

A Survey of Pun Generation: Datasets, Evaluations and Methodologies

Anonymous ACL submission

Abstract

Pun generation seeks to creatively modify linguistic elements in text to produce humour or evoke double meanings. It also aims to preserve coherence and contextual appropriateness, making it useful in creative writing and entertainment across various media and contexts. This field has been widely studied in computational linguistics, while there are currently no surveys that specifically focus on pun generation. To bridge this gap, this paper provides a comprehensive review of pun generation datasets and methods across different stages, including traditional approaches, deep learning techniques, and pre-trained language models. Additionally, we summarise both automated and human evaluation metrics used to assess the quality of pun generation. Finally, we discuss the research challenges and propose promising directions for future work.

1 Introduction

Pun is a kind of rhetorical style that leverages the polysemy or phonetic similarity of words to produce expressions with double or multiple meanings (Delabastita, 2016). Beyond mere wordplay, puns serve as a crucial mechanism of linguistic creativity, enriching communication and making it more engaging (Carter, 2015). For example, the pun sentence “I used to be a banker, but I lost interest” plays on the pun words “interest”, encompassing both a lack of enthusiasm for banking as a profession and the idea of financial loss. This ability to encode multiple layers of meaning fosters cognitive flexibility, encouraging individuals to interpret language in innovative ways (Zheng and Wang, 2023). Due to the unique capacity of puns, they are widely used in advertising (Djafarova, 2008; Van Mulken et al., 2005), literature (Giorgadze, 2014), and various other fields.

Natural language generation (NLG) tasks involve the creation of human-like text by computers

based on given data or input (Gatt and Krahmer, 2018), with pun generation being a notable and challenging aspect of such tasks. There are various approaches utilised in automatic pun generation, including template-based methods (Hong and Ong, 2009), deep neural network approaches (He et al., 2019), and pre-trained language models (PLMs) employing various training and inference styles (Mittal et al., 2022; Xu et al., 2024). These methods are applied to different types of puns, with a particular focus on homophonic (Yu et al., 2020), homographic (Yu et al., 2018; Luo et al., 2019), heterographic puns (Xu et al., 2024) and visual puns (Rebrii et al., 2022).

Despite the long-standing research interest in pun generation, a comprehensive literature review in this field has not been conducted, to the best of our knowledge. Some existing relevant surveys focus on generating creative writing and delve into tasks such as poetry composition (Bena and Kalita, 2020; Elzohbi and Zhao, 2023), storytelling (Gieseke et al., 2021; Alhussain and Azmi, 2021), arts (Shahriar, 2022) and metaphor (Rai and Chakraverty, 2020; Ge et al., 2023). It is noteworthy that Amin and Burghardt (2020) outlined methodologies to humour generation, discussing various systems based on templates and neural networks, along with their respective strengths and weaknesses. However, they did not cover the pun research nor incorporate relevant technologies associated with large language models (LLMs). Therefore, we aim to address this gap by conducting the first comprehensive survey on pun generation.

In this survey, we review the past three decades of research and examine the current state of natural language pun generation, categorising these methods in five groups based on their technological development timeline: (1) Traditional methods, which involve generating puns by manually or automatically constructing templates; (2) Classic Deep Neural networks (DNNs), leveraging architectures,

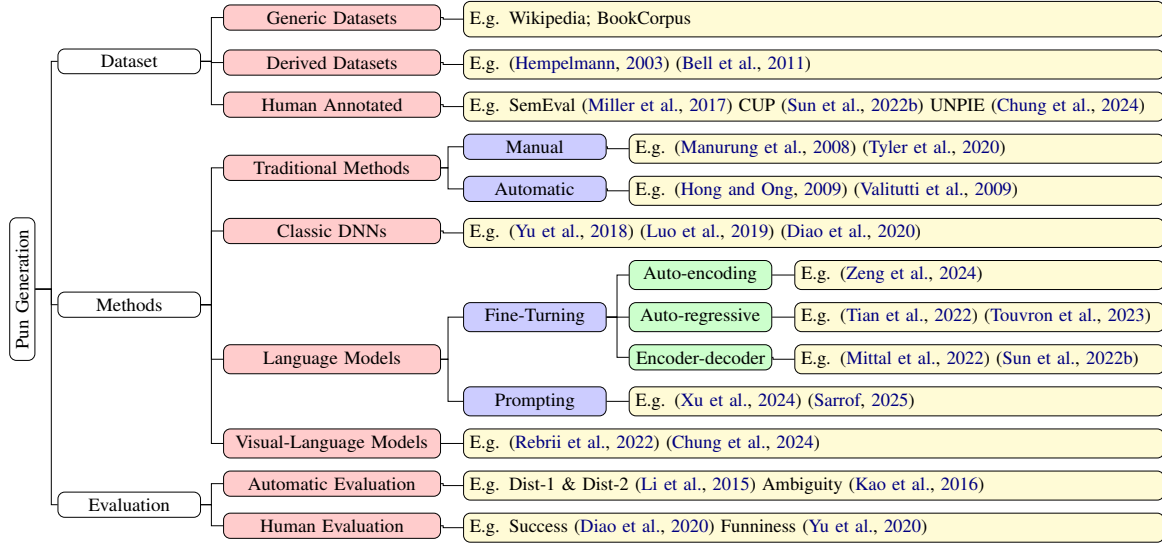


Figure 1: The survey tree for pun generation.

such as RNNs and their variants, to learn pun patterns from data; (3) Fine-tuning of PLMs, where pre-trained models like GPT (Radford, 2018) are adapted with task-specific datasets to improve pun generation, (4) Prompting of PLMs, which utilizes carefully designed prompts to guide models in generating puns without additional training, and (5) Visual-language models, where some preliminary studies on visual pun generation. We further summarise the automatic and human evaluation metrics used to assess the quality of generated puns. Finally, we discuss our findings and propose promising research directions for future work in this field.

The paper is organised as follows: Section 2 reviews the main categories of puns and provides examples for each category. Section 3, 4 and 5 summarise the relevant datasets, methods, and evaluation metrics, as shown in figure 1. We also discuss the challenges and outline future research directions in Section 6, as well as conclude with final remarks in Section 7.

2 Pun Categories

This section outlines the main four types of puns: i) *Homophonic puns*, ii) *Heterographic puns*, iii) *Homographic puns* and iv) *Visual pun*. The main features of these categories are listed in the Appendix A.

2.1 Homophonic Puns

Homophonic puns rely on the dual meanings of homophones, which are words that sound alike but have different meanings (Attardo, 2009), illustrated

in example (a):

- (a) Dentists don’t like a hard day at the orifice (office).

which uses the “orifice” as the pivotal pun word. The term “orifice” refers to the human mouth, while its pronunciation is similar to “office”. This similarity allows it to be interpreted as a dentist working in an office, thereby creating a humorous pun effect.

2.2 Heterographic Puns

Heterographic puns emphasize on differences in spelling with the same pronunciation to achieve their rhetorical effect, which are also classified as homophonic puns in some studies (Sun et al., 2022b; Miller et al., 2017). An example of a heterographic pun is shown as (b):

- (b) Life is a puzzle, look here for the missing peace (piece). (Xu et al., 2024)

The word “peace” can be interpreted as tranquility in life, while it shares the same pronunciation as “piece” which refers to a puzzle piece. Therefore, the pun can be recognized as seeking either peace in life or the missing piece of a puzzle.

2.3 Homographic Puns

Homographic puns exploit words spelled the same homographs but possess different meanings (Attardo, 2009), as shown in example (c):

- (c) Always trust a glue salesman. They tend to stick to their word.



Figure 2: A visual pun example features a white mouse and a mouse trap, where the combination exploits the double meaning of the word “mouse”.

The phrase “stick to their word” refers to the act of keeping a promise in common English expressions. However, the meaning of “stick” is also directly associated with the adhesive properties of “glue”, which artfully plays on the dual meanings of the word “stick”.

2.4 Visual Puns

Visual puns are a form of artistic expression that utilize images or visual elements to create double meanings (Smith et al., 2008). A typical example of a visual pun from Wikipedia¹ is shown in Figure 2. The figure leverages the multiple meanings of the word “mouse” based on the computer device and animal, thereby creating a pun effect by combining the computer mouse and mousetrap.

3 Dataset

In this section, we present the current datasets that have been used and constructed for pun research. We classified the datasets into generic datasets, derived datasets and human-annotated datasets. For the detailed table of the pun dataset, please refer to Appendix C.

3.1 Generic Datasets

In the early days of neural network technology, due to the difficulty of obtaining adequate data to train seq2seq models for some specific tasks (Yu et al., 2018), most research in pun generation relied on general datasets to train conditional language models, enabling them to capture fundamental semantic relationships. For example, some pun generation studies use the English Wikipedia corpus to train the language model (Yu et al., 2018; Luo et al., 2019; Diao et al., 2020), while others rely on Book-Corpus (Zhu, 2015; Yu et al., 2020) as a generic corpus for retrieval and training. Sarrof (2025) analysed the distribution of Hindi words in Latin and

Devanagari scripts using C4 (Raffel et al., 2020) and The Pile (Gao et al., 2020), and then tested on the Dakshina dataset (Roark et al., 2020).

3.2 Derived Datasets

The derived datasets are created for the new datasets by processing, transforming, or extracting specific details from general data. In this section, we present a list of derived datasets and outline the domains used in their creation. Sobkowiak (1991) collected 3850 puns from advertisements and conversation, while Hempelmann (2003) selected a subset for the automatic generation of heterophonic puns. Lucas (2004) proposed a tiny pun corpus that relies on lexical ambiguity from newspaper comics. Bell et al. (2011) created a 373 puns dataset from church marquees and literature to study wordplay in religious advertising. In addition, several studies have created pun datasets by filtering data from specialised joke websites. For example, both Yang et al. (2015) and Kao et al. (2016) curated pun datasets by crawling data from the “Pun of the Day” website. Jaech et al. (2016) compiled a homophonic pun dataset from Tumblr, Reddit, and Twitter to facilitate the automatic recovery of the target word in given puns.

3.3 Human Annotated

This section provides some details of human-annotated pun datasets. **SemEval.** Miller et al. (2017) released two manually annotated pun datasets based on (Miller and Turković, 2016) and (Miller, 2016) including both homophonic and heterogeneous puns, which is one of the most commonly used datasets in the pun generation community. **SemEval Enhancements.** Sun et al. (2022b) augmented the SemEval dataset by adding pun data combined with a given context and provided annotations on the adaptation between context words and their corresponding pun pairs. Furthermore, Sun et al. (2022a) added the fine-grained funniness ratings and natural language explanations based on the SemEval dataset. **ChinesePun.** Chen et al. (2024) introduced the first datasets for Chinese homophonic and homographic puns, specifically designed for pun understanding and generation tasks. **Multimodal Dataset.** Zhang et al. (2024) compiled a large collection of Chinese historical visual puns and provided detailed annotations, including the identification of prominent visual elements, matching of these elements with their symbolic meanings and interpretations. Chung et al. (2024)

¹<https://en.wikipedia.org/wiki/Pun>

selected a subset of homophonic and heterogeneous puns from the SemEval dataset and supplemented it with corresponding explanation images.

4 Methodology

In this section, we provide an overview of existing approaches to pun generation.

4.1 Traditional Models

Early traditional methods are typically through template-based construction. In linguistics, a template refers to a textual structure consisting of pre-defined slots that can be populated with various variables (Amin and Burghardt, 2020). Binsted and Ritchie (1994) developed the simple question-answer system of pun-generator Joke Analysis and Production Engine (JAPE), which was improved in subsequent versions including JAPE-2 (Binsted, 1996) and JAPE-3. The model incorporates two primary structures: schemata, which are used to explore the relationships between different keywords, and templates, which are designed to generate the basic framework for puns. Inspired by JAPE, Manurung et al. (2008) designed the STANDUP system, which expands and variants the elements generated by puns through further semantic and phonological analysis, for children with complex communication needs. Furthermore, Tyler et al. (2020) expanded upon the JAPE system by incorporating more recent knowledge bases and designed the PAUL BOT system, enhancing its capabilities and flexibility in automated pun generation.

Additionally, HCPP (Venour, 2000) and WIS-CRAIC (McKay, 2002) systems both implement models for the specific subclass of puns about homonym common phrase and idiom-based wit-ticisms according to semantic associations, respectively. Hempelmann (2003) studies target recoverability, arguing that a robust model for target alternative words recovery provides the necessary foundation for heterographic pun generation. Ritchie (2005) considered pun generation from the broader perspective of NLG. They analyse the differences in mechanisms between pun generation and traditional NLG, as well as the computational methods that could potentially accomplish this task. As for the research on non-English puns, Dybala et al. (2008) designed a Japanese pun generator as part of a conversational system, while Dehouck and Delaborde (2025) proposed a generator for automatically generating French puns based on a given

name and a word or phrase using rules.

Since building templates manually is a tedious and time-consuming task, Hong and Ong (2009) proposed Template-Based Pun Extractor and Generator (T-PEG) automatically identify, extract and represent the word relationships in a template, and then use these templates as patterns for the computer to generate its own puns. Valitutti et al. (2009) generated funny puns by implementing GraphLaugh to automatically generate different types of lexical associations and visualize them through a dynamic graph. They also explored a method for automatically generating humour through the substitution of words in short texts (Valitutti et al., 2013).

4.2 Classic DNNs

With the development of deep learning, pun generation has increasingly been implemented using deep neural networks, including Sequence-to-Sequence (Seq2Seq) (Sutskever, 2014) and Generative Adversarial Network (GAN) (Goodfellow et al., 2014). In general, Seq2Seq models map input sequences, such as words and phrases, to output the pun sentence, by maximising the conditional log-likelihood of the generated sequence.

Yu et al. (2018) represented the first attempt to apply deep neural networks to generate homographic puns without specific training data by developing a conditional language model (Mou et al., 2015) that creates sentences containing a target word with dual meanings. Building on this generator, Luo et al. (2019) introduced a novel discriminator, which is a word sense classifier with a single-layer bi-directional LSTM, to provide a well-structured ambiguity reward for the generator. Diao et al. (2020) replaced the traditional LSTM network structure with ON-LSTM (Shen et al., 2018) to further enhance performance. Additionally, He et al. (2019) and Yu et al. (2020) used the Seq2Seq model to rewrite the sentence so that it remains grammatically correct after replacing pun words.

In general, classic DNNs can generate puns that are more flexible compared to traditional models by fitting both general and pun datasets. However, existing methods heavily rely on annotated data and limited types of corpora, which restricts further improvement in the quality of pun generation.

4.3 Pre-trained Language Models

Early PLMs, such as Word2Vec (Mikolov, 2013) and GloVe (Pennington et al., 2014), are distributed

Method	Model	Type	Language	Dataset
Classic Deep Neural Networks				
Neural Pun (Yu et al., 2018)	LSTM	hog	English	Wikipedia & (Miller et al., 2017)
Pun-GAN (Luo et al., 2019)	LSTM	hog	English	Wikipedia & (Miller et al., 2017)
SurGen (He et al., 2019)	LSTM	hop	English	BookCorpus & (Miller et al., 2017)
LCR (Yu et al., 2020)	LSTM	hop	English	BookCorpus & (Hu et al., 2019)
AFPun-GAN (Diao et al., 2020)	ON-LSTM	hog	English	Wikipedia & (Miller et al., 2017)
Pre-trained Language Models				
Ext Ambipun(Mittal et al., 2022)	T5	hog	English	(Annamoradnejad and Zoghi, 2020)
Sim Ambipun(Mittal et al., 2022)	T5	hog	English	(Annamoradnejad and Zoghi, 2020)
Gen Ambipun(Mittal et al., 2022)	T5	hog	English	(Annamoradnejad and Zoghi, 2020)
UnifiedPun(Tian et al., 2022)	GPT-2 & BERT	hog&hog	English	(Annamoradnejad and Zoghi, 2020)
Context-pun(Sun et al., 2022b)	T5	hog&heg	English	(Sun et al., 2022b)
PunIntended (Zeng et al., 2024)	BERT	hop&hog	English	(Sun et al., 2022a)
PGCL (Chen et al., 2024)	LLaMA2-7B	hop&hog	English	(Miller et al., 2017)
PGCL (Chen et al., 2024)	Baichuan2-7B	hop&hog	Chinese	(Chen et al., 2024)
Hinglish (Sarraf, 2025)	GPT-3.5	hop	Multi-language	C4 & The Pile & Dakshina

Table 1: Methods of neural network models and pre-trained language models for pun generation task. Hog, hop and heg denote the types of homographic puns, homophonic puns and heterographic puns, respectively.

word representation methods trained on large-scale unlabeled text data, capable of capturing both the semantic and contextual information of words. These models are utilised to address various sub-tasks involved in pun generation. For example, Mittal et al. (2022) proposed to get the context words from Word2Vec based on pun words. Yu et al. (2020) designed a constraint selection algorithm based on lexical semantic relevance and obtained the word embeddings from Continuous Bag of Words model (CBOW) (Mikolov, 2013).

Most contemporary PLMs are built upon the Transformer architecture (Vaswani, 2017), which has shown outstanding performance across various natural language processing tasks (Min et al., 2023). The main model categories are classified into: (1) auto-encoding models, such as BERT (Devlin et al., 2019), (2) auto-regressive models, such as the GPT-2 (Radford et al., 2019), and (3) encoder-decoder models, such as T5 (Raffel et al., 2020). Pun generation tasks are primarily implemented through fine-tuning and prompting strategies.

4.3.1 PLMs with Fine-Tuning

Fine-tuning PLMs is to further train the model on a specific dataset to make it better suited to the needs of a specific task. For auto-encoding models, since the bidirectional encoding characteristics of the model are not suitable for generation tasks, most current work on pun generation employs it as the discriminator in GANs. For example, Zeng et al. (2024) and Tian et al. (2022) both used the BERT-base model, leveraging the [CLS] token rep-

resentation for classification.

In auto-regressive models, Tian et al. (2022) fine-tuned the GPT-2 model based on the combination dataset of Gutenberg BookCorpus and jokes (Annamoradnejad and Zoghi, 2020) and proposed a unified framework for generating both homophonic and homographic puns. Chen et al. (2024) fine-tuned both LLaMA2-7B (Touvron et al., 2023) and Baichuan2-7B (Yang et al., 2023) for generating English and Chinese puns respectively through the standard Direct Preference Optimization (Rafailov et al., 2024) and multistage curriculum learning framework.

For encoder-decoder models, Mittal et al. (2022) explored the generation of puns based on context words associated with pun words and finetuned a keyword-to-sentence model using the T5 model. Similarly, Sun et al. (2022b) proposed the context-situated pun generation, which involves identifying pun words for a given set of contextual keywords and then generating puns based on these keywords and the associated pun words. Zeng et al. (2024) used T5 as a generator, taking the pun semantic trees as input and generating pun text as output.

4.3.2 PLMs with Prompting

Prompting (Liu et al., 2021) refers to a specially designed input mode intended to guide PLMs, especially for LLMs, in performing specific tasks (Alhazmi et al., 2024). However, there are few studies exploring pun generation specifically from the perspective of prompting. Mittal et al. (2022) provides examples of the target pun along with its

two interpretations and instructions for generating the pun in GPT-3 (Brown et al., 2020) to serve as a baseline comparison model. Based on the Chain-of-Thought prompting approach (Wei et al., 2022), Sarrof (2025) designed a novel method that integrates homophone and transliteration modules to enhance the quality of pun generation.

In addition, Xu et al. (2024) selected a range of prominent LLMs to evaluate their capabilities on pun generation, including both open-source models in Llama2-7B-Chat (Touvron et al., 2023), Mistral-7B (Jiang et al., 2023), Vicuna-7B (Zheng et al., 2023), and OpenChat-7B (Wang et al., 2023), and closed-source models in Gemini-Pro (Google, 2023), GPT-3.5-Turbo (OpenAI, 2023a), Claude3-Opus (Anthropic, 2024), and GPT-4-Turbo (OpenAI, 2023b). These studies reveal that although LLMs still exhibit limitations in generating creative and humorous puns, their demonstrated potential highlights a developmental trend in this field. Future research can further optimize existing LLMs to enhance their performance in pun generation tasks.

4.4 Visual-Language Models

There are currently some preliminary studies on visual puns. Rebrii et al. (2022) explored the cross-lingual translation of puns combined with visual elements. Chung et al. (2024) employed the DALL-E-3 (Betker et al., 2023) to generate images that illustrated the meanings of puns based on textual puns. Zhang et al. (2024) leveraged their established dataset to conduct a comprehensive evaluation of large vision-language models in visual pun comprehension. However, to the best of our knowledge, there are no dedicated studies on visual pun generation, which is a potential future research direction.

5 Evaluation Strategies

In this section, we examine both automatic and human evaluation methods for pun generation. Table 2 summarizes the primary metrics for evaluation.

5.1 Automatic Evaluation

The automatic evaluation metrics can be categorized based on the intention and definition. We classify the metrics into funniness, diversity and fluency.

5.1.1 Funniness

Ambiguity & Distinctiveness. Kao et al. (2016) introduced the metrics of *ambiguity* and *distinctive-*

ness based on information theory. These metrics integrate computational models of general language understanding and pun features to quantitatively predict humour with fine-grained precision (Kao et al., 2016). Specifically, ambiguity refers to the uncertainty arising from multiple possible meanings within a sentence, which is formulated as:

$$Amb(M) = - \sum_{k \in \{a,b\}} P(m_k | \vec{w}) \log P(m_k | \vec{w}) \quad (1)$$

where $\vec{w}$ is a vector of observed content words in a sentence and m_k is the latent sentence meaning. Higher ambiguity allows the sentence to better support both the pun and its alternative meanings.

Distinctiveness evaluates the differences between word sets that support distinct meanings within a sentence using the symmetrized Kullback-Leibler divergence D_{KL} , defined as follows:

$$Dist(F_a, F_b) = D_{KL}(F_a || F_b) + D_{KL}(F_b || F_a) \quad (2)$$

where F_a and F_b represent the set of words in a sentence that support two different meanings along with their probability distributions. The high distinctiveness indicates that the distributions of the two-word groups differ significantly, which enhances the sense of humour.

Surprisal. Surprisal is a quantitative metric for surprise based on the pun word and the alternative word given local and global contexts (He et al., 2019). The formulation of local surprisal and global surprisal are defined as follows:

$$\begin{aligned} S_{\text{local}} &= S(x_{p-d:p-1}, x_{p+1:p+d}), \\ S_{\text{global}} &= S(x_{1:p-1}, x_{p+1:n}), \end{aligned} \quad (3)$$

where S is the log-likelihood ratio of two events, $x_1, \dots, x_n$ is a sequence of tokens, p is the pun word and d is the local window size. Finally, a unified metric is defined as a ratio of local-global surprisal to quantify the success of the pun generation. Some details of the formulas are provided in the Appendix B.

5.1.2 Diversity

Unusualness. Given the uniqueness of puns, *unusualness* measures based on the normalised log probabilities from language models are also utilised for pun evaluation (He et al., 2019; Pauls and Klein, 2012), which is formulated as follows:

$$Unusualness \stackrel{\text{def}}{=} -\frac{1}{n} \log \left(\frac{p(x_1, \dots, x_n)}{\prod_{i=1}^n p(x_i)} \right) \quad (4)$$

Paper	Automatic Evaluation							Human Evaluation					
	PPLs.	D1&2.	Succ.	Ambi.	Dist	Surp.	Unus.	Succ.	Funn.	Flun.	Info.	Cohe.	Read.
(Yu et al., 2018)	✓	✓	×	×	×	×	×	×	×	✓	×	✓	✓
(He et al., 2019)	×	×	×	✓	✓	✓	✓	✓	✓	✓	×	×	×
(Luo et al., 2019)	×	✓	×	×	×	×	✓	✓	×	✓	×	×	×
(Yu et al., 2020)	×	✓	×	×	×	×	×	✓	✓	✓	×	×	×
(Diao et al., 2020)	×	✓	×	×	×	×	✓	✓	×	✓	×	×	×
(Mittal et al., 2022)	×	✓	×	×	×	×	×	✓	✓	×	×	✓	×
(Tian et al., 2022)	×	×	×	✓	✓	✓	×	✓	✓	×	✓	×	×
(Sun et al., 2022b)	×	×	✓	×	×	×	×	✓	×	×	×	×	×
(Zeng et al., 2024)	×	✓	×	✓	×	×	×	-	-	-	-	-	-
(Chen et al., 2024)	×	✓	✓	×	×	×	×	✓	×	×	×	×	×

Table 2: Main methods for automatic and human evaluation of pun generation. PPLs., D1&2., Succ., Ambi., Dist., Surp., and Unus. denote the metrics of Perplexity Score, Dist-1 & Dist-2, Structure Succ., Ambiguity, Distinctiveness, Surprisal, and Unusualness, respectively. Similarly, Succ., Funn., Gram., Flun., Info., Cohe., and Read. represent Success, Funniness, Grammar, Fluency, Informativeness, Coherence, and Readability. ✓ indicates metrics that are used, while × indicates metrics that are not used. The symbol “-” signifies that the method is not applicable to this evaluation.

where $p(x_1, \dots, x_n)$ and $p(x_i)$ are the joint and independent probabilities, respectively. A higher metric result suggests the presence of uncommon collocations, innovative sentence structures, and other linguistic features, aligning with the characteristics of puns.

Dist-1 & Dist-2. Dist-1 and Dist-2 focus on the diversity of words and phrases in the generated text (Li et al., 2015), which calculates the proportion of unique n-grams to the total number of n-grams, as formulated Dist-1, for example:

$$\text{Dist-1} = \frac{\text{unique unigrams}}{\text{total generated words}} \quad (5)$$

where a higher Dist-1 score indicates greater diversity in the generated sentences, whereas a lower score suggests more generic and repetitive text. The Dist-2 formulation is shown as Appendix B.

5.1.3 Fluency

Perplexity score (Jelinek et al., 1977). This score evaluates whether the generated puns are natural and fluent. In practice, some studies (Yu et al., 2018) quantified by using the generative language model, formally described as follows:

$$\text{perplexity} = \exp \left(-\frac{1}{N} \sum_{i=1}^N \log P(x_i | x_{<i}) \right) \quad (6)$$

where $P(x_i | x_{<i})$ is the probability of the i -th token of a pun, given the sequence of tokens ahead.

Structure Succ. The evaluation measures the rate of contextual word and pun word integration, specifically the proportion of successful inclusion

of pun words in the generated puns, formally shown as follows:

$$\text{Succ} = \frac{t_{\text{correct}}}{T} \times 100\% \quad (7)$$

where t_{correct} is the number of generated puns with correctly included pun words and T is the total number of generated puns.

5.2 Human Evaluation

In the task of pun generation, since puns are a creative form of language (Yu et al., 2020), human evaluation is essential and intuitively assesses the quality of the generated puns. The primary evaluation metrics are: **Success** recognises whether the generated sentence qualifies as a successful pun based on the definition from (Miller et al., 2017); **Funniness** evaluates the humour and comedic quality of the generated sentences; **Fluency** shows whether the sentence is grammatically correct and flows naturally; **Informativeness** rates whether the generated sentences effectively convey meaningful and specific information; **Coherence** indicates whether the generated sentence is coherency and given senses are suitable in a sentence; **Readability** indicates whether the sentence is easy to understand semantically.

Most studies utilize the Likert Scale (Joshi et al., 2015) to assess the metrics. This commonly used psychological measurement method and relies numerical scales within a specific range to evaluate a given objective (Alhazmi et al., 2024). For example, Mittal et al. (2022) utilized a Likert scale ranging from 1 (not at all) to 5 (extremely) to rate the funniness and coherence of puns. In particular, for

success metrics, some studies adopt a binary classification method in which evaluators determine whether the generated pun is successful by selecting *True* or *False* (Tian et al., 2022; Sun et al., 2022b; Chen et al., 2024).

With the development of LLMs, Chen et al. (2024) conducted a human A/B test, asking annotators to compare paired puns generated by their methods and ChatGPT and select more humorous puns. Since GPT-4’s evaluations aligned closely with those of human reviewers (Liang et al., 2024), Zeng et al. (2024) replaced human reviewers with GPT-4 to assess the metrics of readability, funniness, and coherence.

6 Challenges and Future Directions

This section outlines the challenges and explores potential directions for future work.

6.1 Multilingual Research

With advancements in pun generation research, the majority of studies focus primarily on English, as shown in Table 1, while studies on puns in other languages remain limited. Linguistically, different languages employ distinct mechanisms to create puns. For example, ideographic or mixed languages, such as Chinese and Japanese, tend to emphasize multi-layered pun construction (Shao et al., 2013). Therefore, cross-language pun generation can also serve as a potential future work. Building on previous cross-linguistic research, using parallel data, including word-parallel (Zhao et al., 2020; Alqahtani et al., 2021) and sentence-parallel (Reimers and Gurevych, 2020; Heffernan et al., 2022), can be utilized to achieve targeted alignment of pun words. Additionally, some pioneering works can capture phonological and semantic puns through advanced learning approaches such as contrastive learning (Hu et al., 2024), modify pre-training schemes (Clark, 2020) and adapter tuning (Pfeiffer et al., 2020; Parović et al., 2022).

6.2 Multi-Modal Information

Multimodal information enables a more reliable understanding of the world (Stein, 1993), and incorporating multiple modalities into tasks can enhance the quality of pun generation. Although previous studies have introduced some multimodal evaluations and datasets (Zhang et al., 2024; Chung et al., 2024), few have specifically focused on the generation of multimodal puns. One potential method is shared representation (Ngiam et al.,

2011), which involves integrating complementary information from different modalities to learn higher-performance representations (Lahat et al., 2015). For example, automatic speech recognition (Malik et al., 2021) can be leveraged to enhance homophonic puns. Another direction is to translate puns between modalities, i.e., cross-modal generation (Suzuki and Matsuo, 2022), including text-to-image (Zhang et al., 2023a), image-to-text (He and Deng, 2017), text-to-speech (Zhang et al., 2023b) and speech-to-text (Shadiev et al., 2014; Fortuna and Nunes, 2018)

6.3 PLMs Prompting Design

While prompt engineering has proven effective in enhancing text generation capabilities of LLMs (Liu et al., 2023), current research still faces significant limitations in generating puns, such as an over-reliance on overly simplistic or single-faceted prompts. Chain-of-thought prompting is a powerful technique that significantly improves the reasoning capabilities of LLMs (Wei et al., 2022). Therefore, the quality of pun generation can be enhanced by transferring CoT technique from other fields, such as using iterative bootstrapping (Sun et al., 2023), knowledge enhancement (Dhuliawala et al., 2023; He et al., 2024), question decomposition (Trivedi et al., 2022) and self-ensemble (Yin et al., 2024). Furthermore, the result can be improved by optimizing CoT’s prompt construction, encompassing semi-automatic prompting (Shum et al., 2023) and automatic prompting (Zhang et al., 2022), as well as exploring diverse topological variants (Chu et al., 2024), such as chain structures (Olausson et al., 2023), tree structures (Ning et al., 2023), and graph structures (Besta et al., 2024).

7 Conclusion

In this paper, we present the first comprehensive survey on pun generation tasks, including phonetic, graphic and visual puns. We classify and conduct a thorough analysis of the datasets used in pun research, review previous approaches to pun generation, discuss existing methods, as well as summarize the evaluation metrics for pun generation. Furthermore, we highlight the challenges and future directions, offering valuable insights for researchers interested in pun generation. To enhance research in pun generation, this paper also plans to provide a continuously updated reading list available on a GitHub repository.

Limitations

Although we have attempted to extensively analyse the existing literature on pun generation, some works may still be missed due to variations in search keywords. Furthermore, our exploration of other categories of puns is limited, such as recursive puns and antanaclasis, as we encountered challenges while searching for them, which may be influenced by the relatively low attention they have received in the research community. Finally, due to the rapid development of the research field, this study does not cover the entire historical scope nor the latest advancements following the survey. However, our work represents the first comprehensive survey on pun generation, including datasets, methods, evaluation, challenges and potential directions, making it a valuable resource for scholars in this field.

References

- Elaf Alhazmi, Quan Z. Sheng, W. Zhang, Munazza Zaib, and Ahoud Abdulrahmn F. Alhazmi. 2024. [Distractor generation in multiple-choice tasks: A survey of methods, datasets, and evaluation](#). In *Conference on Empirical Methods in Natural Language Processing*.
- Arwa I Alhussain and Aqil M Azmi. 2021. Automatic story generation: A survey of approaches. *ACM Computing Surveys (CSUR)*, 54(5):1–38.
- Sawsan Alqahtani, Garima Lalwani, Yi Zhang, Salvatore Romeo, and Saab Mansour. 2021. Using optimal transport as alignment objective for fine-tuning multilingual contextualized embeddings. *arXiv preprint arXiv:2110.02887*.
- 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.
- Issa Annamradnejad and Gohar Zoghi. 2020. Colbert: Using bert sentence embedding in parallel neural networks for computational humor. *arXiv preprint arXiv:2004.12765*.
- Anthropic. 2024. [The claude 3 model family: Opus, sonnet, haiku](#).
- Salvatore Attardo. 2009. *Linguistic theories of humor*. Walter de Gruyter.
- Nancy D Bell, Scott Crossley, and Christian F Hempelmann. 2011. Wordplay in church marquees.
- Brendan Bena and Jugal Kalita. 2020. Introducing aspects of creativity in automatic poetry generation. *arXiv preprint arXiv:2002.02511*.

- Maciej Besta, Nils Blach, Ales Kubicek, Robert Gerstenberger, Michal Podstawski, Lukas Gianinazzi, Joanna Gajda, Tomasz Lehmann, Hubert Niewiadomski, Piotr Nyczyk, et al. 2024. Graph of thoughts: Solving elaborate problems with large language models. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 17682–17690.
- James Betker, Gabriel Goh, Li Jing, Tim Brooks, Jianfeng Wang, Linjie Li, Long Ouyang, Juntang Zhuang, Joyce Lee, Yufei Guo, et al. 2023. Improving image generation with better captions. *Computer Science*. <https://cdn.openai.com/papers/dall-e-3.pdf>, 2(3):8.
- Kim Binsted. 1996. Machine humour: An implemented model of puns.
- Kim Binsted and Graeme Ritchie. 1994. *An implemented model of punning riddles*. University of Edinburgh, Department of Artificial Intelligence.
- Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared 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, Jeff Wu, Clemens Winter, Christopher Hesse, Mark Chen, Eric Sigler, Ma teusz 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](#). *ArXiv*, abs/2005.14165.
- Ronald Carter. 2015. *Language and creativity: The art of common talk*. Routledge.
- Yang Chen, Chong Yang, Tu Hu, Xinhao Chen, Man Lan, Li Cai, Xinlin Zhuang, Xuan Lin, Xin Lu, and Aimin Zhou. 2024. Are u a joke master? pun generation via multi-stage curriculum learning towards a humor llm. In *Findings of the Association for Computational Linguistics ACL 2024*, pages 878–890.
- Zheng Chu, Jingchang Chen, Qianglong Chen, Weijiang Yu, Tao He, Haotian Wang, Weihua Peng, Ming Liu, Bing Qin, and Ting Liu. 2024. Navigate through enigmatic labyrinth a survey of chain of thought reasoning: Advances, frontiers and future. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1173–1203.
- Jiwan Chung, Seungwon Lim, Jaehyun Jeon, Seungbeen Lee, and Youngjae Yu. 2024. Can visual language models resolve textual ambiguity with visual cues? let visual puns tell you! *arXiv preprint arXiv:2410.01023*.
- K Clark. 2020. Electra: Pre-training text encoders as discriminators rather than generators. *arXiv preprint arXiv:2003.10555*.
- Mathieu Dehouck and Marine Delaborde. 2025. Rule-based approaches to the automatic generation of puns based on given names in french. In *Proceedings of*

746	<i>the 1st Workshop on Computational Humor (CHum)</i> ,	Albert Gatt and Emiel Krahmer. 2018. Survey of the	799
747	pages 18–22.	state of the art in natural language generation: Core	800
748	Dirk Delabastita. 2016. <i>Traductio: Essays on punning</i>	tasks, applications and evaluation. <i>Journal of Artificial</i>	801
749	<i>and translation</i> . Routledge.	<i>Intelligence Research</i> , 61:65–170.	802
750	Jacob Devlin, Ming-Wei Chang, Kenton Lee, and	Mengshi Ge, Rui Mao, and Erik Cambria. 2023. A	803
751	Kristina Toutanova. 2019. BERT: Pre-training of	survey on computational metaphor processing tech-	804
752	deep bidirectional transformers for language under-	niques: From identification, interpretation, genera-	805
753	standing . In <i>Proceedings of the 2019 Conference of</i>	tion to application. <i>Artificial Intelligence Review</i> ,	806
754	<i>the North American Chapter of the Association for</i>	56(Suppl 2):1829–1895.	807
755	<i>Computational Linguistics: Human Language Tech-</i>	Lena Gieseke, Paul Asente, Radomír Měch, Bedrich	808
756	<i>nologies, Volume 1 (Long and Short Papers)</i> , pages	Benes, and Martin Fuchs. 2021. A survey of control	809
757	4171–4186, Minneapolis, Minnesota. Association for	mechanisms for creative pattern generation. In <i>Com-</i>	810
758	Computational Linguistics.	<i>puter Graphics Forum</i> , volume 40, pages 585–609.	811
759	Shehzaad Dhuliawala, Mojtaba Komeili, Jing Xu,	Wiley Online Library.	812
760	Roberta Raileanu, Xian Li, Asli Celikyilmaz, and Ja-	Meri Giorgadze. 2014. Linguistic features of pun, its	813
761	son Weston. 2023. Chain-of-verification reduces hal-	typology and classification. <i>European Scientific Jour-</i>	814
762	lucination in large language models. <i>arXiv preprint</i>	<i>nal</i> .	815
763	<i>arXiv:2309.11495</i> .	Ian Goodfellow, Jean Pouget-Abadie, Mehdi Mirza,	816
764	Yufeng Diao, Liang Yang, Xiaochao Fan, Yonghe Chu,	Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron	817
765	Di Wu, Shaowu Zhang, and Hongfei Lin. 2020.	Courville, and Yoshua Bengio. 2014. Generative	818
766	Afpun-gan: Ambiguity-fluency generative adversar-	adversarial nets. <i>Advances in neural information</i>	819
767	ial network for pun generation. In <i>Natural Language</i>	<i>processing systems</i> , 27.	820
768	<i>Processing and Chinese Computing: 9th CCF In-</i>	Gemini Team Google. 2023. Gemini: A Family of	821
769	<i>ternational Conference, NLPCC 2020, Zhengzhou,</i>	Highly Capable Multimodal Models . <i>arXiv e-prints</i> ,	822
770	<i>China, October 14–18, 2020, Proceedings, Part I 9</i> ,	arXiv:2312.11805.	823
771	pages 604–616. Springer.	He He, Nanyun Peng, and Percy Liang. 2019.	824
772	Elmira Djafarova. 2008. Why do advertisers use puns?	Pun generation with surprise. <i>arXiv preprint</i>	825
773	a linguistic perspective. <i>Journal of Advertising Re-</i>	<i>arXiv:1904.06828</i> .	826
774	<i>search</i> , 48(2):267–275.	Xiaodong He and Li Deng. 2017. Deep learning	827
775	Pawel Dybala, Michal Ptaszynski, Shinsuke Higuchi,	for image-to-text generation: A technical overview.	828
776	Rafal Rzepka, and Kenji Araki. 2008. Humor	<i>IEEE Signal Processing Magazine</i> , 34(6):109–116.	829
777	prevails!-implementing a joke generator into a con-	Zhiwei He, Tian Liang, Wenxiang Jiao, Zhuosheng	830
778	versational system. In <i>AI 2008: Advances in Artificial</i>	Zhang, Yujia Yang, Rui Wang, Zhaopeng Tu, Shum-	831
779	<i>Intelligence: 21st Australasian Joint Conference</i>	ing Shi, and Xing Wang. 2024. Exploring human-	832
780	<i>on Artificial Intelligence Auckland, New Zealand, De-</i>	like translation strategy with large language models.	833
781	<i>cember 1-5, 2008. Proceedings 21</i> , pages 214–225.	<i>Transactions of the Association for Computational</i>	834
782	Springer.	<i>Linguistics</i> , 12:229–246.	835
783	Mohamad Elzohbi and Richard Zhao. 2023. Creative	Kevin Heffernan, Onur Çelebi, and Holger Schwenk.	836
784	data generation: A review focusing on text and poetry.	2022. Bitext mining using distilled sentence repre-	837
785	<i>arXiv preprint arXiv:2305.08493</i> .	sentations for low-resource languages. <i>arXiv preprint</i>	838
786	Paula Fortuna and Sérgio Nunes. 2018. A survey on	<i>arXiv:2205.12654</i> .	839
787	automatic detection of hate speech in text. <i>ACM</i>	Christian F Hempelmann. 2003. <i>Paronomasic puns:</i>	840
788	<i>Computing Surveys (CSUR)</i> , 51(4):1–30.	<i>Target recoverability towards automatic generation</i> .	841
789	Vaishali Ganganwar, Manvinder, Mohit Singh, Priyank	Ph.D. thesis, Purdue University.	842
790	Patil, and Saurabh Joshi. 2024. Sarcasm and humor	Bryan Anthony Hong and Ethel Ong. 2009. Automati-	843
791	detection in code-mixed hindi data: A survey. In	cally extracting word relationships as templates for	844
792	<i>International Conference on Computing and Machine</i>	pun generation. In <i>Proceedings of the Workshop on</i>	845
793	<i>Learning</i> , pages 453–469. Springer.	<i>Computational Approaches to Linguistic Creativity</i> ,	846
794	Leo Gao, Stella Biderman, Sid Black, Laurence Gold-	pages 24–31.	847
795	ing, Travis Hoppe, Charles Foster, Jason Phang, Ho-	Haigen Hu, Xiaoyuan Wang, Yan Zhang, Qi Chen, and	848
796	race He, Anish Thite, Noa Nabeshima, et al. 2020.	Qiu Guan. 2024. A comprehensive survey on con-	849
797	The pile: An 800gb dataset of diverse text for lan-	trastive learning. <i>Neurocomputing</i> , page 128645.	850
798	guage modeling. <i>arXiv preprint arXiv:2101.00027</i> .		

957	Jiquan Ngiam, Aditya Khosla, Mingyu Kim, Juhan	Rafael Rafailov, Archit Sharma, Eric Mitchell, Christo-	1009
958	Nam, Honglak Lee, Andrew Y Ng, et al. 2011. Mul-	pher D Manning, Stefano Ermon, and Chelsea Finn.	1010
959	timodal deep learning. In <i>ICML</i> , volume 11, pages	2024. Direct preference optimization: Your language	1011
960	689–696.	model is secretly a reward model. <i>Advances in Neu-</i>	1012
		<i>ral Information Processing Systems</i> , 36.	1013
961	Anton Nijholt, Andreea Niculescu, Alessandro Vali-	Colin Raffel, Noam Shazeer, Adam Roberts, Katherine	1014
962	tutti, and Rafael E Banchs. 2017. Humor in human-	Lee, Sharan Narang, Michael Matena, Yanqi Zhou,	1015
963	computer interaction: a short survey. In <i>16th IFIP</i>	Wei Li, and Peter J Liu. 2020. Exploring the lim-	1016
964	<i>TC13 International Conference on Human-Computer</i>	its of transfer learning with a unified text-to-text	1017
965	<i>Interaction, INTERACT 2017</i> , pages 192–214. Indian	transformer. <i>Journal of machine learning research</i> ,	1018
966	Institute of Technology Madras.	21(140):1–67.	1019
967	Xuefei Ning, Zinan Lin, Zixuan Zhou, Zifu Wang,	Sunny Rai and Shampa Chakraverty. 2020. A survey on	1020
968	Huazhong Yang, and Yu Wang. 2023. Skeleton-of-	computational metaphor processing. <i>ACM Comput-</i>	1021
969	thought: Large language models can do parallel de-	<i>ing Surveys (CSUR)</i> , 53(2):1–37.	1022
970	coding. <i>Proceedings ENLSP-III</i> .		
971	Theo X Olausson, Alex Gu, Benjamin Lipkin, Cede-	C Ramakristanaiah, P Namratha, Rajendra Kumar	1023
972	gao E Zhang, Armando Solar-Lezama, Joshua B	Ganiya, and Midde Ranjit Reddy. 2021. A survey	1024
973	Tenenbaum, and Roger Levy. 2023. Linc: A neu-	on humor detection methods in communications. In	1025
974	rosymbolic approach for logical reasoning by com-	<i>2021 Fifth International Conference on I-SMAC (IoT</i>	1026
975	binning language models with first-order logic provers.	<i>in Social, Mobile, Analytics and Cloud)(I-SMAC)</i> ,	1027
976	<i>arXiv preprint arXiv:2310.15164</i> .	pages 668–674. IEEE.	1028
977	OpenAI. 2023a. Gpt-3.5-turbo. https://platform.	Oleksandr Rebrii, Inna Rebrii, and Olha Pieshkova.	1029
978	openai.com/docs/models/gpt-3-5-turbo . Ac-	2022. When words and images play together in	1030
979	cessed: 2025-01-05.	a multimodal pun: From creation to translation.	1031
		<i>Lublin Studies in Modern Languages and Literature</i> ,	1032
980	OpenAI. 2023b. Gpt-4 and gpt-4 turbo.	46(2):85–97.	1033
981	https://platform.openai.com/docs/models/	Nils Reimers and Iryna Gurevych. 2020. Mak-	1034
982	gpt-4-and-gpt-4-turbo . Accessed: 2025-01-05.	ing monolingual sentence embeddings multilin-	1035
		gual using knowledge distillation. <i>arXiv preprint</i>	1036
983	Marinela Parović, Goran Glavaš, Ivan Vulić, and Anna	<i>arXiv:2004.09813</i> .	1037
984	Korhonen. 2022. Bad-x: Bilingual adapters improve	Graeme Ritchie. 2005. Computational mechanisms for	1038
985	zero-shot cross-lingual transfer. In <i>Proceedings of</i>	pun generation. In <i>Proceedings of the Tenth Euro-</i>	1039
986	<i>the 2022 Conference of the North American Chap-</i>	<i>pean Workshop on Natural Language Generation</i>	1040
987	<i>ter of the Association for Computational Linguistics:</i>	<i>(ENLG-05)</i> .	1041
988	<i>Human Language Technologies</i> , pages 1791–1799.		
989	Adam Pauls and Dan Klein. 2012. Large-scale syntactic	Brian Roark, Lawrence Wolf-Sonkin, Christo Kirov,	1042
990	language modeling with treelets. In <i>Proceedings</i>	Sabrina J Mielke, Cibu Johny, Isin Demirsahin, and	1043
991	<i>of the 50th Annual Meeting of the Association for</i>	Keith Hall. 2020. Processing south asian languages	1044
992	<i>Computational Linguistics (Volume 1: Long Papers)</i> ,	written in the latin script: the dakshina dataset. <i>arXiv</i>	1045
993	pages 959–968.	<i>preprint arXiv:2007.01176</i> .	1046
994	Jeffrey Pennington, Richard Socher, and Christopher D	Yash Raj Sarrof. 2025. Homophonic pun generation in	1047
995	Manning. 2014. Glove: Global vectors for word rep-	code mixed hindi english. In <i>Proceedings of the 1st</i>	1048
996	resentation. In <i>Proceedings of the 2014 conference</i>	<i>Workshop on Computational Humor (CHum)</i> , pages	1049
997	<i>on empirical methods in natural language processing</i>	23–31.	1050
998	<i>(EMNLP)</i> , pages 1532–1543.		
999	Jonas Pfeiffer, Ivan Vulić, Iryna Gurevych, and Sebas-	Rustam Shadiev, Wu-Yuin Hwang, Nian-Shing Chen,	1051
1000	tian Ruder. 2020. Mad-x: An adapter-based frame-	and Yueh-Min Huang. 2014. Review of speech-	1052
1001	work for multi-task cross-lingual transfer. <i>arXiv</i>	to-text recognition technology for enhancing learn-	1053
1002	<i>preprint arXiv:2005.00052</i> .	ing. <i>Journal of Educational Technology & Society</i> ,	1054
		17(4):65–84.	1055
1003	Alec Radford. 2018. Improving language understanding	Sakib Shahriar. 2022. Gan computers generate arts? a	1056
1004	by generative pre-training.	survey on visual arts, music, and literary text genera-	1057
		tion using generative adversarial network. <i>Displays</i> ,	1058
1005	Alec Radford, Jeffrey Wu, Rewon Child, David Luan,	73:102237.	1059
1006	Dario Amodei, Ilya Sutskever, et al. 2019. Language	Qing Chen Shao, Zhen Zhen Wang, and Zhi Jie Hao.	1060
1007	models are unsupervised multitask learners. <i>OpenAI</i>	2013. Contrastive studies of pun in figures of speech.	1061
1008	<i>blog</i> , 1(8):9.	<i>Advanced Materials Research</i> , 756:4721–4727.	1062

1063	Yikang Shen, Shawn Tan, Alessandro Sordoni, and Aaron Courville. 2018. Ordered neurons: Integrating tree structures into recurrent neural networks. <i>arXiv preprint arXiv:1810.09536</i> .	1113
1064		1114
1065		1115
1066		1116
1067	KaShun Shum, Shizhe Diao, and Tong Zhang. 2023. Automatic prompt augmentation and selection with chain-of-thought from labeled data. <i>arXiv preprint arXiv:2302.12822</i> .	1117
1068		1118
1069		1119
1070		1120
1071	Robert E Smith, Jiemiao Chen, and Xiaojing Yang. 2008. The impact of advertising creativity on the hierarchy of effects. <i>Journal of advertising</i> , 37(4):47–62.	1121
1072		1122
1073		1123
1074	Włodzimierz Sobkowiak. 1991. <i>Metaphonology of English paronomasic puns</i> . Lang.	1124
1075		1125
1076	BE Stein. 1993. <i>The Merging of the Senses</i> . MIT Press.	1126
1077	Jiao Sun, Anjali Narayan-Chen, Shereen Oraby, Alessandra Cervone, Tagyoung Chung, Jing Huang, Yang Liu, and Nanyun Peng. 2022a. Expunations: Augmenting puns with keywords and explanations. <i>arXiv preprint arXiv:2210.13513</i> .	1127
1078		1128
1079		1129
1080		1130
1081		1131
1082	Jiao Sun, Anjali Narayan-Chen, Shereen Oraby, Shuyang Gao, Tagyoung Chung, Jing Huang, Yang Liu, and Nanyun Peng. 2022b. Context-situated pun generation. <i>arXiv preprint arXiv:2210.13522</i> .	1132
1083		1133
1084		1134
1085		1135
1086	Jiashuo Sun, Yi Luo, Yeyun Gong, Chen Lin, Yelong Shen, Jian Guo, and Nan Duan. 2023. Enhancing chain-of-thoughts prompting with iterative bootstrapping in large language models. <i>arXiv preprint arXiv:2304.11657</i> .	1136
1087		1137
1088		1138
1089		1139
1090		1140
1091	I Sutskever. 2014. Sequence to sequence learning with neural networks. <i>arXiv preprint arXiv:1409.3215</i> .	1141
1092		1142
1093	Masahiro Suzuki and Yutaka Matsuo. 2022. A survey of multimodal deep generative models. <i>Advanced Robotics</i> , 36(5-6):261–278.	1143
1094		1144
1095		1145
1096	Yufei Tian, Divyanshu Sheth, and Nanyun Peng. 2022. A unified framework for pun generation with humor principles. <i>arXiv preprint arXiv:2210.13055</i> .	1146
1097		1147
1098		1148
1099	Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. 2023. Llama 2: Open foundation and fine-tuned chat models. <i>arXiv preprint arXiv:2307.09288</i> .	1149
1100		1150
1101		1151
1102		1152
1103		1153
1104		1154
1105	Harsh Trivedi, Niranjan Balasubramanian, Tushar Khot, and Ashish Sabharwal. 2022. Interleaving retrieval with chain-of-thought reasoning for knowledge-intensive multi-step questions. <i>arXiv preprint arXiv:2212.10509</i> .	1155
1106		1156
1107		1157
1108		1158
1109		1159
1110	Bradley Tyler, Katherine Wilsdon, and Paul M Bodily. 2020. Computational humor: Automated pun generation. In <i>ICCC</i> , pages 181–184.	1160
1111		1161
1112		1162
	Alessandro Valitutti, Oliviero Stock, and Carlo Strapparava. 2009. Graphlaugh: a tool for the interactive generation of humorous puns. In <i>2009 3rd International Conference on Affective Computing and Intelligent Interaction and Workshops</i> , pages 1–2. IEEE.	1163
		1164
	Alessandro Valitutti, Hannu Toivonen, Antoine Doucet, and Jukka M Toivanen. 2013. “let everything turn well in your wife”: generation of adult humor using lexical constraints. In <i>Proceedings of the 51st Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)</i> , pages 243–248.	1165
		1166
	Margot Van Mulken, Renske Van Enschoot-van Dijk, and Hans Hoeken. 2005. Puns, relevance and appreciation in advertisements. <i>Journal of pragmatics</i> , 37(5):707–721.	1167
		1168
	A Vaswani. 2017. Attention is all you need. <i>Advances in Neural Information Processing Systems</i> .	1169
		1170
	Christopher Venour. 2000. <i>The computational generation of a class of pun</i> . Queen’s University.	1171
		1172
	Guan Wang, Sijie Cheng, Xianyu Zhan, Xiangang Li, Sen Song, and Yang Liu. 2023. <i>Openchat: Advancing open-source language models with mixed-quality data</i> . <i>ArXiv</i> , abs/2309.11235.	1173
		1174
	Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, et al. 2022. Chain-of-thought prompting elicits reasoning in large language models. <i>Advances in neural information processing systems</i> , 35:24824–24837.	1175
		1176
	Zhijun Xu, Siyu Yuan, Lingjie Chen, and Deqing Yang. 2024. “a good pun is its own reword”: Can large language models understand puns? <i>arXiv preprint arXiv:2404.13599</i> .	1177
		1178
	Aiyuan Yang, Bin Xiao, Bingning Wang, Borong Zhang, Ce Bian, Chao Yin, Chenxu Lv, Da Pan, Dian Wang, Dong Yan, et al. 2023. Baichuan 2: Open large-scale language models. <i>arXiv preprint arXiv:2309.10305</i> .	1179
		1180
	Di Yi Yang, Alon Lavie, Chris Dyer, and Eduard Hovy. 2015. Humor recognition and humor anchor extraction. In <i>Proceedings of the 2015 conference on empirical methods in natural language processing</i> , pages 2367–2376.	1181
		1182
	Zhangyue Yin, Qiushi Sun, Qipeng Guo, Zhiyuan Zeng, Xiaonan Li, Tianxiang Sun, Cheng Chang, Qinyuan Cheng, Ding Wang, Xiaofeng Mou, et al. 2024. Aggregation of reasoning: A hierarchical framework for enhancing answer selection in large language models. <i>arXiv preprint arXiv:2405.12939</i> .	1183
		1184
	Zhiwei Yu, Jiwei Tan, and Xiaojun Wan. 2018. A neural approach to pun generation. In <i>Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 1650–1660.	1185
		1186

Zhiwei Yu, Hongyu Zang, and Xiaojun Wan. 2020. **Homophonic pun generation with lexically constrained rewriting**. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 2870–2876, Online. Association for Computational Linguistics.

Jingjie Zeng, Liang Yang, Jiahao Kang, Yufeng Diao, Zhihao Yang, and Hongfei Lin. 2024. “barking up the right tree”, a gan-based pun generation model through semantic pruning. In *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, pages 2119–2131.

Chenshuang Zhang, Chaoning Zhang, Mengchun Zhang, and In So Kweon. 2023a. Text-to-image diffusion models in generative ai: A survey. *arXiv preprint arXiv:2303.07909*.

Chenshuang Zhang, Chaoning Zhang, Sheng Zheng, Mengchun Zhang, Maryam Qamar, Sung-Ho Bae, and In So Kweon. 2023b. A survey on audio diffusion models: Text to speech synthesis and enhancement in generative ai. *arXiv preprint arXiv:2303.13336*.

Tuo Zhang, Tiantian Feng, Yibin Ni, Mengqin Cao, Ruying Liu, Katharine Butler, Yanjun Weng, Mi Zhang, Shrikanth S Narayanan, and Salman Avestimehr. 2024. Creating a lens of chinese culture: A multimodal dataset for chinese pun rebus art understanding. *arXiv preprint arXiv:2406.10318*.

Zhuosheng Zhang, Aston Zhang, Mu Li, and Alex Smola. 2022. Automatic chain of thought prompting in large language models. *arXiv preprint arXiv:2210.03493*.

Wei Zhao, Steffen Eger, Johannes Bjerva, and Isabelle Augenstein. 2020. Inducing language-agnostic multilingual representations. *arXiv preprint arXiv:2008.09112*.

Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhonghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric P. Xing, Haotong Zhang, Joseph E. Gonzalez, and Ion Stoica. 2023. **Judging llm-as-a-judge with mt-bench and chatbot arena**. *ArXiv*, abs/2306.05685.

Wei Zheng and Xiaolu Wang. 2023. Humor experience facilitates ongoing cognitive tasks: Evidence from pun comprehension. *Frontiers in Psychology*, 14:1127275.

Yukun Zhu. 2015. Aligning books and movies: Towards story-like visual explanations by watching movies and reading books. *arXiv preprint arXiv:1506.06724*.

A Pun Categories

We outline the characteristics of different types of puns for clearer differentiation, including phonetic,

graphic, meaning, and example, as shown in Table 3. "Same", "similar" and "different" respectively indicate whether the pun word and its substitute word same, similar, or different in phonic, graphic and meaning.

B Additional Evaluation

In this section, we supplemented additional details regarding the evaluation metrics.

Suprisal. Based on (He et al., 2019), the pun word w^p is more surprising relative to its alternative word w^a in the local context, while is less in the global context. Therefore, S_{ratio} is defined as a ratio to balance the metric:

$$S_{ratio} \begin{cases} -1, & S_{local} < 0 \text{ or } S_{global} < 0, \\ S_{local}/S_{global}, & \text{otherwise.} \end{cases} \quad (8)$$

where S_{local} and S_{global} are local surprisal and global surprisal, respectively. A higher value of S_{ratio} indicates a better-quality pun.

Dist-1 & Dist-2. We supplemented the formulation of Dist-2 here:

$$\text{Dist-2} = \frac{\text{unique bigrams}}{\text{total generated bigrams}} \quad (9)$$

where a higher Dist-2 score indicates greater diversity in the generated sentences, whereas a lower score suggests more generic and repetitive text.

C Dataset

The pun dataset for different types are summarized in Table 4. We list the datasets in five dimensions:

- The type of puns.
- The source of the datasets.
- The total number of the datasets.
- The language of the datasets.
- Is the dataset publicly available?

D Paper Collection

This section outlines the approach that we used to collect relevant papers in this survey. We initially searched for the keywords "pun research", "computational humour", and "pun dataset" on arXiv and Google Scholar, identifying a total of around 150 publications. Then, we filtered papers that specifically focused on pun generation, resulting in approximately 30 papers. Subsequently, we applied

Type	Phonic	Graphic	Meaning	Example
Homophonic Puns	Similar	Different	Different	Dentists don't like a hard day at the <u>orifice</u> (office).
Heterographic Puns	Same	Different	Different	Life is a puzzle, look here for the missing <u>peace</u> (piece).
Homographic Puns	Same	Same	Different	Always trust a glue salesman. They tend to <u>stick</u> to their word.
Visual Puns	N/A	N/A	Different	

Table 3: List of pun categories. N/A indicates that the element is not applicable.

Dataset	Type	Source	Corpus (C)	Language	Availability
Paron(Sobkowiak, 1991)	heg	Advertisements	3,850	English	✓
Paron-edit(Jaech et al., 2016)	heg	(Sobkowiak, 1991)	1,182	English	✗
Church(Bell et al., 2011)	hog	Church	373	English	✗
Pun-Yang(Yang et al., 2015)	N/A	Website	2,423	English	✓
Pun-Kao(Kao et al., 2016)	hop	Website	435	English	✓
Puns (Jaech et al., 2016)	N/A	Website	75	English	✗
SemEval (Miller et al., 2017)	hog&heg	Experts	2,878	English	✓
SemEval-P (Miller et al., 2017)	hog	Experts	1,607	English	✓
SemEval-G (Miller et al., 2017)	heg	Experts	1,271	English	✓
ExPUNations (Sun et al., 2022a)	hog&heg	(Miller et al., 2017)	1,999	English	✓
CUP (Sun et al., 2022b)	hog&heg	(Miller et al., 2017)	2,396	English	✓
ChinesePun (Chen et al., 2024)	hop&hog	Website	2,106	Chinese	✓
ChinesePun-P (Chen et al., 2024)	hop	Website	1,049	Chinese	✓
ChinesePun-G (Chen et al., 2024)	hog	Website	1,057	Chinese	✓
Pun Rebus Art (Zhang et al., 2024)	visual	Museum	1,011	Multi-language	✓
UNPIE (Chung et al., 2024)	hog&heg	(Miller et al., 2017)	1,000	Multi-language	✓
UNPIE-P (Chung et al., 2024)	hog	(Miller et al., 2017)	500	Multi-language	✓
UNPIE-G (Chung et al., 2024)	heg	(Miller et al., 2017)	500	Multi-language	✓

Table 4: List of pun datasets. Hog, hop, heg and visual denote the types of homographic puns, homophonic puns, heterographic puns and visual puns, respectively. N/A indicates that the elements are not mentioned in the original paper.

System	Type	Task	Language
JAPE (Binsted and Ritchie, 1994)	heg & hog	Question-Answer	English
HCPP (Venour, 2000)	hop	Text Generation	English
WISCRAIC (McKay, 2002)	heg	Text Generation	English
PUNDA (Dybala et al., 2008)	heg & hog	Dialogue	Japanese
STANDUP (Manurung et al., 2008)	hop	Dialogue	English
T-PEG (Hong and Ong, 2009)	hop & hog	Text Generation	English
PAUL BOT (Tyler et al., 2020)	hop & hog	Dialogue	English
AliGator (Dehouck and Delaborde, 2025)	hop	Text Generation	French

Table 5: System of pun generation using traditional methods. Hog, hop and heg denote the types of homographic puns, homophonic puns and heterographic puns, respectively.

the forward and backward snowballing technique by examining the references and citations of these seed papers to identify additional relevant studies. We carefully reviewed all identified papers and ultimately compiled the findings into this survey.

E Traditional Systems

In this section, we summarized the pun generation systems with traditional methods in Section 4.1, as shown in table 5. We here listed the types of puns, task scenarios and languages corresponding to the system’s applications.

F Related Surveys

To our knowledge, there are currently only surveys on computational humour research, while no focusing exclusively on puns. [Amin and Burghardt \(2020\)](#) provides a survey on humour generation, including generation systems, evaluation methods, and datasets. However, it does not specifically analyze the category of puns and only summarizes papers published prior to 2020. [Nijholt et al. \(2017\)](#) concluded a survey on designing humour and interacting with social media, virtual agents, social robots and smart environments. In addition, other humour studies have been examined from the perspectives of detection ([Ramakrishnaiah et al., 2021](#); [Ganganwar et al., 2024](#)) and recognition ([Kalloniatis and Adamidis, 2024](#)). Furthermore, there are some relevant surveys in creating writing such as poetry composition ([Bena and Kalita, 2020](#); [Elzohbi and Zhao, 2023](#)), storytelling ([Gieseke et al., 2021](#); [Alhussain and Azmi, 2021](#)), arts ([Shahriar, 2022](#)) and metaphor ([Rai and Chakraverty, 2020](#); [Ge et al., 2023](#)). our survey provides a comprehensive overview of various methods focused on pun generation, including those published in recent few years.