

## SURVEY

# Arabic Text Formality Modification: A Review and Future Research Directions

SHADI L. ABUDALFA<sup>1</sup>, TAHAD J. ABDO<sup>2</sup>, AND MAAD M. ALOWAIFEER<sup>3</sup>

<sup>1</sup>DATA&TECH Center for Artificial Intelligence (DAITA)-Jeddah, King Fahd University of Petroleum & Minerals, Dhahran, Kingdom of Saudi Arabia

Corresponding author: Shadi L. Abudalfa (shadi.abudalfa@kfupm.edu.sa)

This work was supported by the Saudi Data & AI Authority (SDAIA) and KFUPM through the SDMA-KFUPM Joint Research Center for Artificial Intelligence under Grant (IRC-AI-EXT-05).

**ABSTRACT** Formality transfer seeks to adjust text formality without altering its core meaning, which carries substantial implications across diverse domains like machine translation, dialogue systems, and social media content creation. This study provides an extensive overview of formality transfer specifically within Arabic text, an emerging domain within natural language processing. Particularly, we carried out a comprehensive review of literature on text formality transfer, focusing on studies published between July 2019 and April 2024. Our focus lies in treating formality transfer in Arabic as akin to a machine translation task, presenting synthesized insights. Despite advancements in formality transfer for English and other languages, Arabic's distinct linguistic features present unique challenges and opportunities. Our investigation uncovers several research gaps necessitating further exploration, emphasizing persistent limitations. Moreover, we delve into text formality transfer as a promising avenue for forthcoming research initiatives in the realm of Arabic text processing.

**INDEX TERMS** Style transfer, formality, Arabic text, dialect, machine translation.

## I. INTRODUCTION

In the era of digital communication, the manipulation of textual content for diverse objectives has seen a notable rise [1]. Of specific interest is the modulation of styles [2] imbedded within text. Text style transfer [3] encompasses the conversion of textual content from one stylistic format to another with the objective of preserving its semantic essence while adjusting its expressive elements. This procedure involves adjustments to linguistic aspects like word choice, sentence structure, and mood to achieve the intended style with minimal deviation in the text. In this study, we concentrate on choosing formality as the attribute for text style transfer.

The process of formality transfer entails modifying the stylistic rendition of a sentence, transitioning it from an informal tone to a formal one, while preserving its inherent meaning. This procedure bears substantial implications for various applications, including machine translation (MT) [4],

dialogue systems [5], [6], [7] through online social media [8], and review systems, emphasizing the critical importance of comprehending and producing text with suitable formality levels.

There is an abundance of Arabic dialects (ADs) [9] in the Arabic world. Nevertheless, a formal dialect exists referred to Modern Standard Arabic (MSA) [10]. Whereas, Arabic dialect (AD) is an informal form of MSA language spoken across all Arabic countries. MSA serves as the academic discourse and communicative variation of the Arabic language, acknowledged by the United Nations. The high prevalence of ADs in Arabic-speaking countries stimulates the need for research focused on creating automated systems capable of implementing text formality transfer.

The primary objective of Arabic text formality transfer entails transforming text from one of ADs to the MSA. The concept of Arabic text formality transfer is depicted in a schematic illustration in Fig. 1. The diagram additionally highlights the key stages utilized in creating a system for formality transfer in Arabic text. The initial critical phase involves gathering Arabic sentences from various sources,

The associate editor coordinating the review of this manuscript and approving it for publication was Marco Chiarini Cavallaro.

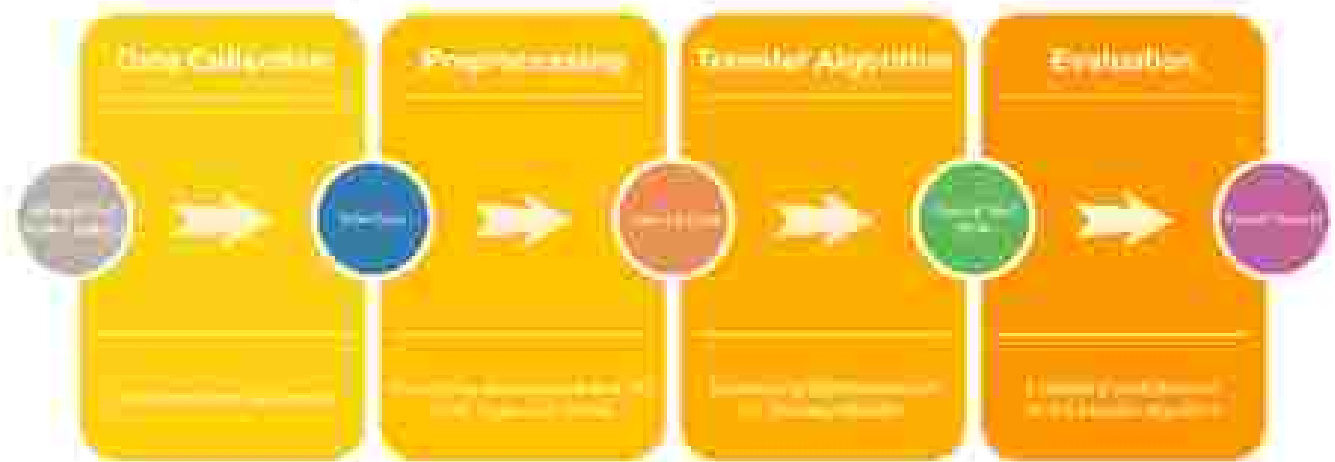


FIGURE 1. Canonical representation of the Arabic text formality transfer.

including social media platforms. This phase also involves annotating data to create a parallel dataset comprising ADs and their corresponding MSA. The second step involves preprocessing the raw data, which includes tokenization and cleaning up special characters by removing unnecessary elements like URLs and HTML tags. The model or technique is then designed to convert informal text into formal MSA. Lastly, the proposed technique's performance is assessed to finalize the development life cycle.

Neural machine translation (NMT) has achieved considerable progress in translating languages with ample resources, such as English, thanks to the availability of large-scale corpora that supply extensive training data for NMT models. However, languages with fewer resources, including Arabic and its various dialects, encounter significant challenges in MT due to the limited amount of text data necessary for developing effective translation models [11]. Over the past two decades, the field of MT has experienced remarkable advancements, resulting in improved translation quality, though it has not yet reached the level of human translation. Despite this progress, the effectiveness of MT systems is further hampered when translating ADs, and research into MT between ADs and MSA has been relatively sparse.

Within this manuscript, we delve into the realm of Arabic text formality transfer, analyzing current approaches, hurdles, and prospective pathways. Through a comprehensive analysis of used techniques and resources, we aim to provide insights and frameworks that advance the field of Arabic natural language processing (NLP).

This study does not consider the transformation in the reverse direction, from MSA to AD [12]. We also exclude in this study the transliteration of Arabic dialects written in Latin script (known as *Arabiized*) into MSA [17]. We excluded as well code-switching style transfer [14].

As far as we are aware, no prior systematic literature review has been conducted in this research direction, which drives the motivation for our current study. Our

study enriches the broader conversation regarding linguistic diversity, intercultural dialogue, and the manipulation of digital text, emphasizing the role of formality adjustment in improving the richness and variety of Arabic textual materials. The contributions of this research are listed as:

- Covering the latest research findings from 2010 to 2024 and a substantial volume of papers derived from a well-defined literature review procedure.
- Presenting a taxonomy framework for organizing the body of literature in AD to MSA machine translation, along with categorizing various methods employed for Arabic text formality transfer.
- Classifying 145 chosen articles based on the suggested taxonomy.
- Summarizing the latest formality transfer models, emphasizing machine learning techniques.
- Proposing new avenues for exploration in future research aimed at enhancing approaches for Arabic text formality transfer.

The structure of the paper unfolds as follows: Section II discusses earlier studies on literature reviews pertaining to formality transfer and associated research trends. In Section III, we describe the methodology employed in conducting this survey. Section IV presents our findings with this study. Additionally, Section V furnishes discussions and suggests avenues for future research. Finally, in Section VI, we draw conclusions.

## II. RELATED WORKS

Arabic text formality transfer can be integrated with other research avenues to advance larger systems for Arabic NLP. The following subsections discuss related works that align with our research focus.

### A. RELATED REVIEWS

Some survey studies can be associated with developing systems for translating Arabic formality from various

perspectives. In this subsection, we present several reviews that are relevant to our research direction from various perspectives.

For example, Elmaghr et al. [15] encompass computational approaches to dialectal Arabic detection. The findings indicate an uneven distribution of research efforts between speech and text and among the dialects, with a bias towards text over speech, regional varieties over individual dialects, and a preference for Egyptian dialects. The review also highlights significant research gaps related to the complete neglect of some Arabic dialects, a lack of resources for city dialects, and a deficiency in deep machine learning experimentation. This study diverges from our work as it concentrates exclusively on detecting dialects. Nonetheless, dialect detection can be utilized to create an automated formality system that identifies the specific Arabic dialect and subsequently translates the text into MSA. Additionally, our study is concentrated solely on the text modality.

Hamza et al. [16] center on MT within the realm of Arabic dialects. It offers an overview of contemporary research in this field, presenting a comprehensive account of each study's methodology and highlighting its key contributions. It is important to note that this study is quite dated, having been submitted to the journal on August 23, 2017. Therefore, our current research is highly significant as it demonstrates the advancements in this field following the recent renaissance of large language models (LLMs).

Alqaide et al. [17] aim to highlight various MT techniques found in the literature to inspire further research in this area. They examine specific linguistic features of the Arabic MT and addressing potential challenges. Additionally, they provide a detailed overview of the primary methods used for translating Arabic into English, evaluating their respective strengths and weaknesses. This study has a limitation in that it does not provide detailed information on translating Arabic dialects. Moreover, it is an older study, having been published online on July 24, 2012. Similarly, Sidya et al. [18] explore a range of translation models, such as Convolutional Neural Networks (CNNs), LSTM, NMT, BERT, and novel hybrid structures like the Transformer-CNN. The study highlights the advantages and drawbacks of each model, revealing their capacities to tackle translation issues.

In addition, Elsharif and Scornro [19] aim to examine the approaches, challenges, and proposed solutions in Arabic MT from approximately seven years ago. In the same context, Almaraytah and Alzobidy [20] identified sixteen challenges from 56 research papers. Their study delves into the four most significant issues and explores the solutions proposed by other researchers. The primary challenges include word sense disambiguation, Arabic named entities, intricate morphology, and limited resources.

Moreover, Zakroui et al. [21] conduct a thorough examination and contrast of various NMT methodologies, primarily focusing on Arabic-English MT endeavors. The methodologies under discussion tackle diverse linguistic and technical hurdles, showcasing significant advancements over

TABLE 1. Comparison with related reviews.

| Reference | Year | Machine translation | Dialect identification | Dialects | Translation systems | Formality transfer | Comparisons | Source language | Target language |
|-----------|------|---------------------|------------------------|----------|---------------------|--------------------|-------------|-----------------|-----------------|
| [15]      | 2018 | ✓                   | ✓                      | ✓        | ✓                   | ✓                  | ✓           | MSA             | Eng             |
| [16]      | 2017 | ✓                   | ✓                      | ✓        | ✓                   | ✓                  | ✓           | MSA             | Eng             |
| [10]      | 2018 | ✓                   | ✓                      | ✓        | ✓                   | ✓                  | ✓           | Ar              | Eng             |
| [22]      | 2020 | ✓                   | ✓                      | ✓        | ✓                   | ✓                  | ✓           | MSA             | Eng             |
| [11]      | 2019 | ✓                   | ✓                      | ✓        | ✓                   | ✓                  | ✓           | Ar              | Ar              |
| [23]      | 2020 | ✓                   | ✓                      | ✓        | ✓                   | ✓                  | ✓           | Ar              | MSA             |
| [13]      | 2021 | ✓                   | ✓                      | ✓        | ✓                   | ✓                  | ✓           | Ar              | Ar              |
| [14]      | 2021 | ✓                   | ✓                      | ✓        | ✓                   | ✓                  | ✓           | MSA             | Ar              |
| [10]      | 2022 | ✓                   | ✓                      | ✓        | ✓                   | ✓                  | ✓           | MSA             | Eng             |
| [18]      | 2019 | ✓                   | ✓                      | ✓        | ✓                   | ✓                  | ✓           | MSA             | Eng             |
| [24]      | 2023 | ✓                   | ✓                      | ✓        | ✓                   | ✓                  | ✓           | Ar              | Ar              |

conventional approaches. The objective of their inquiry is to ensure that researchers are informed about the latest research advancements until 2021, providing them with vital materials to improve Arabic MT. In a similar context, Amour et al. [22] furnish an overview of significant research endeavors in Arabic MT, highlighting key studies and available resources for developing and evaluating Arabic MT systems.

Zampieri et al. [2] explore pertinent applications like language and dialect identification, as well as MT, particularly concerning closely related languages, language variations, and dialects. Their study delves into computational techniques for handling languages with similarities but does not specifically concentrate on Arabic dialects. Within this context, Almansor et al. [23] focus is on Arabic and its dialects, considered low-resource languages, particularly in the context of converting non-standard text through normalization and translation techniques.

According to the results of this subsection, it is evident that our survey diverges from the related literature in several respects (refer to Table 1). Hence, according to our understanding, our study marks the inaugural systematic literature review concentrating on the transfer of formality in Arabic text. This study aims to shed light on contemporary methodologies for enhancing the efficacy of conversational systems to informal Arabic discourse, encouraging further investigation among scholars.

## B. RELATED TASKS

Text formality transfer is generally associated with other tasks as it falls under the broader category of text style transfer. One closely related task is MT, which involves converting text from one language to another. In the following subsections, we explore how various NLP tasks relate to this study, which centers on formality transfer in Arabic.

### 1) MACHINE TRANSLATION OF ARABIC AND NON-ARABIC TEXTS

Many studies have been conducted over the years translating between Arabic text and other languages [24], [25], [26].

Typically, MT relies on parallel datasets [27]. The studies cut across the classical methods (i.e. rule-based and statistical based techniques) and the recent NMT techniques (i.e. sequence to sequence models). Of course, the most recent LLMs generative models have now been employed in this domain for achieving exceptional translation quality. In this subsection, we discuss several studies on Arabic to other language translations to illustrate the principles of MT. The chosen studies in this subsection mainly present Arabic-English translation [28]. Therefore, they are not regarded as Arabic formality transfer. We highlight the distinction between Arabic-to-Arabic translation and AD/MSA translation as elucidated in the subsequent sections of this paper.

For instance, Bertucci and Mazroui [29] developed an English to Arabic NMT model to mitigate the effects of limited vocabulary and long sentences constraint encountered in translating languages with contrasting structures (i.e., English and Arabic). Similarly, Aljohany et al. [30] designed a two-way NMT-based model between Arabic and English texts. The NMT system is based on LSTM encoder-decoder model that employed attention network. Experiments confirmed that the proposed approach led to increased in translation quality, as evidenced by a reduction in translation loss.

Sajjail et al. [31] developed an evaluation suite for translation between Arabic dialects and English (by using some public datasets such as QAraC [32]), encompassing a diverse array of dialects, spanning multiple genres, and varying levels of dialectal characteristics. The suite is the first of its kind, bringing together a wide range of dialects, spanning multiple domains, and exhibiting varying degrees of dialectal characteristics. This research group employed an industrial-scale MSA to English systems to develop robust baselines through fine-tuning, back-translation, and data augmentation techniques.

Additionally, Sidiya et al. [18] developed two distinct models for translating English to Arabic: one utilizing LSTM and the other employing BERT. An extensive performance evaluation of these models is conducted, followed by a detailed examination of the outcomes. The comparative analysis yields valuable insights into the current state of Arabic-English translation models, offering guidance for future studies to enhance model accuracy.

Hamed et al. [33] proposed a NMT system with sole objective of improving the translation quality of the Holy Quran Arabic text to Italian language. The developed model employed two sequencers to sequence-based deep learning algorithms including LSTM and GRU with attention mechanism. Their experimental evaluations demonstrate that LSTM standout with an average BLEU and ROUGE performance of 0.18 and 0.17, respectively. However, their method is evaluated using small training samples, thus extending the experiments using huge amount of data needs to be conducted to ascertain the effectiveness of the developed model.

Bensalah et al. [34] proposed a NMT system between English and Arabic text that employed their novel word

embedding scheme generated by combining word2vec (i.e., CBOW and SG) and FastText embeddings models. The developed embedding method was used to train CNNs and various versions of RNNs to achieve the MT. They examined the efficacy of the developed approach using UN dataset and compared it with word2vec and FastText standalone embedding models.

Abrach [35] presents a NMT system to translate between Arabic and English using different parallel corpora and compared them with conventional approaches such as statistical machine translation (SMT) model (i.e., phrase-based systems). The author also explored the significance of specific preprocessing steps tailored to the Arabic script. After extensive experiments, the author has established that indeed preprocessing steps help in improving the translation quality of both NMT and the traditional methods. The findings indicate that NMT surpassed the phrase-based systems in all settings. Their study also revealed that training a neural network on a limited dataset produced results competitive with those achieved by conventional systems.

Furthermore, Oudah et al. [36] performed comprehensive comparison between SMT and NMT systems for translation between Arabic and English using training samples preprocessed by different notable tokenization approaches. Moreover, they examined various data and vocabulary sizes and assessed their impact on both methodologies. After experimental evaluations, results obtained with their work indicate that selecting the optimal tokenization scheme depends significantly on the model type and dataset size. Specifically, performance achieved from the developed method surpassed those documented in previous studies when evaluated on in-domain tests. In addition, the authors demonstrate that significant advancements can be achieved through a system selection strategy for amalgamating the outputs from both neural and statistical MT.

Bensalah et al. [37] developed a sequence-to-sequence (Seq2Seq) deep learning algorithm to translate between Arabic and English languages. In particular, the authors conducted extensive comparison of four different encoder-decoder based model architectures including LSTM, BiLSTM, GRU and BiGRU integrated with attention mechanisms. They also examined the influence of various preprocessing approaches on the developed MT system. The experimental findings revealed that the combination of Bi-GRU as an encoder, BiLSTM as a decoder, and coupled with attention network achieved the most optimal performance.

Recently, models based on transformers have demonstrated notable efficacy in language understanding and have achieved leading-edge performance across various NLP tasks, particularly MT. In this respect, Bensalah et al. [38] present NMT-based model to enhance the translating quality achieved on the Arabic-English-based translation systems. The developed approach is based on CNN and transformer architectures. The authors proposed novel preprocessing schemes that are based on FARASA and Arabert. Results

obtained using UN Arabic-English corpus indicate that their developed model achieved better translation quality surpassing the current advanced Arabic based NMT methods. In the same context, Bensalah et al. [39] developed a hybrid CNN and RNN attention-based algorithm to translate between Arabic and English text achieving outstanding translation quality.

Moreover, some of related works translate mixed Arabic dialects (i.e., combination of more than one dialect in the same text) into MSA. Nagoudi et al. [40] translate MSA mixed with Egyptian Arabic dialect to English using sequence to sequence transformers (S2ST) trained from scratch and pre-trained large language model. They use Open Parallel Corpus (OPUS) [41] for conducting the experiment work.

MT task can also be achieved by exploiting or adapting some of the pre-trained LLMs such as BERT, T5 and the recently introduced Arabic-based pre-trained transformer models like AraBert and AraT5. In the same context, Nagoudi et al. [42] developed a novel evaluation benchmark for Arabic MT.

In spite of the success recorded by the NMT techniques in translating between Arabic and English text and vice versa, attention has now turned to the most recent LLMs paradigm. Over the last few years different researchers have started deploying recent LLMs like ChatGPT to ascertain their performance in various natural language generation (NLG) tasks such as MT. Specifically, the performance of these LLMs is being investigated in respect to under-resourced languages such as Arabic. For example, Al-Khalifa et al. [43] examined how the most advanced pre-trained language models (PLMs) handle MT tasks from English to Arabic. Evaluations based on various performance metrics illustrate that GPT-4 and GPT-3.5-turbo exhibit better proficiency.

Allderode et al. [44] present the initial research endeavor investigating translation between Syrian Arabic and Turkish. They assess translation accuracy into Turkish from different Arabic dialects and contrast these findings with translations from Syrian Arabic. To establish their MADAR-Turk dataset, they translate the aforementioned 2,000 sentences from the Damascus dialect of Syrian Arabic into Turkish, aided by two native Syrian Arabic speakers proficient in Turkish.

Although LLMs like ChatGPT and Bard are believed to possess multilingual proficiency after being fine-tuned on instructional data, their level of linguistic has not been adequately investigated. LLMs are usually evaluated with zero-shot (ZS) and few-shot (FS) [45] generation strategies. To explore this constraint, Kadraoui et al. [46] provide an extensive assessment of LLMs in MT tasks. Specifically, they assess ChatGPT, GPT-4, and Bard by translating between 10 different Arabic dialects and English. Similarly, Khoshdel et al. [47] performed a comprehensive performance evaluation of LLMs with primary goal of assessing the capabilities of ChatGPT's on different NLP tasks including Arabic language and its associated dialects. The findings of their study suggest that while Chat GPT demonstrates

impressive results in English, it consistently falls behind smaller models that have been fine-tuned for Arabic (e.g., MARBERTv2 and AraT5).

Moreover, Moftem et al. [48] investigate the functionalities of various LLMs models including GPT-3.5, GPT-4, and BLOOM for adaptive machine translation in real-time via in-context learning. Extensive experiments were conducted across five different language combinations (i.e., English-to-Arabic, Chinese, French, Kinyarwanda and to Spanish). Their findings indicate that the translation performance achieved through FS based LLMs (GPT-3.5) can exceed that of robust encoder-decoder-based MT models, particularly for English-to-French and English-to-Spanish. While, low performance was realized for English-to-Arabic.

Alkharaja [49] assesses the suitability of ChatGPT for Arabic-English MT task. The primary aim of the developed approach is to examine the quality of ChatGPT's translations and to compare its performance with dedicated commercial MT systems like Google translate. To achieve this goal, a standardized dataset consisting of 1000 English sentences and their respective Arabic translations was utilized to assess the translation results generated by both MT systems, as well as a human translation reference. Their findings suggest a slight edge of ChatGPT over Google translate in producing high-quality translations. Nevertheless, despite these encouraging translations quality, it is important to recognize that even the most advanced MT technology, including ChatGPT, currently cannot achieve the level of proficiency exhibited by human translation. Table 2 presents a summary of MT approaches used with Arabic and other languages.

Fig. 2 displays the publication trends between 2018 and 2024. The data reveals a steady rise in publication numbers over time. This indicates interest in the field, particularly as it intensified in the latter part of the period under review.

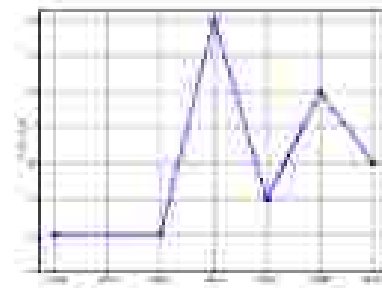


FIGURE 2. Publication distribution by year for Arabic and non-Arabic machine translation.

## 2) MACHINE TRANSLATION FOR NON-ARABIC LANGUAGES

Numerous studies have been conducted on utilizing machine translation across various languages. In this subsection, we discuss several recent studies that address machine translation involving languages other than Arabic. Recently, Lys et al. [30] evaluated different MT tasks such as long document translation, styled MT, and interactive MT using



LLMs. Furthermore, they tackle privacy concerns in MT using LLMs and suggest fundamental privacy-preserving techniques to alleviate potential risks.

Additionally, Zhang et al. [31] explore the abilities of LLMs in performing MT tasks based on QLoRA [52] to translate between English and French. They employed three important techniques including prompting, in-context learning and fine-tuning to accomplish the evaluation. Their study shows that the effectiveness of LLMs for MT is not always consistent. LLaMA-2 surpasses its competitors. However, models that depend completely on FS learning often underperform compared to models that are trained from scratch. The process of fine-tuning consistently yields performance improvements, particularly for models encountering difficulties with FS learning paradigms and document-level translation tasks.

Similarly, Jian et al. [33] evaluated how well ChatGPT performs MT tasks compared to popular commercial services like Google Translate. They looked at how these systems handle different aspects of translation, including following user instructions, translation between many languages (multilingual), and adapting to different text styles. After extensive experiments they were able to establish that ChatGPT performed competitively with Google, DeepL and Tencent translation translators on high resource languages but struggle on low resource languages based on BLEU metric. The authors investigate also an intriguing approach known as pivot prompting for distant languages, wherein ChatGPT translates the candidate sentence into a high-resource pivot language prior to the conversion to target language, thereby notably enhancing translation performance. In respect to translation robustness, ChatGPT does not match the performance of commercial systems on biomedical abstracts or Reddit comments. However, it demonstrates good performance on language translation.

Moreover, Hendy et al. [36] present an extensive assessment of GPT models for MT, addressing different facets including the performance of various GPT models compared to cutting-edge research and popular commercial services, the impact of prompting techniques, resilience to changes in domain, and document level translation. Their experiments involve translation tasks in 18 diverse directions, covering both well-resourced and limited-resource languages. This assessment extends beyond English translation alone, evaluating the performance of three GPT models: ChatGPT and two different versions of GPT-3.5. The achieved performance indicates that GPT models attain translation quality that compares effectively, particularly for high resource languages. However, they exhibit limited capabilities for low resource languages. The authors also demonstrate that integrating GPT models with other translation algorithms has the potential to further improve translation accuracy.

Most of the text style transfer studies in the literature concentrate on the well-resourced languages such as English abandoning limited-resource languages such as Chinese. To address this disparity, Tao et al. [37] developed a Chinese

article style transfer model, harnessing the potential of LLMs. This model integrates a Text Style Definition (TSD) module designed to thoroughly analyze textual attributes in articles, facilitating the seamless transfer of Chinese article-style by LLMs. The TSD module leverages various machine learning techniques in examining the article style. Extensive experiments confirm that the developed approach surpasses previous one in terms of transfer accuracy and maintaining content integrity.

In the same context, Durr et al. [38] explored unsupervised translation system to transform between Mandarin Chinese and Cantonese. A comprehensive comparison was conducted across various model structures, tokenization methods, and embedding techniques at a large scale. Their most successful model attained a BLEU score of 25.1 in translating Mandarin to Cantonese.

### 3) OTHER ATTRIBUTES FOR TEXT STYLE TRANSFER

Our research focuses on examining the techniques employed for formality transfer. Nonetheless, various methods used in related sub-tasks like politeness transfer and emotion modification can also be applied to formality transfer. For example, the impressive ability of diffusion probabilistic-based models in generating high-quality images with a high degree of control has inspired researchers to investigate the potential of applying similar control mechanisms to the domain of text generation. Prior investigations into diffusion-based models have revealed their capability to be trained without the pre-trained weights.

In this regard, Lys et al. [39] performed style transfer by adopting diffusion-based algorithm and training it from scratch (i.e., without pre-trained weights) with very small training samples. They evaluated their approach by using StylePTB dataset (i.e., a benchmark for fine-grained text style transfers). Additionally, their developed approach, trained on a small sample of StylePTB, achieves superior performance compared to prior methods that depend on pre-trained weights, embedding layers, and external linguistic analyzers.

## III. RESEARCH DESIGN

As previously stated, this study seeks to enhance our comprehension of the existing advancements in automated support for Arabic text formality transfer. It explores the literature, examines available datasets, and reviews the techniques employed in text formality transfer. Additionally, it aims to understand how researchers evaluate the quality of formality transfer.

It is important to point out that our focus in this work revolves around elucidating the challenges, datasets, and existing methodologies pertaining to the formality transfer of Arabic text. Additionally, we investigate different methodologies and present empirical results to contribute to the advancement of Arabic natural language processing research. Fig. 3 shows the main themes addressed in this study.

We formulated three Research Questions (RQs) to align with our study objectives. Table 1 presents the RQs



FIGURE 3. A framework to review studies concerning the transformation of formality in Arabic text.

TABLE 3. Research questions and the reasoning behind them.

| Q   | Research question                                                                                 | Reasoning                                                                                                                                                                                                                                                                                                                              |
|-----|---------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| RQ1 | What types of Arabic text are used for model test formalization?                                  | Understanding the types of Arabic text used for model test formalization is crucial for the model test formalization process. It is important to know the types of Arabic text used for model test formalization to ensure the quality of the formalization process.                                                                   |
| RQ2 | How is the quality of text formality transfer measured?                                           | Understanding the quality of text formality transfer is crucial for the model test formalization process. It is important to know how the quality of text formality transfer is measured to ensure the quality of the formalization process.                                                                                           |
| RQ3 | What subjects are related to Arabic text formality transfer in both Arabic and English languages? | Understanding the subjects related to Arabic text formality transfer in both Arabic and English languages is crucial for the model test formalization process. It is important to know the subjects related to Arabic text formality transfer in both Arabic and English languages to ensure the quality of the formalization process. |

considered in our study, along with an explanation of the rationale behind each one.

In this section, we describe the review protocol steps implemented in our study. In Subsection III-A, we outline our approach to searching for relevant studies. Subsection III-B details the criteria used to systematically identify and choose research articles on text formality transfer, as well as the quality evaluation method applied to assess the selected studies. Lastly, in Subsection III-C, we describe our process for data extraction and synthesis to address the posed RQs.

### A. SEARCH STRATEGY

We employed a two-stage hybrid search strategy that combined database searches with the snowballing method. Initially, we conducted a search within databases, followed by a snowballing search. A database search involves creating and executing a search query within digital libraries. A snowballing search involves examining the references (backward) and citations (forward) of chosen studies to identify additional relevant research. According to Mouria et al. [60], an effective approach for locating evidence in systematic reviews is to utilize a hybrid search strategy that combines a representative digital library with the snowballing method.

TABLE 4. Query strings used to identify primary studies.

| # | Query string                           |
|---|----------------------------------------|
| 1 | "Arabic text" AND "formality transfer" |
| 2 | "text formality transfer"              |
| 3 | "Arabic text" AND "formality transfer" |

To conduct the database search, we pinpointed appropriate keywords. By examining our RQs, we identified key concepts and terms linked to our primary topic, which helped to compile a list of relevant domain keywords. We then assessed various combinations of these keywords to determine their effectiveness in retrieving a collection of studies. Consequently, three query strings were employed to identify primary studies, as illustrated in Table 4.

We chose several digital libraries for our research, including ACM Digital Library,<sup>1</sup> Scopus,<sup>2</sup> Springer,<sup>3</sup> ProQuest,<sup>4</sup> IEEE Xplore,<sup>5</sup> and ScienceDirect,<sup>6</sup> due to their comprehensive coverage of leading journals and conference proceedings. Additionally, we conducted a search on Google Scholar<sup>7</sup> to enhance our research scope. In total, we used seven different databases in this study.

We conducted our final search using a timeframe restricted to the period from July 2010 to April 2024. The timeframe was set due to a significant increase in research papers about MT and formality transfer after 2010, compared to the period before. Based on our search method, we found a total of 145 primary studies. The detailed search process and the count of papers identified at each stage are illustrated in Fig. 4.

### B. STUDY SELECTION

The preliminary search conducted through database queries yielded 2,630 possible papers (see Fig. 4). Table 5 presents the outcomes of this stage for each database along with the respective query strings. Then, we filtered out articles that were not relevant by reviewing the titles. To address result overlap from multiple digital libraries, we implemented a structured deduplication process to ensure each article was uniquely represented. This process involved two main steps: first, we computed titles for filtering out duplicates. Then, we manually reviewed a sample of the deduplicated dataset to verify accuracy and confirm that no important articles were excluded. As a result of this phase, 141 papers were obtained.

After that, we utilized the inclusion and exclusion criteria to select the most pertinent papers for our study. Papers that fulfilled all inclusion criteria were included, while those that met any exclusion criteria were not considered. In this step, we selected 77 papers by reviewing their abstracts and entire

<sup>1</sup><https://dl.acm.org>

<sup>2</sup><https://www.scopus.com>

<sup>3</sup><https://link.springer.com>

<sup>4</sup><https://www.proquest.com>

<sup>5</sup><https://ieeexplore.ieee.org>

<sup>6</sup><https://www.sciencedirect.com>

<sup>7</sup><https://scholar.google.com>

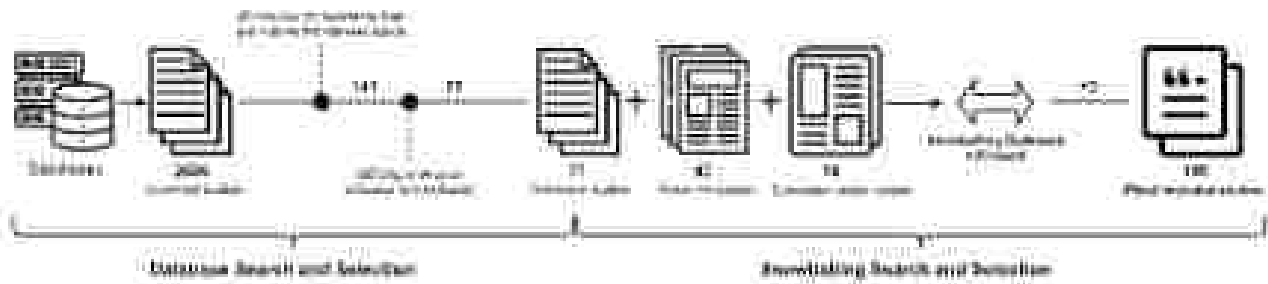


FIGURE 4. Search strategy and selection process.

TABLE 3. Results of preliminary database query search.

| Database       | Query 1 | Query 2 | Query 3 |
|----------------|---------|---------|---------|
| Scopus         | 584     | 4       | 0       |
| Web of Science | 79      | 16      | 0       |
| Pubmed         | 69      | 17      | 0       |
| PsycInfo       | 109     | 28      | 0       |
| Lex Ngram      | 1       | 1       | 0       |
| Scisearch      | 62      | 7       | 0       |
| Google Scholar | 1403    | 142     | 1       |
| Total          | 2365    | 209     | 1       |

texts. The following inclusion and exclusion criteria were used in our study.

• Inclusion criteria

- 1) The paper should be written in English, as it is the common language used within the natural language processing community.
- 2) The paper introduces a corpus that can be utilized for examining the formality transfer in Arabic text.
- 3) The paper introduces an automatic evaluation metric that could be applied to Arabic text formality transfer.
- 4) The paper describes a model or technique employed for MT or text formality transfer.

• Exclusion criteria

- 1) Papers that fail to meet any of the outlined inclusion criteria.
- 2) Extended abstracts, posters, books, patents, and tutorials.
- 3) Papers do not concentrate on textual modality.
- 4) Papers do not employ formality attribute in text style transfer.
- 5) Papers provide dataset that is not organized at the level of individual sentences.

For implementing the second phase of our hybrid search approach, we utilized the snowballing search to ensure that no pertinent studies were overlooked. As a result of applying snowballing search, we selected 77 papers that present techniques for MT or text formality transfer. Additionally, 47 papers have been selected for presenting datasets that can be utilized with Arabic text formality transfer. Moreover, we have chosen 18 papers that introduce

substantiated evaluation metrics for assessing the performance of Arabic text formality transfer. It is worth mention here that two more papers are chosen because they provide both methodology and resources for text formality transfer. While, the final paper selected presents methodology and a new evaluation metric for text formality transfer.

Therefore, the chosen 145 papers introduce methodologies, resources, and assessment criteria for implementing text style transfer and MT. Fig. 5 visually represents the scope of our work by displaying the top 10 most frequent words found in the titles of the selected papers. Whereas, Fig. 6 provides additional insight into our research by visually presenting the 10 most common 2-gram combinations (adjacent word pairs) extracted from the titles of the selected papers.

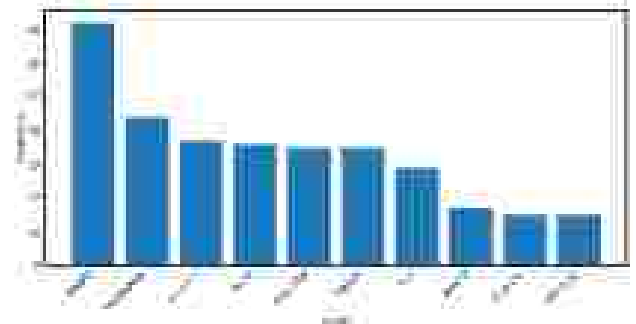


FIGURE 5. The 10 most commonly occurring words in the titles of the 145 papers.

The final studies chosen were based on the quality assessment criteria we developed to evaluate the relevance and strength of the primary research. The criteria for quality assessment can be found in Table 6. We focused exclusively on studies that had a quality score exceeding 50% of the highest possible score, and these were the ones selected for data extraction.

### C. DATA EXTRACTION AND ANALYSIS

We gathered pertinent data from each chosen paper to address the RQs. Additionally, we compiled metadata for each paper to support further statistical analysis. This metadata comprised the title, year of publication, authors,

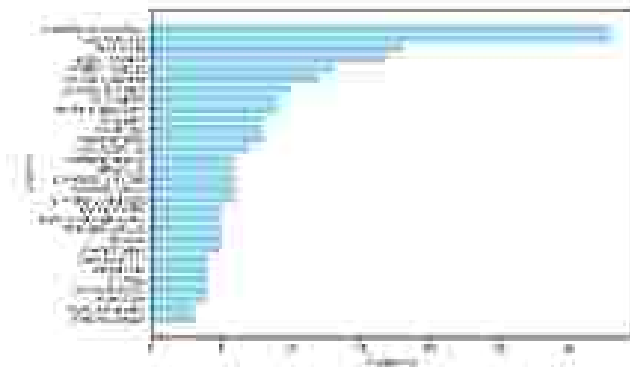


FIGURE 6. The 30 most 3-gram combinations found in the titles of the 145 papers.

TABLE 3. Quality assessment questions.

| Q# | Quality question                                              |
|----|---------------------------------------------------------------|
| 1  | Is the methodology of the paper well defined?                 |
| 2  | Has the dataset size and data source been clearly identified? |
| 3  | Are the machine learning techniques adequately described?     |
| 4  | Does the research make a meaningful contribution?             |
| 5  | Can the proposed methodology be applied and evaluated?        |

and publication type. The collected data were systematically arranged in Excel spreadsheets.

The main objective of data synthesis is to gather and integrate information and statistics from chosen studies to address the RQs and formulate a response. By grouping studies with similar and comparable results, it was possible to provide definitive answers to the RQs through the presentation of research evidence.

We analyzed a range of data, including numerical performance metrics, categorization of approaches, machine learning methods, languages, domains, and datasets. To integrate information from the main studies and tackle the RQs, we employed different methods, including visualization tools. Additionally, tables were utilized to summarize and convey the results.

#### IV. RESULTS

In this section, we provide an in-depth summary of the results from our literature review, addressing each RQ in detail.

##### A. RESOURCES (RQ1)

Several resources [61], [62] are available for gathering data used in Arabic text formality transfer. Some of them are designed for particular Arabic dialects [33], while others are suitable for use across multiple dialects and MSA. The following subsections detail these resources.

##### 1) PARALLEL DATASETS

In this subsection, we offer a survey of prevalent parallel datasets [64] for Arabic dialectal MT as documented in existing literature. A summarized view of these datasets, specifically tailored for translation between ADs and MSA,

TABLE 4. Summary of parallel public datasets for ADs and MSA.

| Dataset   | Domain             | Sample Size                                                                                    | Year |
|-----------|--------------------|------------------------------------------------------------------------------------------------|------|
| OPT       | Middle Eastern     | 10K Egyptian sentences, 1.1M Levantine Egyptian sentences, 10K Algerian dialect sentences, 14K | 2017 |
| AAE       | Algerian dialects  | 2K sentences                                                                                   | 2014 |
| QAD       | Free dialects      | 10K sentences                                                                                  | 2014 |
| MSAE      | Multiple dialects  | 2K to 5K Arabic MS/MSeg                                                                        | 2019 |
| MAAD      | 13 Arabic dialects | 10 sentences, 100,000, 100,000 sentences with 1,000 and 10,000                                 | 2019 |
| MSDOCT    | Multiple dialects  | 50,000 Modern Standard Arabic, 50,000 Levantine Arabic                                         | 2019 |
| MSDT      | Multiple dialects  | 100K words                                                                                     | 2019 |
| QADAT     | Free dialects      | 10 Arabic dialects                                                                             | 2019 |
| MSD       | Multiple dialects  | 100,000 sentences                                                                              | 2019 |
| LIAD      | Lebanese dialects  | 200K sentences                                                                                 | 2020 |
| DS        | Free dialects      | 4.5 sentences                                                                                  | 2020 |
| Arabic-DT | Egyptian, Saudi    | 1.5K sentences                                                                                 | 2021 |
| Leak      | Multiple dialects  | 1.2 million words                                                                              | 2021 |
| Arabic    | Arabic dialects    | 5K sentences                                                                                   | 2021 |
| MSD       | Multiple dialects  | 1.2M sentences                                                                                 | 2021 |

is presented in Table 7. The following subsections provide detailed explanation into the particularities of every dataset presented.

##### a. ARABIC-DIALECT/ENGLISH PARALLEL TEXT (APT) CORPUS

Zhib et al. [65] introduced a dataset containing Arabic-English translations, including MSA and two ADs. The dataset is made up of over 8 million sentences in MSA and English, along with 138,500 sentences in Levantine Arabic and English, and 18,000 sentences in Egyptian Arabic and English. The corpus was sourced from Arabic weblogs, and the conversion process was conducted using Amazon Mechanical Turk. The dataset was collected to be a valuable asset for training and evaluating MT models, particularly for handling ADs and MSA.

##### b. ALGERIAN ARABIC DIALECTS (AAD)

Smadi et al. [66] aim to create parallel corpora for Algerian dialects for the first time. Their ultimate goal is to facilitate MT between MSA and Algerian dialects. Additionally, they introduce language tools designed to handle these dialects. Initially, they adapted HAMA, a renowned MSA analyzer, to develop a morphological analysis model for the dialects. Following this, they propose a diatization system based on an MT process to reinsert vowels into dialect corpora.

##### c. MULTI-DIALECTAL PARALLEL CORPUS (MDPC)

Bouamar et al. [67] developed this corpus by extracting two thousand Egyptian-English sentences out of APT dataset, originally created by Zhib et al. [65]. Following the selection process, indigenous speakers originating from Palestine, Syria, Jordan, and Tunisia were tasked with translating the sentences into their respective native dialects. They also translated an additional 8000 sentences from Egyptian-English data specifically into Syrian dialect. This

dataset provides a valuable resource for comparing identical sentences translated into different dialects.

#### D. PARALLEL ARABIC DIALECT CORPUS (PAIDC)

The PAIDC is a notable resource in the realm of Arabic dialectal studies and NLP. Mefrouh et al. [68] present this corpus, which is a multi-dialect dataset comprising of MSA, Algerian, Tunisian, Palestinian, and Syrian dialects. This corpus encompasses 6.4K parallel sentences for MSA and its dialects. It was generated from manually transcribed recordings of daily discussions, as well as scripts from movies and TV dialogues, which were converted into six varieties. It consists of over 37,000 tokens, with approximately 10,000 word types in both MSA and the five dialects. The corpus is meticulously assembled to incorporate a broad spectrum of linguistic phenomena, expressions, and idiomatic usage unique to each dialect. Its accessibility has expedited [69] the creation of algorithms, models, and utilities designed to accommodate the distinctive features of Arabic dialects.

#### E. MULTI-ARABIC DIALECT APPLICATIONS AND RESOURCES CORPUS (MADAR)

The MADAR dataset constitutes an extensive compilation of Arabic text data representing a range of dialects prevalent across the Arab region. It encompasses texts sourced from diverse platforms like news articles, social media content, blogs, among others. The corpus is meticulously annotated and categorized, facilitating research endeavors and advancements in NLP, machine learning, and computational linguistics pertaining to Arabic dialects.

Bouamor et al. [70] curated 2k English sentences from the IITEC corpus, initially a Japanese/English repository for parallel phrases, and engaged native speakers from 25 Arab cities to translate these phrases into their respective dialects. They further translated 10k sentences for five specific cities: Doha, Beirut, Cairo, Tunisia, and Rabat. The initial resource comprises a substantial parallel dataset encompassing 25 Arabic city dialects, along with the prior parallel sets of English, French, and MSA dialects. While the second corpus consists of a lexicon with entries in the dialects of 25 cities, totaling 1,045 entries each, accompanied by equivalents in MSA, French, and English.

#### F. DIA2MSA

Mahruk [71] introduced a parallel dataset for translating Arabic dialect to MSA. This corpus includes 5,000 tweets written in Egyptian/Maghribi dialects and 5,000 tweets written in Levantine/Gulf dialects. Each tweet was translated to MSA by someone who speaks that specific dialect as their native language.

#### G. SAUDI DIALECT CORPUS (SDAC)

Al-Twairesh et al. [72] developed a dataset consisting of over hundred thousand words generated from various social

media platforms including, X (formerly Twitter), YouTube, WhatsApp, blogs, Instagram and forum.

#### H. CITY-LEVEL DATASET OF ARABIC DIALECTS (CLDAD)

Maged et al. [73] developed this massive Arabic dialect corpus collected from 10 Arabic countries, representing 19 cities containing varying dialects such as Egyptian, Gulf, KSA, Levantine, and Yemeni. The data was gathered using the X API, employing multiple bounding boxes across several Arab countries.

#### I. THE BIBLE

The Bible corpus [74] contains Bible texts translated into multiple languages, serving as a valuable asset for MT endeavors. With translations available in numerous languages, it enables researchers and developers to train and assess MT models across various language combinations. The Bible has been translated into over 330 languages. While there are numerous translations available in Arabic, translations into Arabic dialects remain extremely limited.

#### J. SADID

The SADID dataset [75] encompasses English, Egyptian, Levantine, and MSA. The dialectal materials were sourced from three different origins: 1) Wikipedia, chosen for its wide-ranging topics and straightforward language, 2) Aesop's Fables, selected for their narrative format, and 3) particular dialogues extracted from movie subtitles.

#### K. TUNISIAN ARABIC DIALECT (TAD)

Kobouri et al. [76] introduce a data augmentation method designed to generate a parallel corpus, aligning Tunisian Arabic dialect as expressed on social media with standard Arabic. This corpus aims to support the development of a MT system.

#### L. ARABIC SEMANTIC TEXTUAL SIMILARITY (STS)

The aim of Arabic STS dataset [77] is to ascertain the semantic similarity between two Arabic sentences. Each English phrase was translated into three distinct target languages: MSA, Egyptian Arabic, and Saudi dialect.

#### M. YEMENI, IRAQI, LIBYAN, AND SUDANESE ARABIC DIALECT CORPORA WITH MORPHOLOGICAL ANNOTATIONS (LIAN)

Jarur et al. [78] present Lian, a morphological annotated corpora collected from five Arabic cities dialects, including Yemeni, Iraqi, Libyan and Sudanese dialects consisting of about 1.2 million tokens. Specifically, Yemeni data which is about 1.05 million tokens was generated spontaneously from X platform, while the rest of the three cities (i.e. about 30K token each) was curated manually from Facebook and YouTube.

TABLE 3. Non-parallel corpora for AD sentences.

| Reference | Year | Task                | Dialect          |
|-----------|------|---------------------|------------------|
| [81]      | 2014 | Machine             | Palestinian      |
| [81]      | 2014 | Machine             | Moroccan dialect |
| [83]      | 2017 | Machine translation | Tunisian         |
| [84]      | 2017 | Machine             | Palestinian      |
| [84]      | 2017 | Sentence analysis   | Nigerian         |
| [84]      | 2017 | Machine             | Moroccan dialect |
| [87]      | 2018 | Machine             | Moroccan dialect |
| [88]      | 2018 | Machine             | Moroccan dialect |
| [87]      | 2019 | Machine             | Maghribi         |
| [88]      | 2021 | Machine             | Moroccan         |
| [90]      | 2021 | Machine             | Palestinian      |
| [91]      | 2022 | Sentence analysis   | Nigerian         |
| [91]      | 2022 | Machine             | Moroccan dialect |
| [91]      | 2024 | Machine             | Moroccan dialect |
| [91]      | 2024 | Machine             | Moroccan         |

#### iv. SYRIAN ARABIC DIALECTS WITH MORPHOLOGICAL ANNOTATIONS (NABIRA)

Nayyal et al. [79] developed a Syrian Arabic dialect corpus with morphological annotations called Nabira. Nabira was curated by group of Syrian natives consisting of 6K sentences, totaling 60K words, from various outlets such as social media, movie and series scripts, song lyrics, and local proverbs.

#### iv. THE KING DATASET

Yaman et al. [80] presented a method of gathering data that is crafted to enable participants to translate sentences from MSA into their respective regional dialects. The MSA sentences are drawn from a subset of 11,570 sentences sourced from a widely recognized MSA dataset, specifically the MADAR dataset. Participants are tasked with translating these sentences into the dialect they identified as their own upon entering the competition. Another dataset for question answering is presented as well with 796 samples.

#### 3) NON-PARALLEL CORPORA

Various non-parallel datasets are available for different NLP tasks, which can also serve as a source for implementing Arabic text formality transfer. Table 3 provides several examples and emphasizes that datasets from various tasks can also be leveraged in this line of research.

#### 3) BENCHMARK DATASETS

A benchmark dataset is a standardized collection of data employed to assess and contrast the performance of models within a particular field. These datasets are critical for gathering data applicable to the task of Arabic text formality transfer.

For example, Nagouli et al. [95] introduce Dolphin, a groundbreaking benchmark designed to meet the requirements for evaluating NLG across various Arabic languages and dialects. This innovative benchmark covers an extensive array of 13 distinct NLG tasks, such as dialogue generation, question answering, machine translation, and summarization, among others. Dolphin includes a significant collection of

40 diverse and representative public datasets, organized into 50 test splits, meticulously selected to represent real-world contexts and the linguistic diversity of Arabic.

Similarly, Elmadafy et al. [96] present ORCA, a publicly accessible benchmark designed for evaluating Arabic language understanding. ORCA has been meticulously developed to encompass various Arabic dialects and a broad spectrum of challenging comprehension tasks, utilizing 60 distinct datasets across seven natural language understanding (NLU) tasks.

#### 4) SHARED TASKS

In a research community, a shared task usually involves a collaborative endeavor where various teams or individuals tackle the same problem or dataset to evaluate and compare their methods and outcomes. These collaborative challenges are particularly prevalent in the field of NLP. Such shared tasks can be employed to gather data for Arabic text formality transfer.

For example, Al-Khalifa et al. [97] proposed a shared task<sup>4</sup> for AD to MSA MT. Additionally, Abdel-Mageed et al. [98] proposed a shared tasks for AD-MSA MT. Moreover, we can also collect data from some related task such as AD identification [99], [100].

#### 5) TOOLS

There are various tools and systems that can be used for collecting data and improving performance of Arabic text formality transfer, as explained in the following subsections. In this work, we specifically choose tools and systems that can be applied to both Arabic dialects and MSA.

#### iv. MAGHAD

MAGHAD [101] is a tool designed for analyzing and generating morphological structures within the Arabic language family. It supports both MSA and various spoken dialects, initially focusing on the Levantine dialect. The system is adaptable to new dialects even without an existing lexicon, requiring only minimal manual knowledge engineering.

#### iv. A HYBRID APPROACH FOR CONVERTING WRITTEN EGYPTIAN COLLOQUIAL DIALECT INTO DIACRITIZED ARABIC (HACWECDIA)

The HACWECDIA system [102] transforms a written sentence in Egyptian colloquial Arabic into a diacritized MSA sentence. Egyptian Colloquial Arabic was selected due to its widespread usage in blogs and articles online. Resources were gathered using a rule-based approach from a vast amount of data on the internet. The system's lexicon consists of 41,708 words, including 9,085 non-MSA words, 3,000 distinct colloquial words, with the remainder being spelling errors and non-Arabic names. The system is designed to differentiate between MSA and colloquial words.

<sup>4</sup><https://arabx.com/arabx-github>

## C. COLABA

Diab et al. [107] introduced COLABA, an extensive initiative aimed at developing resources and tools for processing Dialectal Arabic. Their efforts focus on enhancing the existing resources and tools, expanding them to cover additional dialects and varieties of dialectal data.

## D. ARABIC DIALECT HANDLING IN HYBRID MACHINE TRANSLATION (ADHMMT)

The ADHMMT system [104] expands on a hybrid MT system designed for managing Arabic dialects. It incorporates a statistical decoder featuring four rule categories: lexical, syntactic, argument structure, and functional structure rules, alongside semantic disambiguation data, a statistical bilingual lexicon, bilingual phrase table, and target language models.

## E. DIALECTAL TO STANDARD ARABIC PARAPHRASING TO IMPROVE ARABIC-ENGLISH STATISTICAL MACHINE TRANSLATION (DSAPARASMT)

The DSAPARASMT system [105] generates standardized paraphrases in MSA. It incorporates two dialect varieties: Levantine and Egyptian. A simple rule-based method is applied with this system. An existing MSA analyzer is enhanced by including dialect-specific out-of-vocabulary (OOV) words and low-frequency words.

## F. THE ARABIC ONLINE COMMENTARY DATASET (AOCD)

This AOCD system [106] utilizes a vast monolingual dataset of 32 million words, emphasizing diverse dialectal content. It is specifically trained to detect and quantify dialectal elements within sentences. The data is sourced from three newspapers renowned for their extensive use of Levantine, Gulf, and Egyptian dialects. Moreover, the system can differentiate between dialectal content and MSA, as well as between various dialectal forms.

## G. SENTENCE LEVEL DIALECT IDENTIFICATION FOR MACHINE TRANSLATION SYSTEM SELECTION (SLDMMTS)

The SLDMMTS system [107] can identify whether a sentence is purely MSA or contains dialectal elements, enabling the use of the most appropriate MT system for better accuracy.

## H. MULTI-DIALECT MACHINE TRANSLATION (MUDMAT)

The MuDMaT project [108] seeks to promote the advancement of MT systems tailored for under-resourced languages, including various dialects and variants. Specifically, it focuses on enhancing MT capabilities for three North African Arabic dialects—Tunisian, Algerian, and Moroccan—which face severe resource limitations. The project's methodologies and insights are transferable to other less-resourced language variants and dialects. The MuDMaT project focuses on developing hybrid MT systems that combine statistical and rule-based approaches. These systems

aim to translate between various Arabic dialects, MSA, and French bidirectionally.

## I. A NOUJ TUNISIAN DIALECT TRANSLATOR

Turjman et al. [109] have developed a translator to convert the Tunisian dialect into MSA. The suggested translation approach relies on a bilingual dictionary derived from the studied corpus, supplemented by a comprehensive set of local grammars. These local grammars are then converted into finite state transducers using the latest technologies of the NonJ linguistic platform.

## J. TOWARDS A PORTABLE SPOKEN LANGUAGE UNDERSTANDING SYSTEM APPLIED TO MSA AND LOW-RESOURCED ALGERIAN DIALECTS

To circumvent the need for extensive parallel corpora in MT, Lachouri et al. [110] address the challenge of interpreting user intents from natural language queries into a system's semantic format across various languages and dialects. Specifically, they focus on English, MSA, and four distinct Algerian vernacular dialects originating from Blida, Djella, Tenez, and Tizi-Ouzou regions.

## K. A Seq2Seq NEURAL NETWORK BASED CONVERSATIONAL AGENT FOR GULF ARABIC DIALECT

Abuhameed and Sukhija [111] present an initial promising effort towards developing an open-domain conversational agent in the Gulf Arabic dialect. They employed a Seq2Seq neural network to construct a conversational agent using the Arabic Gulf dialect.

## L. CAMEL TOOLS

The CAMEL tools [112] comprises a set of open-source Python tools designed for Arabic NLP. CAMEL Tools currently offers functionalities including text preprocessing, morphological analysis, dialect identification,<sup>4</sup> named entity recognition, and sentiment analysis.

## M. BUILDING THE EMIRATI ARABIC Framework (EAPN)

The EAPN project [113] seeks to establish a FrameNet specifically for Emirati Arabic, leveraging the Emirati Arabic Corpus. The objective is to develop a resource akin to the early stages of the Berkeley FrameNet. The project is structured into manual and automatic tracks, reflecting the main methods employed to gather frames in each track.

## N. ONLINE TOOLS

There are several online tools available that can be used for transforming the formality of Arabic text. For example, the AEADMMT system is an AI-driven tool<sup>10</sup> crafted for precise dialectal translation, cultural understanding, and educational assistance.

<sup>4</sup><https://camel-toolkit.github.io/en/latest/usage/dialect-identification.html>

<sup>10</sup><https://www.yeschat.ai/types-86217/MuMPQ-Arabic-Explorer>

## B. EVALUATION METRICS (MQ3)

It is difficult to design reliable scoring systems [114] for MT because there is not always a single “correct” way to translate a sentence. One method for assessing MT involves human evaluation [115]. Nonetheless, the need for automatic evaluation metrics remains significant in this area of research. Existing metrics [116] mostly judge translations by how closely the words match a reference translation, with the main difference being how they handle word order changes and synonyms. This subsection discusses common metrics employed in literature to automatically evaluate the performance of MT systems.

In this study, we present the automation evaluation since it standard with this research work. It is important to note that while some research includes human evaluation, we have chosen not to incorporate it in our study. This is due to the lack of standardization and the susceptibility to human errors.

### 1) METEOR

METEOR, a popular tool used in MT and other NLP tasks, compares draft translations to corrected ones. It considers individual words (unigrams) and uses three (alpha, beta, and gamma) to adjust its scoring based on the specific language and type of translation being done. Unlike simpler methods, METEOR can handle situations where words are rearranged or synonyms are used, making it valuable for assessing MT quality across various languages and situations.

METEOR developed by Lavie and Denkowski [117], is a performance measure employed to assess the efficiency of MT systems by comparing individual words in a generated translation to their corresponding words in a reference translation. When a word in the candidate translation has multiple matches in the reference translation, METEOR favors alignments that require the least shuffling of words. The extent of rearrangement is determined by tallying the number of crossed alignments. The final alignment is used to compute the unigram precision and recall of the machine translation given by equations 1 and 2.

$$Precision = Matches / (Length trans) \quad (1)$$

$$Recall = Matches / (Length ref) \quad (2)$$

The precision and recall are exploited to obtain the weighted unigram harmonic mean illustrated by equation 3. Equation 4 presents the final equation for METEOR, which is a multiplication between the unigram harmonic mean  $F_{\alpha}$  and the penalty obtain from equation 5. METEOR includes three settings ( $\alpha$ ,  $\beta$ , and  $\gamma$ ) that can be adjusted to better match human evaluation of translations for a particular language and type of translation task.

$$F_{\alpha} = Precision / (\alpha Precision + (1 - \alpha) Recall) \quad (3)$$

$$METEOR_{\alpha, \beta, \gamma} = F_{\alpha} \cdot P_{\beta, \gamma} \quad (4)$$

$$P_{\beta, \gamma} = 1 - \gamma \left( \frac{Chance}{Matches} \right)^{\beta} \quad (5)$$

### 2) BILINGUAL EVALUATION UNDERSTUDY (BLEU)

The BLEU metric was introduced by Papineni et al. [118]. It is a common metric for checking the quality of MT, especially in the world of language processing. It works by looking for how many words match between the machine-translated output and the human-made generated translation. BLEU considers groups of words (n-grams) and rewards translations with more overlap. However, it also discourages translations that are too short, ensuring they capture the full meaning of the original text. It penalizes translations that are not longer than the reference by adjusting the final score based on the length difference.

Equation 6 shows how BLEU calculates a score for machine translation using matching word groups (n-grams) up to a certain span (N). This equation considers a collection of machine-generated translations (T) and their corresponding human-made translations (R). Whereas, BP value is calculated by using equation 7.

$$BLEU = (N) = \left( \prod_{n=1}^N \frac{n\text{-grams}(T \cap R)}{n\text{-grams}(T)} \right)^{\frac{1}{N}} BP \quad (6)$$

$$BP = \min \left( 1, \exp^{(1 - \frac{len(R)}{len(T)})} \right) \quad (7)$$

BLEU is easy to use, but it has some weaknesses. For example, BLEU does not care about the order of words. As long as the translation includes the right groups of words (n-grams), it can get a good score even if the sentence structure is messy. BLEU also does not check if any important information is missing from the translation. It only looks for matching n-grams, not whether the translation captures everything from the original text. Finally, BLEU can not tell the difference between words or phrases with similar meanings. Despite this limitation, BLEU remained one of the most employed metric for assessing the precision of different MT systems.

It is worth mention here that some studies assess the BLEU score by comparing transferred sequences to the original sequences [119], whereas others evaluate the BLEU score by comparing the transferred sequences to human [120].

### 3) TRANSLATION EDIT RATE (TER)

TER was proposed by Snover et al. [121] to establish the quickest way to convert candidate translation by correlating it to a perfect translation done by a human. It can insert, delete, or swap words, or even switch short phrases around. This “swap” ability is what makes TER different from a trivial performance measure called word error rate (WER), which can only insert, delete, or swap individual words. By allowing swaps, TER avoids unfairly penalizing translations where the word order is slightly different. Specifically, TER score is obtained based on equation 8. From a MT standpoint, a lower TER score means the MT is closer to the human-made

version:

$$TER = \frac{(Mean edit)}{(Average ref. len)} \quad (8)$$

Snaver et al. [122] proposed TER-Plus (TERp), which is basically an improvement on TER. It lets researchers adjust the importance of different editing actions (like insertions or deletions) to better match how humans judge translation quality. Additionally, TER-Plus can consider extra ways to fix translations, including word stem matches, WordNet synonym and multiword matches.

#### C) CROSS-LINGUAL OPUS METRIC MEASURES TRANSLATION (COMET)

Rei et al. [123] developed COMET, which is among the most popular performance metrics for evaluating MT over the past few years. It focuses on how well the translation matches the original text word-for-word, across different languages. It does this by comparing the words in the original text and the translation one by one. This approach helps assess how accurate and faithful the translation is, especially when working with different languages.

#### 5) CHARACTER n-GRAM F-SCORE (ChrF AND ChrF++)

ChrF and ChrF++, proposed by Popovic [124], are performance measures developed to assess the performance of MT algorithms. They fall into a class of metrics that concentrate on similarities at the character level between the candidate translation and perfect translations performed by human. ChrF (see equation 9) checks how similar the character n-grams between machine-generated translations and reference translations by employing both precision and recall for calculating an F-score that harmonizes those inputs.

This metric is valuable for assessing translations in scenarios with ambiguous word boundaries or languages with intricate morphology, where conventional word-based metrics encounter limitations. ChrF considers how many of these character sequences match exactly (precision) and how many do not get missed (recall) and combines these aspects into a single score. ChrF++ extends the concept of ChrF by considering word order up to 2.

$$ChrF\beta = \left(1 + \beta^2\right) \frac{CHRP \times CHRR}{\beta^2 \times CHRP + CHRR} \quad (9)$$

where *CHRP* and *CHRR* represent the mean precision and recall values of character n-grams, respectively, encompassing all n-grams. Similar  $\beta$  denotes the weights.

#### 6) BERTScore

Zhang et al. [125] introduced BERTScore as a method that goes beyond merely analyzing single words. It uses a powerful language model called BERT and its contextual embedding to grasp the overall interpretation of the sentence. By comparing the generated MT to a perfect translation, BERTScore can give a similarity score that reflects how well

the translation captures the original meaning. Mathematically, BERTScore can be calculated based on equation 10:

$$BERTScore = \frac{1}{|C|} \sum_{c \in C} \max_{r \in R} \cos(e_r, e_c) \quad (10)$$

where,  $|C|$  denotes the amount of tokens from the machine-generated translations, while  $C$  is the sets of tokens of the machine-generated translation. Similarly  $R$  denotes the token embeddings of the human perfect translation, while  $e_r$  and  $e_c$  denotes the embedding of token  $c$  and  $r$ , respectively.

#### 7) RECALL-ORIENTED UNDERSTUDY FOR GETTING EVALUATION (ROUGE)

ROUGE score is relevant in assessing the extent of content similarity between the machine-generated translation and reference translations. This evaluation measure can be calculated by using equation 11 [126].

$$ROUGE = \frac{\sum_{\text{system-grams}} Recall_{\text{system-grams}}}{\text{Total n-grams in ref. translations}} \quad (11)$$

#### 8) PERPLEXITY (PPL)

Perplexity (PPL) score [127] gauges the level of astonishment the model experiences when confronted with unfamiliar data. It is used for a fluency evaluation. Lower perplexity values indicate more effective training. This metric can be calculated by using equation 12.

$$PPL = \exp\left(-\frac{1}{N} \sum_{i=1}^N \log p(x_i)\right) \quad (12)$$

where  $N$  represents the count of tokens within the same translated sentence,  $x_i$  denotes every token within the same translated sentence, and  $p(x_i)$  represents the likelihood attributed to each token  $x_i$  by the probability distribution.

#### 9) OTHER METRICS

Other metrics are also utilized in some related studies, though they are not as commonly employed. For example, the Generalized Language Evaluation Understanding (GLEU) metric [128], inspired by BLEU, serves as a tool to assess the preservation of content. There is also a measure for style accuracy (SAcc) by evaluating the prediction accuracy of a pre-trained style classifier on the generated sentences. Additionally, the context coherence can be evaluated using a method described by Cheng et al. [129]. This approach involves using a pre-trained coherence classifier to predict how well a generated sentence aligns with its surrounding context, based on the prediction accuracy.

Additional metric entitled as Paraphrase In N-gram Changes (PINC) [130] is used also with this research direction. PINC indicates the degree of difference from the original informal sentences, but this does not necessarily correspond to the effectiveness of the style transfer. Alternative metrics encompass Cosine Similarity (CS) [131], which assesses similarity between the embeddings of original and

style-transferred sentences to evaluate content preservation. Similarly, the degree of Meaning Similarity (MS) [132] between the original and modified texts is evaluated using the Sentence Transformers embedding distance. In a similar context, Hamming distance (Hh) calculates the variance between the source and generated texts. Moreover, Word Overlap (WO) [133] measures the unigram overlap rate between the original and style-altered sentences, while BARTScore [134] evaluates the generated text from multiple perspectives, including informativeness and fluency. In the same context, Sellam et al. [135] introduced BLEURT as a score for assessing the similarity between translation candidates and reference texts. Unlike methods that rely on N-gram overlap, BLEURT is capable of recognizing semantic differences.

Certain studies substitute PPL with the precision of a RoBERTa-large model that has been trained on the CoLA corpus [136] for evaluating grammatical test (GT). Similarly, Mir et al. [137] suggested an alternative approach to assessing fluency, proposing the use of classifiers trained to differentiate between machine-generated and human-written sentences, rather than relying on perplexity.

Moreover, certain studies assess semantic similarity (SIM) by employing the unbiased embedding-based SIM model introduced by Wotting et al. [138]. This metric shows strong performance on semantic textual similarity (STS) benchmarks in SemEval evaluations [139].

Some studies employ the Geometric Mean (GM) to compute the geometric mean of accuracy, BLEU, and 1/PPL. Other research presents findings by using the harmonic mean (HM) of BLEU and style accuracy as a comprehensive score. However, we do not include these metrics in our analysis since they only reflect general performance.

### C. COMPARISON BETWEEN EVALUATION METRICS

The metrics can be categorized according to the technique they use into three groups: "N-gram-based Metrics", "Semantic-based Metrics", and "Edit-based along with Miscellaneous Metrics". Fig. 7 shows the category to which each metric belongs.



FIGURE 7. Categorization of evaluation metrics by technique.

It is clear that BLEU, Chrf, and TER mainly focus on comparing individual words or characters between the candidate translation and perfect translation done by human. COMET, on the other hand, takes a different approach. It uses a special technique to grasp the meaning of the words (summaries) and check if the translation captures the same meaning as the original text. Therefore, we can also categorize evaluation metrics based on their intended purpose into three groups: "Fluency and Adequacy", "Semantic Similarity and Alignment", and "Model and Data Properties" as illustrated in Fig. 8.

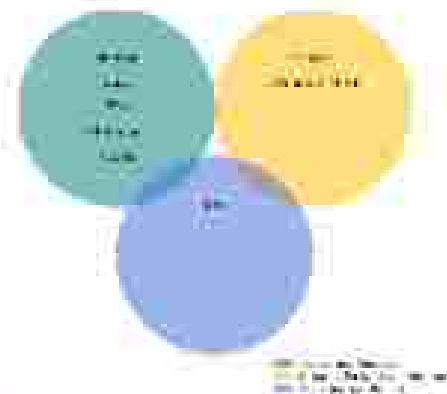


FIGURE 8. Classification of evaluation metrics based on purpose of use.

### D. TAXONOMY OF APPROACHES FOR FORMALITY TRANSFER

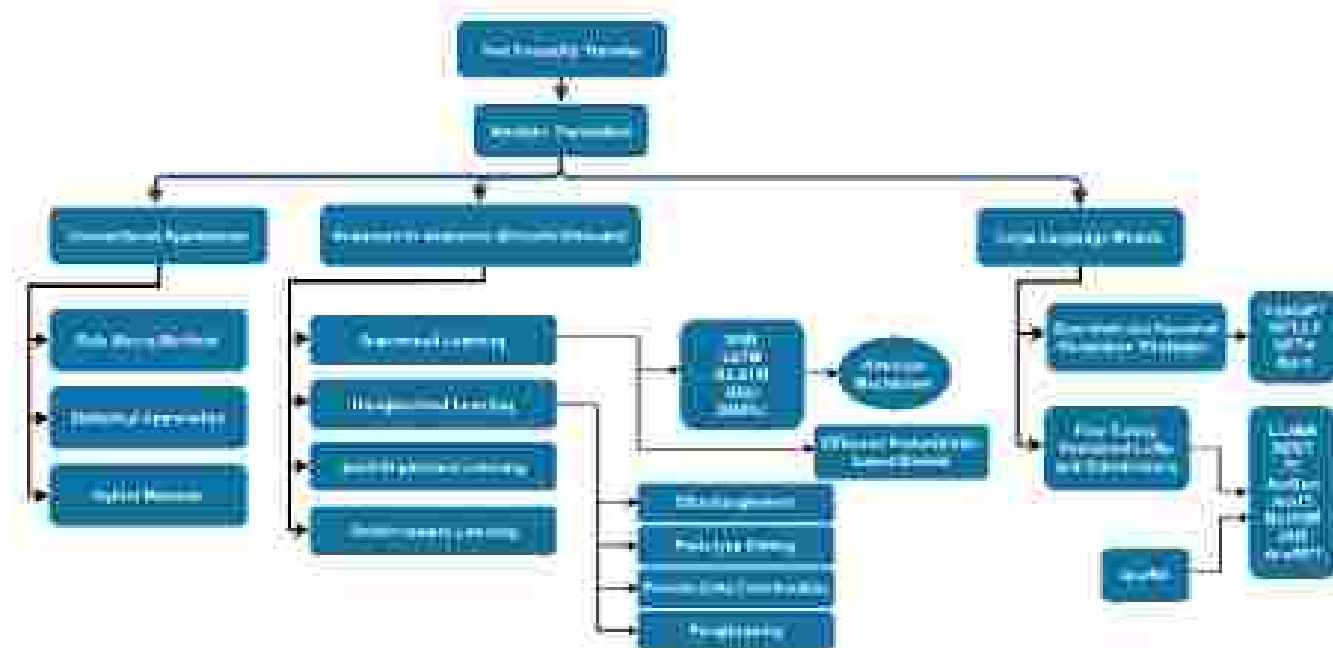
In this study, the various methodologies employed in the chosen papers are classified to evaluate the value of the research for future scholars. Fig. 9 presents a taxonomy of reviewed approaches alongside selected papers. This classification is developed according to the methods used, the data configurations, and the significant results achieved through different approaches.

The figure illustrates the classification of approaches into three main categories: conventional, seq2seq, and LLMs. We distinguish LLMs separately to highlight their significant role due to recent advancements in NLP. Conventional methods are further categorized into rule-based, statistical, and hybrid techniques. The seq2seq approaches are classified into supervised, unsupervised, semi-supervised, and reinforcement learning techniques. Within unsupervised learning (Non-parallel Data), there are four techniques employed: disentanglement, prototype editing, back-translation/pseudo data creation, and paraphrasing [340].

### E. ARABIC TEXT FORMALITY TRANSFER (NQ1)

In this section we explored various studies in the context of Arabic formality transfer ranging from classical approaches, neural machine translation techniques (i.e., sequence to sequence method) and the recent much anticipated LLMs.

We explored previous studies concentrated on the translation from ADs to MSA and vice versa, as well as between



**FIGURE 8** A taxonomic analysis of associations in life of selected species

MSA and English in both directions [21]. Additionally, recent significant works have been discussed as a NMT, ranging from the utilization of basic sequence-to-sequence models to the integration of self-attention mechanisms, thereby establishing transformers and extensive language models as the prevailing approach in NMT. Furthermore, recent advancements in text style transfer are discussed, emphasizing their utility in enhancing the Arabic text formality transfer.

### 15 COMMUTATIONAL APPLICATIONS

There has been increased interest in research on translating low-resource languages like Arabic (i.e., translation between AD and MSA) over the last few years. Thereby, a lot of research has been observed in the context of ADs translation [14].

A lot of the reviewed papers focus on the conventional approaches [142], [143], [144] such as rule based [60], [145], [146], [147], [148], [149], [150], [151], [152], statistical approaches [153], [154], [155], [156] and hybrid methods [104], [157], [158]. For rule-based MT, a set of rules based on the grammar and vocabulary of both the source and target languages is used to translate text. Specifically, it offers variety of approaches including direct translation, transfer system and interlingual system. While, statistical translation computes the probability  $P(S/T)$  for any given source candidate  $S$  and target candidate  $T$ , selecting the translation that maximizes the probability.

Al-Gapfari and Al-Yasouni [150] employed rule-based approach for translation between Sana'a dialect and MSA with achieving 77.42 % performance. Similarly, Tachouri

and Bouazouan [157] employed rule-based technique that depends on language models for transforming the Moroccan dialect to MSA.

The rule-based MT techniques have significant drawbacks. For example, constructing rule-based demands considerable amount of time. Whereas, statistical methods require high computational devices, and they cannot address syntactic challenge present in Arabic dialect such as word ordering. In general, the conventional methods often produce conflicting sentences that are incompatible or linguistically inconsistent. Fig. 10 provides an overview of the MT systems employed for translating Arabic dialects to MS using conventional approaches. It is evident that a rule-based approach is more commonly employed than a statistical approach in the reviewed studies.

## 25. SEQUENCE TO SEQUENCE APPROACHES

Despite the fact that research on translating AD to MSA and vice versa is dominated by the conventional approach like statistical, many researchers have started exploring the translation of AD to MSA based on NMT methods [159], [160], [161], [162], [163], [164], [165].

In particular, Baniam et al. [162] introduced the first NMT architecture designed to transform from ADs to MSA using RNN-based encoder-decoder architecture. Their method introduced a multi-task learning algorithm where the decoder is shared among every language combination, while each source language employs its own individual encoders. The results obtained have demonstrated that the developed model achieved high translation performance against the standalone models given a small training sample. Hence, multi-task

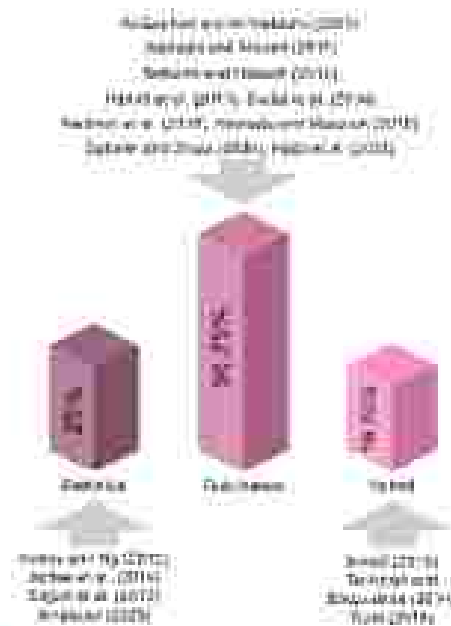


FIGURE 10. Summary of MT systems used with Arabic dialect to MSA translation based on conventional approaches.

learning can handle the challenge of insufficient training samples. Even though the developed system introduces an innovative concept of employing distinct encoders for every single language or dialect, along with a shared decoder. The proposed Arabic dialect-based translation system relies on the basic RNN encoder-decoder architecture.

In the same context, Guehli et al. [159] proposed a NMT model designed to process Arabic dialects, employing basic RNN encoder-decoder-based architecture. The system is specifically designed to translate between Algerian Arabic, a dialect that blends Arabic script with elements of Amazig, and MSA. Similarly, Enna et al. [166] developed a NMT model to convert the spoken Tunisian Dialect (TD) into MSA. Their study has been expanded [167] by utilizing additional models.

Farhan et al. [160] introduce an unsupervised NMT-based approach for translating Arabic dialect. They developed two models for translating Arabic dialect to MSA. The first model is developed using standard attentional sequence-to-sequence model. While, the second model is based on Google NMT system. They firstly evaluated their approach in supervised setting using parallel corpus for Jordan to MSA translation before implementing the unsupervised settings. In the unsupervised setting, the highest BLEU score achieved is 32.14 by using the first architecture on Saudi-MSA corpus. However, their developed approach was evaluated using single Arabic dialect and a single evaluation metric. To effectively ascertain the efficacy of their developed approach it should be subjected to multiple ADs and more than one evaluation measure. Similarly, Al-Ibrahim and Deyoumi [161] developed a machine translation approach for

converting Jordanian Arabic vernaculars into MSA based on encoder-decoder deep learning algorithm that employs RNN model. However, they use a very small size corpus which ultimately hindered the performance of their model. In addition, they only focus on one dialect, leading to a shortage of NMT systems that encompass the complete range of dialects.

Baniata et al. [163] presented an algorithm for a unified multitask-based neural MT with Arabic dialect. The developed NMT algorithm exploit RNN encoder-decoder-based framework and integrates part-of-speech tags linguistic resource. The encoder is shared between Arabic dialect and MSA translation task achieving an improved BLEU score. Similarly, Baniata et al. [164] proposed a novel reverse positional encoding NMT scheme to convert Arabic dialect to MSA. The developed model scheme employed multihead attention (MHA) technique. The integration of the novel encoding scheme and application of sub-word units in the MHA block enhances the encoder's sublayer in the developed approach. These enhancements effectively capture all dependencies among the words in Arabic dialect input text. The outcomes obtained illustrate the ability of the developed approach to mitigate the open grammar issue of Arabic dialect sentences.

Baniata et al. [165] developed the first transformer based NMT algorithm for Arabic dialect exploiting subunits. The developed approach relies on the subunits and the shared vocabulary among the source (i.e. Arabic dialects) and the target candidate (i.e., MSA) to upgrade the conduct of the multi-head attention mechanism of the encoder. Experiments translating Maghrebi, Levantine, Nile, Gulf, and Iraqi dialects to MSA revealed an improvement in the developed method's performance when evaluated using the BLEU score metric. The experimental findings have confirmed that the transformer based NMT model can effectively capitalize on the subword units to enhance the translation performance.

Fahem et al. [168] developed a supervised, semi-supervised, and unsupervised machine translation architecture for translating Egyptian ADs to MSA. They employed both parallel and monolingual datasets to evaluate the developed schemes. Specifically, they developed three neural translations algorithms. Their first model utilizes a sequence-to-sequence architecture capitalizing on the significant overlap in vocabulary between Egyptian dialect and MSA, for effective learning of the word embeddings. The second translation model employed transformer framework trained based on only monolingual dataset, while the third model utilizes both parallel and monolingual datasets in the initial supervised learning stage and training phase, respectively. Results obtained demonstrate that semi-supervised approach produced the best BLEU score against the supervised and unsupervised approaches.

To investigate the spread of misinformation about Covid-19 and its impact on communities in Algeria, Slim et al. [159] explored the issue of translating Algerian vernacular in the

context of Covid-19-related social media interactions. The primary goal of the study is to enhance the comprehension of Algerian vernacular messages concerning Covid-19. The developed approach starts by employing an LSTM algorithm to filter messages and detect comments related to Covid-19. Afterwards, the identified Covid-texts are translated using embedded GRU model from Algerian vernaculars to MSA, achieving a BLEU score of 22.10.

In addition, Slim et al. [170] developed sequence-to-sequence algorithm using attention network to transform Algerian Arabic dialect to MSA. The developed algorithm is based on translductive transfer learning approach. The main objective for integrating translductive learning scheme is to address the learning problem encountered for under-resourced languages such as Arabic dialects (i.e., Algerian dialect). The training process began with a parallel dataset comprising multiple Arabic dialects. Subsequently, the models were fine-tuned on a under-resourced data focused exclusively on the Algerian dialect. The impact of the translductive learning technique was highly noticed. Specifically, it greatly enhances the BLEU score accuracy of the sequence-to-sequence algorithm, increasing it from 6.3 to above 34. With the attentional based method, it increases the BLEU score from less than 17 to more than 35, thus demonstrating the capability of their proposed method.

Moreover, Al-Kharusi and AAlAbdulnaim [171] addressed the lack of specific research on MT between Omani Arabic dialect and English by developing an Omani dialect parallel dataset generated from social media source X. We chose to include this study in our research because the authors use transfer learning techniques for the Omani Arabic dialect based on a pre-existing MSA-English model. Consequently, this research can be viewed as a form of indirect translation between AD and MSA.

Moukaffih et al. [172] have also tried various multitasking learning methods to leverage the connections among Arabic dialects. Table 9 shows a summary of sequence to sequence based NMT approaches for translation between AD and MSA.

### 3) LARGE LANGUAGE MODELS

LLMs are advanced pre-trained neural networks developed to comprehend and produce human-like text. These models constitute a substantial progress in NLP as they possess the capability to understand and generate text across different tasks like MT, question answering, text generation, etc. LLMs usually trained on extensive text datasets containing billions of words or even more, aim to grasp the underlying patterns and structures of human language. They utilize deep learning architectures like transformers, enabling them to capture extensive contextual information and dependencies within text sequences.

LLMs possess the remarkable capability to undertake diverse NLP tasks, eliminating the need for task-specific fine-tuning. Once trained on extensive text corpora, LLMs

can be readily deployed for various downstream tasks with minimal additional training. Nonetheless, obstacles remain, including the shortage of annotated data for dialects, managing diverse dialectal variations, and preserving the cultural nuances inherent in each dialect. Prominent examples of LLMs include OpenAI's GPT series, encompassing models like GPT-2, chat-GPT, GPT-3, and their succeeding versions, alongside models like bidirectional encoder representations from transformers (BERT).

Recently, researchers have started investigating the performance of LLMs in various NLP problems such as NLG subtasks. For instance, in the context of MT (such as translation of ADs to MSA), the process typically involves fine-tuning LLMs on dialectal data and MSA parallel corpora to improve their performance in understanding and generating accurate translations. The LLMs models can be exploited directly using the dialectal/MSA data via *N-shot* (also called in-context learning) settings. Table 10 illustrates a summary of MT systems used to convert Arabic dialect to MSA based on LLMs with ZS generation strategy and fine-tuning.

Recent works show that pre-trained LLaMA models, primarily trained on English-dominated corpora, inherently lack proficiency in handling non-English languages. To investigate this issue, Zhao et al. [180] seek to effectively adapt language generation and instructions-following skills for a language with limited resources. In particular, they conducted an extensive empirical investigation based on LLaMA to examine the need for expanding vocabulary and the necessary scale of training (i.e., additional pre-training, instruction tuning) for successful transfer. Their findings reveal that extending the vocabulary is unnecessary, and comparable transfer accuracy to advanced models can be attained with small additional pre-training samples in the majority of these tasks.

Alyafei et al. [33] comprehensively evaluate the capability of two LLMs series (i.e., GPT3.5 and GPT-4) across seven different Arabic NLP tasks. The evaluation is carried out to gauge the abilities of these emerging foundational models and to compare their performance with advanced approaches in the literature. The extensive experiments indicate that GPT-4 surpassed GPT-3.5 in the majority of the 7 problems analyzed. The findings underscore a notable performance gap between the ChatGPT models and their Arabic-specific counterparts across most tasks. Both ChatGPT models exhibit exceptional performance when compared to current leading models. Nevertheless, certain limitations were noted in particular scenarios, such as summarization, MT, mostly for low resource languages, and reasoning capabilities.

Ghaddar et al. [181] take a fresh look at how to pre-train and evaluate Arabic pre-trained language models (PLMs) by developing five new models based on BERT and T5 pre-training approach, achieving state-of-the-art performance against existing Arabic models on text generation task. Regarding pre-training, the authors investigate enhancing Arabic language models based on three features including,



quality of the data, the model's size, and the integration of character-level details. Consequently, they introduce three Arabic BERT-based and two T5-based architectures. They compared their models to the latest ones on two key tasks for Arabic language processing: NLU and NLG. For NLU, they used a standard benchmark called ALUE, and for NLG, they used a part of another benchmark called ARGON. Their models performed significantly better than existing Arabic PLMs models on both tasks.

Derouich et al. [173] performed an extensive comparison of seven different pre-trained models for sentence-level translation of Arabic varieties such as, Egyptian, Emirati, Jordanian, and Palestinian vernaculars to MSA. Their proposed approach is part of the results obtained in the NAD-2023 Shared competition that exploit MADAR corpus. Specifically, their approach involved fine-tuning of the pre-trained transformer-based models including, AraT5, AraT5v2-base-1024, Sahar-ArabicT5, AraT5-MSA-Small and AraT5-MSA, AraT5-Tweet-Small and AraT5-Tweet-based models. The top-performing model (AraT5v2-base-1024) attained BLEU scores of 10.02% by using the test set.

As part of the NAD-2023 shared tasks, Khared et al. [174] evaluated different T5-based networks for converting various Arabic vernaculars to MSA. Specifically, they employed three different transformer models including AraT5, multilingual T5 (mT5), and multi-task fine-tuned mT5 (mt0), which were trained jointly (with various dialects) and independently (for each dialect). Their results indicate that AraT5 model trained independently reached an outstanding BLEU score of 14.76. While, the jointly trained AraT5 model attained an outstanding BLEU point of 21.10. Their findings indicate that fine-tuning AraT5 and integrating beam search procedure results in world class translation accuracy. Similarly, [175] present their findings and results in the same shared task. The models they proposed attained a BLEU score of 11.57, earning them fourth place in the second subtask and third place in the third subtask. By employing parameter-efficient training methods, their model outperformed traditional fine-tuning approaches in the experimentation phase.

Very recently, Naciri et al. [179] introduce a technique that incorporates data augmentation techniques using generative pre-trained transformer models (GPT-3.5 and GPT-4) combined with the fine-tuning of AraT5 V2, a model specifically designed for Arabic translation. Their work is part of the OSACT 2024 Dialect to MSA Translation Shared Task. Similarly, Arwani et al. [178] reveal the findings of the same shared task. The task encompasses Gulf, Egyptian, Levantine, Iraqi, and Maghrebi dialects, providing 1001 sentences in both MSA and the dialects for fine-tuning, along with 1888 test sentences. Further experiments are conducted by fine-tuning AraT5 and No Language Left Behind (NLLB) models using the MADAR Dataset. Additionally, Fares [176] used AraT5 which achieved a BLEU score of 21.79% on the test set related to the shared task. Moreover, Alshamir [177] fine-tuned four AraT5 models for this shared task.

## F. NON-ARABIC TEXT FORMALITY TRANSFER (RQ3)

Many techniques can be used for applying Arabic text formality transfer. In this subsection, we present some techniques that have not been employed in this research direction to fill more gaps through future works. We reviewed many approaches of text formality transfer that can be adapted in the domain of Arabic text formality transfer. This subsection aims to provide insight into various methodologies that can enhance the efficacy of Arabic text formality transfer.

For example, Lai et al. [182] investigated the capacity of ChatGPT as a multifaceted evaluator for the text formality transfer tasks based on existing metrics in the literature and human judgments using ZS setting. Additionally, Potkova et al. [183] adapt the advanced unsupervised-NMT model to translate dialects based on unilingual data from Croatian dialects to standardize Croatian language.

In conventional style transfer, models are limited to transferring text from one specific source style to a designated target style, facing challenges when handling texts from different domains (i.e., out-of-domain text). To address this challenge Hallinan et al. [184] present unified style transfer with an expert model that can transfer text from any given source style to numerous target styles. Their model was also developed to address the issue of limited parallel corpora. This is achieved through expert-guided decoding and a two-stage reinforcement learning approach. Their developed model surpasses the GPT-3, even though it is 226 times smaller in size.

Table 11 provides an overview of all reviewed papers on non-Arabic formality transfer according to our chosen criteria. In this table we report the best performance achieved by the presented works. The table shows the technique presented with each reviewed paper. The methodology discussed in the reviewed papers is categorized into supervised learning (S), unsupervised learning (U), semi-supervised learning (SS), reinforcement learning (R), and LLMs. We note that Grammarly's Yahoo Answers Formality Corpus (GYAFC) [185] is the most commonly used dataset for English text.

Fig. 11 displays statistics that show the percentage breakdown of the various approaches employed in the reviewed studies. It is evident that many recent studies utilize unsupervised learning to address the issue of limited resources. It is clear also that most of reviewed papers present encoder-decoder (ED) technique as a sequence-to-sequence text generation. It is worth noting that the conventional approaches were not used in any of the studies reviewed for Non-Arabic text formality transfer.

It is important to highlight that many methods utilized for formality transfer in non-Arabic texts have yet to be applied to Arabic texts. This observation paves the way for researchers to explore various future opportunities in the Arabic domain. The upcoming section delves into this observation and explores various other ways in which future research could address different challenges related to formality transfer in Arabic text.

TABLE 15. Summary of MT systems used with Arabic dialect to MSA translation based on LLMs.

| Ref.  | Model                                                                                                                                | Dialects                                                | Terms  | Dataset                                                                   | Results                                                                                                                                     | Setting                                                                                                               |
|-------|--------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------|--------|---------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------|
| [173] | ArCT5 base,<br>ArCT52-base-002A,<br>ArCT5-ArabicT5,<br>ArCT5-MSA-Small<br>and ArCT5-MSA-Basic,<br>ArCT5-TensorFlow<br>ArCT5-DeepBlue | Mixing dialects                                         | 10,000 | MAJMAH                                                                    | Training performance:<br>ArCT52-base-002A<br>(1.0) (with dev set)<br>0.03 (with test set)                                                   | Pre-training:<br>Sequence length: 20 steps,<br>Learning rate: 0.001, 2e-7<br>Weight decay: 0.01                       |
| [174] | multilingual T5 v0.7,<br>multilingual base model<br>T5 v0.7 (v0.7b)<br>and ArCT5                                                     | Egyptian, Jordanian,<br>Lebanese, and<br>Persian/Arabic | 10,000 | MAJMAH, MADRC,<br>DUALMSA, Arabic NLP<br>Dev-GENOT                        | Joint Training:<br>ArCT5: 11.12,<br>T5T5: 11.13<br>ArCT5v2: 11.36<br>Independent Training:<br>ArCT5: 12.07<br>T5T5: 11.27<br>ArCT5v2: 15.14 | Pre-training:<br>Joint Training<br>using various datasets<br>Independent Training<br>using single dialect             |
| [175] | ArCT5                                                                                                                                | Mixing dialects                                         | 10,000 | MAJMAH                                                                    | 11.27                                                                                                                                       | Pre-training:<br>Learning rate: 4e-5<br>Weight decay: 1e-2<br>Scheduler: Cosine Annealing<br>10 epochs, batch size: 2 |
| [176] | ArCT5                                                                                                                                | Egyptian, Gulf,<br>Iraqi, Levantine,<br>Maghrebis       | 10,000 | Combination of<br>MAJMAH, MADRC,<br>North-Lebanese<br>Syria Arabic (DPT3) | 11.76                                                                                                                                       | Pre-training:<br>Learning rate: 3e-3<br>Batch size: 32                                                                |
| [177] | ArCT5 base,<br>ArCT52-base-002A,<br>ArCT5-MSA-Small,<br>ArCT5-MSA-Basic                                                              | Egyptian, Gulf,<br>Iraqi, Levantine,<br>Maghrebis       | 10,000 | MAJMAH, MADRC,<br>DUALMSA, Arabic NLP,<br>KAGNU                           | 2T5 (with dev set)<br>4.33 (with test set)                                                                                                  | Pre-training:<br>Sequence length: 128 tokens<br>Learning rate: 5e-5<br>2 epochs, batch size: 16                       |
| [178] | GPT-3.5, ArCT5, MTEB                                                                                                                 | Egyptian, Gulf,<br>Iraqi, Levantine,<br>Maghrebis       | 10,000 | MAJMAH                                                                    | ArCT5: 50.6<br>ArCT5: 10.41<br>MTEB: 11.56                                                                                                  | 25-<br>few-shot                                                                                                       |
| [179] | GPT-3.5, GPT-4<br>and ArCT5 V2                                                                                                       | Egyptian, Gulf,<br>Iraqi, Levantine,<br>Maghrebis       | 10,000 | MAJMAH                                                                    | 30.7 (with dev set)<br>12.6 (with test set)                                                                                                 | Pre-training:<br>Sequence length: 128 tokens<br>LR: 4e-5, batch size: 16                                              |



FIGURE 15. Percentage distribution of approaches used for text formality transfer.

V. DISCUSSION

Our study outlines several recent studies for text style transfer. Specifically, we considered Arabic text style transfer as a special case of MT, where informal Arabic text style (e.g., ADs) are converted (i.e., transformed) to formal Arabic text style (i.e., MSA). Thus, the study focuses on MT models between ADs and MSA, with special emphasis on the recent approaches like sequence-to-sequence and LLMs. Below, we consolidate the key findings, contributions, limitations,

and future research gaps envisaged from these diverse approaches studied.

In this section, we address the challenges encountered in this research area and identify some research gaps that could be explored in future work. We highlight several research gaps that emerged from our analysis of the reviewed papers. The discussion primarily focuses on models that utilize sequence-to-sequence architectures and LLMs.

A. DISCUSSING THE Seq2Seq APPROACH EXPLORED FOR ARABIC TEXT FORMALITY TRANSFER

Numerous sequence-to-sequence techniques utilize transformer models to achieve exceptional translation quality. This clearly illustrates the strength of modern deep learning models, which perform MT tasks with greater accuracy compared to traditional methods like rule-based or statistical approaches.

In the context of NMT, the small range of the covered languages indicates that MT for AD is in its early stages. Certainly, all the efforts presented with the reviewed studies were focused on translating between various ADs, MSA, and English. It was noted that there is a scarcity of translations from ADs to other foreign languages.

The dialects frequently associated with this research area (as indicated in Table 9) are mainly those spoken in Algeria,

TABLE 11. Summary of reviewed papers for non-finite-time formality studies.

| Reference | Year | Q <sub>1</sub> | Q <sub>2</sub> | Q <sub>3</sub> | Q <sub>4</sub> | Q <sub>5</sub> | Q <sub>6</sub> | Q <sub>7</sub> | Q <sub>8</sub> | Q <sub>9</sub> | Q <sub>10</sub> | Q <sub>11</sub> | Q <sub>12</sub> | Q <sub>13</sub> | Q <sub>14</sub> | Q <sub>15</sub> | Q <sub>16</sub> | Q <sub>17</sub> | Q <sub>18</sub> | Q <sub>19</sub> | Q <sub>20</sub> | Q <sub>21</sub> | Q <sub>22</sub> | Q <sub>23</sub> | Q <sub>24</sub> | Q <sub>25</sub> | Q <sub>26</sub> | Q <sub>27</sub> | Q <sub>28</sub> | Q <sub>29</sub> | Q <sub>30</sub> | Q <sub>31</sub> | Q <sub>32</sub> | Q <sub>33</sub> | Q <sub>34</sub> | Q <sub>35</sub> | Q <sub>36</sub> | Q <sub>37</sub> | Q <sub>38</sub> | Q <sub>39</sub> | Q <sub>40</sub> | Q <sub>41</sub> | Q <sub>42</sub> | Q <sub>43</sub> | Q <sub>44</sub> | Q <sub>45</sub> | Q <sub>46</sub> | Q <sub>47</sub> | Q <sub>48</sub> | Q <sub>49</sub> | Q <sub>50</sub> | Q <sub>51</sub> | Q <sub>52</sub> | Q <sub>53</sub> | Q <sub>54</sub> | Q <sub>55</sub> | Q <sub>56</sub> | Q <sub>57</sub> | Q <sub>58</sub> | Q <sub>59</sub> | Q <sub>60</sub> | Q <sub>61</sub> | Q <sub>62</sub> | Q <sub>63</sub> | Q <sub>64</sub> | Q <sub>65</sub> | Q <sub>66</sub> | Q <sub>67</sub> | Q <sub>68</sub> | Q <sub>69</sub> | Q <sub>70</sub> | Q <sub>71</sub> | Q <sub>72</sub> | Q <sub>73</sub> | Q <sub>74</sub> | Q <sub>75</sub> | Q <sub>76</sub> | Q <sub>77</sub> | Q <sub>78</sub> | Q <sub>79</sub> | Q <sub>80</sub> | Q <sub>81</sub> | Q <sub>82</sub> | Q <sub>83</sub> | Q <sub>84</sub> | Q <sub>85</sub> | Q <sub>86</sub> | Q <sub>87</sub> | Q <sub>88</sub> | Q <sub>89</sub> | Q <sub>90</sub> | Q <sub>91</sub> | Q <sub>92</sub> | Q <sub>93</sub> | Q <sub>94</sub> | Q <sub>95</sub> | Q <sub>96</sub> | Q <sub>97</sub> | Q <sub>98</sub> | Q <sub>99</sub> | Q <sub>100</sub> | Q <sub>101</sub> | Q <sub>102</sub> | Q <sub>103</sub> | Q <sub>104</sub> | Q <sub>105</sub> | Q <sub>106</sub> | Q <sub>107</sub> | Q <sub>108</sub> | Q <sub>109</sub> | Q <sub>110</sub> | Q <sub>111</sub> | Q <sub>112</sub> | Q <sub>113</sub> | Q <sub>114</sub> | Q <sub>115</sub> | Q <sub>116</sub> | Q <sub>117</sub> | Q <sub>118</sub> | Q <sub>119</sub> | Q <sub>120</sub> | Q <sub>121</sub> | Q <sub>122</sub> | Q <sub>123</sub> | Q <sub>124</sub> | Q <sub>125</sub> | Q <sub>126</sub> | Q <sub>127</sub> | Q <sub>128</sub> | Q <sub>129</sub> | Q <sub>130</sub> | Q <sub>131</sub> | Q <sub>132</sub> | Q <sub>133</sub> | Q <sub>134</sub> | Q <sub>135</sub> | Q <sub>136</sub> | Q <sub>137</sub> | Q <sub>138</sub> | Q <sub>139</sub> | Q <sub>140</sub> | Q <sub>141</sub> | Q <sub>142</sub> | Q <sub>143</sub> | Q <sub>144</sub> | Q <sub>145</sub> | Q <sub>146</sub> | Q <sub>147</sub> | Q <sub>148</sub> | Q <sub>149</sub> | Q <sub>150</sub> | Q <sub>151</sub> | Q <sub>152</sub> | Q <sub>153</sub> | Q <sub>154</sub> | Q <sub>155</sub> | Q <sub>156</sub> | Q <sub>157</sub> | Q <sub>158</sub> | Q <sub>159</sub> | Q <sub>160</sub> | Q <sub>161</sub> | Q <sub>162</sub> | Q <sub>163</sub> | Q <sub>164</sub> | Q <sub>165</sub> | Q <sub>166</sub> | Q <sub>167</sub> | Q <sub>168</sub> | Q <sub>169</sub> | Q <sub>170</sub> | Q <sub>171</sub> | Q <sub>172</sub> | Q <sub>173</sub> | Q <sub>174</sub> | Q <sub>175</sub> | Q <sub>176</sub> | Q <sub>177</sub> | Q <sub>178</sub> | Q <sub>179</sub> | Q <sub>180</sub> | Q <sub>181</sub> | Q <sub>182</sub> | Q <sub>183</sub> | Q <sub>184</sub> | Q <sub>185</sub> | Q <sub>186</sub> | Q <sub>187</sub> | Q <sub>188</sub> | Q <sub>189</sub> | Q <sub>190</sub> | Q <sub>191</sub> | Q <sub>192</sub> | Q <sub>193</sub> | Q <sub>194</sub> | Q <sub>195</sub> | Q <sub>196</sub> | Q <sub>197</sub> | Q <sub>198</sub> | Q <sub>199</sub> | Q <sub>200</sub> | Q <sub>201</sub> | Q <sub>202</sub> | Q <sub>203</sub> | Q <sub>204</sub> | Q <sub>205</sub> | Q <sub>206</sub> | Q <sub>207</sub> | Q <sub>208</sub> | Q <sub>209</sub> | Q <sub>210</sub> | Q <sub>211</sub> | Q <sub>212</sub> | Q <sub>213</sub> | Q <sub>214</sub> | Q <sub>215</sub> | Q <sub>216</sub> | Q <sub>217</sub> | Q <sub>218</sub> | Q <sub>219</sub> | Q <sub>220</sub> | Q <sub>221</sub> | Q <sub>222</sub> | Q <sub>223</sub> | Q <sub>224</sub> | Q <sub>225</sub> | Q <sub>226</sub> | Q <sub>227</sub> | Q <sub>228</sub> | Q <sub>229</sub> | Q <sub>230</sub> | Q <sub>231</sub> | Q <sub>232</sub> | Q <sub>233</sub> | Q <sub>234</sub> | Q <sub>235</sub> | Q <sub>236</sub> | Q <sub>237</sub> | Q <sub>238</sub> | Q <sub>239</sub> | Q <sub>240</sub> | Q <sub>241</sub> | Q <sub>242</sub> | Q <sub>243</sub> | Q <sub>244</sub> | Q <sub>245</sub> | Q <sub>246</sub> | Q <sub>247</sub> | Q <sub>248</sub> | Q <sub>249</sub> | Q <sub>250</sub> | Q <sub>251</sub> | Q <sub>252</sub> | Q <sub>253</sub> | Q <sub>254</sub> | Q <sub>255</sub> | Q <sub>256</sub> | Q <sub>257</sub> | Q <sub>258</sub> | Q <sub>259</sub> | Q <sub>260</sub> | Q <sub>261</sub> | Q <sub>262</sub> | Q <sub>263</sub> | Q <sub>264</sub> | Q <sub>265</sub> | Q <sub>266</sub> | Q <sub>267</sub> | Q <sub>268</sub> | Q <sub>269</sub> | Q <sub>270</sub> | Q <sub>271</sub> | Q <sub>272</sub> | Q <sub>273</sub> | Q <sub>274</sub> | Q <sub>275</sub> | Q <sub>276</sub> | Q <sub>277</sub> | Q <sub>278</sub> | Q <sub>279</sub> | Q <sub>280</sub> | Q <sub>281</sub> | Q <sub>282</sub> | Q <sub>283</sub> | Q <sub>284</sub> | Q <sub>285</sub> | Q <sub>286</sub> | Q <sub>287</sub> | Q <sub>288</sub> | Q <sub>289</sub> | Q <sub>290</sub> | Q <sub>291</sub> | Q <sub>292</sub> | Q <sub>293</sub> | Q <sub>294</sub> | Q <sub>295</sub> | Q <sub>296</sub> | Q <sub>297</sub> | Q <sub>298</sub> | Q <sub>299</sub> | Q <sub>300</sub> | Q <sub>301</sub> | Q <sub>302</sub> | Q <sub>303</sub> | Q <sub>304</sub> | Q <sub>305</sub> | Q <sub>306</sub> | Q <sub>307</sub> | Q <sub>308</sub> | Q <sub>309</sub> | Q <sub>310</sub> | Q <sub>311</sub> | Q <sub>312</sub> | Q <sub>313</sub> | Q <sub>314</sub> | Q <sub>315</sub> | Q <sub>316</sub> | Q <sub>317</sub> | Q <sub>318</sub> | Q <sub>319</sub> | Q <sub>320</sub> | Q <sub>321</sub> | Q <sub>322</sub> | Q <sub>323</sub> | Q <sub>324</sub> | Q <sub>325</sub> | Q <sub>326</sub> | Q <sub>327</sub> | Q <sub>328</sub> | Q <sub>329</sub> | Q <sub>330</sub> | Q <sub>331</sub> | Q <sub>332</sub> | Q <sub>333</sub> | Q <sub>334</sub> | Q <sub>335</sub> | Q <sub>336</sub> | Q <sub>337</sub> | Q <sub>338</sub> | Q <sub>339</sub> | Q <sub>340</sub> | Q <sub>341</sub> | Q <sub>342</sub> | Q <sub>343</sub> | Q <sub>344</sub> | Q <sub>345</sub> | Q <sub>346</sub> | Q <sub>347</sub> | Q <sub>348</sub> | Q <sub>349</sub> | Q <sub>350</sub> | Q <sub>351</sub> | Q <sub>352</sub> | Q <sub>353</sub> | Q <sub>354</sub> | Q <sub>355</sub> | Q <sub>356</sub> | Q <sub>357</sub> | Q <sub>358</sub> | Q <sub>359</sub> | Q <sub>360</sub> | Q <sub>361</sub> | Q <sub>362</sub> | Q <sub>363</sub> | Q <sub>364</sub> | Q <sub>365</sub> | Q <sub>366</sub> | Q <sub>367</sub> | Q <sub>368</sub> | Q <sub>369</sub> | Q <sub>370</sub> | Q <sub>371</sub> | Q <sub>372</sub> | Q <sub>373</sub> | Q <sub>374</sub> | Q <sub>375</sub> | Q <sub>376</sub> | Q <sub>377</sub> | Q <sub>378</sub> | Q <sub>379</sub> | Q <sub>380</sub> | Q <sub>381</sub> | Q <sub>382</sub> | Q <sub>383</sub> | Q <sub>384</sub> | Q <sub>385</sub> | Q <sub>386</sub> | Q <sub>387</sub> | Q <sub>388</sub> | Q <sub>389</sub> | Q <sub>390</sub> | Q <sub>391</sub> | Q <sub>392</sub> | Q <sub>393</sub> | Q <sub>394</sub> | Q <sub>395</sub> | Q <sub>396</sub> | Q <sub>397</sub> | Q <sub>398</sub> | Q <sub>399</sub> | Q <sub>400</sub> | Q <sub>401</sub> | Q <sub>402</sub> | Q <sub>403</sub> | Q <sub>404</sub> | Q <sub>405</sub> | Q <sub>406</sub> | Q <sub>407</sub> | Q <sub>408</sub> | Q <sub>409</sub> | Q <sub>410</sub> | Q <sub>411</sub> | Q <sub>412</sub> | Q <sub>413</sub> | Q <sub>414</sub> | Q <sub>415</sub> | Q <sub>416</sub> | Q <sub>417</sub> | Q <sub>418</sub> | Q <sub>419</sub> | Q <sub>420</sub> | Q <sub>421</sub> | Q <sub>422</sub> | Q <sub>423</sub> | Q <sub>424</sub> | Q <sub>425</sub> | Q <sub>426</sub> | Q <sub>427</sub> | Q <sub>428</sub> | Q <sub>429</sub> | Q <sub>430</sub> | Q <sub>431</sub> | Q <sub>432</sub> | Q <sub>433</sub> | Q <sub>434</sub> | Q <sub>435</sub> | Q <sub>436</sub> | Q <sub>437</sub> | Q <sub>438</sub> | Q <sub>439</sub> | Q <sub>440</sub> | Q <sub>441</sub> | Q <sub>442</sub> | Q <sub>443</sub> | Q <sub>444</sub> | Q <sub>445</sub> | Q <sub>446</sub> | Q <sub>447</sub> | Q <sub>448</sub> | Q <sub>449</sub> | Q <sub>450</sub> | Q <sub>451</sub> | Q <sub>452</sub> | Q <sub>453</sub> | Q <sub>454</sub> | Q <sub>455</sub> | Q <sub>456</sub> | Q <sub>457</sub> | Q <sub>458</sub> | Q <sub>459</sub> | Q <sub>460</sub> | Q <sub>461</sub> | Q <sub>462</sub> | Q <sub>463</sub> | Q <sub>464</sub> | Q <sub>465</sub> | Q <sub>466</sub> | Q <sub>467</sub> | Q <sub>468</sub> | Q <sub>469</sub> | Q <sub>470</sub> | Q <sub>471</sub> | Q <sub>472</sub> | Q <sub>473</sub> | Q <sub>474</sub> | Q <sub>475</sub> | Q <sub>476</sub> | Q <sub>477</sub> | Q <sub>478</sub> | Q <sub>479</sub> | Q <sub>480</sub> | Q <sub>481</sub> | Q <sub>482</sub> | Q <sub>483</sub> | Q <sub>484</sub> | Q <sub>485</sub> | Q <sub>486</sub> | Q <sub>487</sub> | Q <sub>488</sub> | Q <sub>489</sub> | Q <sub>490</sub> | Q <sub>491</sub> | Q <sub>492</sub> | Q <sub>493</sub> | Q <sub>494</sub> | Q <sub>495</sub> | Q <sub>496</sub> | Q <sub>497</sub> | Q <sub>498</sub> | Q <sub>499</sub> | Q <sub>500</sub> | Q <sub>501</sub> | Q <sub>502</sub> | Q <sub>503</sub> | Q <sub>504</sub> | Q <sub>505</sub> | Q <sub>506</sub> | Q <sub>507</sub> | Q <sub>508</sub> | Q <sub>509</sub> | Q <sub>510</sub> | Q <sub>511</sub> | Q <sub>512</sub> | Q <sub>513</sub> | Q <sub>514</sub> | Q <sub>515</sub> | Q <sub>516</sub> | Q <sub>517</sub> | Q <sub>518</sub> | Q <sub>519</sub> | Q <sub>520</sub> | Q <sub>521</sub> | Q <sub>522</sub> | Q <sub>523</sub> | Q <sub>524</sub> | Q <sub>525</sub> | Q <sub>526</sub> | Q <sub>527</sub> | Q <sub>528</sub> | Q <sub>529</sub> | Q <sub>530</sub> | Q <sub>531</sub> | Q <sub>532</sub> | Q <sub>533</sub> | Q <sub>534</sub> | Q <sub>535</sub> | Q <sub>536</sub> | Q <sub>537</sub> | Q <sub>538</sub> | Q <sub>539</sub> | Q <sub>540</sub> | Q <sub>541</sub> | Q <sub>542</sub> | Q <sub>543</sub> | Q <sub>544</sub> | Q <sub>545</sub> | Q <sub>546</sub> | Q <sub>547</sub> | Q <sub>548</sub> | Q <sub>549</sub> | Q <sub>550</sub> | Q <sub>551</sub> | Q <sub>552</sub> | Q <sub>553</sub> | Q <sub>554</sub> | Q <sub>555</sub> | Q <sub>556</sub> | Q <sub>557</sub> | Q <sub>558</sub> | Q <sub>559</sub> | Q <sub>560</sub> | Q <sub>561</sub> | Q <sub>562</sub> | Q <sub>563</sub> | Q <sub>564</sub> | Q <sub>565</sub> | Q <sub>566</sub> | Q <sub>567</sub> | Q <sub>568</sub> | Q <sub>569</sub> | Q <sub>570</sub> | Q <sub>571</sub> | Q <sub>572</sub> | Q <sub>573</sub> | Q <sub>574</sub> | Q <sub>575</sub> | Q <sub>576</sub> | Q <sub>577</sub> | Q <sub>578</sub> | Q <sub>579</sub> | Q <sub>580</sub> | Q <sub>581</sub> | Q <sub>582</sub> | Q <sub>583</sub> | Q <sub>584</sub> | Q <sub>585</sub> | Q <sub>586</sub> | Q <sub>587</sub> | Q <sub>588</sub> | Q <sub>589</sub> | Q <sub>590</sub> | Q <sub>591</sub> | Q <sub>592</sub> | Q <sub>593</sub> | Q <sub>594</sub> | Q <sub>595</sub> | Q <sub>596</sub> | Q <sub>597</sub> | Q <sub>598</sub> | Q <sub>599</sub> | Q <sub>600</sub> | Q <sub>601</sub> | Q <sub>602</sub> | Q <sub>603</sub> | Q <sub>604</sub> | Q <sub>605</sub> | Q <sub>606</sub> | Q <sub>607</sub> | Q <sub>608</sub> | Q <sub>609</sub> | Q <sub>610</sub> | Q <sub>611</sub> | Q <sub>612</sub> | Q <sub>613</sub> | Q <sub>614</sub> | Q <sub>615</sub> | Q <sub>616</sub> | Q <sub>617</sub> | Q <sub>618</sub> | Q <sub>619</sub> | Q <sub>620</sub> | Q <sub>621</sub> | Q <sub>622</sub> | Q <sub>623</sub> | Q <sub>624</sub> | Q <sub>625</sub> | Q <sub>626</sub> | Q <sub>627</sub> | Q <sub>628</sub> | Q <sub>629</sub> | Q <sub>630</sub> | Q <sub>631</sub> | Q <sub>632</sub> | Q <sub>633</sub> | Q <sub>634</sub> | Q <sub>635</sub> | Q <sub>636</sub> | Q <sub>637</sub> | Q <sub>638</sub> | Q <sub>639</sub> | Q <sub>640</sub> | Q <sub>641</sub> | Q <sub>642</sub> | Q <sub>643</sub> | Q <sub>644</sub> | Q <sub>645</sub> | Q <sub>646</sub> | Q <sub>647</sub> | Q <sub>648</sub> | Q <sub>649</sub> | Q <sub>650</sub> | Q <sub>651</sub> | Q <sub>652</sub> | Q <sub>653</sub> | Q <sub>654</sub> | Q <sub>655</sub> | Q <sub>656</sub> | Q <sub>657</sub> | Q <sub>658</sub> | Q <sub>659</sub> | Q <sub>660</sub> | Q <sub>661</sub> | Q <sub>662</sub> | Q <sub>663</sub> | Q <sub>664</sub> | Q <sub>665</sub> | Q <sub>666</sub> | Q <sub>667</sub> | Q <sub>668</sub> | Q <sub>669</sub> | Q <sub>670</sub> | Q <sub>671</sub> | Q <sub>672</sub> | Q <sub>673</sub> | Q <sub>674</sub> | Q <sub>675</sub> | Q <sub>676</sub> | Q <sub>677</sub> | Q <sub>678</sub> | Q <sub>679</sub> | Q <sub>680</sub> | Q <sub>681</sub> | Q <sub>682</sub> | Q <sub>683</sub> | Q <sub>684</sub> | Q <sub>685</sub> | Q <sub>686</sub> | Q <sub>687</sub> | Q <sub>688</sub> | Q <sub>689</sub> | Q <sub>690</sub> | Q <sub>691</sub> | Q <sub>692</sub> | Q <sub>693</sub> | Q <sub>694</sub> | Q <sub>695</sub> | Q <sub>696</sub> | Q <sub>697</sub> | Q <sub>698</sub> | Q <sub>699</sub> | Q <sub>700</sub> | Q <sub>701</sub> | Q <sub>702</sub> | Q <sub>703</sub> | Q <sub>704</sub> | Q <sub>705</sub> | Q <sub>706</sub> | Q <sub>707</sub> | Q <sub>708</sub> | Q <sub>709</sub> | Q <sub>710</sub> | Q <sub>711</sub> | Q <sub>712</sub> | Q <sub>713</sub> | Q <sub>714</sub> | Q <sub>715</sub> | Q <sub>716</sub> | Q <sub>717</sub> | Q <sub>718</sub> | Q <sub>719</sub> | Q <sub>720</sub> | Q <sub>721</sub> | Q <sub>722</sub> | Q <sub>723</sub> | Q <sub>724</sub> | Q <sub>725</sub> | Q <sub>726</sub> | Q <sub>727</sub> | Q <sub>728</sub> | Q <sub>729</sub> | Q <sub>730</sub> | Q <sub>731</sub> | Q <sub>732</sub> | Q <sub>733</sub> | Q <sub>734</sub> | Q <sub>735</sub> | Q <sub>736</sub> | Q <sub>737</sub> | Q <sub>738</sub> | Q <sub>739</sub> | Q <sub>740</sub> | Q <sub>741</sub> | Q <sub>742</sub> | Q <sub>743</sub> | Q <sub>744</sub> | Q <sub>745</sub> | Q <sub>746</sub> | Q <sub>747</sub> | Q <sub>748</sub> | Q <sub>749</sub> | Q <sub>750</sub> | Q <sub>751</sub> | Q <sub>752</sub> | Q <sub>753</sub> | Q <sub>754</sub> | Q <sub>755</sub> | Q <sub>756</sub> | Q <sub>757</sub> | Q <sub>758</sub> | Q <sub>759</sub> | Q <sub>760</sub> | Q <sub>761</sub> | Q <sub>762</sub> | Q <sub>763</sub> | Q <sub>764</sub> | Q <sub>765</sub> | Q <sub>766</sub> | Q <sub>767</sub> | Q <sub>768</sub> | Q <sub>769</sub> | Q <sub>770</sub> | Q <sub>771</sub> | Q <sub>772</sub> | Q <sub>773</sub> | Q <sub>774</sub> | Q <sub>775</sub> | Q <sub>776</sub> | Q <sub>777</sub> | Q <sub>778</sub> | Q <sub>779</sub> | Q <sub>780</sub> | Q <sub>781</sub> | Q <sub>782</sub> | Q <sub>783</sub> | Q <sub>784</sub> | Q <sub>785</sub> | Q <sub>786</sub> | Q <sub>787</sub> | Q <sub>788</sub> | Q <sub>789</sub> | Q <sub>790</sub> | Q <sub>791</sub> | Q <sub>792</sub> | Q <sub>793</sub> | Q <sub>794</sub> | Q <sub>795</sub> | Q <sub>796</sub> | Q <sub>797</sub> | Q <sub>798</sub> | Q <sub>799</sub> | Q <sub>800</sub> | Q <sub>801</sub> | Q <sub>802</sub> | Q <sub>803</sub> | Q <sub>804</sub> | Q <sub>805</sub> | Q <sub>806</sub> | Q <sub>807</sub> | Q <sub>808</sub> | Q <sub>809</sub> | Q <sub>810</sub> | Q <sub>811</sub> | Q <sub>812</sub> | Q <sub>813</sub> | Q <sub>814</sub> | Q <sub>815</sub> | Q <sub>816</sub> | Q <sub>817</sub> | Q <sub>818</sub> | Q <sub>819</sub> | Q <sub>820</sub> | Q <sub>821</sub> | Q <sub>822</sub> | Q <sub>823</sub> | Q <sub>824</sub> | Q <sub>825</sub> | Q <sub>826</sub> | Q <sub>827</sub> | Q <sub>828</sub> | Q <sub>829</sub> | Q <sub>830</sub> | Q <sub>831</sub> | Q <sub>832</sub> | Q <sub>833</sub> | Q <sub>834</sub> | Q <sub>835</sub> | Q <sub>836</sub> | Q <sub>837</sub> | Q <sub>838</sub> | Q <sub>839</sub> | Q <sub>840</sub> | Q <sub>841</sub> | Q <sub>842</sub> | Q <sub>843</sub> | Q <sub>844</sub> | Q <sub>845</sub> | Q <sub>846</sub> | Q <sub>847</sub> | Q <sub>848</sub> | Q <sub>849</sub> | Q <sub>850</sub> | Q <sub>851</sub> | Q <sub>852</sub> | Q <sub>853</sub> | Q <sub>854</sub> | Q <sub>855</sub> | Q <sub>856</sub> | Q <sub>857</sub> | Q <sub>858</sub> | Q <sub>859</sub> | Q <sub>860</sub> | Q <sub>861</sub> | Q <sub>862</sub> | Q <sub>863</sub> | Q <sub>864</sub> | Q <sub>865</sub> | Q <sub>866</sub> | Q <sub>867</sub> | Q <sub>868</sub> | Q <sub>869</sub> | Q <sub>870</sub> | Q <sub>871</sub> | Q <sub>872</sub> | Q <sub>873</sub> | Q <sub>874</sub> | Q <sub>875</sub> | Q <sub>876</sub> | Q <sub>877</sub> | Q <sub>878</sub> | Q <sub></sub> |
|-----------|------|----------------|----------------|----------------|----------------|----------------|----------------|----------------|----------------|----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|---------------|
|-----------|------|----------------|----------------|----------------|----------------|----------------|----------------|----------------|----------------|----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|---------------|

Jordan, and Egypt. In contrast, dialects from regions like Saudi Arabia, Oman, Kuwait, Bahrain, and Sudan appear less frequently. Therefore, exploring performance of utilizing more Arabic dialects is a prominent research direction.

We noticed also that some of the reviewed papers focus on single dialects [159], [160], [161], leading to a shortage of models that encompass the complete range of dialects. Investigating the translation of various ADs into MSA can greatly contribute to assessing the robustness of the developed model.

From other side, most of the NMT-based AD to MSA systems reviewed in Subsection IV-B2 (see Table 9) employed single performance metric [162], [163], [164], [165]. However, relying on single evaluation measure would provide limited perspective and difficulty in comparison of the neural translation quality. Hence, for comprehensive evaluation of the NMT systems it is advisable to employ multiple evaluation metrics such as Rouge score, METEOR score, perplexity (See Subsection IV-B1).

Regarding the methodology used for translating AD to MSA, RNN-based uncode-decoder is the most technique explored with seq2seq approach followed by transformer with attention mechanism. Therefore, there are more recent uncode-decoder models can be explored in the future for improving performance of Arabic text formality transfer.

As a result of our analysis, we can now clearly that NMT still grapples with several weaknesses: challenges in handling extensive vocabularies, inaccuracies, or errors in translating source words, and the necessity for a substantial amount of data for training.

## B. DISCUSSING THE LLMs APPROACH EXPLORED FOR ARABIC TEXT FORMALITY TRANSFER

As shown in Section IV-E3 (Table 10), there are limited studies investigating the capabilities of recent LLMs for translating ADs to MSA. In particular, the selected studies employed Arabian version of T5 model (AraT5) and GPT model (GPT-3.5 and GPT-4) to translate between ADs and MSA. Therefore, this indicates that MT between ADs and MSA based on LLMs and pre-trained transformers has just started and there is a need to explore this direction further.

In terms of the performance evaluation, all the selected studies with this category employed a single performance measure. Specifically, BLEU score is used with these selected papers. Therefore, as stated earlier depending on single performance metric could lead to bias assessment and lack of robustness in assessing the MT quality.

It is clear that there are fewer parallel corpora available for MSA and ADs compared to those available for English (see Table 7). In this regard, most researchers built their own private datasets to examine their proposed approach.

## C. RESEARCH CHALLENGES AND FUTURE DIRECTIONS

This section presents a comprehensive overview of the challenges and potential future directions in Arabic text formality transfer.

### 1) CHALLENGES IN ARABIC TEXT FORMALITY TRANSFER

Arabic is a linguistically intricate language [9] with a diverse inflectional structure, showcased through templates and affixes, alongside various classes of attachable clitics. The Arabic language showcases intricate nuances in formality, intricately woven into its linguistic fabric and cultural landscapes. Navigating the fluidity of Arabic text poses both obstacles and openings for scholars to adjust its levels of formality.

Formality transfer encounters challenges when the initial formality of the Arabic text diverges from the desired audience or communicative setting. This discrepancy is evident across various platforms such as social media, news articles, and formal documents, where the suitable degree of formality may vary substantially.

To effectively transfer formality in Arabic text, it is crucial to go beyond just syntactic and semantic changes. Researchers must also consider cultural norms and the intended communicative purpose of the text, as ADs represent a wide range of regional dialects spoken throughout Arab countries. The emergence of social networks has substantially boosted the spread of these complicated dialects, integrating them into the fabric of daily communication. A significant amount of text generated on major social media networks such as X, Facebook, WhatsApp, and Instagram appears in these dialects.

ADs exhibit significant variation across geographic regions as (refer to Fig. 12), leading to challenges in comprehension and interpretation for those unfamiliar with a particular dialect. The language-related variations could be so substantial to the extent that in a singular nation, the same words may possess distinct interpretations. Consequently, because of the diversity among these ADs, it becomes extremely difficult to develop models that can effectively process Arabic social network platforms and other related contents.

A key approach to addressing these challenges involves capitalizing on the extensive lexicon, vast vocabulary, and well-defined structure of MSA by translating the various ADs into MSA. Whereas, the available data for Arabic text typically contains optional diacritics, spelling inconsistencies, and words written in AD is subjected to Arabization (al-Tarib) [224]. As a result, there is a scarcity of research concerning MT for these dialects.

The complexity of the Arabic language [225] renders the process of formality transfer very intricate, consequently diminishing the performance of such task. Navigating the formal attributes of Arabic text necessitates a nuanced grasp of its grammatical structures, lexical preferences, and stylistic norms. Arabic stands apart from numerous other languages by incorporating diverse linguistic markers and honorifics to signal formality. Consequently, numerous ADs [16] are accessible throughout the Arab world which is contributing intricate layers of complexity to the formality transfer endeavor.

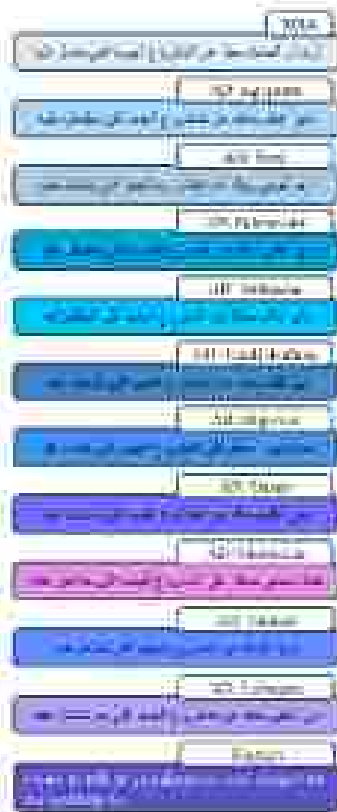


FIGURE 12. Sample sentences in MSA along with their conversions into various ADs.



FIGURE 13. Examples of HQ words and their corresponding MSA forms.

Moreover, there are various dialectal Arabic varieties at the word level. Dialectal Arabic significantly diverges from MSA phonologically, morphologically, and to a lesser extent, syntactically. Fig. 13 presents various Arabic words from different dialects alongside their MSA counterparts. It is evident from the figure that some words undergo a complete transformation when adapted to MSA. Thereby, converting AD to MSA poses numerous challenges, including the need to address these phonological shifts and morphological changes effectively. Additionally, regional idiomatic expressions and colloquialisms further complicate the conversion process, as they often lack direct MSA equivalents.

## 2) FUTURE WORK

As shown in Section IVF, we reviewed very recent approaches for text formality transfer. Such approaches can

be explored to improve the research prospects of Arabic text formality transfer domain. For instance, developing an innovative unified style transfer method to shift Arabic text from MSA to various ADs through out-of-domain [184] transfer presents a significant advancement, offering a promising solution to address challenges in Arabic text formality transfer.

We also observed that none of the available works in the literature studied Arabic text formality transfer using diffusion-based models. Whereas, Horvitz et al. [221] employ diffusion probabilistic models for English text formality transfer. Thus, their approach could be extended to tackle Arabic text formality transfer and check whether diffusion-based models hold an edge in respect to Arabic domain.

There remain several limitations in the field of Arabic text formality transfer that could pave the way for future research. The identified shortcomings include the following:

- Our analysis indicates that the studies examined predominantly utilize AraT5 and GPT models. Consequently, it is important to assess the effectiveness of incorporating additional LLMs, particularly those designed for Arabic text, like ArabicLlama3,<sup>14</sup> Jais,<sup>15</sup> and AceGPT.<sup>16</sup>
- Previous studies on Arabic text formality transfer have predominantly relied on parallel datasets for model training. Consequently, it is important to investigate the effectiveness of non-parallel datasets [168] in Arabic text formality transfer, particularly given the significant challenges associated with creating parallel datasets [226], [227].
- The earlier research on Arabic text formality transfer primarily relies on supervised learning techniques. Consequently, it is important to investigate how unsupervised [160], [238], semi-supervised, and reinforcement [215] learning methods could perform in this area of study. It is worth mentioning here that reinforcement learning can be effectively applied to the process of text formality transfer in an unsupervised learning [220]. There is also a need to explore unsupervised learning with Arabic text formality transfer by using a pseudo-parallel dataset construction [216].
- The existing research primarily relies on the BLEU score to assess the performance of Arabic text formality transfer. Therefore, it is important to consider additional evaluation metrics [229], [230], [231], [232] for a more comprehensive assessment [237]. Additionally, introducing new automatic evaluation metrics [230], [235], [236] tailored specifically for Arabic text would be highly beneficial. One approach to creating such metrics involves examining how named entities contribute to content preservation [237]. Additionally, there is a requirement to investigate how the content is

<sup>14</sup><https://huggingface.co/HaMozArabic/Arabic-Llama3>

<sup>15</sup><https://huggingface.co/jais-ai/Jais-Inception-71b-V1.1>

<sup>16</sup><https://huggingface.co/PredictionIntelligence/AceGPT-33b-chat>

retained [239] when applying Arabic text formality transfer. Moreover, it is important to automatically assess the coherency of the generated Arabic text.

- Given the extensive use of word embeddings in Arabic text formality transfer, it is essential to assess how different word embeddings [239], [240], [241] influence the effectiveness of this research direction.
- To improve the performance of Arabic text formality transfer, it is crucial to thoroughly examine the relationship between dialectal words and MSA words [242].
- Assessing the effectiveness of automatic word-alignment tools [243] is crucial for efficiently generating large quantities of aligned parallel texts for ADs/MSA.
- Assessing how effectively different word augmentation techniques perform in the context of Arabic text formality transfer is crucial [244].
- Assessing how various Arabic parsers handle the Arabic text formality transfer is crucial [245].
- It is essential to assess how the similarities between ADs [246] influence the performance of Arabic text formality transfer.
- It is important to evaluate the impact of transliteration normalization [207] on Arabic text formality transfer.
- It is essential to develop new parallel datasets [248] for the formality transfer in Arabic text. To expand the current parallel datasets, it is important to investigate techniques for extracting parallel sentences [249].
- It is important to investigate how formality transfer relates to other style transfer tasks, like sentiment transfer [189], [190], [208], [250] and code-switching style transfer [14]. Such research may uncover innovative approaches to enhance the performance of Arabic text formality transfer.
- There is a requirement to develop additional directories [251] for Arabic machine translation.
- Improving the performance of Arabic text formality transfer requires more attention toward the Arabic script [252].
- It is important to investigate the formality transfer in Arabic texts at the document level instead of solely focusing on the sentence level [253].

## VI. CONCLUSION

This paper provides insight into research gaps within the domain of Arabic text mining, specifically focusing on formality transfer. Our research underscores the significance of exploring Arabic text formality transfer from various angles. We conducted a survey of state-of-the-art techniques in machine translation/formality transfer and devised a comparison framework to facilitate the evaluation of different methodologies. Through our analysis, we identified several gaps in the field which we have outlined in preceding sections. Additionally, our literature review reveals limitations in current Arabic formality transfer systems, many of which rely on traditional techniques with limited utilization of LLMs.

For future work, our findings suggest a shift towards evaluating newer LLMs that offer Arabic language support, such as Lla and AceGPT. In addition, addressing the scarcity of parallel datasets through non-parallel dataset analysis may offer solutions to a major research limitation. Moreover, expanding beyond supervised learning to include unsupervised, semi-supervised, and reinforcement learning methods could lead to more robust model performance. It is also valuable to employ additional evaluation metrics for a more effective assessment of content preservation and coherence, which can offer a deeper insight into the subtleties of Arabic text. Further investigation into the role of Arabic word embeddings, dialect-to-MSA relationships, and other techniques like word alignment and transliteration normalization is also crucial to advance Arabic formality transfer research.

## REFERENCES

- [1] A. Alsalhi, H. Mubalak, S. Choudhury, M. Hameed, B. Mousi, T. Hameed, J. Alshafie, N.E. Khatib, D. Ismail, F. Dahi, M. Hossain, K. Hoss, Y. El-Hamady, A. Al-H. Damm, N. Mulla-Faraj, and F. Akim, "ArabicBERT: Benchmarking Arabic at 100+ language models," in *Proc. 20th Conf. Emp. Comput. Asian. Comput. Linguistics*, 2024, pp. 447–520.
- [2] M. Sempuri, P. Nakov, and T. Schuster, "Natural language processing for small languages: methods and datasets," *ArXiv*, National Lang. Exp., vol. 28, no. 4, pp. 293–612, Nov. 2020.
- [3] Z. Hu, B. B. Wu, L. Chen, C. C. Apperloo, and A. Zhang, "Text style transfer: A review and experimental evaluation," *ACM SIGMIS Database*, vol. 54, no. 1, pp. 34–45, Jan. 2022.
- [4] D. Zhang, S. Yu, F. Yu, A. Ganesan, and S. Campbell, "Improving machine translation formality control with weakly labeled data augmentation and post-aligning strategies," in *Proc. 10th Int. Conf. Spoken Lang. Transl. (IWALT)*, 2022, pp. 151–160.
- [5] Y. Sami and M. Gammari, "Trends and challenges of Arabic dialects: Linguistic review," *Arabian J. Comput. Inf. Technol.*, vol. 5, no. 1, pp. 1–28, 2023.
- [6] H. Alfaris, N. Alharbi, M. Alharbi, M. Alharbi, N. Alharbi, S. Alfaris, M. Kari, and H. Alfaris, "Arabic language systems: A survey," in *Proc. 20th Conf. Emp. Comput. Asian. Comput. Linguistics*, vol. 4, London, U.K. Cham, Switzerland: Springer, 2022, pp. 173–184.
- [7] D. A. Ali and H. Haddad, "From Arabic dialect to MSA," in *Proc. COLING, 20th Int. Conf. Comput. Linguistics East Asia*, 2018, pp. 208–212.
- [8] E. Pechen and J. Tervahauta, "An empirical analysis of formality in online communications," *Trans. Assoc. Comput. Linguistics*, vol. 4, pp. 61–74, Dec. 2016.
- [9] M. Alharbi and K. Swaidan, "The key challenges for Arabic machine translation," in *Intelligent Natural Language Processing: Trends and Applications*, Cham, Switzerland: Springer, 2018, pp. 128–136.
- [10] W. Salameh and N. Bahati, "Dialectal Arabic to English machine translation: Passing through dialect standard Arabic," in *Proc. Conf. North Amer. Chapter Asian. Comput. Linguistics: Emp. Lang. Technol.*, Jan. 2013, pp. 348–354.
- [11] E. M. Alharbi and A. Al-Ali, "A hybrid neural machine translation technique for translating low resource languages," in *Proc. 10th Int. Conf. Mach. Learn. Data Mining Pattern Recognit. (ICDM)*, New York, NY, USA: Cham, Switzerland: Springer, Jan. 2018, pp. 147–156.
- [12] E. M. Alharbi, B. Mousi, and K. Alharbi, "Translating standard Arabic to colloquial Arabic," in *Proc. 5th Asian. Machine Learn. Comput. Linguistics*, 2012, pp. 176–186.
- [13] I. Ghallil, F. Asmar, M. Alharbi, and F. Sahli, "Arabic transliteration of Algerian Arabic dialect into modern standard Arabic," *Arabic Inf.*, vol. 2017, pp. 1–8, May 2017.
- [14] H. Hu, S. Lu, Z. Sun, C. Ma, and C. Guo, "XNL-based text style transfer with post word coherence learning," 2021, arXiv:2112.07019.

- [15] A. Elmaghr, S. M. Yagi, A. A. Saad, J. Ibrahim, and S. A. Salameh, "Systematic literature review of dialectal Arabic: Identification and detection," *IEEE Access*, vol. 9, pp. 11010–11042, 2021.
- [16] F. Hameed, K. Hleibani, and K. Amari, "Machine translation for Arabic dialects (arabic)," *Int. Process. Manag.*, vol. 76, no. 2, pp. 262–273, Mar. 2019.
- [17] A. Alqahtani, H. Qureshi, and R. Thaler, "Arabic machine translation: A survey," *Arabic World Sci.*, vol. 42, no. 4, pp. 540–572, Dec. 2019.
- [18] A. M. Salameh, H. Alshabaf, K. Almadhoun, and P. Elshahawy, "Term analysis in formalization: A comprehensive review for enhancing Arabic-English machine translation," in *Proc. 21st Int. Symp. Technol. Conf. (IATC)*, Jan. 2024, pp. 108–114.
- [19] M. M. Elshabaf and T. R. Scornen, "Perspectives of Arabic machine translation," *J. Eng. Sci. Technol.*, vol. 12, no. 9, pp. 1115–1117, 2017.
- [20] S. A. Alshabaf and S. A. Alshabaf, "Challenges in modernizing Arabic (arab) English using machine translation: A systematic literature review," *IEEE Access*, vol. 11, pp. 94772–94779, 2023.
- [21] J. Tahaoui, M. Fathi, S. Al-Madhrif, and J. M. Alqahtani, "Arabic machine translation: A survey with challenges and future directions," *IEEE Access*, vol. 9, pp. 101443–101448, 2021.
- [22] M. S. H. Alwan, J. Meriane, and A. Ghomriou, "Arabic machine translation: A survey of the latest trends and challenges," *Comput. Sci. Res.*, vol. 78, Nov. 2020, ArXiv: 2001018.
- [23] E. H. Aljazeera, A. Al-Naw, and F. K. Elshabaf, "Transforming informal text in Arabic to less informal language: Effect of the age and formal speech detection," in *Proc. 21th Int. Conf. Comput. Intell. Intell. Informatics Syst. (CICIS)*, Cham, Switzerland: Springer, 2020, pp. 136–143.
- [24] M. Qasbi and M. M. Alshabaf, "On machine translation dialogue and evaluation from representing dialogue from English to Arabic," *Ar. Lang. Linguistics Literature*, vol. 29, no. 4, pp. 85–91, Dec. 2023.
- [25] E. H. Aljazeera, "Translating Arabic to less informal language using distribution representation and neural machine translation models," Ph.D. dissertation, School of Natural Eng. Univ. Technical Sydney, Sydney, NSW, Australia, 2019.
- [26] W. Amara, "Statistical machine translation of the Arabic language," Ph.D. dissertation, Comput. Sci. Lab. Univ. Le Mans, Le Mans, France, 2015.
- [27] M. Krawczyk, H. Arjoo, S. Balci, A. Prognal, P. Zembek, and P. Poczta, "Multi-parallel corpus of north lebanese Arabic," in *Proc. ArabicNLP*, 2023, pp. 413–417.
- [28] L. E. Hadia, T. M. Hadia, and M. M. Al-Kadi, "Evaluating Arabic to English machine translation," *Int. J. Adv. Comput. Sci. Appl.*, vol. 5, no. 11, pp. 1–6, 2014.
- [29] S. Benrich and A. Hameiri, "Addressing lexical ambiguity and long sentences in English-Arabic neural machine translation," *Arabian J. Sci. Eng.*, vol. 44, no. 5, pp. 8245–8270, Sep. 2021.
- [30] D. A. Aljazeera, H. M. Al-Rubashah, and F. K. Alshabaf, "Arabic machine translation (ArMT) based on LSTM with attention mechanism architecture," in *Proc. 20th Int. Conf. Lang. Eng. (ICOLET)*, vol. 20, Oct. 2022, pp. 78–82.
- [31] H. Sajjad, A. Alshabaf, H. Durrani, and F. Dalvi, "Android-based speech marking dialectal Arabic-English machine translation," in *Proc. 20th Int. Conf. Comput. Engineering*, 2021, pp. 3094–3107.
- [32] M. Elmadhoun, W. Hameiri, H. Durrani, and E. Moudouf, "Development of a TV broadcast speech recognition system for Qatari Arabic," *IEEE*, vol. 14, pp. 1077–1089, Jan. 2014.
- [33] H. Elmadhoun, A. M. Bilel, and A. Moudouf, "Deep learning approach for translating Arabic into Qatari and Italian languages," in *Proc. Int. Symp. Intell. Informatics Comput. Conf. (ISICCI)*, Mar. 2021, pp. 181–190.
- [34] N. Benabdellah, H. Ayadi, A. Adh, and A. J. E. Faruqi, "Arabic machine translation based on the combination of word embedding techniques," in *Intelligence Systems in Big Data, Analytics, Web and Machine Learning*, Cham, Switzerland: Springer, 2021, pp. 243–260.
- [35] A. Alshabaf, "A corpus for Arabic-English neural machine translation," 2018, arXiv: 1808.06116.
- [36] M. Dabbal, A. Alshabaf, and N. Babbat, "The impact of preprocessing on Arabic-English statistical and neural machine translation," 2019, arXiv: 1908.11721.
- [37] N. Benabdellah, H. Ayadi, A. Adh, and A. J. E. Faruqi, "LSTM vs. GRU for Arabic machine translation," in *Proc. 15th Int. Conf. Comput. Pattern Recognit. Conf. (ICPR)*, Cham, Switzerland: Springer, Jan. 2021, pp. 158–167.
- [38] N. Benabdellah, H. Ayadi, A. Adh, and A. J. E. Faruqi, "Transformer model and convolutional neural networks (CNNs) for Arabic to English machine translation," in *Proc. 1st Conf. Big Data Informatics (Big Data Informatics)*, Cham, Switzerland: Springer, 2021, pp. 399–410.
- [39] N. Benabdellah, H. Ayadi, A. Adh, and A. J. E. Faruqi, "TRANS: A hybrid CNN-RNN attention-based model for Arabic machine translation," in *Proc. Arab. World Sci. Conf. (AWSC)*, Cham, Switzerland: Springer, 2022, pp. 47–60.
- [40] E. Moudouf, H. M. Naguib, A. Elmadhoun, and M. Abdel-Maguid, "Integrating word co-occurrence matrix with attention to English machine translation," 2021, arXiv: 2105.11977.
- [41] J. Tahaoui, "Parallel data, words and sentences in Arabic," *IEEE*, vol. 2012, pp. 2214–2218, May 2012.
- [42] E. M. S. Naguib, A. Elmadhoun, and M. Abdel-Maguid, "ArMT: Transformer model for Arabic language generation," 2021, arXiv: 2106.12040.
- [43] H. Al-Khateeb, K. Al-Khatib, and H. Hameiri, "Term analysis of pretrained language models (PLMs) in English-to-Arabic machine translation," *Hum. Comput. Stud.*, vol. 4, no. 2, pp. 236–248, Feb. 2024.
- [44] H. Alshabaf, H. Hameiri, N. Babbat, and A. Zengin, "Reinforcing dialectal Arabic-Turkish machine translation," in *Proc. Arab. World Sci. Conf.*, 2023, pp. 201–279.
- [45] F. Bilel, Z. Djabri, J. A. Bilel, Y. Guez, D. Guez, J. Bilel, C. Fera, and N. Durrani, "TRMT: A benchmark for low-resource Arabic machine translation," *Trans. Assoc. Comput. Linguistics*, vol. 11, pp. 471–482, Jan. 2023.
- [46] K. Kallam, L. M. Hagey, A. Wiland, M. T. J. Khattak, A. O. B. Bilel, E. M. S. Naguib, and M. Abdel-Maguid, "TABIA: MT Evaluation of word and ChatGPT on machine translation of an Arabic sentence," 2023, arXiv: 2306.10001.
- [47] M. T. J. Khattak, A. Wiland, L. M. Hagey, E. M. S. Naguib, and M. Abdel-Maguid, "OPTArabic: A comprehensive evaluation of ChatGPT on Arabic NLP," 2023, arXiv: 2305.14076.
- [48] E. Moudouf, K. Hagey, J. D. Kallam, and A. Wiland, "Adaptive machine translation with large language models," 2023, arXiv: 2301.11294.
- [49] L. Alshabaf, "Uncovering the new frontier: ChatGPT's powered translation for Arabic-English language pairs," *Theory Pract. Lang. Stud.*, vol. 14, no. 2, pp. 347–357, Feb. 2024.
- [50] C. Liu, J. Xu, and L. Wang, "New trends in machine translation using large language models: Case examples with ChatGPT," 2023, arXiv: 2303.01134.
- [51] X. Zhang, N. Kishita, K. Ueh, and P. Kachit, "Machine translation with large language models: Prompting, few-shot learning, and fine-tuning with GPT-3," in *Proc. 1st Conf. Mach. Learn.*, 2023, pp. 465–474.
- [52] T. Durrani, A. Pappas, A. Hameiri, and L. Zaitouny, "Quickly Efficient Factoring of quantized LLMs," in *Proc. Adv. Neural Inf. Process. Syst.*, vol. 36, 2024, pp. 1–12.
- [53] Z. Aljazeera, M. S. Alshabaf, B. Alshabaf, H. Durrani, E. Aljazeera, and A. Fathi, "Dagum: Evaluating Arabic NLP tasks using ChatGPT models," 2023, arXiv: 2308.14022.
- [54] M. Zeng, M. Hameiri, Durrani, and H. Pappas, "The neural network parallel corpus (NLP)," in *Proc. 16th Int. Conf. Lang. Resources Eval. (LREC)*, 2016, pp. 2570–2574.
- [55] W. Liu, W. Wang, J. T. Huang, X. Wang, J. Shi, and J. Yu, "Do ChatGPT's good translation with GPT-4 as a prompt," 2023, arXiv: 2307.07747.
- [56] A. Hameiri, M. Alshabaf, A. Elmadhoun, V. Hameiri, M. Guez, B. Moudouf, T. J. Khat, M. Aljazeera, and H. Hameiri, "How good are GPT models in machine translation? A comprehensive evaluation," 2023, arXiv: 2302.09116.
- [57] Z. Tan, D. Xu, Z. Li, L. Tang, and W. Su, "CATLLM: Prompting large language models with text style adaptation for Chinese article style transfer," 2024, arXiv: 2401.15207.
- [58] M. Guez, V. Hameiri, Durrani, A. Hameiri, E. Wang, and J. Zhang, "Unsupervised Multilingual-Continuous machine translation," 2023, arXiv: 2307.09977.
- [59] T. Liu, T. Liu, T. Su, T. C. Hameiri, and H. Liu, "Fine-grained text style transfer with diffusion-based language models," 2023, arXiv: 2305.09112.
- [60] E. Moudouf, J. F. Pappas, L. Moudouf, M. Kallam, E. Moudouf, and C. Wille, "On the performance of hybrid neural networks for automatic literature reviews in software engineering," *Inf. Systems Technol.*, vol. 11, Jul. 2020, no. 104294.

- 1111-44

- [107] W. Salameh, H. Elharazi, L. Alwan-Salameh, R. Hamark, and M. Tash, "Sentence-level dialect identification for machine translation systems selection," in *Proc. 23rd Asian Meeting Assoc. Comput. Linguistics*, 2019, pp. 772–778.
- [108] F. Sadat, "Multi-Arabic machine translation (MultiMati)," in *Proc. 13th Arab. Conf. Arab. Assoc. Mach. Transl.*, May 2015, p. 226.
- [109] R. Tammes, N. G. Hammond, and K. Heister, "A Novel Tunisian dialect translator," in *Proc. 13th Int. Conf. Frontiers in Recent Lang. Res.*, *Recent Lang. Process. Appl.*, Harbiner, Tunisia, China, Switzerland, Springer, 2019, pp. 123–134.
- [110] M. Tachout, R. Derafi, A. Derafi, and H. Affou, "Tunisian's portable NLP system applied to MSA and low-resources Algerian dialects," in *Proc. 3rd Conf. Adv. Mach. Learn. Technol. Appl. (AMTA)*, Cham, Switzerland: Springer, 2021, pp. 576–581.
- [111] T. Alshamir and M. A. Bakkajin, "A seq2seq neural network based conversational agent for Gulf Arabic dialect," in *Proc. 21st Int. Arab. Conf. Inf. Technol. (ACIT)*, Nov. 2020, pp. 1–7.
- [112] O. Ghaili, N. Zekout, S. Khallaf, D. Taji, M. Oudati, B. Alharbi, O. Ismae, J. Beyrou, A. Boudiaoui, and N. Habam, "C&M&T tools: An open source Python toolkit for Arabic natural language processing," in *Proc. 15th Egypt. Knowledge Eng. Conf.*, May 2021, pp. 2022–2032.
- [113] A. Gargaji and T. T. C. Leung, "Building the Emirati Arabic Framework," in *Proc. 3rd Translating Machine: Towards Global Multilingual Frontiers*, May 2021, pp. 36–38.
- [114] Y. M. N. Salameh, M. E. M. Hassan, R. Elharazi, and A. Omer, "An evaluation of the accuracy of the machine translation systems of social media language," *Int. J. Adv. Comput. Sci. Appl.*, vol. 12, no. 7, pp. 1–16, 2021.
- [115] J. Chen, Q. Anaghi, and W. Zhao, "An integrated participatory platform for formal evaluation of machine translation," in *Proc. 20th ACM SIGSPRING Conf. Emp. Rigor. Methods*, Jan. 2012, pp. 431–442.
- [116] J. Zakariya, M. Ismae, S. Al-Masboud, and J. M. Alfarouni, "Evaluation of Arabic to English machine translation systems," in *Proc. 11th Int. Conf. Inf. Commun. Sys. (ICICS)*, Apr. 2021, pp. 183–192.
- [117] A. Lavee and M. J. Dinkovski, "The wrong metric for automatic evaluation of machine translation," *Mach. Transl.*, vol. 23, nos. 2–3, pp. 105–115, Sep. 2009.
- [118] K. Pejinovic, S. Rindkov, T. Ward, and W.-J. Zhu, "BLEU: A method for automatic evaluation of machine translation," in *Proc. 4th Asian Meeting Assoc. Comput. Linguistics*, 2002, pp. 111–114.
- [119] K. Taji, H. Nakajima, and K. Cho, "Generalizing dialect translation with sentence rules," in *Proc. 17th Asian Meeting Assoc. Comput. Linguistics*, 2019, pp. 1025–1032.
- [120] F. Xiao, J. Peng, Y. Liu, Y. Wang, H. Chen, and K. Cheng, "Translative learning for computerized text style transfer," 2021, [arXiv:2109.07912](https://arxiv.org/abs/2109.07912).
- [121] M. Humeau, R. J. Durr, R. Schuster, L. Marcilla, and J. Mahboud, "A study of machine-to-text rule with targeted machine translation," in *Proc. 26th Conf. Assoc. Mach. Transl. American Inst. Physics*, Aug. 2016, pp. 221–231.
- [122] M. G. Humeau, N. Mahboud, R. Durr, and R. Schuster, "TED plus: Paragraph, semantic, and alignment enhancements to translation rule sets," *Mach. Transl.*, vol. 25, nos. 2–3, pp. 117–127, Sep. 2010.
- [123] R. Ali, U. Saeed, A. C. Farhadi, and A. Lavee, "COMET: A neural framework for MT evaluation," 2021, [arXiv:2009.00113](https://arxiv.org/abs/2009.00113).
- [124] M. Popović, "ChP: Character n-gram F-score for automatic MT evaluation," in *Proc. 10th Workshop Syst. Mach. Transl.*, 2015, pp. 290–295.
- [125] T. Zhang, V. Kishore, T. Wu, K. O. Williams, and T. Arora, "BERTScore: Evaluating text generation with BERT," 2019, [arXiv:1908.08755](https://arxiv.org/abs/1908.08755).
- [126] L. Chen, Y. Yao, "Rouge: A package for automatic evaluation of text transfer," in *Proc. Workshop Text Summarization Research Gen.*, 2002, pp. 1–20.
- [127] D. Colla, M. Dehante, M. Agreus, R. Veldho, and D. P. Raficovic, "Semantic coherence metrics: The contribution of syntactic metrics," *Artif. Intell. Mater.*, vol. 134, Dec. 2022, Art. no. 102195.
- [128] C. Kopylov, K. Nakaguchi, H. Poon, and J. Tenenbaum, "Global rules for grammatical error correction metrics," in *Proc. 13th Asian Meeting Assoc. Comput. Linguistics: 1st Int. Joint Conf. Natural Lang. Process.*, 2015, pp. 388–393.
- [129] F. Cheng, Z. Chen, Y. Zhang, Q. Huo, D. Li, and J. Liu, "Computerized text style transfer," 2020, [arXiv:2005.09102](https://arxiv.org/abs/2005.09102).
- [130] D. Chen and W. B. Dolan, "Collecting highly-parallel sentence paraphrase evaluation," in *Proc. 4th Asian Meeting Assoc. Comput. Linguistics: 1st Int. Joint Conf. Nat. Lang. Process.*, Jan. 2012, pp. 190–200.
- [131] Z. Fu, X. Tai, N. Peng, D. Zhao, and K. Yan, "Style transfer in text: Explanation and evaluation," in *Proc. 44th Conf. Artif. Intell.*, Apr. 2018, vol. 32, no. 1, pp. 1–16.
- [132] N. Berman and J. Goleysch, "Formal-BERT: Sentence embeddings using BERT embeddings," 2020, [arXiv:2004.00044](https://arxiv.org/abs/2004.00044).
- [133] V. Arora, L. Wang, H. Baheti, and O. Vakharenko, "Disentangled representation learning for unpaired text style transfer," 2018, [arXiv:1808.04194](https://arxiv.org/abs/1808.04194).
- [134] W. Yang, Q. Su, and P. Liu, "BERTScore: Evaluating generated text in this generation," in *Proc. 4th Asian Meeting Assoc. Comput. Linguistics*, Jan. 2019, pp. 2720–2727.
- [135] T. Saito, D. Dey, and A. P. Faria, "BLEURT: Learning robust metrics for text generation," 2020, [arXiv:2004.04698](https://arxiv.org/abs/2004.04698).
- [136] A. Wortsch, A. Hupf, and S. B. Hübner, "Neural network embedding algorithms," *Trans. Assoc. Comput. Linguistics*, vol. 7, pp. 625–641, Nov. 2019.
- [137] R. Ma, R. Fokko, N. Ouedraoui, and J. Rabinov, "Evaluating style transfer for text," 2019, [arXiv:1904.01290](https://arxiv.org/abs/1904.01290).
- [138] J. Whiting, T. Song, Kiriljanski, K. Gimpel, and G. Neubig, "Beyond BLEU: Training neural machine translation with semantic similarity," 2019, [arXiv:1907.04884](https://arxiv.org/abs/1907.04884).
- [139] C. Agam, C. Banea, D. Cuk, M. Dahl, A. Gattardo, Agim, R. Mihalcea, C. Riga, and J. Wortsch, "SemEval-2019 task 1: Automatic textual cohesivity, paraphrasing and cross-lingual cohesivity," in *Proc. 10th Int. Workshop Semantic Eval. (SemEval)*, San Diego, CA, USA, 2019, pp. 497–511.
- [140] D. Hu, Z. Hu, Z. Hu, O. Vakharenko, and R. Mihalcea, "Disentangling text style transfer: A survey," *Comput. Linguistics*, vol. 48, no. 1, pp. 115–201, Apr. 2022.
- [141] W. Salameh, *Machine Translation of Arabic Dialects*. New York, NY, USA: Columbia Univ. Press, 2018.
- [142] M. Tachout and M. Salameh, "Machine translation for semi and low-resources dialects using a low resource version of the dialectal parallel corpus (PALRE) v2.0b," in *Proc. 4th Int. Conf. Natural Lang. Speech Process. (ICNLSP)*, 2021, pp. 13–18.
- [143] A. Panchai, K. K. Yogi, and R. Sharma, "A comparison of machine translation methods for natural language processing and their challenges," in *Proc. 3rd Conf. Assoc. Comput. Linguistics: 1st Int. Joint Conf. Natural Lang. Process.*, 2021, pp. 475–487.
- [144] E. H. Alotaibi and A. Al-Ali, "Translating dialectal Arabic to low resource language using word embedding," in *Proc. 3rd Conf. Assoc. Arab. Natural Lang. Process. (ArNLP)*, Nov. 2017, pp. 55–57.
- [145] F. Sadat, F. Mallik, M. Boudiaoui, R. Salameh, and A. Farman, "Collaboratively constructed linguistic resources for language variants and text applications in NLP applications: The case of Tunisian Arabic and its embeddings," in *Proc. Workshop General Computational Resources Lang. Process.*, 2019, pp. 102–110.
- [146] S. Hamada and R. M. Hamark, "Developing a transfer-based system for Arabic dialect translation," in *Intelligent Natural Language Processing: Theory and Applications*, 2018, pp. 121–134.
- [147] A. Hamk, R. Boujelmine, N. Hamk, and A. Naei, "The effect of factoring root and gender mapping in bidirectional Tunisian-standard Arabic machine translation," in *Proc. Mach. Transl. Summit*, 2015, pp. 1–20.
- [148] M. A. Elghamry and M. Zogbi, "Rule-based machine translation from Tunisian dialect to modern standard Arabic," *Proc. Comput. Sci.*, vol. 19, pp. 110–119, Mar. 2020.
- [149] W. Salameh and N. Hamark, "Elbow: A dialect to standard Arabic machine translation system," in *Proc. COLING: International System*, Dec. 2012, pp. 388–392.
- [150] C. Al-Ghathani and M. Al-Tamimi, "Anaphoric coreference in Arabic to modern standard Arabic," *Int. J. Inf. Inf. Manage.*, vol. 8, no. 1, pp. 39–49, 2010.
- [151] E. Alotaibi and F. Alotaibi, "Rule and dialect to modern standard Arabic: Rule-based dialect machine translation," in *Proc. 13th Conf. Artif. Intell. (ICAI)*, 2011, pp. 1–4.
- [152] S. Dabb, O. Amara, N. E. Mouta, R. Roudhwan, and Y. Ghaili, "Machine Arabic-to-Arabic conversion using rule-based transliteration and weighted Levenshtein algorithm," *Sci. Ab.*, vol. 23, Mar. 2024, Art. no. e00075.

- [134] S. Schler, W. Fraz, H. Bannayan, A. Lyye, N. Wilbur, and K. Oflazer, "Diversity and robust adaptivity for machine translation into Egyptian Arabic," in *Proc. EMNLP Workshop Arabic Natural Lang. Process. (ANLP)*, 2014, pp. 196–204.
- [135] H. Tajjal, K. Darwish, and Y. Rifkin, "Translating dialectal Arabic to English," in *Proc. 14th Arab. Meeting Assoc. Comput. Linguistics, Jan. 2014*, pp. 1–4.
- [136] P. Walker and D. T. Ng, "Representing statistical machine translation for a source-target language using related resources such languages," *J. Arab. Acad. Sci.*, vol. 44, pp. 176–222, May 2012.
- [137] A. Khassem, "Automatic classification of Arabic dialects for machine translation," 2023, [arXiv:2301.00447](https://doi.org/10.46471).
- [138] R. Tachikawa and K. Hosonuma, "A hybrid approach to translate Moroccan Arabic dialect," in *Proc. 4th Int. Conf. Arab. World. Comput. Appl. (ICAWA)*, May 2014, pp. 1–4.
- [139] H. Tiedt, "A hybrid approach to statistical machine translation between standard and dialectal varieties," in *Proc. 4th Int. Conf. Arab. World. Comput. Appl. (ICAWA)*, May 2014, pp. 1–4.
- [140] L. Chadi, S. Azzam, and M. Alshai, "Dialect is statistical translation of Algerian Arabic dialect written with Arabic and Arabic letter," in *Proc. 1st Pacific Asia Conf. Lang. Inf. Comput. Process.*, vol. 9, 2017, pp. 1–17.
- [141] W. Farhan, R. Tachikawa, A. Abummar, R. Jahan, M. Al-Ayyash, A. B. Boudil, and A. Tama, "Unsupervised dialectal neural machine translation," *Inf. Process. Manage.*, vol. 57, no. 3, May 2020, Art. no. 102144.
- [142] R. Al-Jarrah and R. M. Oussay, "Neural machine translation from Levantine dialect to modern standard Arabic," in *Proc. 14th Int. Conf. Inf. Commun. Sem. (ICICS)*, Apr. 2020, pp. 173–178.
- [143] L. H. Baniata, S. Park, and Y.-B. Park, "A neural machine translation model for Arabic dialects that utilizes multilingual learning (MTL)," *Comput. Intell. Neurosci.*, vol. 2018, pp. 1–10, Dec. 2018.
- [144] L. H. Baniata, S. Park, and Y.-B. Park, "A multilingual-based neural machine translation model for pan-Arabic speech recognition for Arabic dialects," *Appl. Sci.*, vol. 9, no. 12, p. 5500, Dec. 2019.
- [145] L. H. Baniata, S. Kaya, and Y. B. Arayonah, "A neural positional encoding multi-head attention-based neural machine translation model for Arabic dialects," *Mathematics*, vol. 10, no. 19, p. 3068, Oct. 2022.
- [146] L. H. Baniata, L. K. E. Arayonah, and S. Park, "A transformer-based neural machine translation model for Arabic dialects that utilizes network auto," *Systems*, vol. 21, no. 19, p. 4309, Sep. 2021.
- [147] A. Ene, S. Achary, and R. Bhatnagar, "Neural machine translation of low resource languages: Application to transcriptions of Yiddish dialect," in *Proc. Int. Conf. Arab. World. Comput. Appl. Research, Cham, Switzerland: Springer, Jan. 2022*, pp. 254–267.
- [148] S. Kharin, R. Bhatnagar, and L. Bhatnagar, "Hybrid pipeline for building Arabic Levantine dialect-based Arabic neural machine translation model from scratch," *ACM Trans. Asian Low-Resource Lang. Inf. Process.*, vol. 22, no. 3, pp. 1–21, Mar. 2023.
- [149] M. A. Falestin, K. T. Wasef, H. Barakat, and S. M. Alshai, "Improving neural machine translation for low resource languages through semi-parallel corpora: A case study of Egyptian dialect to modern standard Arabic translation," *Sci. Rep.*, vol. 14, no. 1, p. 1285, Jan. 2024.
- [150] A. Shua, A. Makhoul, Y. Faghihi, and K. Sahli, "Algebraic dialect translation applied on COVID-19 social-media comments," in *AI4Health Intelligence and Generative Towards an Energy Transition, Cham, Switzerland: Springer, 2023*, pp. 714–726.
- [151] A. Shua, A. Makhoul, Y. Faghihi, and K. Sahli, "Improving neural machine translation for low resource Algerian dialect by distributed feature learning strategy," *Arabian J. Sci. Eng.*, vol. 47, no. 8, pp. 10411–10418, Aug. 2023.
- [152] E. Al-Khameis and A. Alalshai, "Machine translation of Quranic Arabic dialect from social media," in *Proc. ArabicNLP*, 2023, pp. 302–308.
- [153] T. Moudathil, R. Shini, M. Chaffee, and K. Ibrahim, "Improving machine translation of Arabic dialects through multi-task learning," in *Proc. Int. Conf. Arab. World. Comput. Appl. Resell, Cham, Switzerland: Springer, 2021*, pp. 580–590.
- [154] W. Drouach, S. Barakat, and R. Boudiabane, "ANLP@NLP 2023 shared task: Machine translation of Arabic dialects: A comparative study of transformer models," in *Proc. ArabicNLP*, 2023, pp. 440–449.
- [155] S. Elmad, I. Abdelkader, R. Abdelkader, A. Islem, and R. Hamed Nasser, "TolMae at NLP 2023 shared task: A comparison of various TS-based models for translating Arabic dialectal text to modern standard Arabic," in *Proc. ArabicNLP*, 2023, pp. 454–464.
- [156] K. Abdelkader, "Reconstruction of NLP 2023 shared task: Parameter efficient tuning for dialect identification and dialect machine translation," in *Proc. ArabicNLP*, 2023, pp. 452–457.
- [157] M. Fara, "ArTS@MMA: Translating dialectal Arabic to MSA," in *Proc. 4th Workshop Open-Source Arabic Corpus Process. Tools (OSACT) Shared Task Arabic LLMs Performance Dialect MSA Arab. Transl. (LREC-COLING)*, 2024, pp. 124–130.
- [158] S. S. Alshai, "Sara, participant in OSACTs 2024 shared task: Fine-tuning ArTS models for translating Arabic dialectal text to modern standard Arabic," in *Proc. 4th Workshop Open-Source Arabic Corpus Process. Tools (OSACT) Shared Task Arabic LLMs Performance Dialect MSA Arab. Transl. (LREC-COLING)*, 2024, pp. 117–123.
- [159] H. Amari, N. Khatib, I. Mohammed, A. Waleed, and H. Ali, "OSACT 2024 task 2: Arabic dialect to MSA translation," in *Proc. 4th Workshop Open-Source Arabic Corpus Process. Tools (OSACT) Shared Task Arabic LLMs Performance Dialect MSA Arab. Transl. (LREC-COLING)*, 2024, pp. 94–101.
- [160] O. Hecar, A. Shari, S. Elmad, I. Alshai, L. Ghali, and A. Khatib, "ArTS at OSACTs shared task: Investigation of data augmentation in Arabic dialect-MSA translation," in *Proc. 4th Workshop Open-Source Arabic Corpus Process. Tools (OSACT) Shared Task Arabic LLMs Performance Dialect MSA Arab. Transl. (LREC-COLING)*, 2024, pp. 104–111.
- [161] J. Zhao, Y. Zhang, L. Guo, Q. Zhang, T. Guo, and X. Huang, "LLAMA beyond English: An empirical study on language capability matrix," 2024, [arXiv:2401.00000](https://arxiv.org/abs/2401.00000).
- [162] A. Ghaffar, X. Wu, Z. Bagg, A. Khatib, K. Dini, M. Hameedkhan, C. Xue, Y. Wang, X. Ding, Z. Wang, S. Hua, X. Jiang, Q. Liu, and F. Lingling, "Evaluating pre-trained language models and their evaluation for Arabic natural language processing," in *Proc. Conf. Empirical Methods Natural Lang. Process.*, 2022, pp. 1175–1191.
- [163] H. Li, A. Saei, and M. Yousefi, "Multidimensional evaluation for text style transfer using ChatGPT," 2023, [arXiv:2304.14962](https://arxiv.org/abs/2304.14962).
- [164] B. Pajouh, M. Mirzaei, K. Ramezani, A. Müller, and M. Hameedkhan, "Learning translation model to transfer written dialects to modern written language," in *Proc. 4th ArabicNLP/ACT Workshop, Cham, Switzerland: Springer, May 2023*, pp. 1083–1088.
- [165] S. Haiman, F. Bealman, X. Li, J. Jiang, S. Wilfong, and Y. Chen, "STEE: Unified style transfer with expert reinforcement," 2023, [arXiv:2311.12767](https://arxiv.org/abs/2311.12767).
- [166] S. Rajendran, T. Tomar, "Dial in a matter, may I introduce the GPT-4 dataset, Corpus, applications and metrics for formality style transfer," 2019, [arXiv:1907.00575](https://arxiv.org/abs/1907.00575).
- [167] X. Wu, S. Bao, and M. Geyrho, "Multi-task neural models for translating between style written and spoken languages," 2019, [arXiv:1907.00575](https://arxiv.org/abs/1907.00575).
- [168] B. Xu, T. Gu, and F. Wei, "Formality style transfer with hybrid neural networks," 2019, [arXiv:1907.00575](https://arxiv.org/abs/1907.00575).
- [169] Z. Gu, D. Gu, J. Ma, S. N. Madhavan, and E. Suresh, "DART: Deep-sequence text style transfer for sentence matching and translation," 2019, [arXiv:1907.11100](https://arxiv.org/abs/1907.11100).
- [170] D. Li, Y. Zhang, Z. Guo, Y. Cheng, C. Buckler, M. T. Hui, and B. Dolan, "Dialectal adaptation and style transfer," 2019, [arXiv:1907.00575](https://arxiv.org/abs/1907.00575).
- [171] H. Guo, S. Wu, L. Wu, J. Kong, and W.-M. Huo, "End-to-end learning based text style transfer without parallel training corpus," 2019, [arXiv:1907.10872](https://arxiv.org/abs/1907.10872).
- [172] F. Kojima, "Negative sentiment classification using the paraphrase generator," in *Proc. 17th Arab. Meeting Assoc. Comput. Linguistics*, 2019, pp. 6007–6012.
- [173] X. Wu, "Formality style transfer within and across languages with limited supervision," Ph.D. dissertation, Dept. Computer Science, Univ. Maryland College Park, MD, USA, 2019.
- [174] X. Wu and M. Geyrho, "Converting neural machine translation formality with syntactic supervision," in *Proc. Arab. Conf. Arab. World. Appl.*, 2020, vol. 14, no. 3, pp. 8564–8573.
- [175] J. He, X. Wang, G. Nang, and T. Beng-Rikgarnik, "A probabilistic translation of unpaired word-text style transfer," 2020, [arXiv:2002.01922](https://arxiv.org/abs/2002.01922).
- [176] Y. Wang, X. Wu, L. Shu, J. Li, and W. Chen, "Formality style transfer with shared latent space," in *Proc. 14th Int. Conf. Comput. Linguistics*, 2020, pp. 2234–2240.

- [196] K. Chaochi and D. Yang, "Semi-supervised formality style transfer using language model maximization and textual information maximization," 2021, [arXiv:2010.04800](#).
- [197] A. Sanchez, K. Krishna, S. Sureshwar, and A. Venugopal, "Reinforced rewards framework for text style transfer," in *Proc. 42nd Int. Conf. Artif. Intell. Syst. (ICAAI)*, Lathem, Portugal, China, Switzerland: Springer, Jan. 2021, pp. 561–566.
- [198] Y. Zhang, T. Gu, and X. Liu, "Parallel data augmentation for formality style transfer," 2021, [arXiv:2007.07523](#).
- [199] H. Gupta, R. V. Sureshwar, A. Venugopal, and A. Sanchez, "Multi-style transfer with discriminator feedback on dataset corpora," 2021, [arXiv:2010.11778](#).
- [200] K. Krishna, I. Witting, and M. Tytlic, "Reinforcement unsupervised style transfer as generative generation," 2021, [arXiv:2009.05790](#).
- [201] H. Lai, A. Tsoi, and M. Nisum, "Generative models for data augmented style transfer tasks without task-specific parallel training data," 2021, [arXiv:2009.04341](#).
- [202] Z. Hu, H. Xu, M. Liu, and C. C. Aggarwal, "Syntax matters: Syntactically controlled text style transfer," 2021, [arXiv:2108.05869](#).
- [203] R. Kulkarni, D. Nathani, S. Chavva, R. Sankaran, and P. Tallekla, "Per-shot controllable style transfer for low-resource pathological settings," 2021, [arXiv:2101.07905](#).
- [204] S. Kumar, E. Malmi, A. Sorensen, and Y. Dwivedi, "Controlled text generation as continuous optimization with multiple constraints," in *Proc. Adv. Neural Inf. Process. Syst.*, vol. 34, Jan. 2021, pp. 14342–14354.
- [205] Y. Ma and Q. Li, "Exploiting non-orthogonality for text style transfer," in *Proc. Conf. Empirical Methods Natural Lang. Process.*, 2021, pp. 4267–4274.
- [206] S. Kulkarni, T. Kulkarni, T. Arora, and M. Omernik, "Rife: Reranked-based continuous learning for generative generation," in *Proc. 39th Annu. Meeting Assoc. Comput. Linguistics (The 1st Joint Conf. Natural Lang. Process., Seattle, WA, Workshop)*, 2021, pp. 226–234.
- [207] Y. Ma, Y. Chen, X. Ma, and Q. Li, "Collaborative learning of bidirectional decoders for unsupervised text style transfer," in *Proc. Conf. Empirical Methods Natural Lang. Process.*, 2021, pp. 4269–4276.
- [208] W. Qian, J. Yang, and S. Hu, "Dissecting the style information in text-to-text transferred hidden layers," in *Proc. Int. Joint Conf. Neural Netw. (IJCNN)*, Jul. 2022, pp. 1–8.
- [209] H. Lai, A. Tsoi, and M. Nisum, "Multilingual pre-training with language and task adaptation for multilingual text style transfer," 2022, [arXiv:2201.08732](#).
- [210] A. Liu, A. Wang, and H. Okazaki, "Semi-supervised formality style transfer with consistency training," 2022, [arXiv:2207.17620](#).
- [211] R. Liu, C. Cao, C. An, D. Xu, and J. Songshi, "New parallel text style transfer with self-supervised supervision," 2022, [arXiv:2204.08112](#).
- [212] Z. Mithrasathani, K. Goyal, and T. Raghavakrishna, "Mts and much: Learning-free controllable text generation using memory language models," 2022, [arXiv:2201.12189](#).
- [213] Y. Xiao, "Tagging without rewriting: A probabilistic model for supervised sentiment and style transfer," in *Proc. 12th Workshop Comput. Approaches to Sentiment (Sentiment: Sentiment Analysis)*, 2022, pp. 281–302.
- [214] S. Yang, "Mask and augmentor: A classifier-based approach for unsupervised sentiment transformation of reviews for electronic commerce websites," in *Proc. 10th Int. Workshop Natural Lang. Process. Social Media*, 2022, pp. 1–10.
- [215] K. G. Rao, H. Sankethi, M. Venkateshwar, and A. V. V. Kumar, "Exposure of natural language processing using transformer learning," *World Scientific*, vol. 25, no. 17, p. 1981, 2022.
- [216] X. Shen, W. Chen, and X. Xu, "Transfer-based sentiment and style transfer from reviews texts," in *Proc. 30th Int. Conf. Database Inf. Technol. Comput. Eng.*, Oct. 2022, pp. 206–213.
- [217] A. Sankaran and A. Rao, "Masked transformer through knowledge distillation for unsupervised text style transfer," *Recent Lang. Eng.*, vol. 35, no. 5, pp. 973–1001, Sep. 2023.
- [218] G. Liu, Y. Ding, B. Li, M. Shi, and M. Huhua, "Prompt-based editing for text style transfer," 2022, [arXiv:2201.11907](#).
- [219] L. Wu, P. Liu, Y. Yang, S. Liu, and Y. Zhang, "Unsupervised style transfer and content reviewer networks for brand style transfer," *Inf. Process. Manage.*, vol. 60, no. 1, May 2023, Art. no. 108265.
- [220] D. Randerpathy and A. Elabb, "Unsupervised text style transfer through differentiable back translation and rewards," in *Proc. Pacific-Asia Conf. Knowl. Discovery Data Mining (PAKDD)*, Switzerland: Springer, Jan. 2023, pp. 210–223.
- [221] Z. Herrera, A. Patel, C. Callison-Burch, Y. Zhou, and K. McKeown, "Paraladder: Gradient-informed paraphrases for play-and-rewind text style transfer," in *Proc. AAAI Conf. Artif. Intell.*, Mar. 2024, vol. 38, no. 16, pp. 18236–18248.
- [222] C. Zhang, H. Cai, Y. Zhang, L. Y. Wu, L. Ren, and M. Abdel-Mottaleb, "Bridging text style transfer with self-supervised from CLMs," 2024, [arXiv:2401.01706](#).
- [223] Y. Cao, Q. Liu, Y. Yang, and K. Wang, "Latent representation discrimination for unsupervised text style generation," *Inf. Process. Manage.*, vol. 61, no. 1, May 2024, Art. no. 109641.
- [224] H. A. Khalil, "Arabic Texts Pre-processor of Dialects with embeddings," in *Proc. 2nd Int. Conf. Culture Lang. Sciences Asia (ICCLAS)*, Amsterdam, The Netherlands: Atlantis Series, 2019, pp. 140–143.
- [225] K. Shaban, S. Siddiqui, M. Alkhateeb, and A. A. Hameed, "Challenges in Arabic natural language processing," in *Computational Linguistics, Speech And Image Processing For Arabic Language*, Singapore: World Scientific, 2019, pp. 29–43.
- [226] G. Nader, S. Garghava, S. Khatib, and B. Ghalib, "Challenges and progress in constructing Arabic dialect corpora and linguistic tools: A focus on Moroccan and Tunisian dialects," in *Proc. 36 IEEE Comp. Inf. Sci. Technol. (CIST)*, Dec. 2023, pp. 293–298.
- [227] M. Omeri and S. Benkhana, "Challenges in building corpora for Algerian Arabic (dialect/MC-arabian)," *J. Soc. Res. Sci.*, vol. 22, no. 4, pp. 594–617, 2021.
- [228] W. Tahaoui and N. Habbas, "Unsupervised Arabic dialect representation for machine translation," *Natural Lang. Eng.*, vol. 28, no. 1, pp. 123–148, Mar. 2022.
- [229] H. Yu, W. Xu, S. Li, and Q. Li, "Machine translation evaluation metrics based on dependency parsing model," *ACM Trans. Asian Low-Resource Lang. Inf. Process.*, vol. 18, no. 4, pp. 1–15, Dec. 2019.
- [230] M. Papayri, "Class-over rules for evaluation of machine translation system," in *Proc. 26 Workshop*, Jan. March Turin, Jan. 2012, pp. 71–75.
- [231] K. Shaban, S. Agreewat, S. Ummat, and M. Caezan, "Evaluating the evaluation metrics for style transfer: A case study in multilingual friendly transfer," 2021, [arXiv:2203.09683](#).
- [232] H. Lai, S. Mao, A. Tsoi, and M. Nisum, "Human judgement is a compass to evaluate automatic metrics for formality transfer," 2022, [arXiv:2201.02546](#).
- [233] S. Gehrmann, S. Clark, and T. Schuster, "Bridging the cracked foundation: A survey of obstacles in evaluation practices for generated text," 2022, [arXiv:2202.06603](#).
- [234] M. Li and M. Wang, "Optimizing automatic evaluation of machine translation with the Lstm2L approach," *ACM Trans. Asian Low-Resource Lang. Inf. Process.*, vol. 18, no. 3, pp. 1–18, Nov. 2019.
- [235] N. Rich, F. Wang, and X. M. Quaresima, "Glossifier: A method for measuring machine translation confidence," in *Proc. 48th Annu. Meeting Assoc. Comput. Linguistics (The 1st Joint Conf. Natural Lang. Process., Seattle, WA, Workshop)*, 2021, pp. 224–234.
- [236] P. Li, C. Chen, W. Zhang, Y. Ding, C. Ye, and Z. Zhang, "ETL: An automatic evaluation metric for machine translation based on word embeddings," *IEEE/ACM Trans. Audio Speech Lang. Process.*, vol. 27, no. 70, pp. 1467–1504, Oct. 2019.
- [237] N. Babiker, D. Gola, Y. Lajouh, S. Kulkarni, and A. Prashanth, "Bridging the gap of mixed content for content generation in text style transfer," in *Proc. Int. Conf. Appl. Natural Lang. Inf. Syst. Comput. Intelligence (Springer)*, Jan. 2023, pp. 437–449.
- [238] N. Babiker, D. Gola, I. Ghom, S. Kulkarni, and A. Prashanth, "Don't lose the message while paraphrasing: A study on content preserving style transfer," in *Proc. Int. Conf. Appl. Natural Lang. Inf. Syst. Comput. Intelligence*, Springer, 2023, pp. 47–61.
- [239] N. Babiker, H. Ayad, A. Adib, and A. I. E. Fawzi, "A comparative study of different dimensionality reduction techniques for Arabic machine translation," *ACM Trans. Asian Low-Resource Lang. Inf. Process.*, vol. 22, no. 12, pp. 1–17, Dec. 2023.
- [240] D. S. Koudan, S. K. Nagmani, and B. Singh, "Impact of word embedding models on text analysis in deep learning environment: A review," *Appl. Artif. Intell.*, vol. 38, no. 5, pp. 10343–10355, Sep. 2024.

- [244] S. Idrissouppara, G. Lécuyer, D. Lohre, and J.-D. Kellens, "Local or global: The variation in the encoding of style across sentences and formality," in *Proc. Int. Conf. Artif. Neural Netw., Cham, Switzerland: Springer, Jan. 2023*, pp. 492–504.
- [245] S. Hameed, K. McIlroy, M. Agha, S. Janssen, M. Seed, and K. Sengul, "Convolutional neural processing," in *Proc. 16th Int. Conf. Comput. Linguistics Intell. Text Process. (COLING)*, Cairo, Egypt, Cham, Switzerland: Springer, Jan. 2013, pp. 629–633.
- [246] S. Gerner, K. Petersen, and X. Li, "Automatic word alignment tools to scale production of manually aligned parallel texts," in *Proc. LREC*, May 2012, pp. 1194–1198.
- [247] M. Fink, A. Fink, and S. Sauer, "Word representations for dialect modulation," in *Proc. Int. Conf. Intell. Text Process. Comput. Linguistics*, Cham, Switzerland: Springer, Jan. 2011, pp. 55–67.
- [248] K. Elmadani, S. Elabbadi, and K. Elshaban, "Evaluating Arabic grammar and sentiment improvement," in *Proc. Int. Conf. Intell. Text Process. Comput. Linguistics*, Cham, Switzerland: Springer, 2016, pp. 417–428.
- [249] A. Alkhatib, W. Alkhatib, and F. Alkhatib, "Simulations between Arabic dialects: Investigating geographical proximity," in *Inf. Process. Manage.*, vol. 58, no. 3, Jan. 2022, Art. no. 110730.
- [250] F. Marouf and I. Zoubeir, "Transformation normalization for information extraction and machine translation," *J. King Saud Univ. Comput. Inf. Sci.*, vol. 26, no. 4, pp. 179–187, Dec. 2014.
- [251] K. Mchalek, H. Bouchehad, and K. Smati, "A study of a non-lexical language: An Algerian dialect," in *Proc. Spoken Lang. Technol. Under. Resource Lang.*, Jan. 2012, pp. 325–332.
- [252] E. Ghannouchi, N. Kapp, and P. Zweigshausen, "Unification of parallel sentences," in *Building and Using Computable Corpora for Multilingual Natural Language Processing*, Cham, Switzerland: Springer, 2023, pp. 35–49.
- [253] S. Idrissouppara, G. Lécuyer, D. Lohre, and J.-D. Kellens, "Style in sentences: corpus style or formality? The answer is different!" in *Proc. 16th Int. Conf. Artif. Neural Netw., Mach. Learn. (ICANN)*, Bratislava, Slovakia, Cham, Switzerland: Springer, Jan. 2023, pp. 487–499.
- [254] H. Alkhatib and K. Elshaban, "Toward building a bilingual dictionary for Lithuanian-Arabic machine-translated Arabic machine translation," in *Proc. LREC*, 2022, pp. 75–81.
- [255] R. Fakhri, A. Alkhatib, and M. Al-Ayyoubi, "Data Adaptation-Based LSTM for Arabic translation," *Int. J. Electr. Comput. Eng.*, vol. 11, no. 3, p. 2022, Jan. 2022.
- [256] M. Wu, Y.-T. Ye, L. Qi, G. Feng, and Q. Hatten, "Adapting large language models for machine-level machine translation," *2024, arXiv:2401.06060*.



**SHADI I. ABUDALIFA** received the B.Sc. and M.Sc. degrees in computer engineering from the Islamic University of Gaza (IUG), Palestine, in 2007 and 2010, respectively, and the Ph.D. degree in computer science and engineering (CSE) from the King Fahd University of Petroleum and Minerals (KFUPM), in 2018. He has a strong teaching experience through his work as an Assistant Professor with the University College of Applied Sciences, Palestine. He is currently a Researcher with the SDAA-KFUPM Joint Research Center for Artificial Intelligence (JRC-AI). From July 2023 to August 2024, he was a Research Assistant with Projects and Research Laboratory, IIS. During the same period, he was a Teaching Assistant with the Faculty of Engineering, IIS. He serves as a reviewer through various international journals and conferences. His current research interests include data mining, activity recognition, Arabic natural language processing, and sentiment analysis.



**FAHAD J. ABDU** was born in Dagebora, Kano, Nigeria. He received the B.Eng. degree in electrical engineering from Bayero University, Kano (BUK), Kano, in 2011; the M.Sc. degree in electrical and electronic engineering from Modares University, Kermanshah, Turkey, in 2014, and the Ph.D. degree in communication and information systems from Xiamen University (XNU), Xiamen, China, in 2022. He is currently a Postdoctoral Researcher with the SDAA-KFUPM Joint Research Center for Artificial Intelligence, King Fahd University of Petroleum and Minerals Research (KFUPM). He is also a Lecturer with the Department of Electronics and Telecommunications Engineering, Ahmadu Bello University (ABU), Zaria, Nigeria. His current research interests include radar signal processing, multibeam radar fusion, computer vision, artificial intelligence (AI), and Arabic natural language processing (NLP).



**MAAD M. ALOWFI** received the B.Sc. and M.Sc. degrees in electrical engineering from King Fahd University of Petroleum and Minerals (KFUPM), Dhahran, Saudi Arabia, in 2013 and 2015, respectively, and the M.Sc. degree in operations research and the Ph.D. degree in electrical and computer engineering from Georgia Institute of Technology, Atlanta, USA, in 2020 and 2021, respectively. He is currently an Assistant Professor with the Electrical Engineering Department, KFUPM. His research interests include optimization, artificial intelligence, and their applications in the energy sector.

44/45