Mi-	3138
3084	for transfer learning . In <i>Proceedings of the 16th Con-</i>	halcea. 2019. MELD: A multimodal multi-party	3139
3085	<i>ference of the European Chapter of the Association</i>	dataset for emotion recognition in conversations . In	3140
3086	<i>for Computational Linguistics</i> , pages 487–503.	<i>Proceedings of the 57th Annual Meeting of the As-</i>	3141
3087	Thang M. Pham, Trung Bui, Long Mai, and Anh	<i>sociation for Computational Linguistics</i> , pages 527–	3142
3088	Nguyen. 2020. Out of order: How important is the	536, Florence, Italy. Association for Computational	3143
3089	sequential order of words in a sentence in natural	Linguistics.	3144
3090	language understanding tasks? <i>arXiv preprint</i> .	Rolandos Alexandros Potamias, Georgios Siolas, and	3145
3091	Steve Piantadosi. 2020. Fleet system .	Andreas-Georgios Stafylopatis. 2020. A transformer-	3146
3092	Mohammad Taher Pilehvar and Jose Camacho-Collados.	based approach to irony and sarcasm detection . <i>Neu-</i>	3147
3093	2019. WiC: The word-in-context dataset for evalu-	<i>ral Computing and Applications</i> , 32:17309–17320.	3148
3094	ating context-sensitive meaning representations . In	Clifton Poth, Hannah Sterz, Indraneil Paul, Sukannya	3149
3095	<i>Proceedings of the 2019 Conference of the North</i>	Purkayastha, Leon Engländer, Timo Imhof, Ivan	3150
3096	<i>American Chapter of the Association for Computa-</i>	Vulić, Sebastian Ruder, Iryna Gurevych, and Jonas	3151
3097	<i>tional Linguistics: Human Language Technologies,</i>	Pfeiffer. 2023. Adapters: A unified library for	3152
3098	<i>Volume 1 (Long and Short Papers)</i> , pages 1267–1273,	parameter-efficient and modular transfer learning .	3153
3099	Minneapolis, Minnesota. Association for Computa-	<i>arXiv preprint arXiv:2311.11077</i> .	3154
3100	tional Linguistics.	Norman M. Prentice and Robert E. Fathman. 1975. Jok-	3155
3101	Tony A. Plate. 1994. <i>Distributed representations and</i>	ing riddles: A developmental index of children’s hu-	3156
3102	<i>nested compositional structure</i> . Ph.D. thesis, Univer-	mor . <i>Developmental Psychology</i> , 11:210–216.	3157
3103	sity of Toronto, Toronto.	Rémi Le Priol, Reza Babanezhad Harikandeh, Yoshua	3158
3104	Tony A. Plate. 2003. <i>Holographic Reduced Represent-</i>	Bengio, and Simon Lacoste-Julien. 2020. An analy-	3159
3105	<i>tations: Distributed Representation for Cognitive</i>	sis of the adaptation speed of causal models . <i>arXiv</i>	3160
3106	<i>Structures</i> . CSLI, Stanford, CA.	<i>preprint</i> .	3161
3107	Robert Plutchik. 1980. A general psychoevolutionary	James Pustejovsky, Robert Ingria, Roser Saurí, José	3162
3108	theory of emotion . In Robert Plutchik and Henry	Castaño, Jessica Littman, Rob Gaizauskas, Andrea	3163
3109	Kellerman, editors, <i>Theories of Emotion</i> , pages 3–33.	Setzer, Graham Katz, and Inderjeet Mani. 2004. The	3164
3110	Academic Press.	specification language TimeML .	3165
3111	Nadia Polikarpova, Ivan Kuraj, and Armando Solar-	Qimingtong. 2016. What are the most popular names	3166
3112	Lezama. 2016. Program synthesis from polymor-	chinese parents give their babies? a perspective	3167
3113	phic refinement types . In <i>Proceedings of the 37th</i>	from big data. https://www.qimingtong.com/	3168
3114	<i>ACM SIGPLAN Conference on Programming Lan-</i>	article/0 (Accessed 3 March 2021).	3169
3115	<i>guage Design and Implementation, PLDI ’16</i> , page	Lianhui Qin, Aditya Gupta, Shyam Upadhyay, Luheng	3170
3116	522–538, New York, NY, USA. Association for Com-	He, Yejin Choi, and Manaal Faruqui. 2021. TIME-	3171
3117	puting Machinery.	DIAL: Temporal commonsense reasoning in dialog .	3172
3118	Stanislas Polu and Ilya Sutskever. 2020. Generative	In <i>Proceedings of the 59th Annual Meeting of the</i>	3173
3119	language modeling for automated theorem proving .	<i>Association for Computational Linguistics and the</i>	3174
3120	<i>arXiv preprint</i> .	<i>11th International Joint Conference on Natural Lan-</i>	3175
3121	Edoardo Maria Ponti, Alessandro Sordani, Yoshua Ben-	<i>guage Processing (Volume 1: Long Papers)</i> , pages	3176
3122	gio, and Siva Reddy. 2023. Combining parameter-	7066–7076, Online. Association for Computational	3177
3123	efficient modules for task-level generalisation. In	Linguistics.	3178
3124	<i>Proceedings of the 17th Conference of the European</i>	Willard V.O. Quine. 1951. Main trends in recent philos-	3179
3125	<i>Chapter of the Association for Computational Lin-</i>	ophy: Two dogmas of empiricism . <i>The Philosophical</i>	3180
3126	<i>guistics</i> , pages 687–702.	<i>Review</i> , 60(1):20–43.	3181

3182	Quora, Inc. 2017. Quora question pairs .	
3183	Dragomir R. Radev, Lori Levin, and Thomas E. Payne.	
3184	2008. The North American computational linguistics olympiad (NACLO) . In <i>Proceedings of the Third Workshop on Issues in Teaching Computational Linguistics</i> , pages 87–96, Columbus, Ohio. Association for Computational Linguistics.	
3185		
3186		
3187		
3188		
3189	Alec Radford, Jeff Wu, Rewon Child, David Luan,	
3190	Dario Amodei, and Ilya Sutskever. 2019. Language models are unsupervised multitask learners . <i>OpenAI blog</i> , 1(8):9.	
3191		
3192		
3193	Kira Radinsky, Eugene Agichtein, Evgeniy Gabrilovich,	
3194	and Shaul Markovitch. 2011. A word at a time: Computing word relatedness using temporal semantic analysis . In <i>WWW '11: Proceedings of the 20th International Conference on World Wide Web</i> , page 337–346, New York, NY, USA. Association for Computing Machinery.	
3195		
3196		
3197		
3198		
3199		
3200	Colin Raffel, Noam Shazeer, Adam Roberts, Katherine	
3201	Lee, Sharan Narang, Michael Matena, Yanqi Zhou,	
3202	Wei Li, and Peter J Liu. 2020. Exploring the limits of transfer learning with a unified text-to-text transformer . <i>Journal of Machine Learning Research</i> , 21:1–67.	
3203		
3204		
3205		
3206	Alessandro Raganato, Jose Camacho-Collados, and	
3207	Roberto Navigli. 2017. Word sense disambiguation: A unified evaluation framework and empirical comparison . In <i>Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics: Volume 1, Long Papers</i> , pages 99–110, Valencia, Spain. Association for Computational Linguistics.	
3208		
3209		
3210		
3211		
3212		
3213		
3214	Altaf Rahman and Vincent Ng. 2012. Resolving complex cases of definite pronouns: The Winograd schema challenge . In <i>Proceedings of the 2012 Joint Conference on Empirical Methods in Natural Language Processing and Computational Natural Language Learning</i> , pages 777–789, Jeju Island, Korea. Association for Computational Linguistics.	
3215		
3216		
3217		
3218		
3219		
3220		
3221	Sunny Rai and Shampa Chakraverty. 2020. A survey on computational metaphor processing . <i>ACM Comput. Surv.</i> , 53(2).	
3222		
3223		
3224	Pranav Rajpurkar, Robin Jia, and Percy Liang. 2018.	
3225	Know what you don’t know: Unanswerable questions for SQuAD . In <i>Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)</i> , pages 784–789, Melbourne, Australia. Association for Computational Linguistics.	
3226		
3227		
3228		
3229		
3230		
3231	Pranav Rajpurkar, Jian Zhang, Konstantin Lopyrev, and	
3232	Percy Liang. 2016. SQuAD: 100,000+ questions for machine comprehension of text . In <i>Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing</i> , pages 2383–2392, Austin, Texas. Association for Computational Linguistics.	
3233		
3234		
3235		
3236		
	Prajit Ramachandran and Quoc V. Le. 2019. Diversity and depth in per-example routing models . In <i>International Conference on Learning Representations</i> .	3237 3238 3239
	Alexandre Ramé, Kartik Ahuja, Jianyu Zhang, Matthieu	3240
	Cord, Léon Bottou, and David Lopez-Paz. 2022. Recycling diverse models for out-of-distribution generalization. <i>arXiv preprint arXiv:2212.10445</i> .	3241 3242 3243
	Abhinav Rastogi, Xiaoxue Zang, Srinivas Sunkara,	3244
	Raghav Gupta, and Pranav Khaitan. 2020. Towards scalable multi-domain conversational agents: The schema-guided dialogue dataset . In <i>Proceedings of the AAAI Conference on Artificial Intelligence</i> , volume 34, pages 8689–8696, Menlo Park, CA. Association for the Advancement of Artificial Intelligence.	3245 3246 3247 3248 3249 3250
	Ian Ravenscroft. 2019. Folk psychology as a theory. In	3251
	Edward N. Zalta, editor, <i>The Stanford Encyclopedia of Philosophy</i> , Summer 2019 edition. Metaphysics Research Lab, Stanford University.	3252 3253 3254
	Siva Reddy, Danqi Chen, and Christopher D. Manning.	3255
	2019. CoQA: A conversational question answering challenge . <i>Transactions of the Association for Computational Linguistics</i> , 7:249–266.	3256 3257 3258
	Scott Reed and Nando de Freitas. 2015. Neural programmer-interpreters . <i>arXiv preprint</i> .	3259 3260
	Marek Rei. 2017. Semi-supervised multitask learning for sequence labeling . <i>arXiv preprint</i> .	3261 3262
	Nils Reimers and Iryna Gurevych. 2019. Sentence-BERT: Sentence embeddings using siamese BERT-networks . <i>arXiv preprint</i> .	3263 3264 3265
	Joseph Reisinger and Raymond J. Mooney. 2010. Multi-prototype vector-space models of word meaning . In <i>Human Language Technologies: The 2010 Annual Conference of the North American Chapter of the Association for Computational Linguistics</i> , pages 109–117, Los Angeles, California. Association for Computational Linguistics.	3266 3267 3268 3269 3270 3271 3272
	He Ren and Quan Yang. 2017. Neural joke generation .	3273
	Philip Resnik. 1995. Using information content to evaluate semantic similarity in a taxonomy . In <i>IJCAI’95: Proceedings of the 14th International Joint Conference on Artificial Intelligence, Volume 1</i> , page 448–453, San Francisco. Morgan Kaufmann.	3274 3275 3276 3277 3278
	Philip Resnik. 1999. Semantic similarity in a taxonomy: An information-based measure and its application to problems of ambiguity in natural language . <i>Journal of Artificial Intelligence Research</i> , 11:95–130.	3279 3280 3281 3282
	Melissa Roemmele, Cosmin Adrian Bejan, and Andrew S Gordon. 2011a. Choice of plausible alternatives: An evaluation of commonsense causal reasoning. In <i>2011 AAAI Spring Symposium Series</i> .	3283 3284 3285 3286
	Melissa Roemmele, Cosmin Adrian Bejan, and Andrew S. Gordon. 2011b. Choice of plausible alternatives: An evaluation of commonsense causal reasoning . <i>AAAI Spring Symposium</i> .	3287 3288 3289 3290

3291	Anna Rogers, Olga Kovaleva, and Anna Rumshisky.	Phillip Rust, Jonas Pfeiffer, Ivan Vulić, Sebastian Ruder,	3344
3292	2020. Getting closer to AI complete question answer-	and Iryna Gurevych. 2020. How good is your tok-	3345
3293	ing: A set of prerequisite real tasks . In <i>Proceedings</i>	enizer? On the monolingual performance of multilin-	3346
3294	<i>of the 2020 Conference on Empirical Methods in Nat-</i>	gual language models . <i>arXiv preprint</i> .	3347
3295	<i>ural Language Processing (EMNLP)</i> , pages 123–150,		
3296	Online. Association for Computational Linguistics.		
3297	Alexis Ross and Ellie Pavlick. 2019. How well do NLI	Amrita Saha, Rahul Aralikatte, Mitesh M. Khapra, and	3348
3298	models capture verb veridicality? In <i>Proceedings of</i>	Karthik Sankaranarayanan. 2018. DuoRC: Towards	3349
3299	<i>the 2019 Conference on Empirical Methods in Nat-</i>	Complex Language Understanding with Paraphrased	3350
3300	<i>ural Language Processing and the 9th International</i>	Reading Comprehension. In <i>Meeting of the Associa-</i>	3351
3301	<i>Joint Conference on Natural Language Processing</i>	<i>tion for Computational Linguistics (ACL)</i> .	3352
3302	<i>(EMNLP-IJCNLP)</i> , pages 2230–2240, Hong Kong,		
3303	China. Association for Computational Linguistics.	Gözde Gül Şahin, Yova Kementchedjhieva, Phillip Rust,	3353
3304	Kenneth J. Rothman and Sander Greenland. 2005. Cau-	and Iryna Gurevych. 2020. PuzzLing Machines: A	3354
3305	sation and causal inference in epidemiology . <i>Ameri-</i>	challenge on learning from small data . In <i>Proceed-</i>	3355
3306	<i>can Journal of Public Health</i> , 95(S1):S144–S150.	<i>ings of the 58th Annual Meeting of the Association</i>	3356
3307	Joshua Rozner, Christopher Potts, and Kyle Mahowald.	<i>for Computational Linguistics</i> , pages 1241–1254, On-	3357
3308	2021. Decrypting cryptic crosswords: Semantically	line. Association for Computational Linguistics.	3358
3309	complex wordplay puzzles as a target for NLP . <i>arXiv</i>	Keisuke Sakaguchi, Kevin Duh, Matt Post, and Ben-	3359
3310	<i>preprint</i> .	jamin Van Durme. 2016. Robsut wrod reocginiton	3360
3311	Sebastian Ruder, Joachim Bingel, Isabelle Augenstein,	via semi-character recurrent neural network . <i>arXiv</i>	3361
3312	and Anders Søgaard. 2017. Latent multi-task archi-	<i>preprint</i> .	3362
3313	tecture learning . In <i>AAAI Conference on Artificial</i>	Keisuke Sakaguchi, Ronan Le Bras, Chandra Bhaga-	3363
3314	<i>Intelligence</i> .	vatula, and Yejin Choi. 2020a. WinoGrande: An	3364
3315	Rachel Rudinger, Jason Naradowsky, Brian Leonard,	adversarial Winograd schema challenge at scale . In	3365
3316	and Benjamin Van Durme. 2018. Gender bias in	<i>Proceedings of the Thirty-Fourth AAAI Conference</i>	3366
3317	coreference resolution . In <i>Proceedings of the 2018</i>	<i>on Artificial Intelligence</i> , pages 8732–8740, New	3367
3318	<i>Conference of the North American Chapter of the</i>	York, NY, USA. Association for the Advancement of	3368
3319	<i>Association for Computational Linguistics: Human</i>	Artificial Intelligence.	3369
3320	<i>Language Technologies, Volume 2 (Short Papers)</i> ,	Keisuke Sakaguchi, Ronan Le Bras, Chandra Bhagavat-	3370
3321	pages 8–14, New Orleans, Louisiana. Association for	ula, and Yejin Choi. 2020b. WINOGRANDE: An	3371
3322	Computational Linguistics.	adversarial Winograd schema challenge at scale . In	3372
3323	Rachel Etta Rudolph and Alexander W. Kocurek. 2020.	<i>The Thirty-Fourth AAAI Conference on Artificial In-</i>	3373
3324	Comparing conventions . In <i>Proceedings of Seman-</i>	<i>teelligence, AAAI-20</i> , pages 8732–8734, Menlo Park,	3374
3325	<i>tics and Linguistic Theory</i> , pages 294–313, Washing-	CA. Association for the Advancement of Artificial	3375
3326	ton, D.C. Linguistic Society of America.	Intelligence.	3376
3327	Rosa Rugani, Giorgio Vallortigara, Konstantinos Priftis,	María del Pilar Salas-Zárate, Mario Andrés Paredes-	3377
3328	and Lucia Regolin. 2015. Number-space mapping in	Valverde, Miguel Ángel Rodríguez-García, Rafael	3378
3329	the newborn chick resembles humans’ mental number	Valencia-García, and Giner Alor-Hernández. 2017.	3379
3330	line . <i>Science</i> , 347(6221):534–536.	Automatic detection of satire in Twitter: A	3380
3331	Joshua S. Rule. 2020. The child as hacker: Building	psycholinguistic-based approach . <i>Knowledge-Based</i>	3381
3332	more human-like models of learning . Ph.D. thesis,	<i>Systems</i> , 128:20–33.	3382
3333	Massachusetts Institute of Technology, Cambridge,	Julian Salazar, Davis Liang, Toan Q. Nguyen, and Ka-	3383
3334	MA.	trrin Kirchhoff. 2020. Masked language model scor-	3384
3335	Joshua S. Rule, Joshua B. Tenenbaum, and Steven T.	ing . In <i>Proceedings of the 58th Annual Meeting of</i>	3385
3336	Piantadosi. 2020. The child as hacker . <i>Trends in</i>	<i>the Association for Computational Linguistics</i> , pages	3386
3337	<i>Cognitive Sciences</i> , 24(11):900–915.	2699–2712, Online. Association for Computational	3387
3338	D. E. Rumelhart, J. L. McClelland, and PDP Research	Linguistics.	3388
3339	Group, editors. 1986. <i>Parallel Distributed Process-</i>	Suresh Kumar Sanampudi and G. Vijaya Kumari. 2010.	3389
3340	<i>ing. Volume 1: Foundations</i> . MIT Press, Cambridge,	Temporal reasoning in natural language processing:	3390
3341	MA.	A survey . <i>International Journal of Computer Appli-</i>	3391
3342	Stuart J. Russell and Peter Norvig. 2002. Artificial In-	<i>cations</i> , 1(4):53–57.	3392
3343	telligence: A Modern Approach . Pearson, Hoboken.	Evan Sandhaus. 2008. The New York Times annotated	3393
		corpus LDC2008T19 . <i>Linguistic Data Consortium</i> .	3394
		Victor Sanh, Albert Webson, Colin Raffel, Stephen H.	3395
		Bach, Lintang Sutawika, Zaid Alyafeai, Antoine	3396
		Chaffin, Arnaud Stiegler, Teven Le Scao, Arun Raja,	3397
		Manan Dey, M. Saiful Bari, Canwen Xu, Urmish	3398
		Thakker, Shanya Sharma, Eliza Szczechla, Taewoon	3399

3400	Kim, Gunjan Chhablani, Nihal V. Nayak, Debajyoti	Megan Scudellari. 2017. Cryopreservation aims to engineer novel ways to freeze, store, and thaw organs.	3458
3401	Datta, Jonathan Chang, Mike Tian-Jian Jiang, Han	Proceedings of the National Academy of Sciences ,	3459
3402	Wang, Matteo Manica, Sheng Shen, Zheng Xin Yong,	114(50):13060–13062.	3460
3403	Harshit Pandey, Rachel Bawden, Thomas Wang, Tr-		3461
3404	ishala Neeraj, Jos Rozen, Abheesht Sharma, Andrea		
3405	Santilli, Thibault Févry, Jason Alan Fries, Ryan Tee-	Abigail See, Peter J. Liu, and Christopher D. Manning.	3462
3406	han, Stella Biderman, Leo Gao, Tali Bers, Thomas	2017a. Get to the point: Summarization with pointer-	3463
3407	Wolf, and Alexander M. Rush. 2022. Multitask	generator networks . In <i>Proceedings of the 55th An-</i>	3464
3408	prompted training enables zero-shot task general-	<i>annual Meeting of the Association for Computational</i>	3465
3409	ization . In <i>The Tenth International Conference on</i>	<i>Linguistics</i> .	3466
3410	<i>Learning Representations</i> .		
3411	Victor Sanh, Albert Webson, Colin Raffel, Stephen H	Abigail See, Peter J. Liu, and Christopher D. Manning.	3467
3412	Bach, Lintang Sutawika, Zaid Alyafeai, Antoine	2017b. Get to the point: Summarization with pointer-	3468
3413	Chaffin, Arnaud Stiegler, Teven Le Scao, Arun	generator networks . In <i>Proceedings of the 55th An-</i>	3469
3414	Raja, et al. 2021. Multitask prompted training en-	<i>annual Meeting of the Association for Computational</i>	3470
3415	ables zero-shot task generalization. <i>arXiv preprint</i>	<i>Linguistics (Volume 1: Long Papers)</i> , pages 1073–	3471
3416	<i>arXiv:2110.08207</i> .	1083, Vancouver, Canada. Association for Computa-	3472
		tional Linguistics.	3473
3417	Adam Santoro, Andrew Lampinen, Kory Mathewson,	Thibault Sellam, Dipanjan Das, and Ankur P. Parikh.	3474
3418	Timothy Lillicrap, and David Raposo. 2021. Sym-	2020. BLEURT: Learning robust metrics for text	3475
3419	bolic behaviour in artificial intelligence . <i>arXiv</i>	generation . <i>arXiv preprint</i> .	3476
3420	<i>preprint</i> .		
3421	Adam Santoro, David Raposo, David G. T. Barrett, Ma-	Daniel Selsam, Matthew Lamm, Benedikt Bünz, Percy	3477
3422	teusz Malinowski, Razvan Pascanu, Peter Battaglia,	Liang, Leonardo de Moura, and David L. Dill. 2018.	3478
3423	and Timothy Lillicrap. 2017. A simple neural net-	Learning a SAT solver from single-bit supervision .	3479
3424	work module for relational reasoning . <i>arXiv preprint</i> .	<i>arXiv preprint</i> .	3480
3425	Maarten Sap, Saadia Gabriel, Lianhui Qin, Dan Juraf-	Lutfi Kerem Senel and Hinrich Schütze. 2021. Does he	3481
3426	sky, Noah A. Smith, and Yejin Choi. 2020. Social	wink or does he nod? A challenging benchmark for	3482
3427	bias frames: Reasoning about social and power im-	evaluating word understanding of language models .	3483
3428	plications of language . In <i>Proceedings of the 58th</i>	<i>arXiv preprint</i> .	3484
3429	<i>Annual Meeting of the Association for Computational</i>		
3430	<i>Linguistics</i> , pages 5477–5490, Online. Association	Luži Sennhauser and Robert C. Berwick. 2018. Eval-	3485
3431	for Computational Linguistics.	uating the ability of LSTMs to learn context-free	3486
		grammars . <i>arXiv preprint</i> .	3487
3432	Maarten Sap, Hannah Rashkin, Derek Chen, Ronan	Rico Sennrich. 2016. How grammatical is character-	3488
3433	Le Bras, and Yejin Choi. 2019. Social IQa: Com-	level neural machine translation? Assessing MT qual-	3489
3434	monsense reasoning about social interactions . In	ity with contrastive translation pairs . <i>arXiv preprint</i> .	3490
3435	<i>Proceedings of the 2019 Conference on Empirical</i>		
3436	<i>Methods in Natural Language Processing and the</i>	Rico Sennrich, Barry Haddow, and Alexandra Birch.	3491
3437	<i>9th International Joint Conference on Natural Lan-</i>	2015. Neural machine translation of rare words with	3492
3438	<i>guage Processing (EMNLP-IJCNLP)</i> , pages 4463–	subword units . <i>arXiv preprint</i> .	3493
3439	4473, Hong Kong, China. Association for Computa-		
3440	tional Linguistics.	Rico Sennrich and Biao Zhang. 2019. Revisiting low-	3494
3441	David Saxton, Edward Grefenstette, Felix Hill, and	resource neural machine translation: A case study .	3495
3442	Pushmeet Kohli. 2019. Analysing mathematical rea-	In <i>Proceedings of the 57th Annual Meeting of the As-</i>	3496
3443	soning abilities of neural models . <i>arXiv preprint</i> .	<i>sociation for Computational Linguistics</i> , pages 211–	3497
3444	Timo Schick, Sahana Udupa, and Hinrich Schütze. 2021.	221, Florence, Italy. Association for Computational	3498
3445	Self-diagnosis and self-debiasing: A proposal for	<i>Linguistics</i> .	3499
3446	reducing corpus-based bias in NLP . <i>arXiv preprint</i> .		
3447	Tal Schuster, Adam Fisch, and Regina Barzilay. 2021.	Min Joon Seo, Hannaneh Hajishirzi, Ali Farhadi, and	3500
3448	Get your vitamin C! Robust fact verification with	Oren Etzioni. 2014. Diagram understanding in ge-	3501
3449	contrastive evidence . In <i>Proceedings of the 2021</i>	ometry questions . In <i>Proceedings of the AAAI Con-</i>	3502
3450	<i>Conference of the North American Chapter of the</i>	<i>ference on Artificial Intelligence</i> , volume 28, Menlo	3503
3451	<i>Association for Computational Linguistics: Human</i>	Park, CA. Association for the Advancement of Artifi-	3504
3452	<i>Language Technologies</i> , pages 624–643, Online. As-	cial Intelligence.	3505
3453	sociation for Computational Linguistics.	Usman Shahid and Elena Zheleva. 2021. Counterfactual	3506
3454	Bernhard Schölkopf, Francesco Locatello, Stefan Bauer,	learning in networks: An empirical study of model	3507
3455	Nan Rosemary Ke, Nal Kalchbrenner, Anirudh	dependence .	3508
3456	Goyal, and Yoshua Bengio. 2021. Towards causal	Janelle Shane. 2020. All your questions answered. AI	3509
3457	representation learning . <i>arXiv preprint</i> .	Weirdness .	3510

3511	David Elliot Shaw, William R. Swartout, and C. Cordell	Ekaterina Shutova. 2010. Automatic metaphor inter-	3565
3512	Green. 1975. Inferring LISP programs from exam-	pretation as a paraphrasing task . In <i>Human Lan-</i>	3566
3513	ples . In <i>IJCAI'75: Proceedings of the 4th Inter-</i>	<i>guage Technologies: The 2010 Annual Conference of</i>	3567
3514	<i>national Joint Conference on Artificial Intelligence</i> ,	<i>the North American Chapter of the Association for</i>	3568
3515	volume 1, pages 260–267. Artificial Intelligence Lab-	<i>Computational Linguistics</i> , pages 1029–1037, Los	3569
3516	oratory, Cambridge, MA.	Angeles, California. Association for Computational	3570
		Linguistics.	3571
3517	Noam Shazeer, Azalia Mirhoseini, Krzysztof Maziarz,	Ekaterina Shutova and Simone Teufel. 2010. Metaphor	3572
3518	Andy Davis, Quoc Le, Geoffrey Hinton, and Jeff	corpus annotated for source-target domain mappings .	3573
3519	Dean. 2016. Outrageously large neural networks:	In <i>Proceedings of the Seventh International Con-</i>	3574
3520	The sparsely-gated mixture-of-experts layer. In <i>Inter-</i>	<i>ference on Language Resources and Evaluation</i>	3575
3521	<i>national Conference on Learning Representations</i> .	(<i>LREC'10</i>), Valletta, Malta. European Language Re-	3576
		sources Association.	3577
3522	Noam Shazeer, Azalia Mirhoseini, Krzysztof Maziarz,	Damien Sileo, Tim Van De Cruys, Camille Pradel, and	3578
3523	Andy Davis, Quoc Le, Geoffrey Hinton, and Jeff	Philippe Muller. 2019. Mining discourse markers	3579
3524	Dean. 2017. Outrageously large neural networks:	for unsupervised sentence representation learning . In	3580
3525	The sparsely-gated mixture-of-experts layer . In <i>Inter-</i>	<i>Proceedings of the 2019 Conference of the North</i>	3581
3526	<i>national Conference on Learning Representations</i> .	<i>American Chapter of the Association for Computa-</i>	3582
		<i>tional Linguistics: Human Language Technologies,</i>	3583
3527	Shannon Zejiang Shen, Hunter Lang, Bailin Wang,	<i>Volume 1 (Long and Short Papers)</i> , pages 3477–3486,	3584
3528	Yoon Kim, and David Sontag. 2024. Learning to	Minneapolis, Minnesota. Association for Computa-	3585
3529	decode collaboratively with multiple language mod-	tional Linguistics.	3586
3530	els. <i>arXiv preprint arXiv:2403.03870</i> .		
3531	Emily Sheng, Kai-Wei Chang, Premkumar Natarajan,	Damien Sileo, Tim Van de Cruys, Camille Pradel, and	3587
3532	and Nanyun Peng. 2019. The woman worked as	Philippe Muller. 2020. DiscSense: Automated se-	3588
3533	a babysitter: On biases in language generation . In	mantic analysis of discourse markers . In <i>Proceedings</i>	3589
3534	<i>Proceedings of the 2019 Conference on Empirical</i>	<i>of the 12th Language Resources and Evaluation Con-</i>	3590
3535	<i>Methods in Natural Language Processing and the</i>	<i>ference</i> , pages 991–999, Marseille, France. European	3591
3536	<i>9th International Joint Conference on Natural Lan-</i>	Language Resources Association.	3592
3537	<i>guage Processing (EMNLP-IJCNLP)</i> , pages 3407–		
3538	3412, Hong Kong, China. Association for Computa-	Damien Sileo, Wout Vossen, and Robbe Raymaekers.	3593
3539	tional Linguistics.	2022. Zero-shot recommendation as language mod-	3594
		eling . In <i>Advances in Information Retrieval: 44th</i>	3595
3540	Shaoyun Shi, Hanxiong Chen, Weizhi Ma, Jiaxin Mao,	<i>European Conference on IR Research, ECIR 2022,</i>	3596
3541	Min Zhang, and Yongfeng Zhang. 2020. Neural logic	<i>Stavanger, Norway, April 10–14, 2022, Proceedings,</i>	3597
3542	reasoning . <i>arXiv preprint</i> .	<i>Part II</i> , page 223–230, Cham. Springer.	3598
3543	Han-Chin Shing, Suraj Nair, Ayah Zirikly, Meir Frieden-	Gurdeep Singh, Kaustubh Dhole, Priyadarshini P. Pai,	3599
3544	berg, Hal Daumé III, and Philip Resnik. 2018. Expert,	and Sukanta Mondal. 2014. SPRINGS: Prediction of	3600
3545	crowdsourced, and machine assessment of suicide	protein-protein interaction sites using artificial neural	3601
3546	risk via online postings . In <i>Proceedings of the Fifth</i>	networks . <i>Journal of Proteomics & Computational</i>	3602
3547	<i>Workshop on Computational Linguistics and Clinical</i>	<i>Biology</i> , 1:7.	3603
3548	<i>Psychology: From Keyboard to Clinic</i> , pages 25–36,		
3549	New Orleans, LA. Association for Computational	Rishabh Singh and Sumit Gulwani. 2015. Predicting	3604
3550	Linguistics.	a correct program in programming by example . In	3605
		<i>Computer Aided Verification</i> , pages 398–414, Cham.	3606
3551	Tal Shnitzer, Anthony Ou, Mírian Silva, Kate Soule,	Springer International Publishing.	3607
3552	Yuekai Sun, Justin Solomon, Neil Thompson, and		
3553	Mikhail Yurochkin. 2023. Large language model	Rishabh Singh and Sumit Gulwani. 2016. Transform-	3608
3554	routing with benchmark datasets. <i>arXiv preprint</i>	ing spreadsheet data types using examples . In	3609
3555	<i>arXiv:2309.15789</i> .	<i>Proceedings of the 43rd Annual ACM SIGPLAN-</i>	3610
		<i>SIGACT Symposium on Principles of Programming</i>	3611
3556	Abu Awal Md Shoeb and Gerard de Melo. 2020. Emo-	<i>Languages, POPL '16</i> , page 343–356, New York,	3612
3557	Tag1200: Understanding the association between	NY, USA. Association for Computing Machinery.	3613
3558	emojis and emotions . In <i>Proceedings of the 2020</i>		
3559	<i>Conference on Empirical Methods in Natural Lan-</i>	Shikhar Singh, Nuan Wen, Yu Hou, Pegah Alipoormo-	3614
3560	<i>guage Processing (EMNLP)</i> , pages 8957–8967, On-	labashi, Te-lin Wu, Xuezhe Ma, and Nanyun Peng.	3615
3561	line. Association for Computational Linguistics.	2021. COM2SENSE: A commonsense reasoning	3616
		benchmark with complementary sentences . In <i>Find-</i>	3617
3562	Kurt Shuster, Spencer Poff, Moya Chen, Douwe Kiela,	<i>ings of the Association for Computational Linguistics:</i>	3618
3563	and Jason Weston. 2021. Retrieval augmentation	<i>ACL-IJCNLP 2021</i> , pages 883–898, Online. Associa-	3619
3564	reduces hallucination in conversation . <i>arXiv preprint</i> .	tion for Computational Linguistics.	3620

3621	Koustuv Sinha, Shagun Sodhani, Jin Dong, Joelle	Animesh Gupta, Anna Gottardi, Antonio Norelli,	3678
3622	Pineau, and William L. Hamilton. 2019. CLUTRR:	Anu Venkatesh, Arash Gholamidavoodi, Arfa Tabas-	3679
3623	A diagnostic benchmark for inductive reasoning from	sum, Arul Menezes, Arun Kirubakaran, Asher Mul-	3680
3624	text. <i>arXiv preprint.</i>	lokandov, Ashish Sabharwal, Austin Herrick, Avia	3681
3625	Natalia Skachkova, Thomas Trost, and Dietrich Klakow.	Efrat, Aykut Erdem, Ayla Karakaş, B. Ryan Roberts,	3682
3626	2018. Closing brackets with recurrent neural net-	Bao Sheng Loe, Barret Zoph, Bartłomiej Bojanowski,	3683
3627	works. In <i>Proceedings of the 2018 EMNLP Work-</i>	Batuhan Özyurt, Behnam Hedayatnia, Behnam	3684
3628	<i>shop BlackboxNLP: Analyzing and Interpreting Neu-</i>	Neyshabur, Benjamin Inden, Benno Stein, Berk	3685
3629	<i>ral Networks for NLP</i> , pages 232–239, Brussels, Bel-	Ekmekci, Bill Yuchen Lin, Blake Howald, Bryan	3686
3630	gium. Association for Computational Linguistics.	Orinion, Cameron Diao, Cameron Dour, Cather-	3687
3631	Douglas R. Smith. 1984. The synthesis of LISP pro-	ine Stinson, Cedrick Argueta, César Ferri Ramírez,	3688
3632	grams from examples: A survey. In Alan W. Bier-	Chandan Singh, Charles Rathkopf, Chenlin Meng,	3689
3633	mann, Gerhard Guiho, and Yves Kodratoff, editors,	Chitta Baral, Chiyu Wu, Chris Callison-Burch, Chris	3690
3634	<i>Automatic Program Construction Techniques</i> , pages	Waites, Christian Voigt, Christopher D. Manning,	3691
3635	307–324. Macmillan, New York.	Christopher Potts, Cindy Ramirez, Clara E. Rivera,	3692
3636	Paul Smolensky. 1988. On the proper treatment of	Clemencia Siro, Colin Raffel, Courtney Ashcraft,	3693
3637	connectionism. <i>Behavioral and Brain Sciences</i> ,	Cristina Garbacea, Damien Sileo, Dan Garrette, Dan	3694
3638	11(1):1–23.	Hendrycks, Dan Kilman, Dan Roth, Daniel Free-	3695
3639	Richard Socher, Alex Perelygin, Jean Wu, Jason	man, Daniel Khashabi, Daniel Levy, Daniel Moseguí	3696
3640	Chuang, Christopher D. Manning, Andrew Ng, and	González, Danielle Perszyk, Danny Hernandez,	3697
3641	Christopher Potts. 2013. Recursive deep models for	Danqi Chen, Daphne Ippolito, Dar Gilboa, David Do-	3698
3642	semantic compositionality over a sentiment treebank.	han, David Drakard, David Jurgens, Debajyoti Datta,	3699
3643	In <i>Proceedings of the 2013 Conference on Empiri-</i>	Deep Ganguli, Denis Emelin, Denis Kleyko, Deniz	3700
3644	<i>cal Methods in Natural Language Processing</i> , pages	Yuret, Derek Chen, Derek Tam, Dieuwke Hupkes,	3701
3645	1631–1642, Seattle, Washington, USA. Association	Diganta Misra, Dilyar Buzan, Dimitri Coelho Mollo,	3702
3646	for Computational Linguistics.	Diyi Yang, Dong-Ho Lee, Dylan Schrader, Ekaterina	3703
3647	Irene Solaiman, Miles Brundage, Jack Clark, Amanda	Shutova, Ekin Dogus Cubuk, Elad Segal, Eleanor	3704
3648	Askill, Ariel Herbert-Voss, Jeff Wu, Alec Radford,	Hagerman, Elizabeth Barnes, Elizabeth Donoway, El-	3705
3649	Gretchen Krueger, Jong Wook Kim, Sarah Kreps,	lie Pavlick, Emanuele Rodola, Emma Lam, Eric Chu,	3706
3650	Miles McCain, Alex Newhouse, Jason Blazakis, Kris	Eric Tang, Erkut Erdem, Ernie Chang, Ethan A. Chi,	3707
3651	McGuffie, and Jasmine Wang. 2019. Release strate-	Ethan Dyer, Ethan Jerzak, Ethan Kim, Eunice En-	3708
3652	gies and the social impacts of language models.	gefu Manyasi, Evgenii Zheltonozhskii, Fanyue Xia,	3709
3653	<i>arXiv preprint.</i>	Fatemeh Siar, Fernando Martínez-Plumed, Francesca	3710
3654	Mahmoud Soltani Firouz, Ali Farahmandi, and	Happé, Francois Chollet, Frieda Rong, Gaurav	3711
3655	Soleiman Hosseinpour. 2021. Early detection of	Mishra, Genta Indra Winata, Gerard de Melo, Ger-	3712
3656	freeze damage in navel orange fruit using nondestructive	mán Kruszewski, Giambattista Parascandolo, Gior-	3713
3657	low intensity ultrasound coupled with machine	gio Mariani, Gloria Wang, Gonzalo Jaimovitch-	3714
3658	learning. <i>Food Analytical Methods</i> , 14:1140–1149.	López, Gregor Betz, Guy Gur-Ari, Hana Galijase-	3715
3659	Dan Sperber and Deirdre Wilson. 2002. Pragmatics,	vic, Hannah Kim, Hannah Rashkin, Hannaneh Ha-	3716
3660	modularity and mind-reading. <i>Mind & Language</i> ,	jishirzi, Harsh Mehta, Hayden Bogar, Henry Shevlin,	3717
3661	17(1-2):3–23.	Hinrich Schütze, Hiromu Yakura, Hongming Zhang,	3718
3662	Peter Spirtes, Clark Glymour, and Richard Scheines.	Hugh Mee Wong, Ian Ng, Isaac Noble, Jaap Jumelet,	3719
3663	2000. <i>Causation, Prediction, and Search</i> . MIT Press,	Jack Geissinger, Jackson Kernion, Jacob Hilton, Jae-	3720
3664	Cambridge, MA.	hooon Lee, Jaime Fernández Fisac, James B. Simon,	3721
3665	Aarohi Srivastava, Abhinav Rastogi, Abhishek Rao,	James Koppel, James Zheng, James Zou, Jan Kocoń,	3722
3666	Abu Awal Md Shoeb, Abubakar Abid, Adam	Jana Thompson, Janelle Wingfield, Jared Kaplan,	3723
3667	Fisch, Adam R. Brown, Adam Santoro, Aditya	Jarema Radom, Jascha Sohl-Dickstein, Jason Phang,	3724
3668	Gupta, Adrià Garriga-Alonso, Agnieszka Kluska,	Jason Wei, Jason Yosinski, Jekaterina Novikova,	3725
3669	Aitor Lewkowycz, Akshat Agarwal, Alethea Power,	Jelle Bosscher, Jennifer Marsh, Jeremy Kim, Jeroen	3726
3670	Alex Ray, Alex Warstadt, Alexander W. Kocurek,	Taal, Jesse Engel, Jesujoba Alabi, Jiacheng Xu, Ji-	3727
3671	Ali Safaya, Ali Tazarv, Alice Xiang, Alicia Par-	aming Song, Jillian Tang, Joan Waweru, John Bur-	3728
3672	rish, Allen Nie, Aman Hussain, Amanda Askill,	den, John Miller, John U. Balis, Jonathan Batchelder,	3729
3673	Amanda Dsouza, Ambrose Slone, Ameet Rahane,	Jonathan Berant, Jörg Froberg, Jos Rozen, Jose	3730
3674	Anantharaman S. Iyer, Anders Andreassen, Andrea	Hernandez-Orallo, Joseph Boudeman, Joseph Guerr,	3731
3675	Madotto, Andrea Santilli, Andreas Stuhlmüller, An-	Joseph Jones, Joshua B. Tenenbaum, Joshua S. Rule,	3732
3676	drew Dai, Andrew La, Andrew Lampinen, Andy	Joyce Chua, Kamil Kanclerz, Karen Livescu, Karl	3733
3677	Zou, Angela Jiang, Angelica Chen, Anh Vuong,	Krauth, Karthik Gopalakrishnan, Katerina Ignatyeva,	3734
		Katja Markert, Kaustubh D. Dhole, Kevin Gimp-	3735
		pel, Kevin Omondi, Kory Mathewson, Kristen Chi-	3736
		afullo, Ksenia Shkaruta, Kumar Shridhar, Kyle Mc-	3737
		Donell, Kyle Richardson, Laria Reynolds, Leo Gao,	3738
		Li Zhang, Liam Dugan, Lianhui Qin, Lidia Contreras-	3739
		Ochando, Louis-Philippe Morency, Luca Moschella,	3740
		Lucas Lam, Lucy Noble, Ludwig Schmidt, Luheng	3741

3742	He, Luis Oliveros Colón, Luke Metz, Lütü Kerem	3806
3743	Şenel, Maarten Bosma, Maarten Sap, Maartje ter	3807
3744	Hoeve, Maheen Farooqi, Manaal Faruqui, Mantas	3808
3745	Mazeika, Marco Baturan, Marco Marelli, Marco	3809
3746	Maru, Maria Jose Ramírez Quintana, Marie Tolkiehn,	3810
3747	Mario Giulianelli, Martha Lewis, Martin Potthast,	3811
3748	Matthew L. Leavitt, Matthias Hagen, Mátyás Schu-	3812
3749	bert, Medina Orduna Baitemirova, Melody Arnaud,	3813
3750	Melvin McElrath, Michael A. Yee, Michael Co-	
3751	hen, Michael Gu, Michael Ivanitskiy, Michael Star-	
3752	ritt, Michael Strube, Michał Śwędrowski, Michele	3814
3753	Bevilacqua, Michihiro Yasunaga, Mihir Kale, Mike	3815
3754	Cain, Mimeo Xu, Mirac Suzgun, Mitch Walker,	3816
3755	Mo Tiwari, Mohit Bansal, Moin Aminnaseri, Mor	3817
3756	Geva, Mozhdah Gheini, Mukund Varma T, Nanyun	3818
3757	Peng, Nathan A. Chi, Nayeon Lee, Neta Gur-Ari	3819
3758	Krakov, Nicholas Cameron, Nicholas Roberts,	
3759	Nick Doiron, Nicole Martinez, Nikita Nangia, Niklas	3820
3760	Deckers, Niklas Muennighoff, Nitish Shirish Keskar,	3821
3761	Niveditha S. Iyer, Noah Constant, Noah Fiedel, Nuan	3822
3762	Wen, Oliver Zhang, Omar Agha, Omar Elbaghdadi,	
3763	Omer Levy, Owain Evans, Pablo Antonio Moreno	
3764	Casares, Parth Doshi, Pascale Fung, Paul Pu Liang,	3823
3765	Paul Vicol, Pegah Alipoormolabashi, Peiyuan Liao,	3824
3766	Percy Liang, Peter Chang, Peter Eckersley, Phu Mon	3825
3767	Htut, Pinyu Hwang, Piotr Miłkowski, Piyush Patil,	3826
3768	Pouya Pezeshkpour, Priti Oli, Qiaozhu Mei, Qing	3827
3769	Lyu, Qinlang Chen, Rabin Banjade, Rachel Etta	
3770	Rudolph, Raefer Gabriel, Rahel Habacker, Ramon	
3771	Risco, Raphaël Millièvre, Rhythm Garg, Richard	3828
3772	Barnes, Rif A. Saurous, Riku Arakawa, Robbe	3829
3773	Raymaekers, Robert Frank, Rohan Sikand, Roman	3830
3774	Novak, Roman Sitelew, Ronan LeBras, Rosanne	3831
3775	Liu, Rowan Jacobs, Rui Zhang, Ruslan Salakhut-	3832
3776	dinov, Ryan Chi, Ryan Lee, Ryan Stovall, Ryan	3833
3777	Teehan, Rylan Yang, Sahib Singh, Saif M. Moham-	
3778	mad, Sajant Anand, Sam Dillavou, Sam Shleifer,	3834
3779	Sam Wiseman, Samuel Gruetter, Samuel R. Bow-	3835
3780	man, Samuel S. Schoenholz, Sanghyun Han, San-	3836
3781	jeev Kwatra, Sarah A. Rous, Sarik Ghazarian, Sayan	3837
3782	Ghosh, Sean Casey, Sebastian Bischoff, Sebastian	3838
3783	Gehrmann, Sebastian Schuster, Sepideh Sadeghi,	3839
3784	Shadi Hamdan, Sharon Zhou, Shashank Srivastava,	
3785	Sherry Shi, Shikhar Singh, Shima Asaadi, Shixi-	
3786	ang Shane Gu, Shubh Pachchigar, Shubham Tosh-	
3787	niwal, Shyam Upadhyay, Shyamolima, Debnath,	
3788	Siamak Shakeri, Simon Thormeyer, Simone Melzi,	
3789	Siva Reddy, Sneha Priscilla Makini, Soo-Hwan Lee,	
3790	Spencer Torene, Sriharsha Hatwar, Stanislas De-	
3791	haene, Stefan Divic, Stefano Ermon, Stella Bider-	
3792	man, Stephanie Lin, Stephen Prasad, Steven T. Pi-	
3793	antadosi, Stuart M. Shieber, Summer Mishnerghi, Svet-	
3794	lana Kiritchenko, Swaroop Mishra, Tal Linzen, Tal	
3795	Schuster, Tao Li, Tao Yu, Tariq Ali, Tatsu Hashimoto,	
3796	Te-Lin Wu, Théo Desbordes, Theodore Rothschild,	
3797	Thomas Phan, Tianle Wang, Tiberius Nkinyili, Timo	
3798	Schick, Timofei Kornev, Titus Tunduny, Tobias Ger-	
3799	stenberg, Trenton Chang, Trishala Neeraj, Tushar	
3800	Khot, Tyler Shultz, Uri Shaham, Vedant Misra, Vera	
3801	Demberg, Victoria Nyamai, Vikas Raunak, Vinay	
3802	Ramasesh, Vinay Uday Prabhu, Vishakh Padmakumar,	
3803	Vivek Srikumar, William Fedus, William Saunders,	
3804	William Zhang, Wout Vossen, Xiang Ren, Xi-	
3805	aoyu Tong, Xinran Zhao, Xinyi Wu, Xudong Shen,	
	Yadollah Yaghoobzadeh, Yair Lakretz, Yangqiu Song,	
	Yasaman Bahri, Yejin Choi, Yichi Yang, Yiding	
	Hao, Yifu Chen, Yonatan Belinkov, Yu Hou, Yufang	
	Hou, Yuntao Bai, Zachary Seid, Zhuoye Zhao, Zijian	
	Wang, Zijie J. Wang, Zirui Wang, and Ziyi Wu. 2023.	
	Beyond the imitation game: Quantifying and extrapolating the capabilities of language models. <i>Preprint</i> , arXiv:2206.04615.	
	Shashank Srivastava, Snigdha Chaturvedi, and Tom	
	Mitchell. 2016. Inferring interpersonal relations in narrative summaries. In <i>Proceedings of the Thirtieth AAAI Conference on Artificial Intelligence, AAAI'16</i> , page 2807–2813, Menlo Park, CA. Association for the Advancement of Artificial Intelligence.	
	Robert Stalnaker. 1978. Assertion. In P. Cole, editor, <i>Pragmatics</i> , Syntax and Semantics 9, pages 315–332. Brill, Leiden.	
	Trevor Standley, Amir Zamir, Dawn Chen, Leonidas	
	Guibas, Jitendra Malik, and Silvio Savarese. 2020. Which tasks should be learned together in multi-task learning? In <i>International conference on machine learning</i> , pages 9120–9132. PMLR.	
	Gabriel Stanovsky, Noah A. Smith, and Luke Zettlemoyer. 2019. Evaluating gender bias in machine translation. In <i>Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics</i> , pages 1679–1684, Florence, Italy. Association for Computational Linguistics.	
	Gerard J. Steen, Aletta G. Dorst, J. Berenike Herrmann, Anna A. Kaal, Tina Krennmayr, and Tryntje Pasma. 2010. <i>A Method for Linguistic Metaphor Identification: From MIP to MIPVU</i> . Converging Evidence in Language and Communication Research 14. John Benjamins, Amsterdam.	
	Bernd Steinbach and Roman Kohut. 2002. Neural networks – a model of boolean functions. <i>5th International Workshop on Boolean Problems, Freiburg, Sept. 2002.</i>	
	Sebastian U. Stich. 2018. Local sgd converges fast and communicates little. <i>arXiv preprint arXiv:1805.09767</i> .	
	Nisan Stiennon, Long Ouyang, Jeff Wu, Daniel M. Ziegler, Ryan Lowe, Chelsea Voss, Alec Radford, Dario Amodei, and Paul Christiano. 2020. Learning to summarize from human feedback. <i>arXiv preprint</i> .	
	Kevin Stowe, Leonardo Ribeiro, and Iryna Gurevych. 2020. Metaphoric paraphrase generation. <i>arXiv preprint</i> .	
	Michael Strube and Simone Paolo Ponzetto. 2006. Wikirelate! Computing semantic relatedness using Wikipedia. In <i>AAAI'06: Proceedings of the 21st National Conference on Artificial Intelligence</i> , volume 2, page 1419–1424. Association for the Advancement of Artificial Intelligence.	

3860	Saku Sugawara, Hikaru Yokono, and Akiko Aizawa.	Colton Swingle, Henry Mellsop, and Alex Lang-	3915
3861	2017. Prerequisite skills for reading comprehension:	shur. 2021. ChePT – applying deep neural trans-	3916
3862	Multi-perspective analysis of MCTest datasets and	former models to chess move prediction and self-	3917
3863	systems . In <i>Proceedings of the AAAI Conference</i>	commentary .	3918
3864	<i>on Artificial Intelligence</i> , volume 31, Menlo Park,		
3865	CA. Association for the Advancement of Artificial	Oyvind Tafjord, Peter Clark, Matt Gardner, Wen-tau	3919
3866	Intelligence.	Yih, and Ashish Sabharwal. 2018. Quarel: A dataset	3920
		and models for answering questions about qualitative	3921
3867	Alane Suhr, Claudia Yan, Jack Schluger, Stanley Yu,	relationships . <i>CoRR</i> , abs/1811.08048.	3922
3868	Hadi Khader, Marwa Mouallem, Iris Zhang, and		
3869	Yoav Artzi. 2019. Executing instructions in situ-	Oyvind Tafjord, Matt Gardner, Kevin Lin, and Peter	3923
3870	ated collaborative interactions . In <i>Proceedings of</i>	Clark. 2019. "quartz: An open-domain dataset of	3924
3871	<i>the 2019 Conference on Empirical Methods in Natu-</i>	qualitative relationship questions". <i>EMNLP</i> .	3925
3872	<i>ral Language Processing and the 9th International</i>		
3873	<i>Joint Conference on Natural Language Processing</i>	Alon Talmor, Yanai Elazar, Yoav Goldberg, and	3926
3874	<i>(EMNLP-IJCNLP)</i> , pages 2119–2130, Hong Kong,	Jonathan Berant. 2019a. oLMpics – on what lan-	3927
3875	China. Association for Computational Linguistics.	guage model pre-training captures . <i>arXiv preprint</i> .	3928
3876	Sainbayar Sukhbaatar, Olga Golovneva, Vasu Sharma,	Alon Talmor, Jonathan Herzig, Nicholas Lourie, and	3929
3877	Hu Xu, Xi Victoria Lin, Baptiste Rozière, Jac-	Jonathan Berant. 2019b. CommonsenseQA: A ques-	3930
3878	cob Kahn, Daniel Li, Wen-tau Yih, Jason We-	tion answering challenge targeting commonsense	3931
3879	ston, et al. 2024. Branch-train-mix: Mixing expert	knowledge . In <i>Proceedings of the 2019 Conference</i>	3932
3880	llms into a mixture-of-experts llm. <i>arXiv preprint</i>	<i>of the North American Chapter of the Association for</i>	3933
3881	<i>arXiv:2403.07816</i> .	<i>Computational Linguistics: Human Language Tech-</i>	3934
		<i>nologies, Volume 1 (Long and Short Papers)</i> , pages	3935
3882	Kai Sun, Dian Yu, Jianshu Chen, Dong Yu, Yejin Choi,	4149–4158, Minneapolis, Minnesota. Association for	3936
3883	and Claire Cardie. 2019a. DREAM: A challenge	Computational Linguistics.	3937
3884	dataset and models for dialogue-based reading com-		
3885	prehension . <i>Transactions of the Association for Com-</i>	Derek Tam, Mohit Bansal, and Colin Raffel. 2023.	3938
3886	<i>putational Linguistics</i> .	Merging by matching models in task subspaces.	3939
		<i>arXiv preprint arXiv:2312.04339</i> .	3940
3887	Wenlong Sun, Olfa Nasraoui, and Patrick Shafto. 2020.	Alex Tamkin, Miles Brundage, Jack Clark, and Deep	3941
3888	Evolution and impact of bias in human and machine	Ganguli. 2021. Understanding the capabilities, limi-	3942
3889	learning algorithm interaction . <i>PLOS ONE</i> , 15(8):1–	tations, and societal impact of large language models .	3943
3890	39.	<i>arXiv preprint</i> .	3944
3891	Ximeng Sun, Rameswar Panda, and Rogério Schmidt	Jiwei Tan, Xiaojun Wan, and Jianguo Xiao. 2015.	3945
3892	Feris. 2019b. Adashare: Learning what to share	Learning to recommend quotes for writing . In <i>Pro-</i>	3946
3893	for efficient deep multi-task learning . <i>ArXiv</i> ,	<i>ceedings of the Twenty-Ninth AAAI Conference on</i>	3947
3894	abs/1911.12423.	<i>Artificial Intelligence, AAAI’15</i> , page 2453–2459,	3948
		Menlo Park, CA. Association for the Advancement	3949
3895	Yi-Lin Sung, Jaemin Cho, and Mohit Bansal. 2022.	of Artificial Intelligence.	3950
3896	Lst: Ladder side-tuning for parameter and memory		
3897	efficient transfer learning. In <i>Advances in Neural</i>	Niket Tandon, Bhavana Dalvi, Keisuke Sakaguchi, Pe-	3951
3898	<i>Information Processing Systems</i> .	ter Clark, and Antoine Bosselut. 2019. WIQA: A	3952
		dataset for “what if...” reasoning over procedural text .	3953
3899	Mirac Suzgun, Yonatan Belinkov, Stuart Shieber, and	In <i>Proceedings of the 2019 Conference on Empirical</i>	3954
3900	Sebastian Gehrmann. 2019a. LSTM networks can	<i>Methods in Natural Language Processing and the</i>	3955
3901	perform dynamic counting . In <i>Proceedings of the</i>	<i>9th International Joint Conference on Natural Lan-</i>	3956
3902	<i>Workshop on Deep Learning and Formal Languages:</i>	<i>guage Processing (EMNLP-IJCNLP)</i> , pages 6076–	3957
3903	<i>Building Bridges</i> , pages 44–54, Florence. Associa-	6085, Hong Kong, China. Association for Computa-	3958
3904	tion for Computational Linguistics.	tional Linguistics.	3959
3905	Mirac Suzgun, Sebastian Gehrmann, Yonatan Belinkov,	Anke Tang, Li Shen, Yong Luo, Nan Yin, Lefei Zhang,	3960
3906	and Stuart M. Shieber. 2019b. Memory-augmented	and Dacheng Tao. 2024. Merging multi-task models	3961
3907	recurrent neural networks can learn generalized Dyck	via weight-ensembling mixture of experts . <i>Preprint</i> ,	3962
3908	languages . <i>arXiv preprint</i> .	arXiv:2402.00433.	3963
3909	Mirac Suzgun, Nathan Scales, Nathanael Schärli, Se-	Jan Arne Telle, José Hernández-Orallo, and Cèsar Ferri.	3964
3910	bastian Gehrmann, Yi Tay, Hyung Won Chung,	2019. The teaching size: Computable teachers and	3965
3911	Aakanksha Chowdhery, Quoc V Le, Ed H Chi, Denny	learners for universal languages . <i>Machine Learning</i> ,	3966
3912	Zhou, and Jason Wei. 2022. Challenging big-bench	108:1653–1675.	3967
3913	tasks and whether chain-of-thought can solve them.		
3914	<i>arXiv preprint arXiv:2210.09261</i> .		

- Damien Teney, Ehsan Abbasnejad, and Anton van den Hengel. 2020. [Learning what makes a difference from counterfactual examples and gradient supervision](#). *arXiv preprint*. 4024
- Paul Thagard, Keith J. Holyoak, Greg Nelson, and David Gochfeld. 1990. [Analog retrieval by constraint satisfaction](#). *Artificial Intelligence*, 46(3):259–310. 4025
- Avijit Thawani, Jay Pujara, Filip Ilievski, and Pedro Szekely. 2021. [Representing numbers in NLP: A survey and a vision](#). In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 644–656, Online. Association for Computational Linguistics. 4026
- Jesse Thomason, Shiqi Zhang, Raymond Mooney, and Peter Stone. 2015. [Learning to interpret natural language commands through human-robot dialog](#). In *Proceedings of the Twenty-Fourth International Joint Conference on Artificial Intelligence, IJCAI-15*, pages 1923–1929. 4027
- Judith Jarvis Thomson. 1976. [Killing, letting die, and the trolley problem](#). *The Monist*, 59(2):204–217. 4028
- Poonam B. Thorat, Rajeshwari M. Goudar, and Sunita Barve. 2015. [Survey on collaborative filtering, content-based filtering and hybrid recommendation system](#). *International Journal of Computer Applications*, 110:31–36. 4029
- James Thorne, Andreas Vlachos, Christos Christodoulopoulos, and Arpit Mittal. 2018. [FEVER: A large-scale dataset for fact extraction and VERification](#). In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)*, pages 809–819, New Orleans, Louisiana. Association for Computational Linguistics. 4030
- Xiaoyu Tong. 2021. [Metaphor paraphrasing and word sense disambiguation: Toward a new approach to automated metaphor](#). 4031
- Xiaoyu Tong, Ekaterina Shutova, and Martha Lewis. 2021. [Recent advances in neural metaphor processing: A linguistic, cognitive and social perspective](#). In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 4673–4686, Online. Association for Computational Linguistics. 4032
- Shubham Toshniwal, Sam Wiseman, Karen Livescu, and Kevin Gimpel. 2021. [Learning chess blindfolded: Evaluating language models on state tracking](#). *arXiv preprint*. 4033
- David Toubiana, Nir Sade, Lifeng Liu, Maria del Mar Rubio Wilhelmi, Yariv Brotman, Urszula Luzarowska, John P. Vogel, and Eduardo Blumwald. 2020. [Correlation-based network analysis combined with machine learning techniques highlight the role of the gaba shunt in brachypodium sylvaticum freezing tolerance](#). *Scientific Reports*, 10:no. 4489. 4034
- Andrew Trask, Felix Hill, Scott Reed, Jack Rae, Chris Dyer, and Phil Blunsom. 2018. [Neural arithmetic logic units](#). *arXiv preprint*. 4035
- George Tsatsaronis, Iraklis Varlamis, and Michalis Vazirgiannis. 2010. [Text relatedness based on a word thesaurus](#). *Journal of Artificial Intelligence Research*, 37(1):1–40. 4036
- George Tsatsaronis, Iraklis Varlamis, Michalis Vazirgiannis, and Kjetil Nørvåg. 2009. Omiotis: A thesaurus-based measure of text relatedness. In *Machine Learning and Knowledge Discovery in Databases*, pages 742–745, Berlin. Springer. 4037
- Yulia Tsvetkov, Leonid Boytsov, Anatole Gershman, Eric Nyberg, and Chris Dyer. 2014. [Metaphor detection with cross-lingual model transfer](#). In *Proceedings of the 52nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 248–258, Baltimore, Maryland. Association for Computational Linguistics. 4038
- Shikhar Vashishth, Shib Sankar Dasgupta, Swayambhu Nath Ray, and Partha Talukdar. 2018. [Dating documents using graph convolution networks](#). In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1605–1615, Melbourne, Australia. Association for Computational Linguistics. 4039
- Siddharth Vashishtha, Adam Poliak, Yash Kumar Lal, Benjamin Van Durme, and Aaron Steven White. 2020. [Temporal reasoning in natural language inference](#). In *Findings of the Association for Computational Linguistics: EMNLP 2020*, pages 4070–4078, Online. Association for Computational Linguistics. 4040
- Siddharth Vashishtha, Benjamin Van Durme, and Aaron Steven White. 2019. [Fine-grained temporal relation extraction](#). *arXiv preprint*. 4041
- Oleg Vasilyev, Vedant Dharnidharka, and John Bohannon. 2020. [Fill in the BLANC: Human-free quality estimation of document summaries](#). In *Proceedings of the First Workshop on Evaluation and Comparison of NLP Systems*, pages 11–20, Online. Association for Computational Linguistics. 4042
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. Attention is all you need. *Advances in neural information processing systems*, 30. 4043
- Jette Viethen and Robert Dale. 2008. [The use of spatial relations in referring expression generation](#). In *Proceedings of the Fifth International Natural Language Generation Conference*, pages 59–67, Salt Fork, Ohio, USA. Association for Computational Linguistics. 4044

4079	Oriol Vinyals, Igor Babuschkin, Wojciech M. Czarnecki,	A multi-task benchmark and analysis platform for nat-	4136
4080	Michaël Mathieu, Andrew Dudzik, Junyoung Chung,	ural language understanding . In <i>Proceedings of the</i>	4137
4081	David H. Choi, Richard Powell, Timo Ewalds, Petko	<i>2018 EMNLP Workshop BlackboxNLP: Analyzing</i>	4138
4082	Georgiev, Junhyuk Oh, Dan Horgan, Manuel Kroiss,	<i>and Interpreting Neural Networks for NLP</i> , pages	4139
4083	Ivo Danihelka, Aja Huang, Laurent Sifre, Trevor	353–355, Brussels, Belgium. Association for Com-	4140
4084	Cai, John P. Agapiou, Max Jaderberg, Alexander S.	putational Linguistics.	4141
4085	Vezhnevets, Rémi Leblond, Tobias Pohlen, Valentin		
4086	Dalibard, David Budden, Yuri Sulsky, James Mol-	Ben Wang and Aran Komatsuzaki. 2021. GPT-J-6B: A	4142
4087	loy, Tom L. Paine, Caglar Gulcehre, Ziyu Wang, To-	6 billion parameter autoregressive language model .	4143
4088	bias Pfaff, Yuhuai Wu, Roman Ring, Dani Yogatama,		
4089	Dario Wünsch, Katrina McKinney, Oliver Smith,	Hanqing Wang, Bowen Ping, Shuo Wang, Xu Han,	4144
4090	Tom Schaul, Timothy Lillicrap, Koray Kavukcuoglu,	Yun Chen, Zhiyuan Liu, and Maosong Sun. 2024.	4145
4091	Demis Hassabis, Chris Apps, and David Silver. 2019.	Lora-flow: Dynamic lora fusion for large lan-	4146
4092	Grandmaster level in StarCraft II using multi-agent	guage models in generative tasks . <i>arXiv preprint</i>	4147
4093	reinforcement learning . <i>Nature</i> , 575:350–354.	<i>arXiv:2402.11455</i> .	4148
4094	Ellen M. Voorhees. 2002. Overview of the trec 2002	Jianfeng Wang, Rong Xiao, Yandong Guo, and Lei	4149
4095	question answering track . In <i>Proceedings of the</i>	Zhang. 2019c. Learning to count objects with few	4150
4096	<i>Eleventh Text REtrieval Conference (TREC 2002)</i> .	exemplar annotations . <i>arXiv preprint</i> .	4151
4097	Tu Vu, Tong Wang, Tsendsuren Munkhdalai, Alessan-	Lingzhi Wang, Jing Li, Xingshan Zeng, Haisong Zhang,	4152
4098	doro Sordoni, Adam Trischler, Andrew Mattarella-	and Kam-Fai Wong. 2020b. Continuity of topic, in-	4153
4099	Micke, Subhransu Maji, and Mohit Iyyer. 2020. Ex-	teraction, and query: Learning to quote in online con-	4154
4100	ploring and predicting transferability across nlp tasks.	versations . In <i>Proceedings of the 2020 Conference on</i>	4155
4101	<i>arXiv preprint arXiv:2005.00770</i> .	<i>Empirical Methods in Natural Language Processing</i>	4156
4102	Henning Wachsmuth, Nona Naderi, Yufang Hou,	(EMNLP), pages 6640–6650, Online. Association for	4157
4103	Yonatan Bilu, Vinodkumar Prabhakaran, Tim Alberd-	Computational Linguistics.	4158
4104	ingk Thijm, Graeme Hirst, and Benno Stein. 2017.	Po-Wei Wang, Priya L. Donti, Bryan Wilder, and Zico	4159
4105	Computational argumentation quality assessment in	Kolter. 2019d. SATNet: Bridging deep learning and	4160
4106	natural language . In <i>Proceedings of the 15th Con-</i>	logical reasoning using a differentiable satisfiability	4161
4107	<i>ference of the European Chapter of the Association</i>	solver . <i>arXiv preprint</i> .	4162
4108	<i>for Computational Linguistics: Volume 1, Long Pa-</i>		
4109	<i>pers</i> , pages 176–187, Valencia, Spain. Association	Sida I. Wang, Percy Liang, and Christopher D. Manning.	4163
4110	for Computational Linguistics.	2016. Learning language games through interaction .	4164
4111	Eric Wallace, Florian Tramèr, Matthew Jagielski, and	In <i>Proceedings of the 54th Annual Meeting of the</i>	4165
4112	Ariel Herbert-Voss. 2020. Does GPT-2 know your	<i>Association for Computational Linguistics (Volume</i>	4166
4113	phone number? <i>Berkeley Artificial Intelligence Re-</i>	<i>1: Long Papers</i>), pages 2368–2378, Berlin, Germany.	4167
4114	<i>search blog</i> .	Association for Computational Linguistics.	4168
4115	Alex Wang, Kyunghyun Cho, and Mike Lewis. 2020a.	Yaqing Wang, Subhabrata Mukherjee, Xiaodong Liu,	4169
4116	Asking and answering questions to evaluate the fac-	Jing Gao, Ahmed Hassan Awadallah, and Jian-	4170
4117	tual consistency of summaries . In <i>Proceedings of the</i>	feng Gao. 2022a. Adamix: Mixture-of-adapters for	4171
4118	<i>58th Annual Meeting of the Association for Compu-</i>	parameter-efficient tuning of large language models.	4172
4119	<i>tational Linguistics</i> , pages 5008–5020, Online. Asso-	<i>arXiv preprint arXiv:2205.12410</i> .	4173
4120	ciation for Computational Linguistics.	Yizhong Wang, Swaroop Mishra, Pegah Alipoor-	4174
4121	Alex Wang, Yada Pruksachatkun, Nikita Nangia, Aman-	molabashi, Yeganeh Kordi, Amirreza Mirzaei,	4175
4122	preet Singh, Julian Michael, Felix Hill, Omer Levy,	Anjana Arunkumar, Arjun Ashok, Arut Selvan	4176
4123	and Samuel Bowman. 2019a. Superglue: A stick-	Dhanasekaran, Atharva Naik, David Stap, et al.	4177
4124	ier benchmark for general-purpose language under-	2022b. Super-naturalinstructions: Generalization via	4178
4125	standing systems. <i>Advances in neural information</i>	declarative instructions on 1600+ nlp tasks. <i>arXiv</i>	4179
4126	<i>processing systems</i> , 32.	<i>preprint arXiv:2204.07705</i> .	4180
4127	Alex Wang, Yada Pruksachatkun, Nikita Nangia, Aman-	Zijian Wang and David Jurgens. 2018. It’s going to be	4181
4128	preet Singh, Julian Michael, Felix Hill, Omer Levy,	okay: Measuring access to support in online com-	4182
4129	and Samuel R. Bowman. 2019b. SuperGLUE: A	munities . In <i>Proceedings of the 2018 Conference</i>	4183
4130	stickier benchmark for general-purpose language un-	<i>on Empirical Methods in Natural Language Process-</i>	4184
4131	derstanding systems . In <i>Advances in Neural Infor-</i>	<i>ing</i> , pages 33–45, Brussels, Belgium. Association for	4185
4132	<i>mation Processing Systems</i> , volume 32. Curran Asso-	Computational Linguistics.	4186
4133	ciates, Inc.	Zijian Wang and Christopher Potts. 2019. TalkDown:	4187
4134	Alex Wang, Amanpreet Singh, Julian Michael, Felix	A corpus for condescension detection in context . In	4188
4135	Hill, Omer Levy, and Samuel Bowman. 2018. GLUE:	<i>Proceedings of the 2019 Conference on Empirical</i>	4189
		<i>Methods in Natural Language Processing and the</i>	4190

4191	9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP), pages 3711–3719, Hong Kong, China. Association for Computational Linguistics.	4247
4192		4248
4193		4249
4194		4250
4195	Ingmar Weber and Alejandro Jaimes. 2011. Who uses web search for what: And how. In <i>Proceedings of the Fourth ACM International Conference on Web Search and Data Mining, WSDM '11</i> , page 15–24, New York, NY, USA. Association for Computing Machinery.	4251
4196		4252
4197		4253
4198		4254
4199		4255
4200		4256
4201	David Wechsler. 2008. <i>Wechsler Adult Intelligence Scale—Fourth Edition (WAIS-IV)</i> . Pearson, San Antonio.	4257
4202		4258
4203		4259
4204	Jason Wei, Maarten Bosma, Vincent Zhao, Kelvin Guu, Adams Wei Yu, Brian Lester, Nan Du, Andrew M. Dai, and Quoc V Le. 2022a. Finetuned language models are zero-shot learners. In <i>International Conference on Learning Representations</i> .	4260
4205		4261
4206		4262
4207		4263
4208		
4209	Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Ed Chi, Quoc V Le, and Denny Zhou. 2022b. Chain of thought prompting elicits reasoning in large language models. <i>arXiv preprint arXiv:2201.11903</i> .	4264
4210		4265
4211		
4212		
4213	Johannes Welbl, Pontus Stenetorp, and Sebastian Riedel. 2018. Constructing datasets for multi-hop reading comprehension across documents.	4266
4214		4267
4215		4268
4216	Orion Weller and Kevin Seppi. 2019. Humor detection: A transformer gets the last laugh. In <i>Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)</i> , pages 3621–3625, Hong Kong, China. Association for Computational Linguistics.	4269
4217		4270
4218		4271
4219		4272
4220		4273
4221		4274
4222		4275
4223	Jason Weston, Antoine Bordes, Sumit Chopra, Alexander M. Rush, Bart van Merriënboer, Armand Joulin, and Tomas Mikolov. 2015. Towards AI-complete question answering: A set of prerequisite toy tasks. <i>arXiv preprint</i> .	4276
4224		4277
4225		
4226		
4227		
4228	Aaron Steven White, Rachel Rudinger, Kyle Rawlins, and Benjamin Van Durme. 2018. Lexicosyntactic inference in neural models. In <i>Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing</i> , pages 4717–4724, Brussels, Belgium. Association for Computational Linguistics.	4278
4229		4279
4230		4280
4231		4281
4232		4282
4233		4283
4234	Sarah White, Elisabeth Hill, Francesca Happé, and Uta Frith. 2009. Revisiting the strange stories: Revealing mentalizing impairments in autism. <i>Child Development</i> , 80(4):1097–1117.	4284
4235		4285
4236		
4237		
4238	Adina Williams, Nikita Nangia, and Samuel Bowman. 2018. A broad-coverage challenge corpus for sentence understanding through inference. In <i>Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)</i> , pages 1112–1122, New Orleans, Louisiana. Association for Computational Linguistics.	4286
4239		4287
4240		4288
4241		4289
4242		4290
4243		
4244		
4245		
4246		
	Ulrike Willinger, Andreas Hergovich, Michaela Schmoeger, Matthias Deckert, Susanne Stoettner, Iris Bunda, Andrea Witting, Melanie Seidler, Reinhilde Moser, Stefanie Kacena, David Jaeckle, Benjamin Loader, Christian Mueller, and Eduard Auff. 2017. Cognitive and emotional demands of black humour processing: The role of intelligence, aggressiveness and mood. <i>Cognitive Processing</i> , 18:159–167.	4291
		4292
		4293
		4294
		4295
		4296
		4297
		4298
	Genta Indra Winata, Andrea Madotto, Zhaojiang Lin, Rosanne Liu, Jason Yosinski, and Pascale Fung. 2021. Language models are few-shot multilingual learners. In <i>Proceedings of the 1st Workshop on Multilingual Representation Learning</i> , pages 1–15, Punta Cana, Dominican Republic. Association for Computational Linguistics.	4299
		4300
		4301
		4302
		4303
		4304
	Terry Winograd. 1972. Understanding natural language. <i>Cognitive Psychology</i> , 3(1):1–191.	
	Ludwig Wittgenstein. 1953. <i>Philosophical investigations</i> . Basil Blackwell, Oxford.	
	Thomas Wolf. 2019. Some additional experiments extending the tech report “assessing BERT’s syntactic abilities” by Yoav Goldberg.	
	Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Rémi Louf, Morgan Funtowicz, Joe Davison, Sam Shleifer, Patrick von Platen, Clara Ma, Yacine Jernite, Julien Plu, Canwen Xu, Teven Le Scao, Sylvain Gugger, Mariama Drame, Quentin Lhoest, and Alexander M. Rush. 2019. Huggingface’s transformers: State-of-the-art natural language processing. <i>arXiv preprint</i> .	
	Min-Sub Won, YunSeok Choi, Samuel Kim, Cheol-Won Na, and Jee-Hyong Lee. 2021. An embedding method for unseen words considering contextual information and morphological information. In <i>Proceedings of the 36th Annual ACM Symposium on Applied Computing, SAC '21</i> , page 1055–1062, New York, NY, USA. Association for Computing Machinery.	
	Mitchell Wortsman, Maxwell C Horton, Carlos Guestrin, Ali Farhadi, and Mohammad Rastegari. 2021. Learning neural network subspaces. In <i>International Conference on Machine Learning</i> , pages 11217–11227. PMLR.	
	Mitchell Wortsman, Gabriel Ilharco, Samir Ya Gadre, Rebecca Roelofs, Raphael Gontijo-Lopes, Ari S Morcos, Hongseok Namkoong, Ali Farhadi, Yair Carmon, Simon Kornblith, et al. 2022. Model soups: averaging weights of multiple fine-tuned models improves accuracy without increasing inference time. In <i>International Conference on Machine Learning</i> , pages 23965–23998. PMLR.	
	Bo Wu, Pedro Szekely, and Craig A. Knoblock. 2012. Learning data transformation rules through examples: Preliminary results. In <i>Proceedings of the Ninth International Workshop on Information Integration on the Web, IIWeb '12</i> , New York, NY, USA. Association for Computing Machinery.	

4305	Chengyue Wu, Teng Wang, Yixiao Ge, Zeyu Lu,	Linting Xue, Noah Constant, Adam Roberts, Mihir Kale,	4361
4306	Ruisong Zhou, Ying Shan, and Ping Luo. 2023. <i>pi-</i>	Rami Al-Rfou, Aditya Siddhant, Aditya Barua, and	4362
4307	tuning: Transferring multimodal foundation mod-	Colin Raffel. 2021b. <i>mT5: A massively multilingual</i>	4363
4308	els with optimal multi-task interpolation. In <i>Inter-</i>	<i>pre-trained text-to-text transformer</i> . In <i>Proceedings</i>	4364
4309	<i>national Conference on Machine Learning</i> , pages	<i>of the 2021 Conference of the North American Chap-</i>	4365
4310	37713–37727. PMLR.	<i>ter of the Association for Computational Linguistics:</i>	4366
4311	Shijie Wu, Ryan Cotterell, and Mans Hulden. 2020. <i>Ap-</i>	<i>Human Language Technologies</i> , pages 483–498, On-	4367
4312	<i>plying the transformer to character-level transduction.</i>	line. Association for Computational Linguistics.	4368
4313	<i>arXiv preprint</i> .		
4314	Xun Wu, Shaohan Huang, and Furu Wei. 2024. <i>Mix-</i>	Prateek Yadav, Leshem Choshen, Colin Raffel, and Mo-	4369
4315	<i>ture of loRA experts</i> . In <i>The Twelfth International</i>	hit Bansal. 2023a. <i>Compeft: Compression for com-</i>	4370
4316	<i>Conference on Learning Representations</i> .	<i>municating parameter efficient updates via sparsifica-</i>	4371
		<i>tion and quantization</i> . <i>Preprint</i> , arXiv:2311.13171.	4372
4317	Yonghui Wu, Mike Schuster, Zhifeng Chen, Quoc V. Le,	Prateek Yadav, Colin Raffel, Mohammed Muqeeth,	4373
4318	Mohammad Norouzi, Wolfgang Macherey, Maxim	Lucas Caccia, Haokun Liu, Tianlong Chen, Mohit	4374
4319	Krikun, Yuan Cao, Qin Gao, Klaus Macherey, Jeff	Bansal, Leshem Choshen, and Alessandro Sordoni.	4375
4320	Klingner, Apurva Shah, Melvin Johnson, Xiaobing	2024. A survey on model moerging: Recycling and	4376
4321	Liu, Łukasz Kaiser, Stephan Gouw, Yoshikiyo Kato,	routing among specialized experts for collaborative	4377
4322	Taku Kudo, Hideto Kazawa, Keith Stevens, George	learning. <i>arXiv preprint arXiv:2408.07057</i> .	4378
4323	Kurian, Nishant Patil, Wei Wang, Cliff Young, Jason		
4324	Smith, Jason Riesa, Alex Rudnick, Oriol Vinyals,	Prateek Yadav, Derek Tam, Leshem Choshen, Colin	4379
4325	Greg Corrado, Macduff Hughes, and Jeffrey Dean.	Raffel, and Mohit Bansal. 2023b. <i>TIES-merging:</i>	4380
4326	2016. <i>Google’s neural machine translation system:</i>	<i>Resolving interference when merging models</i> . In	4381
4327	<i>Bridging the gap between human and machine trans-</i>	<i>Thirty-seventh Conference on Neural Information</i>	4382
4328	<i>lation</i> . <i>arXiv preprint</i> .	<i>Processing Systems</i> .	4383
4329	Kevin Xia, Kai-Zhan Lee, Yoshua Bengio, and Elias	Xinru Yan and Ted Pedersen. 2017. <i>Who’s to say what’s</i>	4384
4330	Bareinboim. 2021. <i>The causal-neural connection:</i>	<i>funny? A computer using language models and deep</i>	4385
4331	<i>Expressiveness, learnability, and inference</i> . <i>arXiv</i>	<i>learning, that’s who!</i> <i>arXiv preprint</i> .	4386
4332	<i>preprint</i> .		
4333	Yijun Xiao and William Yang Wang. 2021. <i>On hal-</i>	Diyi Yang, Alon Lavie, Chris Dyer, and Eduard Hovy.	4387
4334	<i>lucination and predictive uncertainty in conditional</i>	2015. <i>Humor recognition and humor anchor extrac-</i>	4388
4335	<i>language generation</i> . <i>arXiv preprint</i> .	<i>tion</i> . In <i>Proceedings of the 2015 Conference on Em-</i>	4389
4336	Jing Xu, Da Ju, Margaret Li, Y-Lan Boureau, Jason	<i>pirical Methods in Natural Language Processing</i> ,	4390
4337	Weston, and Emily Dinan. 2020a. <i>Recipes for safety</i>	pages 2367–2376, Lisbon, Portugal. Association for	4391
4338	<i>in open-domain chatbots</i> . <i>arXiv preprint</i> .	Computational Linguistics.	4392
4339	Jingwei Xu, Junyu Lai, and Yunpeng Huang. 2024. <i>Me-</i>	Enneng Yang, Zhenyi Wang, Li Shen, Shiwei Liu, Guib-	4393
4340	<i>teora: Multiple-tasks embedded lora for large lan-</i>	ing Guo, Xingwei Wang, and Dacheng Tao. 2023.	4394
4341	<i>guage models</i> . <i>arXiv preprint arXiv:2405.13053</i> .	Adamerging: Adaptive model merging for multi-task	4395
4342	Silei Xu, Sina Semnani, Giovanni Campagna, and Mon-	learning. <i>arXiv preprint arXiv:2310.02575</i> .	4396
4343	ica Lam. 2020b. <i>AutoQA: From databases to QA</i>	Kaiyu Yang and Jia Deng. 2019. <i>Learning to prove</i>	4397
4344	<i>semantic parsers with only synthetic training data</i> . In	<i>theorems via interacting with proof assistants</i> . <i>arXiv</i>	4398
4345	<i>Proceedings of the 2020 Conference on Empirical</i>	<i>preprint</i> .	4399
4346	<i>Methods in Natural Language Processing (EMNLP)</i> ,	Scott Cheng-Hsin Yang and Patrick Shafto. 2017. <i>Ex-</i>	4400
4347	pages 422–434, Online. Association for Computa-	<i>plainable artificial intelligence via Bayesian teaching.</i>	4401
4348	tional Linguistics.	<i>Workshop on Teaching Machines, Robots, and Hu-</i>	4402
4349	Yang Xu, Jiawei Liu, Wei Yang, and Liusheng Huang.	<i>mans, NIPS 2017</i> .	4403
4350	2018. <i>Incorporating latent meanings of morpholog-</i>	Zhilin Yang, Peng Qi, Saizheng Zhang, Yoshua Ben-	4404
4351	<i>ical compositions to enhance word embeddings</i> . In	gio, William W. Cohen, Ruslan Salakhutdinov, and	4405
4352	<i>Proceedings of the 56th Annual Meeting of the As-</i>	Christopher D. Manning. 2018. <i>Hotpotqa: A dataset</i>	4406
4353	<i>sociation for Computational Linguistics (Volume 1:</i>	<i>for diverse, explainable multi-hop question answer-</i>	4407
4354	<i>Long Papers)</i> , pages 1232–1242, Melbourne, Aus-	<i>ing</i> . In <i>Proceedings of the 2018 Conference on Em-</i>	4408
4355	tralia. Association for Computational Linguistics.	<i>pirical Methods in Natural Language Processing</i> .	4409
4356	Linting Xue, Aditya Barua, Noah Constant, Rami Al-	Qinyuan Ye, Juan Zha, and Xiang Ren. 2022. <i>Elic-</i>	4410
4357	Rfou, Sharan Narang, Mihir Kale, Adam Roberts,	<i>iting and understanding cross-task skills with</i>	4411
4358	and Colin Raffel. 2021a. <i>ByT5: Towards a token-free</i>	<i>task-level mixture-of-experts</i> . <i>arXiv preprint</i>	4412
4359	<i>future with pre-trained byte-to-byte models</i> . <i>arXiv</i>	<i>arXiv:2205.12701</i> .	4413
4360	<i>preprint</i> .		

4414	Eric Yeh, Daniel Ramage, Christopher D. Manning,	<i>Workshop BlackboxNLP: Analyzing and Interpreting</i>	4471
4415	Eneko Agirre, and Aitor Soroa. 2009. WikiWalk:	<i>Neural Networks for NLP</i> , pages 138–146, Florence,	4472
4416	Random walks on Wikipedia for semantic related-	Italy. Association for Computational Linguistics.	4473
4417	ness . In <i>Proceedings of the 2009 Workshop on Graph-</i>		
4418	<i>based Methods for Natural Language Processing</i>	Zhou Yu, Dejing Xu, Jun Yu, Ting Yu, Zhou Zhao, Yuet-	4474
4419	(<i>TextGraphs-4</i>), pages 41–49, Suntec, Singapore. As-	ing Zhuang, and Dacheng Tao. 2019d. ActivityNet-	4475
4420	sociation for Computational Linguistics.	QA: A dataset for understanding complex web videos	4476
		via question answering . <i>arXiv preprint</i> .	4477
4421	Yelp, Inc. 2018. Yelp open dataset .		
4422	Yang Yi, Yih Wen-tau, and Christopher Meek. 2015.	Eliezer Yudkowsky. 2008. Artificial intelligence as a	4478
4423	WikiQA: A Challenge Dataset for Open-Domain	positive and negative factor in global risk . In Nick	4479
4424	Question Answering . Association for Computational	Bostrom and Milan M. Ćirković, editors, <i>Global</i>	4480
4425	<i>Linguistics</i> , page 2013–2018.	<i>Catastrophic Risks</i> , pages 308–345. Oxford Univer-	4481
		sity Press, Oxford.	4482
4426	Wenpeng Yin and Yulong Pei. 2015. Optimizing sen-	Ted Zadori, Ahmet Üstün, Arash Ahmadian, Beyza Er-	4483
4427	tence modeling and selection for document summa-	miş, Acyr Locatelli, and Sara Hooker. 2023. Pushing	4484
4428	rization . In <i>Proceedings of the 24th International</i>	mixture of experts to the limit: Extremely parameter	4485
4429	<i>Conference on Artificial Intelligence, IJCAI-15</i> , page	efficient moe for instruction tuning. <i>arXiv preprint</i>	4486
4430	1383–1389.	<i>arXiv:2309.05444</i> .	4487
4431	Tao Yu, Rui Zhang, Heyang Er, Suyi Li, Eric Xue,	Amir Zamir, Alexander Sax, Bokui (William) Shen,	4488
4432	Bo Pang, Xi Victoria Lin, Yi Chern Tan, Tianze	Leonidas J. Guibas, Jitendra Malik, and Silvio	4489
4433	Shi, Zihan Li, Youxuan Jiang, Michihiro Yasunaga,	Savarese. 2018. Taskonomy: Disentangling task	4490
4434	Sungrok Shim, Tao Chen, Alexander Fabbri, Zifan	transfer learning . <i>2018 IEEE/CVF Conference on</i>	4491
4435	Li, Luyao Chen, Yuwen Zhang, Shreya Dixit, Vin-	<i>Computer Vision and Pattern Recognition</i> , pages	4492
4436	cent Zhang, Caiming Xiong, Richard Socher, Walter	3712–3722.	4493
4437	Lasecki, and Dragomir Radev. 2019a. CoSQL: A	Wojciech Zaremba and Ilya Sutskever. 2014. Learning	4494
4438	conversational text-to-SQL challenge towards cross-	to execute . <i>arXiv preprint</i> .	4495
4439	domain natural language interfaces to databases . In		
4440	<i>Proceedings of the 2019 Conference on Empirical</i>	Poorya Zaremoondi, Wray L. Buntine, and Gholamreza	4496
4441	<i>Methods in Natural Language Processing and the</i>	Haffari. 2018. Adaptive knowledge sharing in multi-	4497
4442	<i>9th International Joint Conference on Natural Lan-</i>	task learning: Improving low-resource neural ma-	4498
4443	<i>guage Processing (EMNLP-IJCNLP)</i> , pages 1962–	chine translation . In <i>Annual Meeting of the Associa-</i>	4499
4444	1979, Hong Kong, China. Association for Computa-	<i>tion for Computational Linguistics</i> .	4500
4445	tional Linguistics.		
4446	Tao Yu, Rui Zhang, Kai Yang, Michihiro Yasunaga,	Omnia Zayed, John Philip McCrae, and Paul Buite-	4501
4447	Dongxu Wang, Zifan Li, James Ma, Irene Li, Qingn-	laar. 2020. Figure me out: A gold standard dataset	4502
4448	ing Yao, Shanelle Roman, Zilin Zhang, and Dragomir	for metaphor interpretation . In <i>Proceedings of the</i>	4503
4449	Radev. 2018. Spider: A large-scale human-labeled	<i>12th Language Resources and Evaluation Confer-</i>	4504
4450	dataset for complex and cross-domain semantic pars-	<i>ence</i> , pages 5810–5819, Marseille, France. European	4505
4451	ing and text-to-SQL task . In <i>Proceedings of the 2018</i>	Language Resources Association.	4506
4452	<i>Conference on Empirical Methods in Natural Lan-</i>		
4453	<i>guage Processing</i> , pages 3911–3921, Brussels, Bel-	Rowan Zellers, Yonatan Bisk, Ali Farhadi, and Yejin	4507
4454	gium. Association for Computational Linguistics.	Choi. 2018. From recognition to cognition: Visual	4508
		commonsense reasoning . <i>arXiv preprint</i> .	4509
4455	Tao Yu, Rui Zhang, Michihiro Yasunaga, Yi Chern	Rowan Zellers, Ari Holtzman, Yonatan Bisk, Ali	4510
4456	Tan, Xi Victoria Lin, Suyi Li, Heyang Er, Irene	Farhadi, and Yejin Choi. 2019a. HellaSwag: Can a	4511
4457	Li, Bo Pang, Tao Chen, Emily Ji, Shreya Dixit,	machine really finish your sentence? <i>arXiv preprint</i>	4512
4458	David Proctor, Sungrok Shim, Jonathan Kraft, Vin-	<i>arXiv:1905.07830</i> .	4513
4459	cent Zhang, Caiming Xiong, Richard Socher, and		
4460	Dragomir Radev. 2019b. SPaRC: Cross-domain se-	Rowan Zellers, Ari Holtzman, Hannah Rashkin,	4514
4461	mantic parsing in context . In <i>Proceedings of the</i>	Yonatan Bisk, Ali Farhadi, Franziska Roesner, and	4515
4462	<i>57th Annual Meeting of the Association for Computa-</i>	Yejin Choi. 2019b. Defending against neural fake	4516
4463	<i>tional Linguistics</i> , pages 4511–4523, Florence, Italy.	news . <i>arXiv preprint</i> .	4517
4464	Association for Computational Linguistics.		
4465	Weihaoyu, Zihang Jiang, Yanfei Dong, and Jiashi Feng.	Zihao Zeng, Yibo Miao, Hongcheng Gao, Hao Zhang,	4518
4466	2020. ReClor: A reading comprehension dataset re-	and Zhijie Deng. 2024. Adamoe: Token-adaptive	4519
4467	quiring logical reasoning . <i>CoRR</i> , arXiv:2002.04326.	routing with null experts for mixture-of-experts lan-	4520
		guage models . <i>Preprint</i> , arXiv:2406.13233.	4521
4468	Xiang Yu, Ngoc Thang Vu, and Jonas Kuhn. 2019c.	Dongxiang Zhang, Lei Wang, Luming Zhang, Bing Tian	4522
4469	Learning the Dyck language with attention-based	Dai, and Heng Tao Shen. 2020a. The gap of semantic	4523
4470	Seq2Seq models . In <i>Proceedings of the 2019 ACL</i>	parsing: A survey on automatic math word problem	4524

4525 [solvers](#). *IEEE Transactions on Pattern Analysis & Machine Intelligence*, 42(09):2287–2305. 4580

4526 4581

4527 Haoran Zhang, Amy X. Lu, Mohamed Abdalla, 4582

4528 Matthew McDermott, and Marzyeh Ghassemi. 2020b. 4583

4529 [Hurtful words: Quantifying biases in clinical con-](#)

4530 [textual word embeddings](#). In *Proceedings of the*

4531 *ACM Conference on Health, Inference, and Learn-*

4532 *ing, CHIL '20*, page 110–120, New York, NY, USA. 4584

4533 Association for Computing Machinery. 4585

4534 Hongming Zhang, Xinran Zhao, and Yangqiu Song. 4586

4535 2020c. [WinoWhy: A deep diagnosis of essential](#)

4536 [commonsense knowledge for answering Winograd](#)

4537 [schema challenge](#). In *Proceedings of the 58th An-*

4538 *ual Meeting of the Association for Computational*

4539 *Linguistics*, pages 5736–5745, Online. Association 4587

4540 for Computational Linguistics. 4588

4541 Kaiqing Zhang, Zhuoran Yang, and Tamer Başar. 2019a. 4589

4542 [Multi-agent reinforcement learning: A selective](#)

4543 [overview of theories and algorithms](#). *arXiv preprint*. 4590

4544 Li Zhang, Qing Lyu, and Chris Callison-Burch. 2020d. 4591

4545 [Reasoning about goals, steps, and temporal ordering](#)

4546 [with WikiHow](#). In *Proceedings of the 2020 Con-*

4547 *ference on Empirical Methods in Natural Language*

4548 *Processing (EMNLP)*, pages 4630–4639, Online. As- 4592

4549 sociation for Computational Linguistics. 4593

4550 Meishan Zhang, Yue Zhang, and Guohong Fu. 2016. 4594

4551 [Tweet sarcasm detection using deep neural network](#). 4595

4552 In *Proceedings of COLING 2016, the 26th Inter-*

4553 *national Conference on Computational Linguistics:*

4554 *Technical Papers*, pages 2449–2460, Osaka, Japan. 4596

4555 The COLING 2016 Organizing Committee. 4597

4556 Sheng Zhang, Xiaodong Liu, Jingjing Liu, Jianfeng 4598

4557 Gao, Kevin Duh, and Benjamin Van Durme. 2018a. 4599

4558 Record: Bridging the gap between human and ma- 4600

4559 chine commonsense reading comprehension. *arXiv*

4560 *preprint arXiv:1810.12885*. 4601

4561 Shiwei Zhang, Xiuzhen Zhang, Jeffrey Chan, and Paolo 4602

4562 Rosso. 2019b. [Irony detection via sentiment-based](#)

4563 [transfer learning](#). *Information Processing & Manage-*

4564 *ment*, 56(5):1633–1644. 4603

4565 Xiang Zhang, Junbo Zhao, and Yann LeCun. 2015. 4604

4566 [Character-level convolutional networks for text clas-](#)

4567 [sification](#). In *Advances in Neural Information Pro-*

4568 *cessing Systems*. 4605

4569 Yan Zhang, Jonathon Hare, and Adam Prügel-Bennett. 4606

4570 2018b. [Learning to count objects in natural images](#)

4571 [for visual question answering](#). *arXiv preprint*. 4607

4572 Yuan Zhang, Jason Baldridge, and Luheng He. 2019c. 4608

4573 PAWS: Paraphrase Adversaries from Word Scram- 4609

4574 bling. In *Proc. of NAACL*. 4610

4575 Jieyu Zhao, Tianlu Wang, Mark Yatskar, Vicente Or- 4611

4576 donez, and Kai-Wei Chang. 2018. [Gender bias in](#)

4577 [coreference resolution: Evaluation and debiasing](#)

4578 [methods](#). In *Proceedings of the 2018 Conference*

4579 *of the North American Chapter of the Association for*

Computational Linguistics: Human Language Tech-

nologies, Volume 2 (Short Papers), pages 15–20, New

Orleans, Louisiana. Association for Computational

Linguistics.

Tony Z. Zhao, Eric Wallace, Shi Feng, Dan Klein, and

Sameer Singh. 2021. [Calibrate before use: Improv-](#)

[ing few-shot performance of language models](#). *arXiv*

preprint.

Xinyu Zhao, Guoheng Sun, Ruisi Cai, Yukun Zhou,

Pingzhi Li, Peihao Wang, Bowen Tan, Yexiao He,

Li Chen, Yi Liang, et al. 2024a. Model-glue: Democ-

ratized llm scaling for a large model zoo in the wild.

arXiv preprint arXiv:2410.05357.

Ziyu Zhao, Leilei Gan, Guoyin Wang, Wangchunshu

Zhou, Hongxia Yang, Kun Kuang, and Fei Wu.

2024b. [Loraretriever: Input-aware lora retrieval and](#)

[composition for mixed tasks in the wild](#). *Preprint*,

arXiv:2402.09997.

Ben Zhou, Daniel Khashabi, Qiang Ning, and Dan Roth.

2019. ["Going on a vacation" takes longer than "go-](#)

[ing for a walk": A study of temporal commonsense](#)

[understanding](#). *arXiv preprint*.

Chunting Zhou, Graham Neubig, Jiatao Gu, Mona

Diab, Paco Guzman, Luke Zettlemoyer, and Marjan

Ghazvininejad. 2020. [Detecting hallucinated con-](#)

[tent in conditional neural sequence generation](#). *arXiv*

preprint.

Jing Zhou, Zongyu Lin, Yanan Zheng, Jian Li, and

Zhilin Yang. 2022. Not all tasks are born equal:

Understanding zero-shot generalization. In *The*

Eleventh International Conference on Learning Rep-

resentations.

Haichao Zhu, Li Dong, Furu Wei, Wenhui Wang, Bing

Qin, and Ting Liu. 2019. [Learning to ask unanswer-](#)

[able questions for machine reading comprehension](#).

In *Proceedings of the 57th Annual Meeting of the As-*

sociation for Computational Linguistics, pages 4238–

4248, Florence, Italy. Association for Computational

Linguistics.

Xiaojin Zhu. 2015. [Machine teaching: An inverse prob-](#)

[lem to machine learning and an approach toward](#)

[optimal education](#). In *Proceedings of the AAAI Con-*

ference on Artificial Intelligence, volume 29, Menlo

Park, CA. Association for the Advancement of Arti-

ficial Intelligence.

Zhemina Zhu, Delphine Bernhard, and Iryna Gurevych.

2010. [Sentence simplification by monolingual ma-](#)

[chine translation](#). In *Proceedings of the 5th Interna-*

tional Conference on Natural Language Processing.

Alan Zucconi. 6 Jan. 2016. [The secrets of colour inter-](#)

[polation](#).

Appendix

A Extended Related Work

Model Merging. Model merging (Yadav et al., 2023b; Choshen et al., 2022; Wortsman et al., 2022; Ramé et al., 2022; Matena and Raffel, 2022; Ilharco et al., 2022; Tam et al., 2023; Jin et al., 2022; Yang et al., 2023; Zhao et al., 2024a) consolidates multiple independently trained models with identical architectures into a unified model that preserves multi-model capabilities. While simple parameter averaging suffices for models within a linearly connected low-loss parameter space (McMahan et al., 2017; Stich, 2018; Frankle et al., 2020; Wortsman et al., 2021; Li et al., 2023), more sophisticated techniques are necessary for complex scenarios. For instance, task vectors facilitate merging expert models trained on diverse domains (Ilharco et al., 2022). Additionally, methods like weighted merging using Fisher Importance Matrices (Matena and Raffel, 2022; Tam et al., 2023) and TIES-Merging, which addresses sign disagreements and redundancy (Yadav et al., 2023b) offers improved performance. As a non-adaptive expert aggregation method, merging serves as a fundamental baseline for numerous Model Editing with Regularization (MoErging) techniques.

Multitask Learning (MTL) research offers valuable insights for decentralized development. Notably, investigations into task-relatedness (Standley et al., 2020; Bingel and Søgaard, 2017; Achille et al., 2019; Vu et al., 2020; Zamir et al., 2018; Mou et al., 2016) provide guidance for designing routing mechanisms, while MTL architectures addressing the balance between shared and task-specific knowledge (Misra et al., 2016; Ruder et al., 2017; Meyerson and Miikkulainen, 2017; Zareemoodi et al., 2018; Sun et al., 2019b) offer strategies for combining expert contributions in a decentralized manner.

MoE for Multitask Learning. Recent research has extensively investigated mixture-of-experts (MoE) models for multitask learning, achieving promising results in unseen task generalization. These approaches generally fall into two categories: (1) Example Routing: Studies like Muqeeth et al. (2023); Zadouri et al. (2023); Wang et al. (2022a) train routers to dynamically select experts for each input, while Caccia et al. (2023) demonstrate the efficacy of routing at a finer granularity by splitting expert parameters into blocks. (2) Task Routing:

Ponti et al. (2023) employs a trainable skill matrix to assign tasks to specific parameter-efficient modules, while Gupta et al. (2022) leverages task-specific routers selected based on domain knowledge. Ye et al. (2022) proposes a layer-wise expert selection mechanism informed by task representations derived from input embeddings. Such approaches leverage task-specific representation to allow the router to effectively select the most suitable experts for unseen tasks. While these studies differ from our setting by assuming simultaneous data access, they offer valuable insights applicable to our exploration of creating routing mechanisms over expert models.

B LLM for Task Instruction Generation.

B.1 Prompt Template

We use the following prompt with 3 randomly selected samples for each task to generate its description. The prompt is then fed into the gpt-4-turbo OpenAI API to get the generated task descriptions.

The following are three pairs of input-output examples from one task. Generate the task instruction in one sentence that is most possibly used to command a language model to produce them. In the instruction, remember to point out the skill or knowledge required for the task to guide the language model.

- Input:
- Output:

- Input:
- Output:

- Input:
- Output:

B.2 Examples of the Generated Instructions

We provide several examples of LLM-generated instructions in this section.

WikiBio (Lebret et al., 2016a) (T0 Held-In):

- *Create a short biography using the provided facts, demonstrating knowledge in historical and biographical writing.*
- *Write a short biography based on the given factual bullet points, demonstrating profi-*

4711	ciency in summarizing and transforming struc-		
4712	tured data into coherent narrative text.		
4713	CommonGen (Lin et al., 2020b) (T0 Held-In):		
4714	• Generate a coherent sentence using all the		
4715	given abstract concepts, requiring the skill		
4716	of concept integration to form a meaningful		
4717	sentence.		
4718	• Generate a coherent sentence by creatively		
4719	combining a given set of abstract concepts.		
4720	COPA (Huang et al., 2024b) (T0 Held-Out):		
4721	• Identify the most logically consistent sentence		
4722	from two given options based on the provided		
4723	context, demonstrating reasoning and causal		
4724	relationship skills.		
4725	• Generate the most likely outcome for a given		
4726	scenario by choosing between two provided		
4727	options based on contextual clues and causal		
4728	reasoning.		
4729	Date Understanding (Srivastava et al., 2023)		
4730	(BigBench-Hard):		
4731	• Calculate the date based on the given informa-		
4732	tion and present it in MM/DD/YYYY format,		
4733	ensuring that you accurately account for day,		
4734	month, and year changes.		
4735	Hindu Mythology Trivia (Srivastava et al.,		
4736	2023) (BigBench-Lite):		
4737	• Generate the correct answer by making use		
4738	of your knowledge in Hindu mythology and		
4739	culture.		
4740	C Demonstrating Compositional		
4741	Generation		
4742	In addition to significant improvements on held-in		
4743	tasks, GLIDER demonstrates strong performance		
4744	on held-out tasks, showcasing its generalization		
4745	capability. To further examine this ability to handle		
4746	unseen tasks by composing experts, we provide		
4747	specific task examples illustrating the association		
4748	between selected experts and the evaluated task.		
4749	As Figure 2 shows, GLIDER primarily selects two		
4750	experts for the COPA (T0 held-out) task, corre-		
4751	sponding to CosmosQA and QuaRel. The follow-		
4752	ing three examples from these tasks demonstrate		
4753	their close semantic relationship:		
4754	• COPA:		
	– Question: Everyone in the class turned	4755	
	to stare at the student. Select the most	4756	
	plausible cause: - The student’s phone	4757	
	rang. - The student took notes.	4758	
	– Answer: The student’s phone rang.	4759	
	• CosmosQA:	4760	
	– Question: That idea still weirds me out .	4761	
	I made a blanket for the baby ’s older sis-	4762	
	ter before she was born but I completely	4763	
	spaced that this one was on the way ,	4764	
	caught up in my own dramas and what-	4765	
	not . Luckily , I had started a few rows	4766	
	in white just to learn a stitch ages ago ,	4767	
	and continuing that stitch will make an	4768	
	acceptable woobie , I think . Accord-	4769	
	ing to the above context, choose the best	4770	
	option to answer the following question.	4771	
	Question: What did I make for the baby .	4772	
	Options: A. I made a carseat . B. None	4773	
	of the above choices . C. I made a crb .	4774	
	D. I finished a pair of booties .	4775	
	– Answer: D.	4776	
	• QuaRel:	4777	
	– Question: Here’s a short story: A piece	4778	
	of thread is much thinner than a tree	4779	
	so it is (A) less strong (B) more strong.	4780	
	What is the most sensical answer be-	4781	
	tween "Thread" and "Tree"?	4782	
	– Answer: Thread.	4783	
	D Datasets and Metric	4784	
	The specific details of all the datasets we use in this	4785	
	work are provided in this section.	4786	
	D.1 T0 Held-In Datasets	4787	
	• CommonsenseQA (Talmor et al., 2019b) un-	4788	
	der MIT License, evaluated by accuracy.	4789	
	• DREAM (Sun et al., 2019a) under MIT Li-	4790	
	cense, evaluated by accuracy.	4791	
	• QUAIL (Rogers et al., 2020) under CC BY-SA	4792	
	4.0, evaluated by accuracy.	4793	
	• QuaRTz (Tafjord et al., 2019) under Apache	4794	
	2.0 License, evaluated by accuracy.	4795	
	• Social IQA (Sap et al., 2019) under MIT Li-	4796	
	cense, evaluated by accuracy.	4797	

4798	• WiQA (Tandon et al., 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy.	4838
4799		4839
4800	• Cosmos QA (Huang et al., 2019) under <i>MIT License</i> , evaluated by accuracy.	4840
4801		
4802	• QASC (Khot et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy.	4841
4803		4842
4804	• Quarel (Tafjord et al., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy.	4843
4805		4844
4806	• SciQ (Johannes Welbl, 2017) under <i>MIT License</i> , evaluated by accuracy.	4845
4807		4846
4808	• Wiki Hop (Welbl et al., 2018) under <i>CC BY-SA 3.0</i> , evaluated by accuracy.	4847
4809		4848
4810	• Adversarial QA (Bartolo et al., 2020) under <i>Apache 2.0 License</i> , evaluated by F1 score.	4849
4811		4850
4812	• Quoref (Dasigi et al., 2019) under <i>Apache 2.0 License</i> , evaluated by F1 score.	4851
4813		4852
4814	• DuoRC (Saha et al., 2018) under <i>MIT License</i> , evaluated by F1 score.	4853
4815		4854
4816	• ROPES (Lin et al., 2019b) under <i>Apache 2.0 License</i> , evaluated by F1 score.	4855
4817		4856
4818	• Hotpot QA (Yang et al., 2018) under <i>CC BY-SA 4.0</i> , evaluated by exact match and F1 score.	4857
4819		
4820	• Wiki QA (Yi et al., 2015) under <i>MIT License</i> , evaluated by mean average precision (MAP) and mean reciprocal rank (MRR).	4858
4821		4859
4822		
4823	• Common Gen (Lin et al., 2020c) under <i>MIT License</i> , evaluated by BLEU and ROUGE scores.	4860
4824		4861
4825		4862
4826	• Wiki Bio (Lebret et al., 2016b) under <i>CC BY-SA 3.0</i> , evaluated by BLEU score.	4863
4827		
4828	• Amazon (Blitzer et al., 2007) under <i>Proprietary License</i> , evaluated by accuracy.	4864
4829		
4830	• App Reviews (Maas et al., 2011a) under <i>Proprietary License</i> , evaluated by accuracy.	4865
4831		4866
4832	• IMDB (Maas et al., 2011b) under <i>Proprietary License</i> , evaluated by accuracy.	4867
4833		4868
4834	• Rotten Tomatoes (Zhu et al., 2010) under <i>Proprietary License</i> , evaluated by accuracy.	4869
4835		4870
4836	• Yelp (Yelp, Inc., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy.	4871
4837		4872
		4873
		4874
		4875
		4876

D.2 T0 Held-Out Datasets

Held-out Tasks

• ANLI (Nie et al., 2020) under <i>MIT License</i> , evaluated by accuracy.	4865
• CB (de Marneffe et al., 2017) under <i>CC-BY-SA License</i> , evaluated by accuracy.	4866
• RTE (Dagan et al., 2005; Bar-Haim et al., 2006; Giampiccolo et al., 2007; Bentivogli et al., 2009) under <i>Apache 2.0 License</i> , evaluated by accuracy.	4867
• WSC (Levesque et al., 2012) under <i>Creative Commons License</i> , evaluated by accuracy.	4868
• Winogrande (Sakaguchi et al., 2020a) under <i>Apache License 2.0</i> , evaluated by accuracy.	4869
	4870
	4871
	4872
	4873
	4874
	4875
	4876

- **WiC** (Pilehvar and Camacho-Collados, 2019) under *CC BY-SA 4.0 License*, evaluated by accuracy.
- **COPA** (Roemmele et al., 2011a) under *BSD-2-Clause License*, evaluated by accuracy.
- **HellaSwag** (Zellers et al., 2019a) under *MIT License*, evaluated by accuracy.
- **Story Cloze** (Mostafazadeh et al., 2016b) under *CC-BY 4.0 License*, evaluated by accuracy.

D.3 BigBench-Hard Datasets

- **abstract_narrative_understanding** (Ghosh and Srivastava, 2021; Holyoak, 2012; Nippold et al., 2001; Tan et al., 2015; Wang et al., 2020b; Mostafazadeh et al., 2016a) under *Apache 2.0 License*, evaluated by accuracy
- **abstraction_and_reasoning_corpus** (Chollet, 2019; Brown et al., 2020; Chollet, 2020) under *Apache 2.0 License*, evaluated by accuracy
- **anachronisms** (Otterbacher et al., 2002; Popescu and Strapparava, 2015; Llorens et al., 2015; Meng et al., 2017; Geva et al., 2021) under *Apache 2.0 License*, evaluated by accuracy
- **analogical_similarity** (Plate, 2003, 1994; Thagard et al., 1990; Gentner et al., 1993) under *Apache 2.0 License*, evaluated by accuracy
- **analytic_entailment** (Hume, 1739–1740; Kant, 1781/1787; Wittgenstein, 1953; Quine, 1951; Grice and Strawson, 1956; Bolukbasi et al., 2016; Kocurek et al., 2020; Rudolph and Kocurek, 2020; Kocurek and Jerzak, 2021) under *Apache 2.0 License*, evaluated by accuracy
- **arithmetic** (Brown et al., 2020; Saxton et al., 2019) under *Apache 2.0 License*, evaluated by accuracy
- **ascii_word_recognition** (Child et al., 2019; Chen et al., 2020) under *Apache 2.0 License*, evaluated by accuracy
- **authorship_verification** (Bischoff et al., 2020; Koppel and Schler, 2004) under *Apache 2.0 License*, evaluated by accuracy

- **auto_categorization** under *Apache 2.0 License*, evaluated by accuracy
- **bbq_lite** (Crawford, 2017; Khashabi et al., 2020b; Li et al., 2020a) under *Apache 2.0 License*, evaluated by accuracy
- **bias_from_probabilities** (Bender et al., 2021; Abid et al., 2021) under *Apache 2.0 License*, evaluated by accuracy
- **boolean_expressions** (Habernal et al., 2018; Yu et al., 2020; Dua et al., 2019; Liu et al., 2020a; Sinha et al., 2019; Wang et al., 2019b; Steinbach and Kohut, 2002; Saxton et al., 2019; Payani and Fekri, 2019; Trask et al., 2018; Selsam et al., 2018; Allamanis et al., 2016; Evans et al., 2018; Shi et al., 2020) under *Apache 2.0 License*, evaluated by accuracy
- **bridging_anaphora_resolution_barqa** (Hou, 2020; Hou et al., 2013; Markert et al., 2012; Rajpurkar et al., 2016) under *Apache 2.0 License*, evaluated by accuracy
- **causal_judgment** (Gordon, 2010; Bosselut et al., 2019; Halpern, 2016; Knobe, 2003) under *Apache 2.0 License*, evaluated by accuracy
- **cause_and_effect** (Gordon, 2010) under *Apache 2.0 License*, evaluated by accuracy
- **checkmate_in_one** (Alexander, 2020; Ammanabrolu et al., 2019; Dambekodi et al., 2020; Ammanabrolu et al., 2020) under *Apache 2.0 License*, evaluated by accuracy
- **chess_state_tracking** (Weston et al., 2015; Côté et al., 2018; Toshniwal et al., 2021; Alexander, 2020; Chen, 2020; Noever et al., 2020; Swingle et al., 2021) under *Apache 2.0 License*, evaluated by accuracy
- **chinese_remainder_theorem** under *Apache 2.0 License*, evaluated by accuracy
- **cifar10_classification** under *Apache 2.0 License*, evaluated by accuracy
- **codenames** (Kim et al., 2019) under *Apache 2.0 License*, evaluated by accuracy
- **color** (Gibson et al., 2017; Zucconi, 6 Jan. 2016) under *Apache 2.0 License*, evaluated by accuracy

4963	• com2sense (Singh et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	
4964		
4965	• common_morpheme (Devlin et al., 2018; Wu et al., 2016; Won et al., 2021; Xu et al., 2018; Edmiston and Stratos, 2018; El-Kishky et al., 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	
4966		
4967		
4968		
4969		
4970	• context_definition_alignment (Senel and Schütze, 2021; Reimers and Gurevych, 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	
4971		
4972		
4973		
4974	• convinceme (Lin et al., 2021b; Levy et al., 2021; Clark et al., 2018; Maynez et al., 2020; Wang et al., 2020a; Kenton et al., 2021; Xu et al., 2020a; Tamkin et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	
4975		
4976		
4977		
4978		
4979	• coqa_conversational_question_answering (Reddy et al., 2019; Radford et al., 2019; Brown et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	
4980		
4981		
4982		
4983	• crash_blossom under <i>Apache 2.0 License</i> , evaluated by accuracy	
4984		
4985	• crass_ai (Schölkopf et al., 2021; Teney et al., 2020; Liang et al., 2020b; Pearl, 2000; Shahid and Zheleva, 2021; Xia et al., 2021; Priol et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	
4986		
4987		
4988		
4989		
4990	• cryobiology_spanish (Scudellari, 2017; Soltani Firouz et al., 2021; Toubiana et al., 2020; Mbogba et al., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	
4991		
4992		
4993		
4994	• cryptonite (Efrat et al., 2021; Raganato et al., 2017; Sakaguchi et al., 2020b; Miller and Gurevych, 2015; Miller et al., 2017; Joshi et al., 2017; Oprea and Magdy, 2020; Friedlander and Fine, 2018; Lewis et al., 2020c) under <i>Apache 2.0 License</i> , evaluated by accuracy	
4995		
4996		
4997		
4998		
4999		
5000	• cs_algorithms under <i>Apache 2.0 License</i> , evaluated by accuracy	
5001		
5002	• cycled_letters (Brown et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	
5003		
5004	• dark_humor_detection (Weller and Seppi, 2019; Fan et al., 2020; Willinger et al., 2017; Yang et al., 2015; Mihalcea and Strapparava, 2005) under <i>Apache 2.0 License</i> , evaluated by accuracy	
5005		
5006		
5007		
5008		
	• date_understanding (Vashishth et al., 2018; Chambers, 2012; Kotsakos et al., 2014; Vashishtha et al., 2020, 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	5009 5010 5011 5012
	• disambiguation_qa (Zhao et al., 2018; Rudinger et al., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	5013 5014 5015
	• discourse_marker_prediction (Malmi et al., 2018; Nie et al., 2019; Sileo et al., 2019, 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5016 5017 5018 5019
	• disfl_qa (Gupta et al., 2021; Rajpurkar et al., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	5020 5021 5022
	• diverse_social_bias (Sheng et al., 2019; Nadeem et al., 2020; Hendrycks et al., 2020; Sap et al., 2020; Gehman et al., 2020; Bolukbasi et al., 2016; Caliskan et al., 2017; May et al., 2019; Liang et al., 2020a; Barocas and Selbst, 2016; Cho et al., 2019; Blodgett et al., 2020; Merity et al., 2016; Socher et al., 2013; Poria et al., 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	5023 5024 5025 5026 5027 5028 5029 5030 5031
	• dyck_languages (Chomsky and Schützenberger, 1959; Suzgun et al., 2019b; Hao et al., 2018; Hewitt et al., 2020; Hahn, 2020; Suzgun et al., 2019a; Sennhauser and Berwick, 2018; Skachkova et al., 2018; Bhattamishra et al., 2020a; Yu et al., 2019c; Ebrahimi et al., 2020; Ackerman and Cybenko, 2020; Bhattamishra et al., 2020b) under <i>Apache 2.0 License</i> , evaluated by accuracy	5032 5033 5034 5035 5036 5037 5038 5039 5040
	• dynamic_counting (Suzgun et al., 2019a; Skachkova et al., 2018; Bhattamishra et al., 2020a; Suzgun et al., 2019b; Yu et al., 2019c; Ebrahimi et al., 2020; Ackerman and Cybenko, 2020; Bhattamishra et al., 2020b; Sennhauser and Berwick, 2018; Merrill, 2020; Karpathy, 2015) under <i>Apache 2.0 License</i> , evaluated by accuracy	5041 5042 5043 5044 5045 5046 5047 5048
	• elementary_math_qa (Amini et al., 2019; Ling et al., 2017; Hendrycks et al., 2021c; Patel et al., 2021; Zhang et al., 2020a; Hendrycks et al., 2021b) under <i>Apache 2.0 License</i> , evaluated by accuracy	5049 5050 5051 5052 5053

5054	• emojis_emotion_prediction (Shoeb and de Melo, 2020; Plutchik, 1980) under <i>Apache 2.0 License</i> , evaluated by accuracy	Lourie et al., 2021; Liu et al., 2020a; Saxton et al., 2019; Clark et al., 2018; Yu et al., 2019d; Zhang et al., 2018a; Johnson et al., 2016) under <i>Apache 2.0 License</i> , evaluated by accuracy	5099
5055			5100
5056			5101
5057	• empirical_judgments (Kant, 1781/1787, 1783; Spirtes et al., 2000; Pearl, 1988; Goldberger, 1972; Rothman and Greenland, 2005; Roemmele et al., 2011b; Wang et al., 2019b; Evans et al., 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy		5102
5058			5103
5059		• few_shot_nlg (Rastogi et al., 2020; Kale and Rastogi, 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5104
5060			5105
5061			5106
5062		• figure_of_speech_detection (Potamias et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5107
5063	• english_proverbs (Gyasi Obeng, 1996; Honeck, 1997; Hrisztova-Gotthardt and Aleksa Varga, 2015) under <i>Apache 2.0 License</i> , evaluated by accuracy		5108
5064			5109
5065		• forecasting_subquestions under <i>Apache 2.0 License</i> , evaluated by accuracy	5110
5066			5111
5067	• english_russian_proverbs (Bodrova, 2007; Gvarjalaze and Mchedlishvili, 1971; Wik) under <i>Apache 2.0 License</i> , evaluated by accuracy	• gem (Gehrmann et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	5112
5068			5113
5069		• gender_inclusive_sentences_german under <i>Apache 2.0 License</i> , evaluated by accuracy	5114
5070	• entailed_polarity (Karttunen, 2012) under <i>Apache 2.0 License</i> , evaluated by accuracy		5115
5071		• gender_sensitivity_chinese (on the Revision of the National Standard Occupational Classification, 2015; Household Management Research Center, 2021, 2020, 2019; Qimingtong, 2016; Ministry of the Interior, 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	5116
5072	• entailed_polarity_hindi (Karttunen, 2012) under <i>Apache 2.0 License</i> , evaluated by accuracy		5117
5073			5118
5074			5119
5075	• epistemic_reasoning (Ravenscroft, 2019; Call and Tomasello, 2008; Bugnyar et al., 2016; Stalnaker, 1978; Sperber and Wilson, 2002; Nematzadeh et al., 2018; Le et al., 2019; Jiang and de Marneffe, 2019; Ross and Pavlick, 2019; Bowman et al., 2015; de Marneffe et al., 2012; Jeretic et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	• gender_sensitivity_english (Bordia and Bowman, 2019; Marcus et al., 1994; Caliskan et al., 2017; Bolukbasi et al., 2016; Rudinger et al., 2018; Lu et al., 2020; Gonen and Goldberg, 2019; Hall Maudslay et al., 2019; Fellbaum, 1998) under <i>Apache 2.0 License</i> , evaluated by accuracy	5120
5076			5121
5077		• general_knowledge (Shane, 2020; Dhingra et al., 2017; Rajpurkar et al., 2016, 2018; Lacker, 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5122
5078			5123
5079		• geometric_shapes (Bostock et al., 2011; Marriott et al., 2021; Boillot, 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	5124
5080			5125
5081		• goal_step_wikihow (Zhang et al., 2020d) under <i>Apache 2.0 License</i> , evaluated by accuracy	5126
5082			5127
5083	• evaluating_information_essentiality (Rajpurkar et al., 2018; Hosseini et al., 2014; Levy et al., 2017; Yin and Pei, 2015; de Marneffe et al., 2008; Zhu et al., 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy		5128
5084			5129
5085		• gre_reading_comprehension (Lai et al., 2017) under <i>Apache 2.0 License</i> , evaluated by accuracy	5130
5086			5131
5087		• hhh_alignment under <i>Apache 2.0 License</i> , evaluated by accuracy	5132
5088	• fact_checker (Thorne et al., 2018; Lee et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy		5133
5089			5134
5090			5135
5091	• factuality_of_summary (Eyal et al., 2019; Wang et al., 2020a; Durmus et al., 2020; Vasilyev et al., 2020; See et al., 2017b; Hermann et al., 2015b; Narayan et al., 2018; Pagnoni et al., 2021; Kryscinski et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy		5136
5092			5137
5093			5138
5094			5139
5095			5140
5096			
5097	• fantasy_reasoning (Wang et al., 2019b, 2018; McCann et al., 2018; Bhagavatula et al., 2019;		5141
5098			5142

5143	• high_low_game under <i>Apache 2.0 License</i> , evaluated by accuracy	5187
5144		5188
5145	• hindi_question_answering (Brown et al., 2020; Radford et al., 2019; Jain et al., 2020; Lewis et al., 2020a; Artetxe et al., 2020; Ra- jpurkar et al., 2016) under <i>Apache 2.0 License</i> , evaluated by accuracy	5189
5146		5190
5147		
5148		
5149		
5150	• hinglish_toxicity under <i>Apache 2.0 License</i> , evaluated by accuracy	
5151		
5152	• human_organs_senses under <i>Apache 2.0 Li-</i> <i>cence</i> , evaluated by accuracy	
5153		
5154	• hyperbaton (Forsyth, 2014) under <i>Apache</i> <i>2.0 License</i> , evaluated by accuracy	5191
5155		5192
5156	• identify_math_theorems (Gao et al., 2021; Black et al., 2022; Brown et al., 2020; Radford et al., 2019; Wang and Komatsuzaki, 2021) un- der <i>Apache 2.0 License</i> , evaluated by accuracy	
5157		
5158		
5159		
5160	• identify_odd_metaphor (Lakoff and John- son, 2008; Gao et al., 2018) under <i>Apache 2.0</i> <i>License</i> , evaluated by accuracy	
5161		
5162		
5163	• implicatures (Davis, 2019; George and Mamidi, 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	
5164		
5165		
5166	• implicit_relations (Cain and Oakhill, 1999; Bayat and Çetinkaya, 2020; Srivastava et al., 2016; Lin et al., 2019a; Massey et al., 2015) under <i>Apache 2.0 License</i> , evaluated by accu- racy	
5167		
5168		
5169		
5170		
5171	• intent_recognition (Brown et al., 2020; Winata et al., 2021; Madotto et al., 2020; Coucke et al., 2018; Madotto et al., 2021) un- der <i>Apache 2.0 License</i> , evaluated by accuracy	
5172		
5173		
5174		
5175	• international_phonetic_alphabet_nli (Williams et al., 2018) under <i>Apache 2.0 License</i> , evalu- ated by accuracy	
5176		
5177		
5178	• international_phonetic_alphabet_transliterate (Brown et al., 2020; Liu et al., 2020c; Williams et al., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	
5179		
5180		
5181		
5182	• intersect_geometry (Weston et al., 2015; Agrawal et al., 2015; Trask et al., 2018; Seo et al., 2014; Hosseini et al., 2014; Polu and Sutskever, 2020; Yang and Deng, 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	
5183		
5184		
5185		
5186		
	• irony_identification (Zhang et al., 2019b; Ghanem et al., 2020; Salas-Zárate et al., 2017) under <i>Apache 2.0 License</i> , evaluated by accu- racy	5193
		5194
		5195
		5196
	• kanji_ascii under <i>Apache 2.0 License</i> , evalu- ated by accuracy	5197
		5198
	• kannada (Prentice and Fathman, 1975; Narasimhachar, 1988; Liu et al., 2021b; Lev et al., 2004; Lin et al., 2021a) under <i>Apache</i> <i>2.0 License</i> , evaluated by accuracy	5199
		5200
	• key_value_maps under <i>Apache 2.0 License</i> , evaluated by accuracy	5201
		5202
	• language_games under <i>Apache 2.0 License</i> , evaluated by accuracy	5203
		5204
	• linguistic_mappings (McCoy et al., 2018, 2020; Mulligan et al., 2021; Rumelhart et al., 1986; Kirov and Cotterell, 2018; Berko, 1958; Baayen et al., 1995) under <i>Apache 2.0 License</i> , evaluated by accuracy	5205
		5206
	• list_functions (Rule et al., 2020; Rule, 2020; Green et al., 1974; Shaw et al., 1975; Bier- mann, 1978; Green, 1981; Smith, 1984; Feser et al., 2015; Osera and Zdancewic, 2015; Po- likarpova et al., 2016; Cropper et al., 2020; Graves et al., 2014; Reed and de Freitas, 2015; Joulin and Mikolov, 2015; Balog et al., 2016; Bošnjak et al., 2017; Gaunt et al., 2016; Chen et al., 2019b; Kitzelmann, 2010; Flener and Schmid, 2008; Gulwani et al., 2017a; Devlin et al., 2017; Ellis et al., 2020; Cropper and Muggleton, 2016; Piantadosi, 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5207
		5208
		5209
		5210
		5211
		5212
		5213
		5214
		5215
		5216
		5217
		5218
	• logical_args under <i>Apache 2.0 License</i> , eval- uated by accuracy	5219
		5220
	• logical_fallacy_detection (Brown et al., 2020; Hendrycks et al., 2021b; Bender et al., 2021; Wachsmuth et al., 2017; Yu et al., 2020; Copi et al., 2018; Oberauer et al., 2005; Ober- auer and Wilhelm, 2000) under <i>Apache 2.0</i> <i>License</i> , evaluated by accuracy	5221
		5222
		5223
		5224
		5225
		5226
	• logical_sequence (Saxton et al., 2019; Lin et al., 2020a; Bowman and Dahl, 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	5227
		5228
		5229
	• long_context_integration under <i>Apache 2.0</i> <i>License</i> , evaluated by accuracy	5230
		5231

5232	• mathematical_induction (Hendrycks et al., 2021c; Patel et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	2018, 2020; Petrova-Antonova and Tancheva, 2020; Huynh and Mazzocchi, 2012; Kandel et al., 2011; Bhupatiraju et al., 2017; Ellis and Gulwani, 2017; Gulwani et al., 2012; Gulwani, 2011; Singh and Gulwani, 2016) under <i>Apache 2.0 License</i> , evaluated by accuracy	5277
5233			5278
5234			5279
5235	• matrixshapes under <i>Apache 2.0 License</i> , evaluated by accuracy		5280
5236			5281
5237	• metaphor_boolean (Lakoff and Johnson, 2008; Bizzoni and Lappin, 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	• multiemo (Kocoń et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	5283
5238			5284
5239			
5240	• metaphor_understanding (Paul, 1970; Tong et al., 2021; Radford et al., 2019; Rai and Chakraverty, 2020; Shutova, 2010; Stowe et al., 2020; Mohler et al., 2016; Shutova and Teufel, 2010; Birke and Sarkar, 2006; Zayed et al., 2020; Steen et al., 2010; Tong, 2021; Bizzoni and Lappin, 2018; Tsvetkov et al., 2014) under <i>Apache 2.0 License</i> , evaluated by accuracy	• multistep_arithmetic (Flanagan and Dixon, 2014) under <i>Apache 2.0 License</i> , evaluated by accuracy	5285
5241			5286
5242			5287
5243		• muslim_violence_bias (Abid et al., 2021; Bender et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	5288
5244			5289
5245			5290
5246			
5247		• natural_instructions (Mishra et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	5291
5248			5292
5249	• minute_mysteries_qa (Sugawara et al., 2017; Dunietz et al., 2020; Kočiský et al., 2018; Mostafazadeh et al., 2016a; Frermann et al., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	• navigate (Graves et al., 2016; Henaff et al., 2016; Geva et al., 2020b; Chen et al., 2019a; Kryscinski et al., 2020; Côté et al., 2018; Luketina et al., 2019; Thawani et al., 2021; Lake and Baroni, 2017) under <i>Apache 2.0 License</i> , evaluated by accuracy	5293
5250			5294
5251			5295
5252			5296
5253			5297
5254	• misconceptions (Irving et al., 2018; Atanasova et al., 2020; Boller and George, 1989) under <i>Apache 2.0 License</i> , evaluated by accuracy		5298
5255		• nonsense_words_grammar under <i>Apache 2.0 License</i> , evaluated by accuracy	5299
5256			5300
5257			
5258	• mnist_ascii under <i>Apache 2.0 License</i> , evaluated by accuracy	• object_counting (Rugani et al., 2015; Wang et al., 2019c; Zhang et al., 2018b; Brown et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5301
5259			5302
5260	• modified_arithmetic (Brown et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy		5303
5261			5304
5262	• moral_permissibility (Hendrycks et al., 2020; Lourie et al., 2020; Thomson, 1976) under <i>Apache 2.0 License</i> , evaluated by accuracy	• odd_one_out (Resnik, 1995, 1999; Jiang and Conrath, 1997; Li et al., 2003; Banerjee and Pedersen, 2003; Jarmasz, 2012; Hughes and Ramage, 2007; Tsatsaronis et al., 2010, 2009; Morris and Hirst, 1991; Strube and Ponzetto, 2006; Ponzetto and Strube, 2007; Gabrilovich and Markovitch, 2007; Milne and Witten, 2008; Yeh et al., 2009; Radinsky et al., 2011; Cilibiasi and Vitanyi, 2007; Deerwester et al., 1990; Reisinger and Mooney, 2010; El-Yaniv and Yanay, 2013) under <i>Apache 2.0 License</i> , evaluated by accuracy	5305
5263			5306
5264			5307
5265			5308
5266	• movie_dialog_same_or_different (Park et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy		5309
5267			5310
5268			5311
5269	• movie_recommendation (Sileo et al., 2022; Thorat et al., 2015; Barkan and Koenigstein, 2016; Harper and Konstan, 2015) under <i>Apache 2.0 License</i> , evaluated by accuracy		5312
5270			5313
5271			5314
5272			5315
5273	• mult_data_wrangling (Bender et al., 2021; Tamkin et al., 2021; Singh and Gulwani, 2015; Cropper et al., 2016; Wu et al., 2012; Gulwani et al., 2015; Contreras-Ochando et al.,	• paragraph_segmentation under <i>Apache 2.0 License</i> , evaluated by accuracy	5317
5274			5318
5275		• parsinlu_qa (Khashabi et al., 2020a) under <i>Apache 2.0 License</i> , evaluated by accuracy	5319
5276			5320

5321	• penguins_in_a_table (Herzig et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5364
5322		5365
5323	• periodic_elements under <i>Apache 2.0 License</i> , evaluated by accuracy	5366
5324		5367
5325	• persian_idioms under <i>Apache 2.0 License</i> , evaluated by accuracy	5368
5326		5369
5327	• phrase_relatedness (Asaadi et al., 2019 ; Levy et al., 2015 ; Ein Dor et al., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	5370
5328		5371
5329		5372
5330	• physical_intuition under <i>Apache 2.0 License</i> , evaluated by accuracy	5373
5331		5374
5332	• physics under <i>Apache 2.0 License</i> , evaluated by accuracy	5375
5333		5376
5334	• physics_questions (Ling et al., 2017 ; Amini et al., 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	5377
5335		
5336		
5337	• polish_sequence_labeling (Nguyen and Guo, 2007 ; Rei, 2017 ; Gu et al., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	5378
5338		5379
5339		
5340	• presuppositions_as_nli (Heim, 1983 ; de Marneffe et al., 2019 ; White et al., 2018 ; Jeretic et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5380
5341		5381
5342		5382
5343		
5344	• program_synthesis (Gulwani et al., 2017b) under <i>Apache 2.0 License</i> , evaluated by accuracy	5383
5345		5384
5346		5385
5347	• protein_interacting_sites (Dhole et al., 2014 ; Singh et al., 2014 ; Li et al., 2020b ; Murakami and Mizuguchi, 2010) under <i>Apache 2.0 License</i> , evaluated by accuracy	5386
5348		5387
5349		
5350		
5351	• python_programming_challenge (Allamanis et al., 2018 ; Alon et al., 2020 ; Hendrycks et al., 2021a) under <i>Apache 2.0 License</i> , evaluated by accuracy	5388
5352		5389
5353		5390
5354		5391
5355	• qa_wikidata (Radford et al., 2019 ; Kwiatkowski et al., 2019 ; Weber and Jaimes, 2011) under <i>Apache 2.0 License</i> , evaluated by accuracy	5392
5356		5393
5357		
5358		
5359	• question_answer_creation under <i>Apache 2.0 License</i> , evaluated by accuracy	5394
5360		5395
5361	• question_selection (Rajpurkar et al., 2016) under <i>Apache 2.0 License</i> , evaluated by accuracy	5396
5362		5397
5363		
	• real_or_fake_text (Dugan et al., 2020 ; Ippolito et al., 2020 ; Solaiman et al., 2019 ; Zellers et al., 2019b ; Brown et al., 2020 ; Bakhtin et al., 2019 ; Sandhaus, 2008 ; Fan et al., 2018 ; Marín et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	5398
		5399
	• reasoning_about_colored_objects (Hendricks et al., 2018 ; Hosseini et al., 2014 ; Winograd, 1972 ; Wang et al., 2016 ; Jayanavar et al., 2020 ; Suhr et al., 2019 ; Thomason et al., 2015 ; Mitchell et al., 2010 ; Viethen and Dale, 2008 ; Gatt et al., 2009 ; Mitchell et al., 2013 ; Liang et al., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	5400
		5401
	• rephrase under <i>Apache 2.0 License</i> , evaluated by accuracy	5402
		5403
	• riddle_sense (Lin et al., 2021a ; Talmor et al., 2019b) under <i>Apache 2.0 License</i> , evaluated by accuracy	5404
		5405
	• roots_optimization_and_games (Lample and Charton, 2019 ; Polu and Sutskever, 2020 ; Amos and Kolter, 2017 ; Agrawal et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5406
		5407
	• ruin_names (Attardo, 2017 ; Ren and Yang, 2017 ; Amin and Burghardt, 2020 ; Annamoradnejad and Zoghi, 2020 ; Blinov et al., 2019 ; Yan and Pedersen, 2017 ; Frolovs, 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	
	• salient_translation_error_detection under <i>Apache 2.0 License</i> , evaluated by accuracy	
	• scientific_press_release under <i>Apache 2.0 License</i> , evaluated by accuracy	
	• self_awareness (Yudkowsky, 2008 ; Chella et al., 2020 ; Schick et al., 2021 ; Kounev et al., 2017 ; Huttunen et al., 2017 ; Wallace et al., 2020 ; Clark and Jackson, 1994 ; Horowitz, 2017 ; Branwen, 2020 ; Chu et al., 2017) under <i>Apache 2.0 License</i> , evaluated by accuracy	
	• self_evaluation_courtroom (Hildebrandt, 2018 ; Daley, 2021 ; King and Cook, 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	

5408	• self_evaluation_tutoring (Zhang et al., 2019a; Irving et al., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	5454
5409		5455
5410		
5411	• semantic_parsing_in_context_sparc (Yu et al., 2019b, 2018, 2019a) under <i>Apache 2.0 License</i> , evaluated by accuracy	5456
5412		5457
5413		
5414	• semantic_parsing_spider (Yu et al., 2018, 2019b,a) under <i>Apache 2.0 License</i> , evaluated by accuracy	5458
5415		5459
5416		5460
5417	• sentence_ambiguity under <i>Apache 2.0 License</i> , evaluated by accuracy	5461
5418		5462
5419	• similarities_abstraction (Nasreddine et al., 2005; Wechsler, 2008) under <i>Apache 2.0 License</i> , evaluated by accuracy	5463
5420		5464
5421		
5422	• simp_turing_concept (Böhm, 1964; Sun et al., 2020; Devlin et al., 2018; Radford et al., 2019; Brown et al., 2020; Vaswani et al., 2017; Hendrycks et al., 2021b; Xu et al., 2020b; Izacard and Grave, 2020; Zhu, 2015; Cakmak and Thomaz, 2014; Goodman and Frank, 2016; Degen et al., 2020; Khan et al., 2011; Basu and Christensen, 2013; Yang and Shafto, 2017; Melo et al., 2018; Telle et al., 2019; Chater and Vitányi, 2003; Hupkes et al., 2020; Lakretz et al., 2019; Toshniwal et al., 2021; Bender and Koller, 2020; Köhl et al., 2020; Marcus and Davis, 2020; Sinha et al., 2019; McClelland et al., 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	5465
5423		5466
5424		5467
5425		5468
5426		
5427		
5428		
5429		
5430		
5431		
5432		
5433		
5434		
5435		
5436		
5437	• simple_ethical_questions (Hendrycks et al., 2020; Lourie et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5469
5438		5470
5439		5471
5440	• simple_text_editing (Branwen, 2020; Malmi et al., 2019; Faltings et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5472
5441		5473
5442		5474
5443	• snarks (Brown et al., 2020; Devlin et al., 2018; Lan et al., 2019; Liu et al., 2019; Radford et al., 2019; Annamradnejad and Zoghi, 2020; Chen and Soo, 2018; Mao and Liu, 2019; Weller and Seppi, 2019; Khodak et al., 2017; Ghosh et al., 2020; González-Ibáñez et al., 2011; Joshi et al., 2015; McCoy et al., 2019; Kaushik et al., 2019; Gardner et al., 2020; Sennrich, 2016; Burlot et al., 2018; Naik et al., 2018; Zhang et al., 2016; Felbo et al., 2017; Pant and Dadu, 2020; Pelser	5475
5444		5476
5445		5477
5446		5478
5447		
5448		
5449		
5450		
5451		
5452		
5453		
	and Murrell, 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	5479
	• social_support (Wang and Jurgens, 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	5480
		5481
	• social_iqa (Sap et al., 2019; Bisk et al., 2020; Talmor et al., 2019b; Zellers et al., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	5482
		5483
	• spelling_bee (Ginsberg, 2014) under <i>Apache 2.0 License</i> , evaluated by accuracy	5484
		5485
	• sports_understanding under <i>Apache 2.0 License</i> , evaluated by accuracy	5486
		5487
	• squad_shifts (Miller et al., 2020; Brown et al., 2020; Rajpurkar et al., 2016; Baumgartner et al., 2020; McAuley et al., 2015) under <i>Apache 2.0 License</i> , evaluated by accuracy	5488
		5489
	• subject_verb_agreement (Lakretz et al., 2021b, 2019; Linzen et al., 2016; Gulordava et al., 2018; Marvin and Linzen, 2018; Goldberg, 2019; Lakretz et al., 2021a; Wolf, 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	5490
		5491
	• sudoku (Wang et al., 2019d; Russell and Norvig, 2002; Garcez and Lamb, 2020; Huang et al., 2018; Hendrycks et al., 2021c) under <i>Apache 2.0 License</i> , evaluated by accuracy	5492
		5493
	• sufficient_information under <i>Apache 2.0 License</i> , evaluated by accuracy	5494
		5495
	• suicide_risk (Gaur et al., 2019; Mohammadi et al., 2019; Matero et al., 2019; Shing et al., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	5496
		5497
	• swahili_english_proverbs under <i>Apache 2.0 License</i> , evaluated by accuracy	
	• swedish_to_german_proverbs (Hanzén, 2007; Korhonen, 2009; Meister, 2007; Mieder, 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	
	• taboo (Joshi et al., 2017) under <i>Apache 2.0 License</i> , evaluated by accuracy	
	• talkdown (Wang and Potts, 2019; Mendelsohn et al., 2020; Fiske, 1993; Nolan and Mikami, 2013; Breitfeller et al., 2019; Perez Almendros et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	

5498	• temporal_sequences (Elazar et al., 2021; Pustejovsky et al., 2004; Sanampudi and Kumari, 2010; Han et al., 2020; Ma et al., 2021; Brown et al., 2020; Petroni et al., 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	5544
5499		5545
5500		5546
5501		
5502		
5503	• tense (Logeswaran et al., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	
5504		
5505	• text_navigation_game (Vinyals et al., 2019; Küttler et al., 2020; Kanagawa and Kaneko, 2019; Noever et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	
5506		
5507		
5508		
5509	• timedial (Qin et al., 2021; Li et al., 2017; Zhou et al., 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	
5510		
5511		
5512	• topical_chat (Gopalakrishnan et al., 2019; Mehri and Eskenazi, 2020; Gopalakrishnan et al., 2020; Hedayatnia et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	
5513		
5514		
5515		
5516	• tracking_shuffled_objects (Liu et al., 2020b; Dong et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	
5517		
5518		
5519	• training_on_test_set under <i>Apache 2.0 License</i> , evaluated by accuracy	
5520		
5521	• truthful_qa (Brown et al., 2020; Sellam et al., 2020; Amodei et al., 2016; Leike et al., 2018; Kenton et al., 2021; Clark et al., 2018; Bhakthavatsalam et al., 2021; Hendrycks et al., 2021b; Khashabi et al., 2020b; Kreps et al., 2020; Maynez et al., 2020; Gabriel et al., 2020; Wang et al., 2020a; Stiennon et al., 2020; Lewis et al., 2020b; Krishna et al., 2021; Gehrmann et al., 2021; Xu et al., 2020a; Dinan et al., 2019; Tamkin et al., 2021; Bowman and Dahl, 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	
5522		
5523		
5524		
5525		
5526		
5527		
5528		
5529		
5530		
5531		
5532		
5533	• twenty_questions (Rajpurkar et al., 2016; Choi et al., 2018; Reddy et al., 2019; Aliannejadi et al., 2019; Clark et al., 2019; Zhang et al., 2019a) under <i>Apache 2.0 License</i> , evaluated by accuracy	
5534		
5535		
5536		
5537		
5538	• understanding_fables (Reimers and Gurevych, 2019; Salazar et al., 2020; Wolf et al., 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	
5539		
5540		
5541		
5542	• undo_permutation (Pham et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	
5543		
	• unit_conversion (Hendrycks et al., 2021c; Geva et al., 2020a) under <i>Apache 2.0 License</i> , evaluated by accuracy	5544
		5545
		5546
	• unit_interpretation under <i>Apache 2.0 License</i> , evaluated by accuracy	5547
		5548
	• unnatural_in_context_learning (Brown et al., 2020; Kaplan et al., 2020; Henighan et al., 2020; Hernandez et al., 2021; Bahri et al., 2021; Wang et al., 2019b; Hernandez et al., 2020; Hendrycks et al., 2021c,a; Liu et al., 2021a; Zhao et al., 2021; Perez et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	5549
		5550
		5551
		5552
		5553
		5554
		5555
		5556
	• unqover (Li et al., 2020a; Caliskan et al., 2017; Rudinger et al., 2018; Zhao et al., 2018; Dev et al., 2020; Stanovsky et al., 2019; Nadeem et al., 2020; Sheng et al., 2019; Zhang et al., 2020b) under <i>Apache 2.0 License</i> , evaluated by accuracy	5557
		5558
		5559
		5560
		5561
		5562
	• web_of_lies under <i>Apache 2.0 License</i> , evaluated by accuracy	5563
		5564
	• what_is_the_tao under <i>Apache 2.0 License</i> , evaluated by accuracy	5565
		5566
	• which_wiki_edit under <i>Apache 2.0 License</i> , evaluated by accuracy	5567
		5568
	• word_problems_on_sets_and_graphs (Bency et al., 2019; Mnih et al., 2013; Russell and Norvig, 2002; Besold et al., 2017; Clark et al., 2020; Wang et al., 2018; Lacker, 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5569
		5570
		5571
		5572
		5573
	• word_sorting (Grover et al., 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	5574
		5575
	• word_unscrambling (Nishino et al., 2019; Rozner et al., 2021; Jones et al., 2020; Mays et al., 1991; Edizel et al., 2019; Sakaguchi et al., 2016; Kim et al., 2015; Xue et al., 2021a; Wu et al., 2020; Rust et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5576
		5577
		5578
		5579
		5580
		5581
	• yes_no_black_white under <i>Apache 2.0 License</i> , evaluated by accuracy	5582
		5583
	D.4 BigBench-Lite Datasets	5584
	• auto_debugging (Zaremba and Sutskever, 2014) under <i>Apache 2.0 License</i> , evaluated by accuracy	5585
		5586
		5587

5588	• bbq_lite_json (Crawford, 2017; Khashabi et al., 2020b; Li et al., 2020a) under <i>Apache 2.0 License</i> , evaluated by accuracy	5631
5589		5632
5590		
5591	• code_line_description (Alon et al., 2018) under <i>Apache 2.0 License</i> , evaluated by accuracy	5633
5592		5634
5593		5635
5594	• conceptual_combinations (Fodor, 1975; Fodor and Pylyshyn, 1988; Smolensky, 1988; Lake et al., 2017; Lake and Murphy, 2020; Marcus, 2020; Henrich et al., 2010; Murphy, 1988) under <i>Apache 2.0 License</i> , evaluated by accuracy	5636
5595		
5596		
5597		
5598		
5599	• conlang_translation (Canfield, 2010; Şahin et al., 2020; Sennrich and Zhang, 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	5641
5600		5642
5601		5643
5602		
5603	• emoji_movie (Cruse, 2015; Instagram Engineering, 2015; Chandra Guntuku et al., 2019; Eisner et al., 2016; Mayne, 2020; Boillot, 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	5644
5604		5645
5605		
5606		
5607		
5608	• formal_fallacies_syllogisms_negation (Kassner and Schütze, 2019; Talmor et al., 2019a; Betz et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5646
5609		5647
5610		5648
5611		
5612	• hindu_knowledge under <i>Apache 2.0 License</i> , evaluated by accuracy	5649
5613		5650
5614		
5615	• known_unknowns (Liu et al., 2021c; Xiao and Wang, 2021; Shuster et al., 2021; Zhou et al., 2020; Dziri et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	5651
5616		5652
5617		5653
5618	• language_identification (Brown, 2014) under <i>Apache 2.0 License</i> , evaluated by accuracy	5654
5619		5655
5620		
5621	• linguistics_puzzles (Bozhanov and Derzhanski, 2013; Radev et al., 2008; Sennrich and Zhang, 2019; Clark et al., 2018; Şahin et al., 2020) under <i>Apache 2.0 License</i> , evaluated by accuracy	5656
5622		5657
5623		5658
5624		
5625	• logic_grid_puzzle under <i>Apache 2.0 License</i> , evaluated by accuracy	5659
5626		5660
5627		5661
5628		5662
5629	• logical_deduction under <i>Apache 2.0 License</i> , evaluated by accuracy	5663
5630		
	• novel_concepts (Santoro et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	5664
	• operators (Brown et al., 2020; Kassner et al., 2020; Hendrycks et al., 2021c; Saxton et al., 2019) under <i>Apache 2.0 License</i> , evaluated by accuracy	5665
		5666
		5667
		5668
	• parsinlu_reading_comprehension (Khashabi et al., 2020a; Xue et al., 2021b; Rajpurkar et al., 2016) under <i>Apache 2.0 License</i> , evaluated by accuracy	5669
		5670
		5671
		5672
		5673
		5674
	• play_dialog_same_or_different (Park et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	5675
	• repeat_copy_logic (Graves et al., 2014) under <i>Apache 2.0 License</i> , evaluated by accuracy	5676
	• strange_stories (Happé, 1994; White et al., 2009) under <i>Apache 2.0 License</i> , evaluated by accuracy	5677
	• strategyqa (Geva et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	5678
	• symbol_interpretation (Brown et al., 2020; Johnson et al., 2016; Santoro et al., 2017; Hudson and Manning, 2019; Sennrich et al., 2015) under <i>Apache 2.0 License</i> , evaluated by accuracy	5679
		5680
		5681
		5682
		5683
	• vitaminc_fact_verification (Schuster et al., 2021) under <i>Apache 2.0 License</i> , evaluated by accuracy	5684
	• winowhy (Zhang et al., 2020c; Devlin et al., 2018; Liu et al., 2019; Kocijan et al., 2019; Rahman and Ng, 2012; Sakaguchi et al., 2020b) under <i>Apache 2.0 License</i> , evaluated by accuracy	5685
		5686
		5687
		5688
		5689
		5690
		5691
		5692
		5693
		5694
		5695
		5696
		5697
		5698
		5699
		5700
		5701
		5702
		5703
		5704
		5705
		5706
		5707
		5708
		5709
		5710
		5711
		5712
		5713
		5714
		5715
		5716
		5717
		5718
		5719
		5720
		5721
		5722
		5723
		5724
		5725
		5726
		5727
		5728
		5729
		5730
		5731
		5732
		5733
		5734
		5735
		5736
		5737
		5738
		5739
		5740
		5741
		5742
		5743
		5744
		5745
		5746
		5747
		5748
		5749
		5750
		5751
		5752
		5753
		5754
		5755
		5756
		5757
		5758
		5759
		5760
		5761
		5762
		5763
		5764
		5765
		5766
		5767
		5768
		5769
		5770
		5771
		5772
		5773
		5774
		5775
		5776
		5777
		5778
		5779
		5780
		5781
		5782
		5783
		5784
		5785
		5786
		5787
		5788
		5789
		5790
		5791
		5792
		5793
		5794
		5795
		5796
		5797
		5798
		5799
		5800
		5801
		5802
		5803
		5804
		5805
		5806
		5807
		5808
		5809
		5810
		5811
		5812
		5813
		5814
		5815
		5816
		5817
		5818
		5819
		5820
		5821
		5822
		5823
		5824
		5825
		5826
		5827
		5828
		5829
		5830
		5831
		5832
		5833
		5834
		5835
		5836
		5837
		5838
		5839
		5840
		5841
		5842
		5843
		5844
		5845
		5846
		5847
		5848
		5849
		5850
		5851
		5852
		5853
		5854
		5855
		5856
		5857
		5858
		5859
		5860
		5861
		5862
		5863
		5864
		5865
		5866
		5867
		5868
		5869
		5870
		5871
		5872
		5873
		5874
		5875
		5876
		5877
		5878
		5879
		5880
		5881
		5882
		5883
		5884
		5885
		5886
		5887
		5888
		5889
		5890
		5891
		5892
		5893
		5894
		5895
		5896
		5897
		5898
		5899
		5900
		5901
		5902
		5903
		5904
		5905
		5906
		5907
		5908
		5909
		5910
		5911
		5912
		5913
		5914
		5915
		5916
		5917
		5918
		5919
		5920
		5921
		5922
		5923
		5924
		5925
		5926
		5927
		5928
		5929
		5930
		5931
		5932
		5933
		5934
		5935
		5936
		5937
		5938
		5939
		5940
		5941
		5942
		5943
		5944
		5945
		5946
		5947
		5948
		5949
		5950
		5951
		5952
		5953
		5954
		5955
		5956
		5957
		5958
		5959
		5960
		5961
		5962
		5963
		5964
		5965
		5966
		5967
		5968
		5969
		5970
		5971
		5972
		5973
		5974
		5975
		5976
		5977
		5978
		5979
		5980
		5981
		5982
		5983
		5984
		5985
		5986
		5987
		5988
		5989
		5990
		5991
		5992
		5993
		5994
		5995
		5996
		5997
		5998
		5999
		6000

E Efficiency Analysis

GLIDER introduces minimal computational overhead by requiring only two lightweight operations per LoRA layer: a single cosine similarity calculation between the query’s task embedding and global routing vectors and a simple vector addition to combine this with the local routing score. With typical values for N (experts) and d_g (embedding dimension) in the hundreds, this amounts to just $(N \times d_g + d_g)$ FLOPs per layer, which is negligible compared to the base model’s computation.

F Detailed Performance

We list detailed performance of all baselines for each task in Table 4, 6, 5, 7, and 7.

Method	Avg	RTE	H-Swag	COPA	WIC	Winogrande	CB	StoryCloze	ANLI-R1	ANLI-R2	ANLI-R3	WSC
Multi-Task Fine-Tuning	51.6	60.1	26.9	74.8	51.3	50.9	52.7	85.1	34.7	33	33.5	64.9
Oracle Expert	57.2	66.9	36.8	89.6	52.4	57.6	59.9	96.9	34.5	34.8	36.7	63.6
Merged Experts	45.4	55.7	25.7	61.8	50.3	53.3	45.6	63.8	33.1	33.4	33.4	43.5
LoRA Hub	46.9	52.1	27.2	75.1	50.4	50.6	33.5	84.3	33.2	34.2	33.6	41.1
Glider	57.3	64.8	29.9	91.6	50.6	60.2	69.6	96.2	36.0	34.3	38.1	58.5

Table 4: Complete results on T0 Held-Out datasets.

Expert	Avg	Boolean Expression	Causal Judgment	Date Understanding	Disambiguator QA	Formal Fallacies	Geometric Shapes	Hyperbaton
Multi-Task Fine-Tuning	34.9	49.6	55.3	35.2	55.4	51.5	10.6	50
Oracle Expert	42.2	64.4	59.5	41.7	65.1	51.7	30.1	52.7
Merged Experts	35.3	60.8	58.4	38.5	45.3	50.1	10	50
LoRA Hub	32.0	59.2	49.5	36.6	30.2	50.9	24.2	50.0
Glider	34.9	52.8	59.5	39.6	57.0	50.1	10.3	50.0

Expert	Logical Detection	Movie Recommendation	Multistep Arithmetic	Navigate	Object Counting	Penguins in a Table	Reasoning about Colored Objects	Ruin Names
Multi-Task Fine-Tuning	47.9	34.8	0	50	22.5	32.9	42	19.6
Oracle Expert	45.8	49.2	1.6	50	28.1	36.9	53.2	49.6
Merged Experts	44.3	23	0.4	50	25.4	35.6	47.5	26.6
LoRA Hub	27.5	32.0	1.2	50.0	0.0	30.9	19.3	21.2
Glider	40.8	31.2	1.6	50.0	19.7	34.2	48.4	24.3

Expert	Salient Translation Error Detection	Snarks	Sports Understanding	Temporal Sequences	Track Shuffled Objects	Web of Lies	Word Sorting
Multi-Task Fine-Tuning	27.8	46.4	50.2	16.1	17.4	51.6	0.5
Oracle Expert	27	61.3	50.9	28.2	20.1	59.2	2.8
Merged Experts	24.9	44.8	51.7	19.5	17	51.6	0.9
LoRA Hub	24.9	43.6	50.3	12.8	19.7	55.6	0.0
Glider	25.6	48.1	51.4	12.8	16.1	52.8	0.0

Table 5: BIG-bench Hard (BBH) results of different methods in T0 Held-In setting

Expert	Avg	BBQ Lite Json	Conceptual Combinations	Conlong Translation	Formal Fallacies	Hindu Knowledge
Multi-Task Fine-Tuning	36.6	40.8	44.7	26	51.5	40.6
Oracle Expert	43.5	55.3	62.1	29.8	51.6	46.3
Merged Experts	36	42.5	33	28.9	50.1	40
LoRA Hub	32.4	40.1	39.8	0.2	50.1	29.1
Glider	37.5	48.4	44.7	17.1	50.1	48.6

Expert	Known Unknowns	Linguistic Puzzles	Logic Grid Puzzle	Logical Detection	Novel Concepts	Operators
Multi-Task Fine-Tuning	47.8	0	35.9	48.1	40.6	1
Oracle Expert	65.2	0	41.7	45.4	43.8	8.6
Merged Experts	45.7	0	39.6	44.3	28.1	7.1
LoRA Hub	45.7	0.0	32.8	20.1	28.1	5.7
Glider	52.2	0.0	39.6	40.8	43.8	2.9

Expert	Play Dialog Same or Different	Repeat Copy Logic	Strange Stories	Strategy QA	Vitamin C Fact Verification	Winowhy
Multi-Task Fine-Tuning	45.8	0	47.7	52.5	54.2	44.3
Oracle Expert	63.3	0	68.4	56.1	51.1	50.5
Merged Experts	36.9	0	56.3	54.3	61.3	44.3
LoRA Hub	47.5	0.0	43.7	52.9	49.9	44.3
Glider	36.9	0.0	63.8	53.6	49.8	44.5

Table 6: BIG-bench Lite (BBL) results of different methods in T0 Held-In setting

Expert	Avg	Boolean Expression	Causal Judgment	Date Understanding	Disambiguator QA	Formal Fallacies	Geometric Shapes	Hyperbaton
Multi-Task Fine-Tuning	38.9	50	61.1	36.6	65.9	52.2	9.7	51.1
Oracle Expert	45.5	66	59.5	42.3	65.1	52.9	30.1	69.3
Merged Experts	34.6	53.6	56.8	36.9	45.7	50	12	52.2
LoRA Hub	32.8	55.2	54.2	26.8	30.2	50.0	19.5	50.0
Glider	35.3	50.8	58.4	35.2	50.4	50.5	10.3	48.2

Expert	Logical Detection	Movie Recommendation	Multistep Arithmetic	Navigate	Object Counting	Penguins in a Table	Reasoning about Colored Objects	Ruin Names
Multi-Task Fine-Tuning	49.6	32.8	0	50	35.7	39.6	56.6	19
Oracle Expert	48.8	49.2	1.6	54.6	45.7	37.6	53.5	49.6
Merged Experts	42.9	23.2	0.8	50	24.1	34.2	44.5	28.3
LoRA Hub	31.4	33.8	0.0	50.0	0.0	29.5	25.1	24.6
Glider	40.1	28.2	0.4	50.0	41.1	32.2	46.9	33.0

Expert	Salient Translation Error Detection	Snarks	Sports Understanding	Temporal Sequences	Track Shuffled Objects	Web of Lies	Word Sorting
Multi-Task Fine-Tuning	39.2	59.7	51.2	27.2	15.7	53.6	0
Oracle Expert	31.5	61.3	52.5	45.3	20.3	61.6	2.9
Merged Experts	26.4	40.9	49.9	22.4	16.3	49.6	1.2
LoRA Hub	12.4	48.1	50.3	100.0	19.1	48.8	0.0
Glider	23.8	41.4	49.6	8.8	17.2	52.0	0.1

Table 7: BIG-bench Hard (BBH) results of different methods in the FLAN setting.

Expert	Avg	BBQ Lite Json	Conceptual Combinations	Conlong Translation	Formal Fallacies	Hindu Knowledge
Multi-Task Fine-Tuning	45.4	66.9	72.8	27.9	52.2	40
Oracle Expert	46.5	61.3	62.1	36	52.9	46.3
Merged Experts	34	40.2	27.2	29.1	50	36.6
LoRA Hub	31.8	36.0	29.1	2.6	50.0	25.7
Glider	35.5	49.4	35.9	10.1	50.5	47.4

Expert	Known Unknowns	Linguistic Puzzles	Logic Grid Puzzle	Logical Detection	Novel Concepts	Operators
Multi-Task Fine-Tuning	58.7	0	42.8	49.9	37.5	13.3
Oracle Expert	65.2	0	42.5	48.6	50	12.4
Merged Experts	43.5	0	37.9	42.9	34.4	7.6
LoRA Hub	52.2	0.0	29.3	26.2	34.4	0.0
Glider	41.3	0.0	34.4	40.1	25.0	3.3

Expert	Play Dialog Same or Different	Repeat Copy Logic	Strange Stories	Strategy QA	Vitamin C Fact Verification	Winowhy
Multi-Task Fine-Tuning	44.4	0	75.9	65.7	78.5	45.3
Oracle Expert	63.3	0	74.1	56.1	66.7	53.8
Merged Experts	36.9	0	46.6	52.1	47.9	44.3
LoRA Hub	63.1	0.0	28.2	46.7	51.4	44.3
Glider	37.6	0.0	61.5	52.6	66.5	48.3

Table 8: BIG-bench Lite (BBL) results of different methods in the FLAN setting.