

Multilingual Text-to-Image Person Retrieval via Bidirectional Relation Reasoning and Aligning

Min Cao, Xinyu Zhou, Ding Jiang, Bo Du, *Senior Member, IEEE*, Mang Ye, *Senior Member, IEEE*,
Min Zhang

Abstract—Text-to-image person retrieval (TIPR) aims to identify the target person using textual descriptions, facing challenge in modality heterogeneity. Prior works have attempted to address it by developing cross-modal global or local alignment strategies. However, global methods typically overlook fine-grained cross-modal differences, whereas local methods require prior information to explore explicit part alignments. Additionally, current methods are English-centric, restricting their application in multilingual contexts. To alleviate these issues, we pioneer a multilingual TIPR task by developing a multilingual TIPR benchmark, for which we leverage large language models for initial translations and refine them by integrating domain-specific knowledge. Correspondingly, we propose Bi-IRRA: a Bidirectional Implicit Relation Reasoning and Aligning framework to learn alignment across languages and modalities. Within Bi-IRRA, a bidirectional implicit relation reasoning module enables bidirectional prediction of masked image and text, implicitly enhancing the modeling of local relations across languages and modalities, a multi-dimensional global alignment module is integrated to bridge the modality heterogeneity. The proposed method achieves new state-of-the-art results on all multilingual TIPR datasets. Data and code are presented in <https://github.com/Flame-Chasers/Bi-IRRA>.

Index Terms—Text-to-Image Person Retrieval, multilingual image-text learning, person re-identification.

1 INTRODUCTION

GIVEN a text query, Text-to-Image Person Retrieval (TIPR) [1] aims to identify the most relevant person images from an extensive gallery of such images. The task is similar to the person re-identification task (Re-ID) [2], [3], [4], which involves identifying person images across cameras based on the image query. In contrast to the structured image query in Re-ID, the text query in TIPR takes the form of free, flexible characters, making it more accessible and offering substantial application potential in public safety domains. In recent years, TIPR has garnered increasing attention [5], [6], [7], [8], [9].

A key challenge in TIPR is the inherent modality gap between vision and language, driving research toward robust cross-modal alignment. These efforts can be broadly categorized into global-matching [10], [11], [12], [13], [14], [15], [16] and local-matching methods [17], [18], [19], [20], [21]. The former aligns global text-image representations at the coarse-grained level via cross-modal matching loss functions (Fig. 1(a)), while the latter establishes fine-grained associations between textual entities and image body parts (Fig. 1(b)).

Despite notable progress in this task, two critical issues remain to be addressed. The first issue is the limited focus on language environments. Current methods [5], [6], [7],

[8], [22] center around addressing TIPR with English as the default text query. However, in practical application, there is a growing need for supporting multiple languages, (*e.g.*, Chinese and French) as queries for TIPR. Biased towards a single language (*i.e.*, English), current methods struggle to achieve accurate retrieval when confronted with diverse language requirements in real-world scenarios. Another issue involves the constraint of existing cross-modal alignment strategies. Global-matching methods align the global representations of texts and images, often overlooking the need to bridge modality heterogeneity at a more detailed, fine-grained level, potentially impacting performance. On the other hand, local-matching methods are tailored to build the explicit correspondence between body parts and textual descriptions with the aid of external technologies and predefined rules. As a result, they confine the exploration of cross-modal alignment within the boundaries of these set rules, and also pose resource-intensive demands as it necessitates the extraction and storage of multiple local part representations of images and texts during inference.

In this paper, we pioneer a multilingual TIPR task. Specifically, we build the multilingual TIPR benchmark and propose Bi-IRRA: a cross-modal Bidirectional Implicit Relation Reasoning and Aligning framework to learn alignment across languages and modalities at both coarse-grained and fine-grained levels. This work is centered on addressing both data and framework aspects.

Data. For a novel multilingual TIPR task, a critical obstacle is the scarcity of multilingual TIPR data to support its research. While manual annotation of multilingual TIPR data presents a direct solution, it is labor-intensive and impractical for covering a wide range of languages over extensive image data. An alternative solution

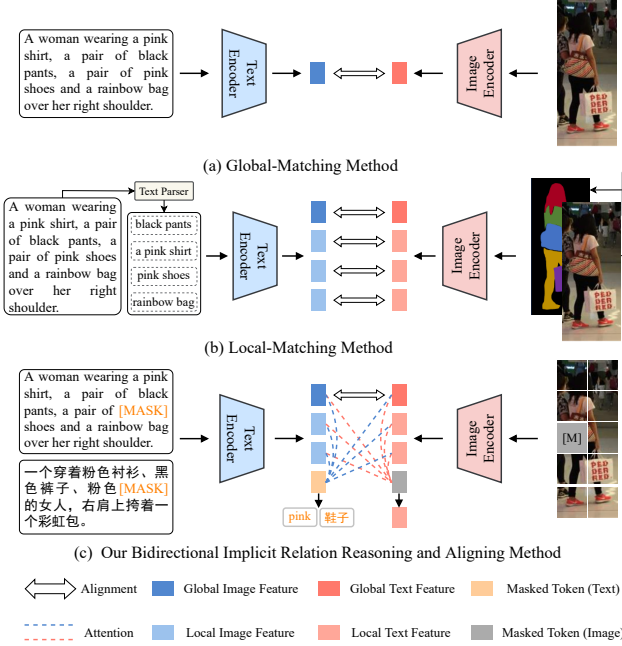


Fig. 1: Illustration of different TIPR methods. (a) Global-matching methods directly align global image and text feature representations. (b) Local-matching methods explicitly extract and align local image and text representations. (c) Our bidirectional implicit relation reasoning and aligning method not only implicitly reasons about the relations among all local tokens but also aligns global image and text representations in a multilingual environment.

[28] for automatic translation on existing TIPR datasets, thereby extending them beyond English. However, directly using LMs for translation often introduces noise due to their lack of domain-specific knowledge. For this, we develop a LMs-driven Domain Adaptive Translation (LDAT) pipeline, consisting of translation, filtering, and rewriting phases. After an initial translation by Large Language Models (LLMs) [23], [24], [29] in the translation phase, we identify clean and noisy translation texts in the filtering phase. Subsequently, the clean texts with corresponding person images serve as the supervision signal to finetune Multimodal Large Language Models (MLLMs) [25], [26], [27]. The finetuned MLLMs, enriched with comprehensive domain-specific knowledge, are employed to rewrite noisy texts in the rewriting phase, thus effectively mitigating the noise issue. The proposed LDAT, as a concise translation pipeline, enables the cost-effective acquisition of the high-quality multilingual TIPR benchmark.

Framework. In contrast to the traditional TIPR task that deals solely with heterogeneity between text and image modalities, the multilingual TIPR encapsulates modality heterogeneity and linguistic diversity challenges. In response, we propose the Bi-IRRA framework, designed to establish robust global alignment and explore implicit fine-grained relations across diverse languages and modalities (as depicted in Fig. 1 (c)). Specifically, Bi-IRRA comprises a Bidirectional Implicit Relation Reasoning (Bi-IRR) module and a Multi-dimensional Global Alignment (Md-GA) module. The Bi-IRR module performs bidirectional prediction

of masked text and image, enabling the implicit modeling of local relations between vision and language. It includes a *bi-lingual Masked Language Modeling (MLM)* pretext task and a *cross-lingual Distillation Masked Image Modeling (D-MIM)* pretext task, tailored for adaptive modeling and interaction across languages. Meanwhile, Md-GA aligns global text and image representations from multiple dimensions, facilitating the extraction of discriminative global text and image representations. The Md-GA module comprises a *bi-lingual Image-Text Contrastive (ITC)* pretext task and a *bi-lingual Asymmetric Image-Text Matching (A-ITM)* pretext task, with the former applied to the unimodal encoders and the latter to the subsequent multimodal interaction encoder. Notably, an asymmetric masking operation on input data is integrated into the *bi-lingual A-ITM* to facilitate noise-robust learning for noisy target text. Finally, by integrating these two modules, Bi-IRRA achieves comprehensive alignment across languages and modalities at both fine-grained and coarse-grained levels.

Our main contributions are as follows. 1) We pioneer a multilingual TIPR task that enables querying in multiple languages to retrieve the target person, offering substantial real-world application potential compared to traditional TIPR. 2) We develop the LDAT pipeline to acquire multilingual TIPR data. By introducing domain-specific knowledge into LMs to mitigate the noise issue in LMs-driven translation, LDAT enables the construction of high-quality multilingual TIPR benchmarks. 3) To address both modality heterogeneity and linguistic diversity in multilingual TIPR, we propose the Bi-IRRA framework, which learns bidirectional implicit relations at a fine-grained level while attaining global alignment at a coarse-grained level across languages and modalities. 4) Extensive experiments on the multilingual TIPR benchmarks demonstrate that the proposed framework surpasses existing SOTA methods.

guages and modalities from multiple dimensions. Notably, we also consider the noisy interference from the generated target texts during global alignment by designing a *bi-lingual A-ITM* pretext task into Md-GA. (3) We expand experiments to encompass a wider range of language environments for TIPR, accompanied by more thorough analyses. To the best of our knowledge, our work is the first effort to address TIPR in a multilingual setting, significantly promoting its practical application value.

2 RELATED WORK

2.1 Text-to-Image Person Retrieval

Since Li *et al.* [1] pioneered the TIPR task, there have been notable advancements in its development [5], [7], [8], [30], [31]. The key challenge in TIPR lies in the significant modality heterogeneity between vision and language. Existing methods address this challenge by focusing on representation learning and cross-modal alignment.

Representation Learning. This category of methods is centered on the development of robust representation learning network designed to extract discriminative feature representations relevant to individuals. Early works [1], [12], [32], [33], [34] employed convolutional neural networks [35], [36] as image encoder and LSTM [37] as text encoder. Later works [10], [38] improved these networks with ViT [39] for images and BERT [40] for texts. More recent advancements have embraced powerful Vision-Language Pre-training (VLP) models [41], [42] (e.g., CLIP [41]) as the representation learning networks. These VLP models are typically pre-trained on large-scale vision-language datasets, enabling them to extract more discriminative person feature representations, thereby garnering considerable attention [5], [6], [7], [9], [43]. For instance, Han *et al.* [44] introduced the CLIP model into TIPR, devising a momentum contrastive learning framework to transfer knowledge from large-scale image-text pairs to the representation learning in TIPR; Cao *et al.* [7] conducted a comprehensive empirical study of CLIP for TIPR, establishing a robust TBPS-CLIP baseline for effective representation learning in this task.

Cross-modal Alignment. The second category of methods is dedicated to designing an effective alignment strategy to achieve favorable cross-modal alignment. Global-matching methods [7], [10], [11], [12], [13] aligned global textual and visual representations directly through designing the rational cross-modal matching loss functions. While these methods are straightforward and intuitive, they often overlook fine-grained information when performing cross-modal alignment. Later, local-matching methods [19], [20], [43], [45], [46], [47], [48] have been introduced to align fine-grained visual and textual information, enhancing cross-modal alignment. Typically, most of these methods [17], [19], [20], [48], [49] relied on external technologies and predefined rules to explicitly extract local textual and visual information, such as text phrases, human segmentation [20], [49], and color information [48], to model fine-grained relations between two modalities. For example, Fujii *et al.* [49] leveraged human parsing models to obtain semantic labels of images, which serve as supervision signals to align cross-modal information. Although incorporating such fine-grained information enhances retrieval performance, these

explicit local-matching methods introduce additional computational complexity during inference when computing the similarity of all these local part representations of images and texts. In comparison, several implicit local-matching methods [5], [6], [8], [50], [51], [52] have been proposed. These methods explore fine-grained cross-modal alignment relations without relying on external explicit dependencies, significantly reducing additional computational overhead. For example, the conference version of this work [5] leveraged the MLM pretext task where visual and unmasked textual information are integrated to predict the masked textual tokens. These masked tokens are treated as anchors to align fine-grained cross-modal information implicitly.

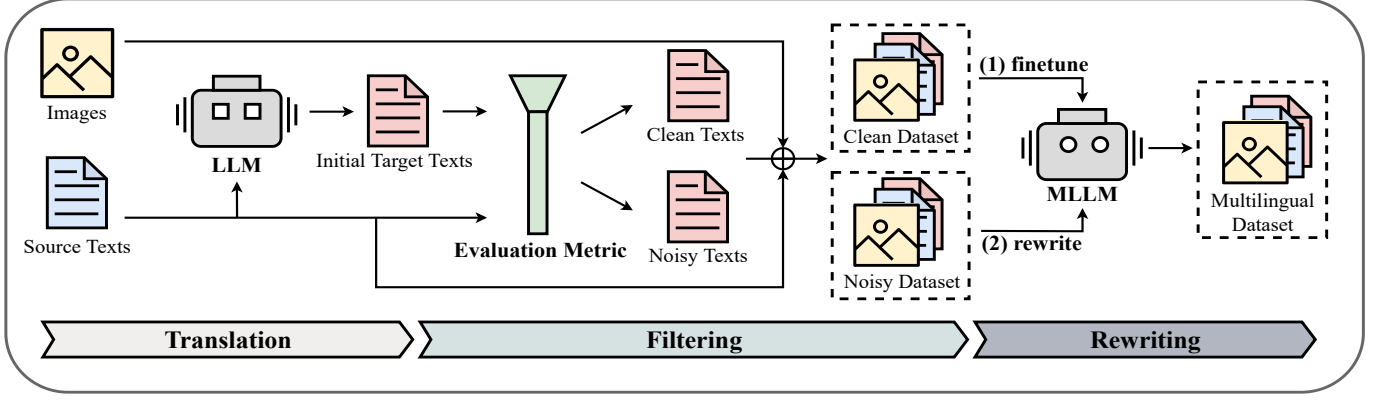


Fig. 2: An Overview of LMs-driven Domain Adaptive Translation (LDAT) pipeline. It consists of three main phases: translation, filtering, and rewriting. In the translation phase, the source texts are translated into the initial target texts by an LLM. During the filtering phase, an evaluation metric is adopted to assess the noise level of each initial target text, based on which we divide the initial target texts into two groups: clean texts and noisy texts. These texts are then paired with their corresponding images and source texts, resulting in a clean dataset and a noisy dataset, respectively. The rewriting phase includes two steps: (1) finetuning an MLLM with the clean dataset, and (2) using the finetuned MLLM to rewrite the initial target texts from the noisy dataset, thereby producing a high-quality multilingual TIPR dataset.

task, we propose the LDAT pipeline, which incorporates domain-specific knowledge to construct high-quality multilingual data. In terms of framework design, we place a greater emphasis on fine-grained cross-modal alignment by introducing the Bi-IRRA framework.

3 LMS-DRIVEN DOMAIN ADAPTIVE TRANSLATION

Beginning with existing TIPR dataset $\mathcal{D} = \{I_n, T_n^s\}_{n=1}^N$, containing the n -th person image I_n and its corresponding English text T_n^s (referred to as the source text), we build its multilingual counterpart $\mathcal{D}_{\mathcal{M}} = \{I_n, T_n^s, T_n^t\}_{n=1}^N$, where T_n^t represents the non-English text (referred to as the target text) and is paired with I_n and T_n^s .

We propose the LDAT pipeline¹ to automatically acquire T^t . First, an LLM is employed to translate each source text T^s into its corresponding initial target text $\tilde{T}^t$. These initial target texts are then filtered into clean and noisy texts. Finally, noisy texts are rewritten by the MLLM finetuned on domain-specific data, which includes clean texts along with the person images, to produce the final high-quality target texts. To sum up, this process, depicted in Fig. 2, comprises three key phases: translation, filtering, and rewriting.

Translation. We first translate the source text T^s in English, to the target text $\tilde{T}^t$ in a non-English language, by leveraging the language understanding capability of LLMs. Specifically, we launch an LLM [23], [29] with a fixed instruction template:

Please translate the English sentence '{source-text}' into {target-language}.

Here, the placeholder {source-text} is filled with the source text T^s , and {target-language} is replaced with the target language name, (e.g., Chinese or French). With this

We introduce a threshold θ to divide these initial target texts into clean and noisy texts. Specifically, if $\phi_n \leq \theta$, the corresponding $\tilde{T}_n^t$ is classified as the clean text and also serves as the final target text T_n^t ; otherwise, it is identified as the noisy text. Subsequently, by pairing these texts with their respective images and source texts, we construct two datasets: a clean dataset $\mathcal{D}_{\mathcal{M}}^C$ and a noisy dataset $\mathcal{D}_{\mathcal{M}}^N$:

$$\mathcal{D}_{\mathcal{M}}^C = \{(I_n, T_n^s, T_n^t) \mid \phi_n \leq \theta\}_{n=1}^{N_{clean}}, \quad (2)$$

$$\mathcal{D}_{\mathcal{M}}^N = \{(I_n, T_n^s, \tilde{T}_n^t) \mid \phi_n > \theta\}_{n=1}^{N_{noisy}}. \quad (3)$$

Here, N_{clean} and N_{noisy} are the numbers of clean texts and noisy texts, respectively, and $N_{clean} + N_{noisy} = N$.

Rewriting. The goal of this phase is to rewrite the noisy text $\tilde{T}_n^t$ from $\mathcal{D}_{\mathcal{M}}^N$. We first utilize $\mathcal{D}_{\mathcal{M}}^C$ as the supervision signal to finetune an MLLM [25], [26] and then employ the finetuned MLLM to execute the rewriting process.

Specifically, for each triplet data $(I, T^s, T^t) \in \mathcal{D}_{\mathcal{M}}^C$, we construct an instruction based on a predefined template:

{image-path} Please combine the image information, translate the English sentence '{source-text}' into {target-language}.

In this template, ** and ** serve as the special tokens to indicate the region of image input. The placeholder *{image-path}* is replaced with the path to the image I , which enables the model to access and process visual information. *{source-text}* and *{target-language}* are substituted with T^s and the target language name, respectively.

We use the constructed instruction as input, with the target text T^t serving as the ground truth, to finetune an MLLM with LoRA [74] in a fully supervised manner. Subsequently, the same template employed during the finetuning process is used by incorporating the image I and the source text T^s from $\mathcal{D}_{\mathcal{M}}^N$, constructing the instruction to guide the finetuned MLLM to re-translate source text into the target text. These re-translated target texts with their corresponding images and source texts are added to $\mathcal{D}_{\mathcal{M}}^C$, ultimately forming a complete multilingual dataset $\mathcal{D}_{\mathcal{M}} = \{I_n, T_n^s, T_n^t\}_{n=1}^N$.

Typically, finetuning the MLLM with high-quality person data from the clean dataset $\mathcal{D}_{\mathcal{M}}^C$ enriches it with the domain

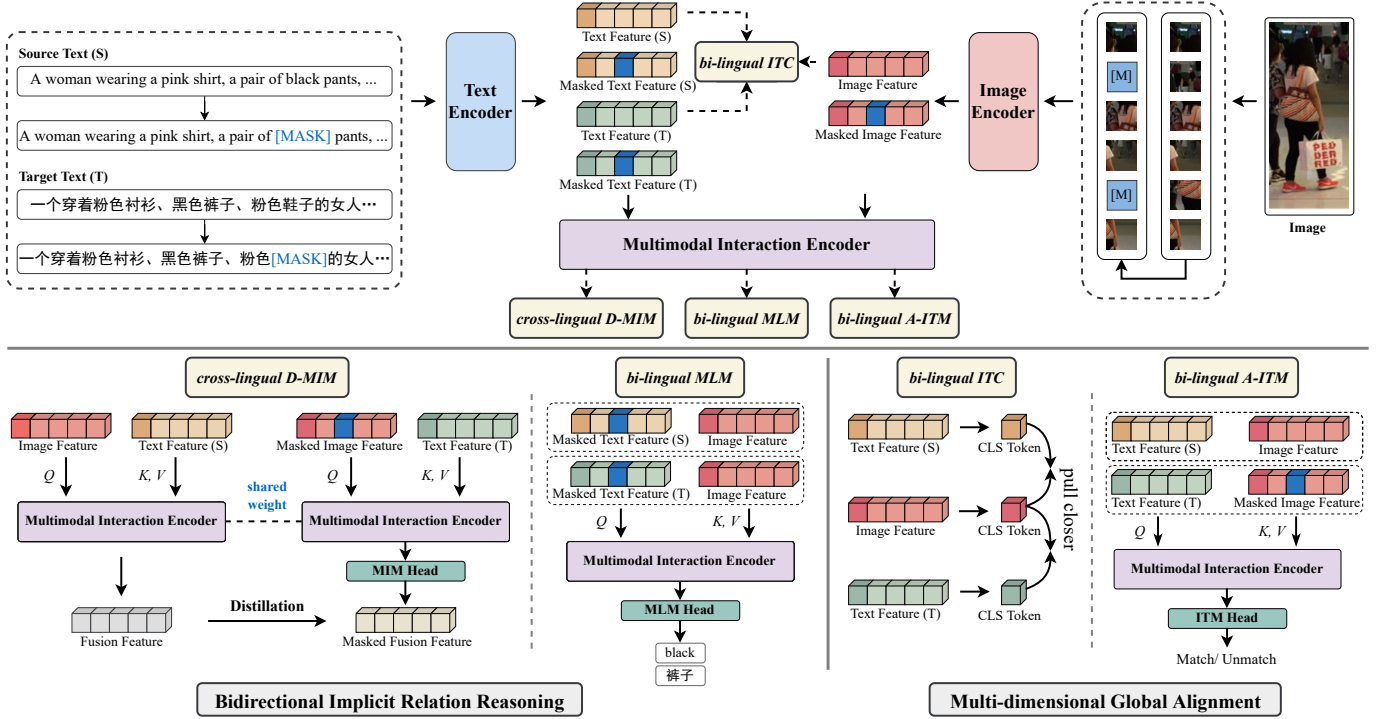


Fig. 3: Overview of the proposed Bidirectional Implicit Relation Reasoning and Aligning (Bi-IRRA) framework. It consists of an image encoder, a text encoder, and a multimodal interaction encoder. With the triplet data as input, Bi-IRRA achieves cross-lingual cross-modal alignment through two key modules. The first is the Bidirectional Implicit Relation Reasoning (Bi-IRR) module, which has two pretext tasks—*cross-lingual D-MIM* and *bi-lingual MLM*. Bi-IRR facilitates the bidirectional modeling of fine-grained relations across languages and modalities. The second is the Multi-dimensional Global Alignment (Md-GA) module, which contains *bi-lingual ITC* and *bi-lingual A-ITM* pretext tasks to align global feature representations of texts and images.

masked and replaced by a learnable masked token, resulting in a masked image denoted as $\hat{I}$. Correspondingly, the masked image feature $F(\hat{I})$ is obtained by feeding $\hat{I}$ into the image encoder. The multimodal interaction encoder with the input of $F(I)$ and $F(T^s)$ is used as the teacher model, while that with the input of $F(\hat{I})$ and $F(T^t)$ is used as the student model. Specifically, taking the teacher model as an example, the multimodal interaction encoder performs the cross-modal fusion via image representation $F(I)$ as query (Q) and text representation $F(T^s)$ as key (K) and value (V) and we obtain the fusion representations:

$$G(I, T^s) = \text{Transformer}(\text{softmax}(\frac{QK^T}{\sqrt{d}})V), \quad (4)$$

$$Q = F(I)W^Q, \quad K = F(T^s)W^K, \quad V = F(T^s)W^V, \quad (5)$$

where W^Q , W^K and W^V are the trainable parameter matrices, and d is the representation dimension. Similarly, we can obtain the fusion representations $G(\hat{I}, T^s)$ from the student model. Notably, the multimodal interaction encoders in the teacher model and student model share weights. Although the encoder processes two distinct inputs—image paired with source text and image paired with target text—both inputs describe the same information and thus share equivalent semantic meaning. By using a shared encoder, we encourage the model to learn a unified multimodal representation space, where semantically similar image-text pairs, regardless of language, are mapped close together.

This design

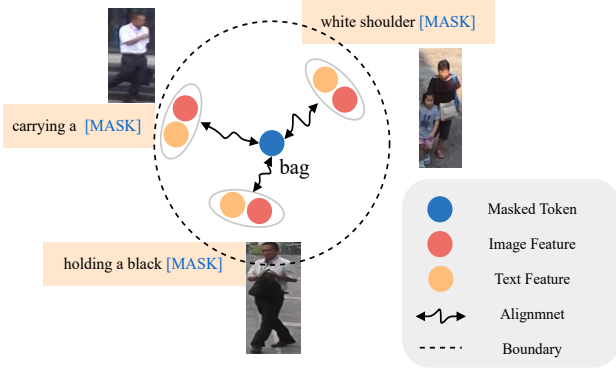


Fig. 4: Illustration of *bi-lingual MLM*, as exemplified with English text-image pairs. It uses masked textual tokens as local fine-grained keys to align visual and textual information.

here employs text representations as queries and text representations as keys/values, and the fusion representations are denoted as $G^*(I, \hat{T}^s)$. Then, a cross-lingual MLM head Ψ_{mlm} , composed of a multi-layer perceptron, predicts the probability distribution $\Psi_{mlm}(I, \hat{T}^s)$ for the masked tokens based on the fusion representations $G^*(I, \hat{T}^s)$. The source text-specific MLM pretext task is formulated as:

$$\mathcal{L}_{mlm}^s = \mathbb{E}_{(I, T^s) \sim \mathcal{D}_{\mathcal{M}}} \mathcal{H}(y_{mlm}, \Psi_{mlm}(I, \hat{T}^s)), \quad (7)$$

where y_{mlm} is a one-hot vocabulary distribution denoting the ground truth and $\mathcal{H}$ represents the cross-entropy.

For the target text T^t , we follow the same procedure to predict its masked textual tokens based on both unmasked target textual and visual information. The target text-specific MLM pretext task is formulated as:

$$\mathcal{L}_{mlm}^t = \mathbb{E}_{(I, T^t) \sim \mathcal{D}_{\mathcal{M}}} \mathcal{H}(y_{mlm}, \Psi_{mlm}(I, \hat{T}^t)). \quad (8)$$

Taken together, the overall *bi-lingual MLM* is defined as:

$$\mathcal{L}_{mlm} = \mathcal{L}_{mlm}^s + \mathcal{L}_{mlm}^t. \quad (9)$$

Notably, in the *bi-lingual MLM*, we share the same multimodal interaction encoder for both the source and target text-specific MLM computations, facilitating indirect cross-lingual relationship modeling.

Discussion. There exists a structural asymmetry between the *cross-lingual D-MIM* and the *bi-lingual MLM*, where the distillation mechanism is integrated into the former but not the latter. Traditional MIM [76], [77], which reconstructs images at the feature level, enables the model to better capture the semantic information of the image. In this process, an additional teacher model (e.g., CLIP) is usually introduced to extract feature representations of the image, providing supervision signals. In our case, given the model's ability to establish robust connections between the source text and image, we utilize the multimodal interaction encoder with the input of source text and the image for supervision signals, thereby leading to *cross-lingual D-MIM*. It avoids extra



Fig. 5: Examples from CUHK-PEDES(M). Each image is paired with text descriptions in four different languages.

Ψ_{itm} to predict the matching probability $\Psi_{itm}(I, T^s)$, and the ITM on (I, T^s) is formulated as:

$$\mathcal{L}_{itm}^s = \mathbb{E}_{(I, T^s) \sim \mathcal{D}_{\mathcal{M}}} \mathcal{H}(y^{itm}, \Psi_{itm}(I, T^s)), \quad (14)$$

where y^{itm} is a two-dimensional one-hot vector representing the ground truth label.

On the other hand, for the image-text pair (I, T^t) , we also compute the global fusion representation g_{cls}^t in $G(\hat{I}, T^t)$. Notably, the target text representations $F(T^s)$ and the masked image representations $F(\hat{I})$ are employed as input to the multimodal interaction encoder. Thus, the ITM on (I, T^t) is formulated as:

$$\mathcal{L}_{itm}^t = \mathbb{E}_{(I, T^t) \sim \mathcal{D}_{\mathcal{M}}} \mathcal{H}(y^{itm}, \Psi_{itm}(\hat{I}, T^t)). \quad (15)$$

Overall, the *bi-lingual A-ITM* pretext task is defined as:

$$\mathcal{L}_{a-itm} = \frac{1}{2}(\mathcal{L}_{itm}^s + \mathcal{L}_{itm}^t). \quad (16)$$

Here, an asymmetric masking design is implemented in computing the two ITM tasks. Specifically, the image representations paired with the target text are masked, whereas those paired with the source text are unmasked. Despite efforts to denoise the generated target texts in LDAT, there is inevitably some noisy correspondence between the target text and image data. By applying masking exclusively on the target text branch, *i.e.*, randomly masking a subset of image tokens paired with the target text, we effectively reduce the model's reliance on potentially noisy image-text alignments. This introduces a form of semantic regularization, encouraging the model to learn robust, holistic representations by reasoning over partial visual input. In this way, masking serves as a regularizer that enhances generalization under noisy supervision.

4.4 Joint Optimization

Finally, we combine the Bi-IRR module with the Md-GA module and formulate the joint optimization objective:

TABLE 1: Performance comparisons on English TIPR with state-of-the-art TIPR methods on CUHK-PEDES, ICFG-PEDES, and RSTPREid. We categorize these methods into two groups based on whether they are pre-trained on large-scale person data. The method indicated by * is trained on the corresponding multilingual dataset.

Methods	Reference	CUHK-PEDES				ICFG-PEDES				RSTPREid			
		R@1	R@5	R@10	mAP	R@1	R@5	R@10	mAP	R@1	R@5	R@10	mAP
<i>Methods without Pre-training on Large-scale Person Data:</i>													
CMKA [33]	TIP21	54.69	73.65	81.86	-	-	-	-	-	-	-	-	-
LapsCore [48]	ICCV21	63.40	-	87.80	-	-	-	-	-	-	-	-	-
SAF [38]	ICASSP22	64.13	82.62	88.40	-	-	-	-	-	-	-	-	-
TIPCB [17]	Neuro22	64.26	83.19	89.10	-	-	-	-	-	-	-	-	-
AXM-Net [79]	MM22	64.44	80.52	86.77	58.73	-	-	-	-	-	-	-	-
MANet [50]	TNNLS23	65.64	83.01	88.78	-	59.44	76.80	82.75	-	-	-	-	-
CFine [80]	TIP23	69.57	85.93	91.15	-	60.83	76.55	82.42	-	50.55	72.50	81.60	-
IRRA [5]	CVPR23	73.38	89.93	93.71	66.13	63.46	80.25	85.82	38.0				

TABLE 3: Performance comparisons with state-of-the-art MITR methods on CUHK-PEDES(M), ICFG-PEDES(M), and RSTPReid(M). The results of other methods are reproduced by running the official code on these multilingual TIPR datasets.

Language	Methods	Reference	CUHK-PEDES(M)				ICFG-PEDES(M)				RSTPReid(M)			
			R@1	R@5	R@10	mAP	R@1	R@5	R@10	mAP	R@1	R@5	R@10	mAP
Chinese	IRRA [5]	CVPR23	66.52	84.58	90.85	60.56	56.83	74.99	81.42	34.66	51.60	74.15	82.45	41.91
	MLLM+IRRA [9]	CVPR24	70.63	87.22	92.24	64.01	59.71	76.58	82.78	37.00	56.40	75.95	83.45	45.81
	CCLM [62]	ACL23	70.83	87.78	92.85	63.14	61.83	78.06	83.98	34.01	68.60	86.15	90.85	52.63
	X ² -VLM [89]	TPAMI23	74.81	90.09	94.17	66.00	66.54	81.08	86.28	39.19	72.00	86.65	91.35	54.67
	CCRK [56]	KDD24	68.31	86.24	91.91	60.69	57.98	75.65	81.99	30.91	64.05	84.70	90.90	49.73
	Bi-IRRA	Ours	76.43	90.72	94.79	67.79	66.79	81.55	86.44	40.26	71.75	87.10	91.65	55.77

TABLE 5: Ablation study on the proposed LDAT on CUHK-PEDES(M). Trans. is the abbreviation of translation.

	English				Chinese			
	R@1	R@5	R@10	mAP	R@1	R@5	R@10	mAP
LLMs-driven Trans.	78.09	91.94	95.52	69.37	75.89	90.58	94.57	67.53
MLLMs-driven Trans.	78.61	92.02	95.47	69.55	76.19	90.85	94.61	67.78
LDAT (w/ LLM)	78.22	91.96	95.42	69.27	76.02	90.46	94.54	67.64
LDAT	78.82	92.02	95.47	69.68	76.43	90.72	94.79	67.79

Moreover, we present the results on UFineBench which involves longer texts with ultra-fine-grained information in Table 2. Bi-IRRA still showcases superior performance on this dataset. Specifically, Bi-IRRA outperforms the current state-of-the-art method CFAM [8] by the significant margins of 1.94%/2.57% at R@1/mAP.

Performance Comparison on non-English TIPR. We compare Bi-IRRA with other methods on non-English TIPR, including some TIPR methods trained on multilingual corpora (IRRA [5], MLLM+IRRA [9]) and state-of-the-art MITR methods (CCRK [56], CCLM [62] and X^2 -VLM [89]). The results are shown in Tables 3 and 4. Bi-IRRA consistently shows superior performance across various non-English environments. For instance, on the CUHK-PEDES(M) dataset, the Bi-IRRA method surpasses the state-of-the-art MITR method X^2 -VLM [89] by 1.62%/1.79%, 1.28%/1.09%, and 0.86%/1.34% in terms of R@1/mAP in Chinese, French, and German, respectively. It indicates that Bi-IRRA achieves effective alignment between different languages and images, demonstrating strong robustness and generalization.

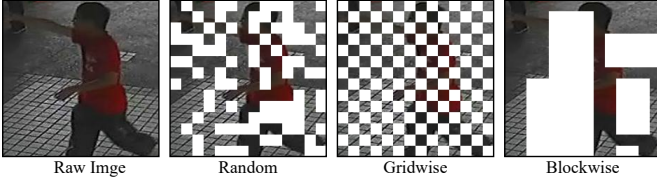
5.4 Ablation Study

We conduct ablation studies on the proposed LDAT and Bi-IRRA to validate their effectiveness. The ablation studies are performed on CUHK-PEDES(M), with English text as the source text and Chinese text as the target text.

Effectiveness of LDAT. We propose the LDAT pipeline, comprising translation, filtering, and rewriting phases, to automatically generate

TABLE 7: Analysis of cross-lingual D-MIM on CUHK-PEDES(M). The cross-lingual D-MIM employs the multimodal interaction encoder with different inputs as both the teacher and student models via distilling the fusion feature representations. Additionally, the attention map generated by the multimodal interaction encoder and the CLS token of the fusion representations can act as alternatives to the fusion feature representations for distillation.

No.	Distill	Teacher		Student		English				Chinese			
		(T^s, I)	(T^t, I)	$(T^s, \hat{I})$	$(T^t, \hat{I})$	R@1	R@5	R@10	mAP	R@1	R@5	R@10	mAP
1	Attention Map	✓			✓	78.72	92.11	95.58	69.60	76.33	90.87	94.70	67.62
2	CLS Token	✓			✓	78.59	92.14	95.66	69.56	75.86	90.61	94.56	67.71
3	Feature		✓	✓		78.59	92.09	95.52	69.55	76.01	90.71	94.59	67.65
4	Feature	✓		✓		78.38	92.33	95.58	69.51	75.97	90.92	94.64	67.63
5	Feature		✓		✓	78.23	92.14	95.66	69.57	75.65	90.79	94.48	67.61
6	Feature	✓		✓		78.41	92.24	95.47	69.67	75.83	90.90	94.51	67.34
7	Feature	✓	✓		✓	78.82	92.02	95.47	69.68	76.43	90.72	94.79	67.79



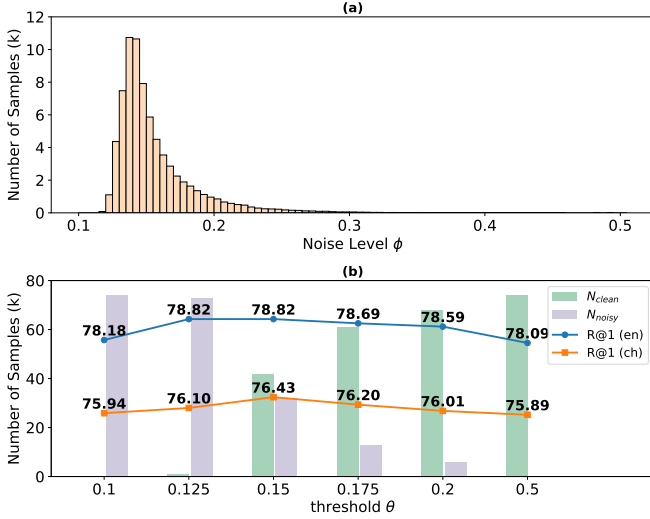


Fig. 7: Illustration of the distribution of noise level (a) and experimental results with varying thresholds (b). In (a), as no data falls within the noise level range of 0.5 ~ 1, so this interval is omitted. In (b), the green and purple bars represent the sample sizes N_{clean} and N_{noisy} of $\mathcal{D}_{\mathcal{M}}^C$ and $\mathcal{D}_{\mathcal{M}}^N$, respectively. The blue and orange line charts illustrate the retrieval performance in English and Chinese at the corresponding θ .

input settings, with results reported in Table 9. (1) We first employ a conventional setup where (I, T^s) and (I, T^t) (without any mask operation) are used to compute ITM (No.1), and yet resulting in limited performance. (2) We experiment with masking input image data during ITM computation. As shown in No.2 ~ No.4, only masking the image paired with the target text yields the most favorable results. It facilitates noise-robust learning for the potentially noisy correspondence between the target text and image. (3) Alternatively, we explore masking input text data during ITM computation. However, as shown in No.5 ~ No.7, the performance of masking text is generally inferior to that of masking image. The image typically exhibits more semantic coherence than the text and retains more semantic information after masking for effective cross-modal alignment.

Analysis of Bi-Lingual ITC. Inspired by the masking



Fig. 10: Examples of Chinese texts before and after the rewriting phase in LDAT. (a)~(c) display clean initial target texts that do not require rewriting, while (d)~(i) show noisy initial target texts along with their corresponding rewritten versions. Errors in source texts and initial target texts are marked in orange and red, respectively.

TABLE 9: Analysis of bi-lingual A-ITM on CUHK-PEDES(M). Input-1 and Input-2 are the inputs for Eq. (14) and Eq. (15), respectively. A special mark $[M]$ is used to indicate that T^s , T^t , and I are masked during computation.

No.	Input-1 I	Input-2 T^s	Input-2 I	Input-2 T^t	English			Chinese		
					R@1	R@10	mAP	R@1	R@10	mAP
1					77.92	95.24	68.71	75.49	94.43	66.83
2	[M]				78.53	95.61	69.60	76.17	94.69	67.59
3		[M]			78.82	95.47	69.68	76.43	94.79	67.79
4	[M]		[M]		77.55	95.47	69.38	75.32	94.41	67.52
5		[M]			76.95	95.11	67.86	75.91	94.57	66.98
6				[M]	78.14	95.22	68.63	75.13	94.22	66.42
7		[M]		[M]	73.86	94				

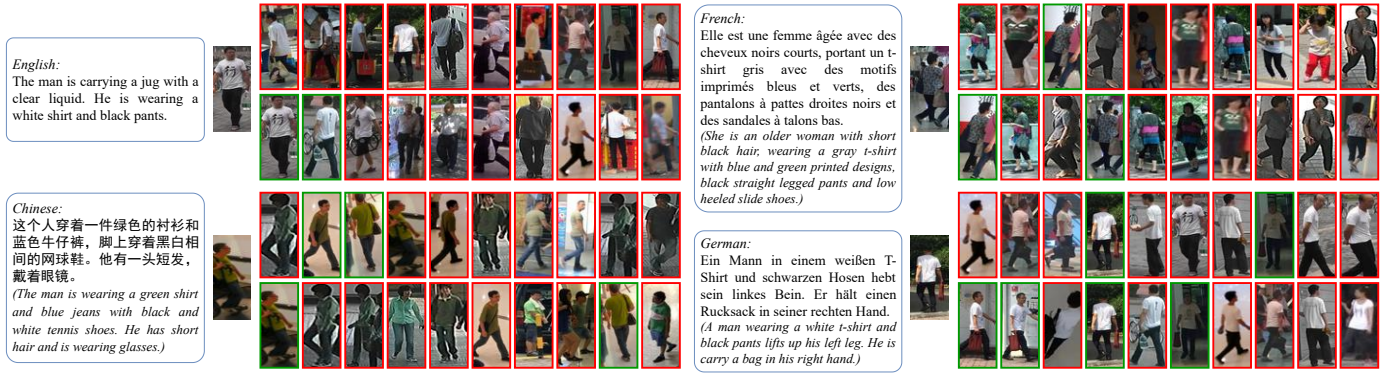


Fig. 11: Comparison of top-10 retrieved results on CUHK-PEDES(M) between IRRA (the first row) and Bi-IRRA (the second row) for each text query. The matched and mismatched images are marked with green and red rectangles, respectively.

clear and concise descriptions, enabling LLMs to achieve satisfactory translation results. Fig. 10 (d)~(i) present initial target texts containing noise and their rewritten versions. Translation errors in the initial target texts are highlighted in red. Notably, in (h)~(i), noise present in the source texts, marked in orange, significantly misleads LLMs for translation. For instance, the misspelling of “sneakers” as “snickers” in the source text results in an inaccurate translation in (h). These initial target texts are effectively revised during the subsequent rewriting phase. When incorporating the domain knowledge, LDAT demonstrates strong performance in translation, consistently producing accurate target texts.

Fig. 11 presents a comparison of the top-10 retrieval results obtained from IRRA and our proposed Bi-IRRA, utilizing text queries in various languages. As illustrated in the figure, Bi-IRRA demonstrates superior performance in capturing fine-grained information such as “carrying a jug with a clear liquid”, “white shirt”, and “black pants”, leading to significantly more accurate retrieval results with queries in multiple languages.

6 CONCLUSION

This paper pioneers a multilingual TIPR task, exploring TIPR in multilingual scenarios. First, we propose the LDAT pipeline to construct a multilingual TIPR benchmark automatically. LDAT alleviates noise issues in large model translations by effectively acquiring and leveraging domain-specific knowledge, enabling the efficient creation of high-quality multilingual TIPR datasets. In addition, we introduce Bi-IRRA: a cross-modal Bidirectional Implicit Relation Reasoning and Aligning framework to achieve comprehensive alignment across different languages and modalities. Extensive experiments demonstrate that the proposed framework consistently achieves superior retrieval performance across various languages. We believe that research on the multilingual TIPR task can further drive the practical application of this field.

REFERENCES

- [1] S. Li, T. Xiao, H. Li, B. Zhou, D. Yue, and X. Wang, “Person search with natural language description,” in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2017, pp. 1970–1979.
- [2] S. He, H. Luo, P. Wang, F. Wang, H. Li, and W. Jiang, “Transreid: Transformer-based object re-identification,” in *Proceedings of the IEEE/CVF international conference on computer vision*, 2021, pp. 15 013–15 022.
- [3] H. Luo, Y. Gu, X. L

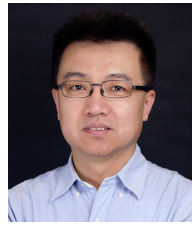
- retrieval," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 39, no. 6, 2025, pp. 5703–5711.
- [17] Y. Chen, G. Zhang, Y. Lu, Z. Wang, and Y. Zheng, "Tipcb: A simple but effective part-based convolutional baseline for text-based person search," *Neurocomputing*, vol. 494, pp. 171–181, 2022.
 - [18] Z. Ding, C. Ding, Z. Shao, and D. Tao, "Semantically self-aligned network for text-to-image part-aware person re-identification," *arXiv preprint arXiv:2107.12666*, 2021.
 - [19] Y. Jing, C. Si, J. Wang, W. Wang, L. Wang, and T. Tan, "Pose-guided multi-granularity attention network for text-based person search," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 34, no. 07, 2020, pp. 11 189–11 196.
 - [20] Z. Wang, Z. Fang, J. Wang, and Y. Yang, "Vitaa: Visual-textual attributes alignment in person search by natural language," in *Computer Vision—ECCV 2020: 16th European Conference, Glasgow, UK, August 23–28, 2020, Proceedings, Part XII 16*. Springer, 2020, pp. 402–420.
 - [21] S. You, C. Chen, Y. Feng, H. Liu, Y. Ji, and M. Ye, "Diverse co-saliency feature learning for text-based person retrieval," *IEEE Transactions on Information Forensics and Security*, 2025.
 - [22] Y. Wu, H. Wang, M. Wu, M. Cao, and M. Zhang, "Laip: Learning local alignment from image-phrase modeling for text-based person search," in *2024 IEEE International Conference on Multimedia and Expo (ICME)*. IEEE, 2024, pp. 1–10.
 - [23] A. Dubey, A. Jauhri, A. Pandey, A. Kadian, A. Al-Dahle, A. Letman, A. Mathur, A. Schelten, A. Yang, A. Fan *et al.*, "The llama 3 herd of models," *arXiv preprint arXiv:2407.21783*, 2024.
 - [24] L. Ouyang, J. Wu, X. Jiang, D. Almeida, C. Wainwright, P. Mishkin, C. Zhang, S. Agarwal, K. Slama, A. Ray *et al.*, "Training language models to follow instructions with human feedback," *Advances in neural information processing systems*, vol. 35, pp. 27 730–27 744, 2022.
 - [25] J. Bai, S. Bai, S. Yang, S. Wang, S. Tan, P. Wang, J. Lin, C. Zhou, and J. Zhou, "Qwen-vl: A frontier large vision-language model with versatile abilities," *arXiv preprint arXiv:2308.12966*, 2023.
 - [26] M. Abdin, S. A. Jacobs, A. A. Awan, J. Aneja, A. Awadallah, H. Awadalla, N. Bach, A. Bahree, A. Bakhtiari, H. Behl *et al.*, "Phi-3 technical report: A highly capable language model locally on your phone," *arXiv preprint arXiv:2404.14219*, 2024.
 - [27] D. Zhu, J. Chen, X. Shen, X. Li, and M. Elhoseiny, "Minigtpt-4: Enhancing vision-language understanding with advanced large language

- [60] Z. Li, Z. Fan, J. Chen, Q. Zhang, X.-J. Huang, and Z. Wei, "Unifying cross-lingual and cross-modal modeling towards weakly supervised multilingual vision-language pre-training," in *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 2023, pp. 5939–5958.
- [61] M. Zhou, L. Zhou, S. Wang, Y. Cheng, L. Li, Z. Yu, and J. Liu, "Uc2: Universal cross-lingual cross-modal vision-and-language pre-training," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2021, pp. 4155–4165.
- [62] Y. Zeng, W. Zhou, A. Luo, Z. Cheng, and X. Zhang, "Cross-view language modeling: Towards unified cross-lingual cross-modal pre-training," *arXiv preprint arXiv:2206.00621*, 2022.
- [63] H. Fei, T. Yu, and P. Li, "Cross-lingual cross-modal pretraining for multimodal retrieval," in *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, 2021, pp. 3644–3650.
- [64] F. Carlsson, P. Eisen, F. Rekathati, and M. Sahlgren, "Cross-lingual and multilingual clip," in *Proceedings of the thirteenth language resources and evaluation conference*, 2022, pp. 6848–6854.
- [65] L. Zhang, A. Hu, and Q. Jin, "Multi-lingual acquisition on multi-modal pre-training for cross-modal retrieval," *Advances in Neural Information Processing Systems*, vol. 35, pp. 29 691–29 704, 2022.
- [66] Y. Wang, J. Dong, T. Liang, M. Zhang, R. Cai, and X. Wang, "Cross-lingual cross-modal retrieval with noise-robust learning," in *Proceedings of the 30th ACM International Conference on Multimedia*, 2022, pp. 422–433.
- [67] Y. Wang, S. Wang, H. Luo, J. Dong, F. Wang, M. Han, X. Wang, and M. Wang, "Dual-view curricular optimal transport for cross-lingual cross-modal retrieval," *IEEE Transactions on Image Processing*, vol. 33, pp. 1522–1533, 2024.
- [68] Y. Wang, F. Wang, J. Dong, and H. Luo, "Cl2cm: Improving cross-lingual cross-modal retrieval via cross-lingual knowledge transfer," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 38, no. 6, 2024, pp. 5651–5659.
- [69] R. Cai, J. Dong, T. Liang, Y. Liang, Y. Wang, X. Yang, X. Wang, and M. Wang, "Cross-lingual cross-modal retrieval with noise-robust fine-tuning," *IEEE Transactions on Knowledge and Data Engineering*, 2024.
- [70] L. Huang, W. Yu, W. Ma, W. Zhong, Z. Feng, H. Wang, Q. Chen, W. Peng, X. Feng, B. Qin *et al.*, "A survey on hallucination in large language models: Principles, taxonomy, challenges, and open questions," *arXiv preprint arXiv:2311.05232*, 2023.
- [71] K. Papineni, S. Roukos, T. Ward, and W.-J. Zhu, "Bleu: a method for automatic evaluation of machine translation," in *Proceedings of the 40th annual meeting of the Association for Computational Linguistics*, 2002, pp. 311–318.
- [72] S. Banerjee and A. Lavie, "Meteor: An automatic metric for mt evaluation with improved correlation with human judgments," in *Proceedings of the acl workshop on intrinsic and extrinsic evaluation measures for machine translation and/or summarization*, 2005, pp. 65–72.
- [73] R. Rei, M. Treviso, N. M.



person re-identification.

Min Cao received her Ph.D. degree in pattern recognition and intelligent systems from the Institute of Automation, Chinese Academy of Sciences, Beijing, China, in 2020. In March 2020, she became a member of the computer science and technology school at Soochow University, where she is currently an Associate Professor. She was a visiting scholar in computer graphics research, Fraunhofer-Gesellschaft, Darmstadt, German, in 2018. Her research interests include cross-modal vision-language learning and per-



Min Zhang (Fellow, ACL) received the bachelor's and PhD degrees from the Harbin Institute of Technology, in 1991 and 1997, respectively. He is a distinguished professor with Soochow University (China). His current research interests include machine translation, natural language processing, large model and AI.



Xinyu Zhou received the BE degree from the School of Computer Science and Technology, Soochow University, Suzhou, China. He is currently work toward the ME degree with the School of Computer Science and Technology, Soochow University. His research interests include cross-modal retrieval and person re-identification.



Ding Jiang received the Master's degree from the School of Computer Science, Wuhan University, Wuhan, China. His research interests include cross-modal retrieval, person re-identification and 3D human pose estimation.



Bo Du (Senior Member, IEEE) received the PhD degree in photogrammetry and remote sensing from the State Key Laboratory of Information Engineering in Surveying, Mapping and Remote Sensing, Wuhan University, Wuhan, China, in 2010. He is a professor with the School of Computer Science, Wuhan University. He has more than 60 research articles published in the IEEE TGRS, TIP, JSTARS, and GRSL. Thirteen of them are ESI hot articles or highly cited articles. His major research interests include pattern recognition, hyperspectral image processing, machine learning, and signal processing. He was a recipient of the Distinguished Paper Award from IJCAI 2018, the Best Paper Award of the IEEE Whispers 2018, the Champion Award of the IEEE Data Fusion Contest 2018, the Best Reviewer Award from the IEEE GRSS for his service to the IEEE Journal of Selected Topics in Earth Observations and Applied Remote Sensing, in 2011, and the ACM rising star awards for his academic progress, in 2015. He was the session chair of the IGARSS 2018/2016 and the 4th IEEE GRSS Workshop on Hyperspectral Image and Signal Processing: Evolution in Remote Sensing.



His research interests focus on computer vision, pattern recognition, and federated learning.

Mang Ye (Senior Member, IEEE) received the PhD degree in computer science from Hong Kong Baptist University, in 2019. He is currently a full professor with the School of Computer Science, Wuhan University, Wuhan, China. He has published more than 100 articles in top-tier venues. He serves as the associate editor for IEEE Transactions on Image Processing, IEEE Transactions on Information Forensics and Security, the Journal of Electronic Imaging, and CAAI Transactions on Intelligence Technology.