

TableVLM: A visual language model with multi-modal joint guided learning for end-to-end image-based table recognition

Anonymous ACL submission

Abstract

Image-based table recognition is one of the important issues in intelligent document processing. Existing solutions usually decompose it into multiple subtasks to solve them separately, but lead to shortcomings like error propagation, weak generalization, etc. Considering multi-modal large language models usually have excellent performance in image captioning and support multiple languages, we propose an innovative end-to-end solution, and construct corresponding datasets, models and evaluation metrics. Specifically, we firstly redefine the HTML representation of the table and remove some unnecessary tags for fair comparison and save limited tokens. Then, we construct a multi-modal dataset containing more than 600k question-answer pairs in total, and each image is annotated only with its HTML representation for training and evaluating the performance of the corresponding methods. In addition, to make the evaluation scheme more comprehensive, we proposed EDSC, Efficiency to evaluate the content recognition ability and cost-effectiveness of various methods. Finally, we construct a multi-modal image-based table recognition model TableVLM, including two different versions, 4B and 14B, focusing on cost-effectiveness and performance respectively. Experimental results show that the proposed TableVLM is able to recognize table images of various styles. Its recognition and generalization capabilities surpass those of existing table-related multi-modal large language models. Therefore, it is an effective and innovative end-to-end solution.

1 Introduction

Table images, as one of the main forms of tables, are extremely common in our daily life. However, since the information in them cannot be read directly, which hinders its widespread application. Therefore, image-based table recognition that can convert table images into readable files is crucial.

To address this problem, existing solutions usually decompose it into four sub-tasks: table structure recognition, table content detection, table content recognition, and table reconstruction, as shown in Fig. 1. Furthermore, according to the differences in the feature representation, existing solutions can be roughly divided into two categories: visual recognition-based methods and sequence generation-based methods. Specifically, the former usually firstly recognizes the visual elements in the table, such as cells, separator lines, rows, columns, etc., and reconstructs the table structure based on their relevant information. Some representative methods include TSRFormer(Lin et al., 2022), SEM(Zhang et al., 2022), LORE(Xing et al., 2023), TGRNet(Xue et al., 2021), RobustTabNet(Ma et al., 2023), TableNet(Paliwal et al., 2019), DeepTabStR(Siddiqui et al., 2019), etc. Then, it is necessary to extract the content in the table based on text detection and text recognition methods, such as DBNet(Liao et al., 2020), DBNet++(Liao et al., 2022), SVTR(Du et al., 2022), etc., and finally embed the content into the correct cells according to pre-set rules. Unlike the former, the latter regards image-based table recognition as a sequence generation task, that is, converting table images into corresponding sequence representations, such as HTML, LaTeX, Markdown, etc. They firstly convert the table structure into the corresponding sequence representation and obtain the pixel coordinates of each cell. Some representative methods include TableMaster(Ye et al., 2021), EDD(Zhong et al., 2020), VAST(Huang et al., 2023), etc. Then, the content in the table is still extracted through text detection and text recognition methods, and embedded into the corresponding cells according to the coordinates of the cells and text blocks obtained previously. According to the above analysis, we can find that most of the existing solutions include multiple steps and require various models to cooperate with each other. However, an inherent defect

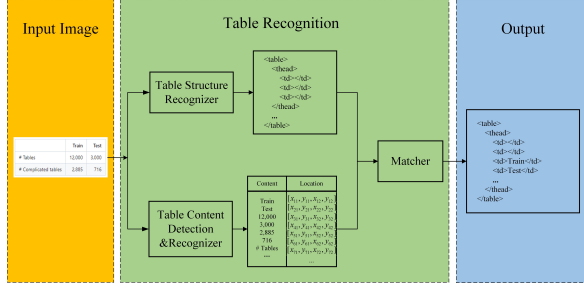


Figure 1: The traditional pipeline of image-based table recognition.

of this pipeline is that the performance degradation of any sub-task will directly lead to the decline in the quality of the final recognition result, such as the loss of table content due to low table structure recognition accuracy, the embedding of table content into the wrong cell due to low cell positioning ability, etc. In addition, the complexity and diversity of table content is also a major challenge. For example, multiple languages, mathematical formulas, subscripts, subscripts, special symbols, etc. In summary, the existing multi-stage solutions still have some limitations that need to be addressed to improve the performance of image-based table recognition.

With the emergence of ChatGPT(Achiam et al., 2023), large language model-related technologies have developed rapidly and have been widely used in various scenarios. Image-based table recognition is no exception. So far, many studies have focused on how to process the tabular data according to the natural language prompts, such as TableLLM(Zhang et al., 2024a), TableGPT(Zha et al., 2023), TableGPT2(Su et al., 2024), etc., or how to use multi-modal large language models to recognize and understand table images, such as Table-LLaVA(Zheng et al., 2024), UniTable(Peng et al., 2024), TabPedia(Zhao et al., 2024), etc., but they have not fully utilized the excellent capabilities of multi-modal large language models. Therefore, we conducts an in-depth exploration and analysis of the image-based table recognition capabilities of multi-modal large language models. Specifically, considering that a table can be represented by the HTML sequence, we regard image-based table recognition as a holistic Image-To-Seq task without having to decompose it into multiple sub-tasks. In addition, since multi-modal large language models usually have excellent image captioning capabilities, we believe that they can understand table images and the correspond-

ing sequence representations. Finally, large language models usually have strong text processing capabilities, and support multiple languages, while OCR-related methods trained on public available datasets do not have such capabilities. Therefore, we believe that it is a good choice to use the excellent capabilities of multi-modal large language models to improve the performance of image-based table recognition. Based on the above analysis, we propose an innovative multi-modal-based end-to-end solution and construct corresponding datasets and models. Specifically, we firstly unified define the HTML representation of table images, that is, remove some unnecessary tags for fair comparison while saving limited tokens. Then, we construct a multi-modal dataset for image-based table recognition, which contains more than 600k question-answer pairs. Each image is annotated only with its HTML representation. Based on this dataset, we comprehensively evaluate the recognition capabilities of existing multi-modal large language models, select the most suitable foundation model, and further construct a multi-modal large language model TableVLM for image-based table recognition, including two different versions, 4B and 14B, focusing on cost-effectiveness and performance, respectively. Experimental results show that the proposed TableVLM is able to recognize the table images with various styles and show competitive performance. In summary, the contributions of this paper are as follows:

1. We proposed an innovative multi-modal-based end-to-end solution, which is more concise than the traditional multi-stage solution and can achieve competitive results with only sequence-level annotation.
2. We unified the HTML representation of table images, that is, deleted some tag pairs that can be added by preset rules and only retained necessary tags, which is conducive to fair comparison and also saves limited tokens.
3. We constructed a multi-modal dataset with more than 600k question-answer pairs. Each table image is annotated only with its HTML annotation.
4. We constructed a multi-modal large language model for image-based table recognition TableVLM, including two different versions, 4B and 14B, which focus on cost-effectiveness and performance respectively.
5. We introduced new evaluation metrics, including EDSC, Efficiency, to more comprehensively evaluate the image-based table recognition ability

of multi-modal large language models.

2 Related Works

In this section, we review the representative works related to image-based table recognition and multi-modal large language models. They are as follows.

2.1 Image-based Table Recognition

Image-based table recognition usually refers to extracting the structure and content in the table image and converting it into a machine-readable format, such as Excel, database, etc. It usually includes four main sub-tasks: table structure recognition, table content detection, table content recognition, and table reconstruction. They work together to achieve image-based table recognition. At present, many excellent solutions have been proposed for this problem, which can be roughly divided into two categories: visual recognition-based methods and sequence generation-based methods. The difference between the two is only in the representation of the table structure. Specifically, the former uses relevant information of visual elements such as cells, separator lines, rows and columns to describe the table structure. Some representative works include TGRNet(Xue et al., 2021), LORE(Xing et al., 2023), LORE++(Long et al., 2025), Robust-TabNet(Ma et al., 2023), TSRFormer(Lin et al., 2022), SEM(Zhang et al., 2022), SEMv2(Zhang et al., 2024b), etc. The latter converts the table structure into the corresponding sequence representation and obtains the pixel coordinates of each cell. The representative methods include Table-Master(Ye et al., 2021), EDD(Zhong et al., 2020), VAST(Huang et al., 2023) etc. After completing the table structure recognition, the contents and their pixel coordinates in the table are extracted through text detection and text recognition methods. And finally, they are embedded into the corresponding cells to obtain the final recognition result. The above is the most common pipeline in image-based table recognition. Although the recognition result can be obtained accurately, it is a multi-stage solution that requires different models and data to solve the above sub-tasks respectively. The performance of any sub-task will directly affect the quality of the final recognition result. For example, due to insufficient accuracy of table content detection, the content is cut off, resulting in recognition errors, and due to low table structure recognition accuracy, the corresponding table content is lost,

etc. Therefore, researchers gradually focus on how to construct a more effective pipeline.

In response to this problem, some corresponding methods have also been proposed, such as MTLTabNet(Ly and Takasu, 2023), which is a multi-task joint optimization model that uses three decoder heads for table structure recognition, cell position recognition, and cell content recognition, respectively, to directly generate HTML sequence representations of table images. In addition, GPT4V-OCR(Shi et al., 2023) also explores the potential in image-based table recognition by fine-tuning GPT-4V. This preliminarily verifies the feasibility of solutions based on multi-modal large language models.

2.2 Multi-Modal Large Language Models

Multi-Modal large language model, its main feature is combine the natural language processing capabilities of the large language model with the ability to understand and generate the data of other modalities, aiming to provide more powerful interactive capabilities by integrating multiple types of input and output such as text, images, sounds, etc. Existing multi-modal large language models mainly include Qwen-VL-Series(Bai et al., 2023), Qwen2-VL-Series(Yang et al., 2024), LLaVA-Series(Liu et al., 2024), mPLUG-Owl-Series(Ye et al., 2024), Phi-Series(Abdin et al., 2024), MiniCPM-V-Series(Hu et al., 2024b), etc. From the perspective of architecture, they can be roughly divided into three parts: pre-trained modal encoder, pre-trained large language model, and adapter. If multi-modal data generation is involved, the generator may also be included. Among them, the pre-trained large language model is very important, and its performance usually directly determines the capabilities of the corresponding multi-modal large language models. At present, many pre-trained large language models have been released one after another. They usually provide multiple different versions. Generally, the larger the number of parameters, the better the performance. The above-mentioned general multi-modal large language models perform well in tasks such as image captioning, visual question answering, OCR, etc., but their performance is not ideal when used directly for image-based table recognition, because image-based table recognition requires accurate content recognition and strict sequential output, which is a more difficult task. However, although it cannot be directly applied, it shows

Attributes		Traditional Solution	Our proposed Solution
Overall Pipeline		Multi-Stage	End-To-End
Supervision	HTML	×	✓
	HTML-Structure	✓	×
	Cell Content	✓	×
	Cell Bbox	✓	×
	Content Bbox	✓	×
Post-processing	Method	Position Matching	String modification
	Complexity	Complex	Simple
Recognition Ability	Table Structure	✓	✓
	Table Content	✓	✓
Generalization Ability	Various Styles	×	✓
	Multi-Language	×	✓

Table 1: The detailed comparison of the traditional multi-stage Image-To-HTML solution and our proposed solution.

competitive performance after fine-tuning, so some table-related multi-modal large language models have been proposed. For example, TabPedia(Zhao et al., 2024) is an innovative table-related visual language model that can seamlessly integrate multiple table-related tasks such as table detection, table structure recognition, and table question answering. A large number of experiments conducted on various public benchmarks have verified its effectiveness and superiority. UniTable(Peng et al., 2024) proposed a novel framework that unify table structure recognition, cell content recognition, and cell bounding box recognition into sequence modeling tasks, and achieved excellent performance on various public available datasets. In addition, Table-LLaVA(Zheng et al., 2024) uses LLaVA-1.5 as the foundation model and trains it on the proposed multi-modal table understanding dataset MMTab. After comprehensive evaluation, Table-LLaVA shows quite good image-based table recognition and understanding capabilities.

Although the above multi-modal large language models have excellent performance, they do not deeply analyze the influencing factors, limitations, etc. between multi-modal large language models and image-based table recognition. We believe that for image-based table recognition, considering that the size of text blocks in table images is usually small, the resolution of the input image should be large to ensure sufficient visual information. In addition, since table content usually contains multiple languages, the pre-trained large language model should also support multiple languages. Considering that the sequence representation of large tables is usually long, the model

should also have a large context length. Based on the above analysis, we comprehensively evaluated existing multi-modal large language models and public available datasets, further constructed an innovative end-to-end solution with corresponding datasets and models. In the following, we introduce them in details.

3 Task Definition

Considering that most existing image-based table recognition methods and public available datasets use HTML to represent tables, we also adopts the same approach. However, there are some differences in the tag pair sets used by different methods and datasets, which in turn affects the fairness of performance comparison. For example, in Table-Master(Ye et al., 2021), the author specifically introduced `<eb></eb>` to represent blank cells, although the fluctuation caused by the difference may not be large. In addition, some tag pairs in HTML have fixed meanings and are not very helpful for recognition. For example, `<html></html>` usually indicates the beginning and end of an HTML sequence, and `<thead></thead>` usually specifically indicates the first row of a table. They all can be added to the HTML sequence when necessary through pre-set rules.

To solve the above problems, we uniformly define the HTML representation of tables to ensure fair comparison. Specifically, we perform the tag pruning to remove some unnecessary tag pairs, such as `<thead></thead>`, `<html></html>`, etc. The specific representation rules are set as follows:

1. We only retain the essential structural tags in HTML, including `<table></table>`, `<tr></tr>`,

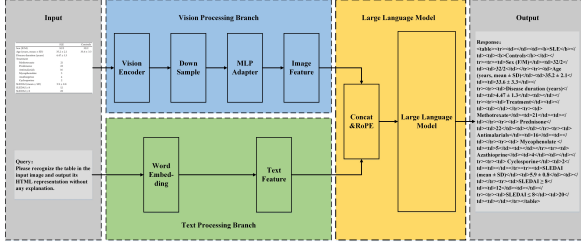


Figure 2: The overall architecture of the innovative multi-modal-based end-to-end solution.

which represent tables, table rows, and cells respectively. If a cell has row-span or column-span, it is achieved by setting the rowspan or colspan attributes of the cell.

2. We retain all content-related special tags, such as , , , etc., which is conducive to enhancing the description ability of HTML.

3. The HTML representation of the table should start with `<table>` and end with `</table>`. It should contain several rows, each starting with `<tr>` and ending with `</tr>`. And each row should contain several cells, each starting with `<td>` and ending with `</td>`. The cell content should be embedded between its start and end tags.

The benefit of following the above rules is that the model can focus on understanding the table structure and content, while making full use of limited tokens, simplifying the task and improving the recognition accuracy.

4 TableVLM

In this section, we introduce the details of the innovative multi-modal-based end-to-end solution and the proposed TableVLM. In the following, we introduce them in details.

4.1 Overall Architecture

The overall architecture of the innovative multi-modal-based end-to-end solution is shown in Fig. 2. It includes several main components, such as visual encoder, adapter, large language models. They are responsible for extracting the visual feature map, alignment, etc., respectively, finally converting the table image into the HTML representation. Compared with traditional solutions, it only requires HTML-level supervision to achieve image-based table recognition, which is a more concise and effective solution. Its detailed comparison with traditional multi-stage Image-To-HTML solution is shown in Table 1.

4.2 TableVLM-4B

In this section, we introduce TableVLM-4B, which includes multiple core modules such as visual encoder, projector, and small language model, with a total of about 4.2 billion parameters. The details of the main components are as follows.

Visual Encoder. Considering that CLIP has excellent image-text alignment capabilities, we use CLIP ViT-L/14 (Cherti et al., 2022) as the visual encoder of TableVLM-4B to extract visual feature maps with stronger description capabilities.

Projector. For TableVLM-4B, two stacked MLP layers are used as the visual language adapters, which mainly for compressing or aligning visual feature maps and text feature maps.

Small Language Model. Due to Phi3.5-mini-Series (Abdin et al., 2024) language models perform well in various public benchmarks, such as DocVQA (Mathew et al., 2020), TextVQA (Singh et al., 2019), OCRbench (Liu et al., 2023), etc., we use Phi3.5-mini-instruct as the small language model in Table VLM-4B.

The above components together constitute TableVLM-4B. It has a small number of parameters while maintaining satisfactory performance. However, in some complex scenarios, its performance is not ideal, so it is necessary to explore the image-based table recognition capabilities of multi-modal large language models with larger parameters.

4.3 TableVLM-14B

In this section, we introduce the details of TableVLM-14B, which is a large visual language model with better performance. It also contains multiple core components such as visual encoder, projector, large language model, etc., with a total of about 13.9 billion parameters. We introduce the main components in details below.

Visual Encoder. Compared with CLIP, EVA-CLIP (Sun et al., 2023) performs better with the same number of parameters and lower training cost. Therefore, we use EVA-02-CLIP-E as the visual encoder of TableVLM-14B, which has more parameters and stronger visual feature map extraction ability. And the resolution of the input image is set to 1120×1120 .

Projector. Similar to the above, we use two layers of stacked MLP as projectors to align visual feature maps and textual feature maps.

Large language model. Since GLM4-9B-

Chat(Zeng et al., 2024) performs well on various public test benchmarks and supports multiple languages, we use it as the large language model of TableVLM-14B, with a total parameter size of about 9.4B. In addition, we still use RoPE to encode the various feature maps, and all visual feature maps share the same position id. The above components together constitute TableVLM-14B.

5 Experiments and Analysis

In this section, we introduce the proposed multi-modal image-based table recognition dataset, the more comprehensive evaluation scheme, and a series of experiments to verify the effectiveness and advancement of TableVLM, including the evaluation of image-based table recognition ability, the evaluation of generalization ability, and the comparison with existing multi-modal large language models and existing image-based table recognition methods. In the following, we introduce the above in details.

5.1 Dataset

Considering the complexity of table content and morphology, our proposed dataset should contain many table images of different styles as much as possible. Therefore, we selected four distinctive datasets from the existing public available datasets, namely PubTabNet(Zhong et al., 2020), FinTabNet(Zheng et al., 2020), SciTSR(Chi et al., 2019) and TableOCR, all of which contain rich table images and accurate annotations. Based on them, we carefully designed the corresponding scripts to generate the HTML representations of table images. Then, according to the preset dialogue format, we constructed a large multi-modal image-based table recognition dataset, which contains more than 600k question-answer pairs for training and evaluating the performance of multi-modal large language models.

5.2 Evaluation Scheme

For the evaluation scheme, we require it to be able to comprehensively evaluate the image-based table recognition capabilities of the corresponding methods from various dimensions, including precision, cost-effectiveness, etc. Therefore, we proposed two innovative evaluation metrics to further supplement the existing evaluation scheme, as follows.

1.EDSC. Its full name is Edit-Distance-based Similarity for table Content, which is mainly used

to evaluate the accuracy of table content recognition. It takes the average of the edit distance between the content prediction and the corresponding label of each cell in the table as the result, as shown in the following formula.

$$EDSC = \frac{1}{N} \sum_{i=0}^{N-1} EditDist(Pre_i, GT_i) \quad (1)$$

2.Efficiency. It is used to measure the cost-effectiveness of the corresponding method, the calculation formula is shown in formula 2. It should be noted that the basic unit of parameter quantity is billion.

$$Efficiency = \frac{TEDS}{Parameters} \quad (2)$$

The above evaluation metrics, together with TEDS, TEDS-Struct, and Acc, constitute the comprehensive evaluation scheme.

5.3 Results and Analysis

The performance evaluation of TableVLM. We evaluate the image-based table recognition capabilities of TableVLM on four different datasets, including PubTabNet, FinTabNet, SciTSR, and TableOCR. The corresponding evaluation metrics are shown in sub-section 5.2. The quantitative results are shown in Table 2.

According to the above, we conclude that the proposed TableVLM can accurately recognize the table images of various styles, whether Chinese tables, English tables, scanned tables, or distorted tables. Specifically, for FinTabNet and SciTSR, the TEDS and TEDS-Struct of TableVLM both exceed 95%, showing excellent performance. And for TableOCR, it also achieves great performance even though there are a large number of distorted images in this dataset. The above phenomenon shows that TableVLM can uniformly model the recognition of table images of different styles as a sequence generation task without designing solutions for each. For PubTabNet, TableVLM performs competitively but not completely satisfactory. The main reasons are as follows. Firstly, its content is more complex, including superscripts, subscripts, special symbols, etc., which increases the difficulty of image-based table recognition. Secondly, its original annotations do not contain the correspondence between the table content and the cells, which requires us

Models	Dataset	TEDS	TEDS-Struct	Acc	EDSC	Efficiency
Qwen2-VL-7B-Instruct	PubTabNet	79.12%	90.94%	54.21%	68.94%	9.53
	FinTabNet	90.93%	95.55%	75.76%	86.78%	10.97
	SciTSR	96.60%	97.36%	82.57%	95.26%	11.65
	TableOCR	93.67%	96.26%	86.03%	93.25%	11.30
LLava1.5-7B-Instruct	PubTabNet	59.31%	74.59%	11.76%	32.66%	8.35
	FinTabNet	35.07%	71.35%	7.79%	25.75%	4.94
	SciTSR	41.19%	73.61%	16.10%	27.64%	5.80
	TableOCR	39.21%	72.47%	20.13%	31.17%	5.52
LLava1.5-13B-Instruct	PubTabNet	60.21%	75.86%	12.86%	33.71%	4.49
	FinTabNet	37.08%	73.20%	10.08%	26.78%	2.77
	SciTSR	42.91%	74.79%	18.17%	28.92%	3.20
	TableOCR	42.72%	76.04%	24.59%	33.21%	3.19
MiniCPM-V.V2.5	PubTabNet	71.91%	88.03%	39.14%	55.24%	8.46
	FinTabNet	83.97%	92.31%	59.95%	77.18%	9.88
	SciTSR	88.94%	93.77%	61.53%	85.14%	10.46
	TableOCR	87.34%	92.66%	72.69%	86.09%	10.28
TableVLM-4B	PubTabNet	81.86%	92.65%	59.61%	73.63%	19.49
	FinTabNet	93.50%	96.35%	78.07%	89.94%	22.26
	SciTSR	96.55%	97.38%	81.03%	95.15%	22.99
	TableOCR	88.22%	95.55%	82.31%	86.97%	21.00
TableVLM-14B	PubTabNet	83.90%	95.51%	71.10%	76.16%	6.03
	FinTabNet	95.82%	97.82%	84.94%	93.53%	6.89
	SciTSR	97.89%	98.33%	87.67%	97.02%	7.04
	TableOCR	97.70%	99.24%	95.56%	97.44%	7.03

Table 2: The performance comparison of the proposed TableVLM with existing multi-modal large language models on different datasets.

to embed them one by one according to their arrangement order when generating the HTML representation, which may cause some errors, thereby reducing the quality of the multi-modal dataset.

In summary, TableVLM is an effective end-to-end image-based table recognition solution. Its advantage is that it does not need to split the image-based table recognition into multiple sub-tasks to be solved separately, only HTML is needed as supervision, and it can adapt to more different scenarios, etc. In addition, although its performance on PubTabNet is not ideal, we still believe that it can have excellent performance by providing enough high-quality datasets.

The performance comparison with existing multi-modal large language models. We compare the performance of the proposed TableVLM with various existing multi-modal large language models. The dataset and evaluation scheme used remain unchanged. The experimental results are shown in Table 1. It should be noted that the above models used for comparison have been fine-tuned

using the proposed dataset.

According to the data shown in Table 1, it can be seen that TableVLM-14B has the best performance on all datasets, its TEDS and TEDS-Struct are both higher than 98%, far exceeding LLava-1.5-13B-Instruct with similar parameters. TableVLM-4B has the highest cost-effectiveness while taking into account performance. We believe that this should be attributed to its 128k context length, because the longer context length enables it to learn the dependencies between tokens that are far away, which is beneficial for image-based table recognition. For Qwen2-VL-7B-Instruct, LLava-1.5-7B-Instruct, and MiniCPM-V-V2.5-Chat, their performance is slightly inferior to or far inferior to TableVLM. We believe that the difference in their performance comes from different configurations, such as the pre-trained language model, the input resolution of the visual encoder, the context length, etc. Therefore, we make the following conclusion: a multi-modal large language model that is beneficial to image-based table recognition should meet

	Methods	Datasets	TEDS	TEDS-Struct	Acc
Image-based Table Recognition Methods	MTLTabNet	PubTabNet	96.67%	97.88%	-
		FinTabNet	-	98.79%	-
	UniTable	PubTabNet	96.50%	97.89%	-
		FinTabNet	-	98.89%	-
	MuTabNet	PubTabNet	96.87%	-	-
		FinTabNet	97.69%	98.87%	-
	TableMaster+PSENet+Master		PubTabNet	96.84%	-
				-	-
Ours	TableVLM-4B	PubTabNet	81.86%	92.65%	59.61%
		FinTabNet	93.50%	96.35%	78.07%
	TableVLM-14B	PubTabNet	83.90%	95.51%	71.10%
		FinTabNet	95.82%	97.82%	84.94%

Table 3: The performance comparison of TableVLM and various existing image-based table recognition methods on different datasets. It is necessary to note that ‘-’ represent there is no related data in the original paper.

the following conditions: support for multiple languages, large context length, large and flexible input image resolution, powerful OCR capabilities and image captioning capabilities. This is also the direction we are working towards.

The performance comparison with existing image-based table recognition methods. We compare the performance of the proposed TableVLM with various existing image-based table recognition methods, including MTLTabNet(Ly and Takasu, 2023), UniTable(Peng et al., 2024), etc. The experimental results are shown in Table 2.

According to the data shown in Table 2, we can find that, for table structure recognition, TableVLM shows performance close to that of existing image-based table recognition methods, although we did not fine-tune it using data related to table structure recognition. In addition, for table content recognition, TableVLM still shows good performance on FinTabNet, but its performance on PubTabNet is still not ideal. This is because the quality of the corresponding generated multi-modal data is not good enough and only HTML is used as supervision. Nevertheless, we still think this is a promising solution that can unify the recognition of table images with different styles into the same task without splitting it into multiple sub-tasks to solve separately, and requires less supervision. This is also the development trend of image-based table recognition and understanding. In the future, we will strive to improve its performance and cover more complex scenarios as much as possible. **Please note that, due to the limited space, more experimental results are shown in the appendix.**

6 Conclusion

In view of the shortcomings of existing image-based table recognition methods, such as cumbersome processes, error propagation, and weak generalization ability, considering the excellent cross-modal recognition ability and image captioning ability of multi-modal large language models, we proposed an innovative end-to-end solution and constructed corresponding datasets, models, and evaluation metrics. Specifically, we unified the HTML representation of the table and removed some unnecessary tags to make the corresponding models focus on the table structure and content. Then, we constructed a large-scale multi-modal image-based table recognition dataset and a comprehensive evaluation scheme for training and evaluating the performance of the corresponding methods. Finally, we constructed an image-based table recognition model TableVLM, including two different versions, 4B and 14B, focusing on cost-effectiveness and performance respectively. Experimental results show that the proposed TableVLM can unify the recognition of table images of different styles into a same task, and its performance is excellent, with strong recognition and generalization capabilities. The TEDS and TEDS-Struct are as high as more than 98%. In addition, we also analyzed the conditions for the adaptation of multi-modal large language models to image-based table recognition, including long context length, large input resolution, richer data, etc., which provided direction for subsequent optimization. In summary, it is an effective, innovative, and promising end-to-end solution that provides a novel and concise pipeline for image-based table recognition.

7 Limitations

Although this work has comprehensively explored the problem of image-based table recognition based on multi-modal large language models, there are still some limitations that need to be addressed in subsequent research. Firstly, the proposed dataset does not contain a large number of more difficult samples such as cross-page tables, blurred tables, rotated tables, etc., which limits the versatility of TableVLM to a certain extent. In the future, we will collect more complex table images and construct corresponding multi-modal data, such as WTW(Long et al., 2021). We believe that this can greatly enhance its application prospects. Then, the performance of the proposed TableVLM is still unsatisfactory. It still makes mistakes in low-quality table images, approximate characters, etc., and its inference process is not controllable. In the future, we will design a chain of thought for image-based table recognition, so that it can gradually generate HTML sequences, improve the thinking ability of the corresponding models, and obtain better recognition results. Thirdly, the proposed TableVLM is only effective for Image-To-HTML, and cannot be generated for other common sequences such as LaTeX and Markdown. In the future, we will further expand the data volume to make the performance of TableVLM more powerful. Finally, with the development of multi-modal large language models, more and more excellent models and fine-tune methods have emerged, such as mPLUG-DocOwl2(Hu et al., 2024a), DocPedia(Feng et al., 2023), GOT-OCR2.0(Wei et al., 2024), LLaVA-OneVision(Li et al., 2024), DeepSeek-VL2-Tiny(Wu et al., 2024), InternVL2.5-1B/2B, LoRA-GA(Wang et al., 2024), LoRA+(Hayou et al., 2024), etc. Therefore, how to improve the cost-effectiveness of TableVLM but maintain the great performance is also an important issue. In the future, we will evaluate more multi-modal large language models and select stronger foundation model to obtain higher cost-effectiveness.

8 Ethical Considerations

The multi-modal image-based table recognition dataset and model proposed in this paper are constructed based on public available datasets and models, which are usually free and open, using the MIT license. Based on the above, we carefully designed Python scripts to generate HTML

sequence representations according to the original annotations in the dataset and further construct the corresponding multi-modal dataset. The multi-modal dataset will also be an open resource related to multi-modal image-based table recognition and understanding. Therefore, the research in this paper does not involve any ethical issues.

References

- Marah Abidin, Jyoti Aneja, Hany Awadalla, Ahmed Awadallah, Ammar Ahmad Awan, Nguyen Bach, Amit Bahree, Arash Bakhtiari, Jianmin Bao, Harkirat Behl, et al. 2024. Phi-3 technical report: A highly capable language model locally on your phone. *arXiv preprint arXiv:2404.14219*.
- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- Jinze Bai, Shuai Bai, Shusheng Yang, Shijie Wang, Sinan Tan, Peng Wang, Junyang Lin, Chang Zhou, and Jingren Zhou. 2023. Qwen-vl: A frontier large vision-language model with versatile abilities. *arXiv preprint arXiv:2308.12966*.
- Mehdi Cherti, Romain Beaumont, Ross Wightman, Mitchell Wortsman, Gabriel Ilharco, Cade Gordon, Christoph Schuhmann, Ludwig Schmidt, and Jenia Jitsev. 2022. [Reproducible scaling laws for contrastive language-image learning](#). *2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 2818–2829.
- Zewen Chi, Heyan Huang, Heng-Da Xu, Houjin Yu, Wanxuan Yin, and Xian-Ling Mao. 2019. Complicated table structure recognition. *arXiv preprint arXiv:1908.04729*.
- Yongkun Du, Zhineng Chen, Caiyan Jia, Xiaoting Yin, Tianlun Zheng, Chenxia Li, Yuning Du, and Yu-Gang Jiang. 2022. Svtr: Scene text recognition with a single visual model. *arXiv preprint arXiv:2205.00159*.
- Hao Feng, Qi Liu, Hao Liu, Wen gang Zhou, Houqiang Li, and Can Huang. 2023. [Docpedia: Unleashing the power of large multimodal model in the frequency domain for versatile document understanding](#). *ArXiv*, abs/2311.11810.
- Soufiane Hayou, Nikhil Ghosh, and Bin Yu. 2024. [Lora+: Efficient low rank adaptation of large models](#). *ArXiv*, abs/2402.12354.
- Anwen Hu, Haiyang Xu, Liang Zhang, Jiabo Ye, Ming-shi Yan, Ji Zhang, Qin Jin, Fei Huang, and Jingren Zhou. 2024a. [mplug-docowl2: High-resolution compressing for ocr-free multi-page document understanding](#). *ArXiv*, abs/2409.03420.

742	Shengding Hu, Yuge Tu, Xu Han, Chaoqun He, Ganqu	Chixiang Ma, Weihong Lin, Lei Sun, and Qiang Huo.	799
743	Cui, Xiang Long, Zhi Zheng, Yewei Fang, Yuxiang	2023. Robust table detection and structure recogni-	800
744	Huang, Weilin Zhao, Xinrong Zhang, Zhen Leng	tion from heterogeneous document images. <i>Pattern</i>	801
745	Thai, Kaihuo Zhang, Chongyi Wang, Yuan Yao,	<i>Recognition</i> , 133:109006.	802
746	Chenyang Zhao, Jie Zhou, Jie Cai, Zhongwu Zhai,		
747	Ning Ding, Chaochao Jia, Guoyang Zeng, Da-	Minesh Mathew, Dimosthenis Karatzas, R. Manmatha,	803
748	hai Li, Zhiyuan Liu, and Maosong Sun. 2024b.	and C. V. Jawahar. 2020. Docvqa: A dataset for vqa	804
749	Minicpm: Unveiling the potential of small language	on document images. <i>2021 IEEE Winter Conference</i>	805
750	models with scalable training strategies . <i>ArXiv</i> ,	on Applications of Computer Vision (WACV), pages	806
751	abs/2404.06395.	2199–2208.	807
752	Yongshuai Huang, Ning Lu, Dapeng Chen, Yibo Li,	Shubham Singh Paliwal, D Vishwanath, Rohit Rahul,	808
753	Zecheng Xie, Shenggao Zhu, Liangcai Gao, and Wei	Monika Sharma, and Lovekesh Vig. 2019. <i>Tablenet:</i>	809
754	Peng. 2023. Improving table structure recognition	Deep learning model for end-to-end table detection	810
755	with visual-alignment sequential coordinate model-	and tabular data extraction from scanned document	811
756	ing. In <i>Proceedings of the IEEE/CVF Conference</i>	images. In <i>2019 International Conference on Docu-</i>	812
757	<i>on Computer Vision and Pattern Recognition</i> , pages	<i>ment Analysis and Recognition (ICDAR)</i> , pages 128–	813
758	11134–11143.	133. IEEE.	814
759	Bo Li, Yuanhan Zhang, Dong Guo, Renrui Zhang, Feng	ShengYun Peng, Aishwarya Chakravarthy, Seong-	815
760	Li, Hao Zhang, Kaichen Zhang, Yanwei Li, Ziwei	min Lee, Xiaojing Wang, Rajarajeswari Balasubra-	816
761	Liu, and Chunyuan Li. 2024. Llava-onevision: Easy	maniyan, and Duen Horng Chau. 2024. <i>Unitable:</i>	817
762	visual task transfer . <i>ArXiv</i> , abs/2408.03326.	Towards a unified framework for table recognition	818
763	Minghui Liao, Zhaoyi Wan, Cong Yao, Kai Chen, and	via self-supervised pretraining. In <i>NeurIPS 2024</i>	819
764	Xiang Bai. 2020. Real-time scene text detection with	<i>Third Table Representation Learning Workshop</i> .	820
765	differentiable binarization. In <i>Proceedings of the</i>		
766	<i>AAAI conference on artificial intelligence</i> , volume 34,	Yongxin Shi, Dezhi Peng, Wenhui Liao, Zening Lin,	821
767	pages 11474–11481.	Xinhong Chen, Chongyu Liu, Yuyi Zhang, and Lian-	822
768	Minghui Liao, Zhisheng Zou, Zhaoyi Wan, Cong Yao,	wen Jin. 2023. Exploring ocr capabilities of gpt-4v	823
769	and Xiang Bai. 2022. Real-time scene text detection	(ision): A quantitative and in-depth evaluation. <i>arXiv</i>	824
770	with differentiable binarization and adaptive scale	<i>preprint arXiv:2310.16809</i> .	825
771	fusion. <i>IEEE transactions on pattern analysis and</i>		
772	<i>machine intelligence</i> , 45(1):919–931.	Shoaib Ahmed Siddiqui, Imran Ali Fateh, Syed Tah-	826
773	Weihong Lin, Zheng Sun, Chixiang Ma, Mingze Li,	seen Raza Rizvi, Andreas Dengel, and Sheraz Ahmed.	827
774	Jiawei Wang, Lei Sun, and Qiang Huo. 2022. <i>Tsr-</i>	2019. <i>Deeptabstr: Deep learning based table struc-</i>	828
775	former: Table structure recognition with transform-	ture recognition. In <i>2019 international conference on</i>	829
776	ers. In <i>Proceedings of the 30th ACM International</i>	<i>document analysis and recognition (ICDAR)</i> , pages	830
777	<i>Conference on Multimedia</i> , pages 6473–6482.	1403–1409. IEEE.	831
778	Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae	Amanpreet Singh, Vivek Natarajan, Meet Shah,	832
779	Lee. 2024. Visual instruction tuning. <i>Advances in</i>	Yu Jiang, Xinlei Chen, Dhruv Batra, Devi Parikh,	833
780	<i>neural information processing systems</i> , 36.	and Marcus Rohrbach. 2019. Towards vqa models	834
781	Yuliang Liu, Zhang Li, Mingxin Huang, Biao Yang,	that can read . <i>2019 IEEE/CVF Conference on Com-</i>	835
782	Wenwen Yu, Chunyuan Li, Xucheng Yin, Cheng lin	<i>puter Vision and Pattern Recognition (CVPR)</i> , pages	836
783	Liu, Lianwen Jin, and Xiang Bai. 2023. Ocrbench:	8309–8318.	837
784	on the hidden mystery of ocr in large multimodal	Aofeng Su, Aowen Wang, Chao Ye, Chen Zhou,	838
785	models . <i>Science China Information Sciences</i> .	Ga Zhang, Guangcheng Zhu, Haobo Wang, Haokai	839
786	Rujiao Long, Wen Wang, Nan Xue, Feiyu Gao, Zhibo	Xu, Hao Chen, Haoze Li, et al. 2024. <i>Tablegpt2: A</i>	840
787	Yang, Yongpan Wang, and Gui-Song Xia. 2021. Pars-	large multimodal model with tabular data integration.	841
788	ing table structures in the wild . <i>2021 IEEE/CVF In-</i>	<i>arXiv preprint arXiv:2411.02059</i> .	842
789	<i>ternational Conference on Computer Vision (ICCV)</i> ,		
790	pages 924–932.	Quan Sun, Yuxin Fang, Ledell Yu Wu, Xinlong Wang,	843
791	Rujiao Long, Hangdi Xing, Zhibo Yang, Qi Zheng, Zhi	and Yue Cao. 2023. Eva-clip: Improved training	844
792	Yu, Fei Huang, and Cong Yao. 2025. <i>Lore++: logi-</i>	techniques for clip at scale . <i>ArXiv</i> , abs/2303.15389.	845
793	cal location regression network for table structure		
794	recognition with pre-training. <i>Pattern Recognition</i> ,	Shaowen Wang, Linxi Yu, and Jian Li. 2024. Lora-ga:	846
795	157:110816.	Low-rank adaptation with gradient approximation .	847
796	Nam Tuan Ly and Atsuhiko Takasu. 2023. An end-to-	<i>ArXiv</i> , abs/2407.05000.	848
797	end multi-task learning model for image-based table		
798	recognition. <i>arXiv preprint arXiv:2303.08648</i> .	Haoran Wei, Chenglong Liu, Jinyue Chen, Jia Wang,	849
		Lingyu Kong, Yanming Xu, Zheng Ge, Liang Zhao,	850
		Jian-Yuan Sun, Yuang Peng, Chunrui Han, and Xi-	851
		angyu Zhang. 2024. General ocr theory: Towards	852
		ocr-2.0 via a unified end-to-end model . <i>ArXiv</i> ,	853
		abs/2409.01704.	854

855	Zhiyu Wu, Xiaokang Chen, Zizheng Pan, Xingchao	Yifan Xu, Yilin Niu, Yuantao Yang, Yueyan Li, Yushi	914
856	Liu, Wen Liu, Damai Dai, Huazuo Gao, Yiyang Ma,	Bai, Yuxiao Dong, Zehan Qi, Zhaoyu Wang, Zhenyi	915
857	Chengyue Wu, Bing-Li Wang, Zhenda Xie, Yu Wu,	Yang, Zhengxiao Du, Zhen-Ping Hou, and Zihan	916
858	Kai Hu, Jiawei Wang, Yaofeng Sun, Yukun Li, Yishi	Wang. 2024. Chatglm: A family of large language	917
859	Piao, Kang Guan, Aixin Liu, Xin Xie, Yu mei	models from glm-130b to glm-4 all tools . <i>ArXiv</i> ,	918
860	You, Kaihong Dong, Xingkai Yu, Haowei Zhang,	abs/2406.12793 .	919
861	Liang Zhao, Yisong Wang, and Chong Ruan. 2024.		
862	Deepseek-vl2: Mixture-of-experts vision-language	Liangyu Zha, Junlin Zhou, Liyao Li, Rui Wang, Qingyi	920
863	models for advanced multimodal understanding .	Huang, Saisai Yang, Jing Yuan, Changbao Su, Xiang	921
864		Li, Aofeng Su, et al. 2023. Tablegpt: Towards unify-	922
865	Hangdi Xing, Feiyu Gao, Rujiao Long, Jiajun Bu,	ing tables, nature language and commands into one	923
866	Qi Zheng, Liangcheng Li, Cong Yao, and Zhi Yu.	gpt. <i>arXiv preprint arXiv:2307.08674</i> .	924
867	2023. Lore: logical location regression network		
868	for table structure recognition. In <i>Proceedings of</i>	Xiaokang Zhang, Jing Zhang, Zeyao Ma, Yang Li,	925
869	<i>the AAAI Conference on Artificial Intelligence</i> , vol-	Bohan Zhang, Guanlin Li, Zijun Yao, Kangli Xu,	926
	ume 37, pages 2992–3000.	Jinchang Zhou, Daniel Zhang-Li, et al. 2024a.	927
		Tablellm: Enabling tabular data manipulation by	928
		llms in real office usage scenarios. <i>arXiv preprint</i>	929
870	Wenyuan Xue, Baosheng Yu, Wen Wang, Dacheng Tao,	<i>arXiv:2403.19318</i> .	930
871	and Qingyong Li. 2021. Tgrnet: A table graph recon-		
872	struction network for table structure recognition. In	Zhenrong Zhang, Pengfei Hu, Jiefeng Ma, Jun Du,	931
873	<i>Proceedings of the IEEE/CVF International Confer-</i>	Jianshu Zhang, Baocai Yin, Bing Yin, and Cong	932
874	<i>ence on Computer Vision (ICCV)</i> , pages 1295–1304.	Liu. 2024b. Semv2: Table separation line detection	933
		based on instance segmentation. <i>Pattern Recognition</i> ,	934
875	An Yang, Baosong Yang, Binyuan Hui, Bo Zheng,	149:110279.	935
876	Bowen Yu, Chang Zhou, Chengpeng Li, Chengyuan		
877	Li, Dayiheng Liu, Fei Huang, Guanting Dong, Hao-	Zhenrong Zhang, Jianshu Zhang, Jun Du, and Fengren	936
878	ran Wei, Huan Lin, Jialong Tang, Jialin Wang, Jian	Wang. 2022. Split, embed and merge: An accu-	937
879	Yang, Jianhong Tu, Jianwei Zhang, Jianxin Ma, Jin	rate table structure recognizer. <i>Pattern Recognition</i> ,	938
880	Xu, Jingren Zhou, Jinze Bai, Jinzheng He, Junyang	126:108565.	939
881	Lin, Kai Dang, Keming Lu, Ke-Yang Chen, Kexin		
882	Yang, Mei Li, Min Xue, Na Ni, Pei Zhang, Peng	Weichao Zhao, Hao Feng, Qi Liu, Jingqun Tang, Shu	940
883	Wang, Ru Peng, Rui Men, Ruize Gao, Runji Lin,	Wei, Binghong Wu, Lei Liao, Yongjie Ye, Hao Liu,	941
884	Shijie Wang, Shuai Bai, Sinan Tan, Tianhang Zhu,	Wengang Zhou, et al. 2024. Tabpedia: Towards com-	942
885	Tianhao Li, Tianyu Liu, Wenbin Ge, Xiaodong Deng,	prehensive visual table understanding with concept	943
886	Xiaohuan Zhou, Xingzhang Ren, Xinyu Zhang, Xipin	synergy. <i>arXiv preprint arXiv:2406.01326</i> .	944
887	Wei, Xuancheng Ren, Yang Fan, Yang Yao, Yichang		
888	Zhang, Yunyang Wan, Yunfei Chu, Zeyu Cui, Zhenru	Mingyu Zheng, Xinwei Feng, Qingyi Si, Qiaoqiao She,	945
889	Zhang, and Zhi-Wei Fan. 2024. Qwen2 technical	Zheng Lin, Wenbin Jiang, and Weiping Wang. 2024.	946
890	report . <i>ArXiv</i> , abs/2407.10671.	Multimodal table understanding. <i>arXiv preprint</i>	947
		<i>arXiv:2406.08100</i> .	948
891	Jiaquan Ye, Xianbiao Qi, Yelin He, Yihao Chen, Dengyi		
892	Gu, Peng Gao, and Rong Xiao. 2021. Pingan-	Xinyi Zheng, Douglas Burdick, Lucian Popa, and Nancy	949
893	vcgroup’s solution for icdar 2021 competition on	Xin Ru Wang. 2020. Global table extractor (gte): A	950
894	scientific literature parsing task b: table recognition	framework for joint table identification and cell struc-	951
895	to html. <i>arXiv preprint arXiv:2105.01848</i> .	ture recognition using visual context . <i>2021 IEEE</i>	952
		<i>Winter Conference on Applications of Computer Vi-</i>	953
896	Qinghao Ye, Haiyang Xu, Jiabo Ye, Ming Yan, An-	<i>sion (WACV)</i> , pages 697–706.	954
897	wen Hu, Haowei Liu, Qi Qian, Ji Zhang, and Fei		
898	Huang. 2024. mplug-owl2: Revolutionizing multi-	Xu Zhong, Elaheh ShafieiBavani, and Antonio Ji-	955
899	modal large language model with modality collabor-	meno Yepes. 2020. Image-based table recognition:	956
900	ation. In <i>Proceedings of the IEEE/CVF Conference</i>	data, model, and evaluation. In <i>European conference</i>	957
901	<i>on Computer Vision and Pattern Recognition</i> , pages	<i>on computer vision</i> , pages 564–580. Springer.	958
902	13040–13051.		
903	Team Glm Aohan Zeng, Bin Xu, Bowen Wang, Chen-		
904	hui Zhang, Da Yin, Diego Rojas, Guanyu Feng, Han-		
905	lin Zhao, Hanyu Lai, Hao Yu, Hongning Wang, Ji-		
906	adai Sun, Jiajie Zhang, Jiale Cheng, Jiayi Gui, Jie		
907	Tang, Jing Zhang, Juanzi Li, Lei Zhao, Lindong Wu,		
908	Lucen Zhong, Ming yue Liu, Minlie Huang, Peng		
909	Zhang, Qinkai Zheng, Rui Lu, Shuaiqi Duan, Shu-		
910	dan Zhang, Shulin Cao, Shuxun Yang, Weng Lam		
911	Tam, Wenyi Zhao, Xiao Liu, Xiaoyu Xia, Xiaohan		
912	Zhang, Xiaotao Gu, Xin Lv, Xinghan Liu, Xinyi Liu,		
913	Xinyue Yang, Xixuan Song, Xunkai Zhang, Yi An,		

A More details of Multi-Modal Dataset

Due to limited space, we present here the relevant details of the proposed multi-modal dataset, including the composition of the dataset, main characteristics, multi-modal dialogue format, some necessary processing, etc., as follows.

A.1 The Multi-Modal Dialogue Format

Fig. 3 shows the dialogue format used in this multi-modal image-based table recognition dataset, which includes two roles: user and assistant. Among them, the former is responsible for inputting the user’s instructions, and the latter is responsible for outputting the HTML representation of the table images. Based on this, we convert table images and the corresponding HTML representations into single-round multi-modal dialogues, and use them to train and evaluate the image-based table recognition capabilities of the multi-modal large language models.

A.2 The detailed description of Multi-Modal Dataset

As mentioned above, the multi-modal image-based table recognition dataset consists of four public available datasets, namely PubTabNet, FinTabNet, SciTSR and TableOCR. They each have their own characteristics. Specifically, the table images in the first three datasets are all flat English tables, but the last dataset contains many distorted table images and its content is in Chinese. In addition, considering that except for TableOCR, the original annotations of the above datasets do not directly contain the HTML representation of the table images, we carefully design the corresponding scripts to generate the corresponding HTML representation according to the original annotations of the dataset, and further removed a small number of erroneous samples or damaged samples to ensure the high quality of the multi-modal dataset. The basic properties of the above datasets are shown in Table 4, some representative samples are shown in Fig. 4, and the methods used to generate the HTML representation of the table images are shown in Fig. 5, 6, and 7.

B Implement Details

In this sub-section, we introduce the details of experimental configuration, including the hardware environment, the hyperparameters used in the training and inference stages, etc. They are shown in

Table 5.

C More Experimental Results and Analysis

Due to limited space, we put more detailed experimental results in the appendix as the supplement, including the performance comparison of more table-related methods, the analysis of samples that failed to be recognized, the evaluation of generalization ability, etc., as shown below.

C.1 The performance comparison with other existing table recognition-related methods

In this sub-section, we provide more detailed comparison of TableVLM with various existing table-related methods, including the performance comparison with existing table structure recognition methods, the performance comparison with existing table-related multi-modal large language models, etc. The quantitative results are shown in Table 6.

According to the data shown in Table 6, we can find that TableVLM has excellent performance in table structure recognition, and TEDS-Struct exceeds 95%, which is only 2% lower than the current best method SEMv2, and exceeds the table-related multi-modal large language models such as TabPedia, etc., highlighting its excellent table structure recognition ability. However, in terms of table content recognition, the performance gap between TableVLM and existing multi-stage methods is quite obvious, and further optimization is needed. However, a careful observation of the evaluation results shows that this gap is particularly obvious on PubTabNet, but not on FinTabNet. This is because the quality of the multi-modal data generated by us based on PubTabNet is not good enough, resulting in a decrease in the learning effect of the corresponding model, while FinTabNet does not have this defect. Based on this, we still believe that after improving the data quality, TableVLM can also have an excellent performance comparable to SOTA. In general, it is feasible to construct an end-to-end solution based on multi-modal large language models because they have some irreplaceable advantages, such as compatibility with more different scenarios, less supervision required, multi-language generalization capabilities, etc.

C.2 The Analysis of Failure Samples

In this sub-section, we analyze the defects of the proposed solution, mainly including the following

three points.

Firstly, considering the principle of the multi-modal large language model is to predict the next token, so the basic unit of the output is also the token. But for HTML, the tag pairs will be decomposed into multiple tokens for output. If a token is predicted incorrectly during this process, meaningless output will be generated, which often occurs when recognizing cells with row-column spans, as shown in Fig. 8a. However, after a comprehensive statistical analysis of all recognition results, we find that its probability is low, usually less than 0.1%.

Secondly, the length of the token that can be output by the multi-modal large language model is often limited. As the number of tokens that can be output increases, the performance of TableVLM can gain certain gains. However, when facing large tables, due to the extremely long length of its HTML representation, incomplete output will occur, as shown in Fig. 8b.

Finally, the hallucination of the multi-modal large language model is still an inevitable problem, such as the recognition errors, the expression preference formed in context learning, etc. But considering image-based table recognition is a very rigorous task, even if two sentences express the same meaning, they cannot be replaced arbitrarily, so hallucination will also have an adverse effect on image-based table recognition, as shown in Fig. 8c.

C.3 The comprehensive evaluation of the generalization ability of TableVLM

In this sub-section, we comprehensively evaluate the out-of-distribution generalization capabilities of TableVLM and existing multi-modal large language models. Specifically, we use the above four datasets alternately as training sets and test sets to evaluate their recognition ability for out-of-distribution samples from different dimensions such as structure, content, etc. The quantitative results are shown in Tables 7, 8, 9, and 10.

According to the above results, we can conclude the follows. Firstly, TableVLM-14B still has the best performance, far exceeding other multi-modal large language models. Specifically, for table structure recognition, after fine-tuning on large-scale datasets such as PubTabNet, FinTabNet, etc., it can still obtain excellent results on other datasets, and most of TEDS-Struct exceeds 85%, as shown in Table 7. And for table content recognition, due

to the large language model has been pre-trained on large-scale textual datasets, so it has Chinese recognition capabilities even if it is fine-tuned on an English dataset, and vice versa. This is one of the important advantages that the current OCR-related methods do not have. For example, in Table 8, after fine-tuning TableVLM-14B on FinTabNet, its TEDS can still reach 74.78% and TEDS-Struct reaches 87.77% on a dataset named TableOCR that is very different from it, which highlights its excellent generalization ability. In addition, TableVLM-4B has the best cost-effectiveness ratio, but its performance is slightly weaker than TableVLM-14B, close to Qwen2-VL-7B-Instruct, and better than MiniCPM-V-V2.5 and LLaVA-1.5.

Secondly, we find that the performance of the model fine-tuned on a large-scale dataset is better than that of the model fine-tuned on a small-scale dataset. For example, in Table 8 and Table 9, we evaluate the performance of the TableVLM-14B after fine-tuned on FinTabNet on SciTSR, its TEDS reaches 92.47% and TEDS-Struct reaches 95.01%. In contrast, its TEDS and TEDS-Struct are only 72.07% and 78.36%. This is because large-scale datasets usually contain richer and more diverse samples, which enables the model to learn more feature maps with stronger descriptive ability.

Finally, we also find that even if the model is fine-tuned on a dataset containing only flat table images, it still has considerable recognition ability for distorted images, as shown in Table 10. This is because the image-based table recognition is modeled as a sequence generation task, which weakens the negative impact of image distortion. This ability is also one of the advantages that existing multi-stage solutions do not have. Because they often rely on the pixel coordinates of table cells and text blocks for table reconstruction, the changes in pixel coordinates will greatly affect the effect of table reconstruction, but the solutions based on multi-modal large language models do not rely on the pixel coordinates of table-related elements, so they are less affected.

In summary, we conclude that the proposed end-to-end image-based table recognition method based on a multi-modal large language model is an excellent solution with strong recognition and generalization capabilities. It can uniformly model the recognition of table images with different languages and styles as sequence generation, so it is an effective and promising solution that deserves further in-depth exploration by researchers.

Dataset	Image Styles	Language	Data Volumn (Train/Test)
PubTabNet	Flat, no distortion	English	500777/9115
FinTabNet	Flat, no distortion	English	91505/10627
SciTSR	Flat, no distortion	English	11970/3000
TableOCR	Some images are distorted	Chinese	12800/3200

Table 4: The details of the public available datasets used to construct the multi-modal dataset, including image style, language, data volume, etc.

Configurations	Hyperparameters	Details
Hardware Environment	GPU	NVIDIA A800 (80GB)
	CPU	Intel(R) Xeon(R) Gold 6348 CPU @ 2.60GHz
Training Stage	LoRA-Rank	8
	LoRA-Scaling factor	32
	Dropout Rate	0.05
	LoRA-Modules	The Language Model and Vision Module
	LoRA-Epoches	1 or 3 Epoches
	Optimizer	AdamW
	Initial Learning Rate	1e-4
	Learning Rate Schedule	Cosine Learning Rate Decay Strategy
	Beta1	0.9
	Beta2	0.999
Inference Stage	Weight Decay Coefficient	0.1
	Max New Tokens	2048

Table 5: The details of experimental configuration, including the hardware environment, the hyperparameters used in the training and inference stages, etc.

Categories	Methods	Datasets	TEDS	TEDS-Struct	Acc
Table Structure Recognition Methods	TSRFormer	PubTabNet	-	97.5%	-
		FinTabNet	-	98.4%	-
	VAST	PubTabNet	96.31%	97.23%	-
		FinTabNet	98.21%	98.63%	-
	TabStructNet	PubTabNet	-	90.1%	-
	SEMr2	PubTabNet	-	97.5%	-
	GTE	PubTabNet	-	93.0%	-
Table-Related Multi-Modal Large Language Models	TabPedia	PubTabNet	-	95.41%	-
		FinTabNet	-	95.11%	-
	GPT4V-OCR	SciTSR	-	99.19%	-
Ours	TableVLM-4B	PubTabNet	81.86%	92.65%	59.61%
		FinTabNet	93.50%	96.35%	78.07%
		SciTSR	96.55%	97.38%	81.03%
	TableVLM-14B	PubTabNet	83.90%	95.51%	71.10%
		FinTabNet	95.82%	97.82%	84.94%
		SciTSR	97.89%	98.33%	87.67%

Table 6: The performance comparison of TableVLM and various existing image-related methods. It is necessary to note that '-' represent there is no related data in the original paper.

	Years Ended December 31,			% Change	
	2015	2014	2013	2015 vs. 2014	2014 vs. 2013
<i>(dollars in thousands)</i>					
Revenue:					
Drilling & Production	\$ 1,009,416	\$ 1,459,514	\$ 1,378,225	(30.8)%	5.9 %
Bearings & Composites	306,387	347,470	341,628	(11.8)%	1.7 %
Automation	167,877	210,255	134,000	(20.2)%	56.9 %
Total	\$ 1,483,680	\$ 2,017,239	\$ 1,853,853	(26.4)%	8.8 %
Segment earnings	\$ 173,190	\$ 461,815	\$ 459,649	(62.5)%	0.5 %
Operating margin	11.7%	22.9%	24.8%		
Segment EBITDA	\$ 314,960	\$ 573,771	\$ 558,724	(45.1)%	2.7 %
Segment EBITDA margin	21.2%	28.4%	30.1%		
Other measures:					
Depreciation and amortization	\$ 141,779	\$ 111,956	\$ 99,075	26.6 %	13.0 %
Goodwill impairment	1,429,260	2,016,411	1,853,562	(29.1)%	8.8 %
Backlog	155,586	233,347	206,700	(33.3)%	12.8 %
Components of revenue growth:				2015 vs. 2014	2014 vs. 2013
Organic (decline) growth				(34.3)%	3.1 %
Acquisitions				9.3 %	6.6 %
Foreign currency translation				(11.4)%	(0.9)%
				(36.4)%	8.8 %

[illegible]

Figure 3: The dialogue template in the proposed multi-modal image-based table recognition dataset.

Table
Images

HTML Representations

HTML Representations

HTML Representations

Table
Images

HTML Representations

Parameter name	Meaning
N.B.F	Number of fact tables
N.DIM(f)	Number of dimensions describing fact table #f
TOT.NB.DIM	Total number of dimensions
N.B.MEAS(f)	Number of measures in fact table #f
DENSITY(f)	Density rate in fact table #f
N.B.LEVELS(f)	Number of hierarchy levels in dimension #d
N.B.ATT(d)	Number of attributes in hierarchy level #d of dimension #d
N.B.LEVELS(d)	Number of tuples in the latest hierarchy level of dimension #d
DIM.SF.WTOR(d)	Size scale factor in the hierarchy levels of dimension #d

Parameter name	Meaning
N.B.F	Number of fact tables
N.DIM(f)	Number of dimensions describing fact table #f
TOT.NB.DIM	Total number of dimensions
N.B.MEAS(f)	Number of measures in fact table #f
DEN.SF(f)	Density rate in fact table #f
N.B.LEVELS(f)	Number of hierarchy levels in dimension #d
N.B.ATT(d)	Number of attributes in hierarchy level #d of dimension #d
N.B.LEVELS(d)	Number of tuples in the latest hierarchy level of dimension #d
DIM.SF.WTOR(d)	Size scale factor in the hierarchy levels of dimension #d

实验序号	重量/mg	体积/cm ³	下落高度/cm	弹的直径/mm
1	20	7.5	50	2.6
2	20	7.5	60	3.5
3	20	7.5	70	5.0
4	20	3.5	70	11.3
5	20	1.8	70	12.4
6	60	7.5	70	14.6
7	85	7.5	70	15.9

实验序号	质量/mg	体积/cm ³	下落高度/cm	弹的直径/mm
1	20	7.5	50	2.6
2	20	7.5	60	3.5
3	20	7.5	70	5.0
4	20	2.5	70	11.3
5	20	1.8	70	12.4
6	60	7.5	70	14.6
7	85	7.5	70	15.9

		Revenue	
		December 29, 2017	December 29, 2018
		(in thousands)	(in thousands)
U.S.		\$ 890.36	\$ 826.91
International:			
China		276.20	254.38
United Kingdom		134.64	133.61
Germany		101.81	99.15
Japan		93.81	95.76
Italy		85.43	76.28
France		84.94	81.79
Other international		667.58	587.65
Total international		1,387.88	1,330.95
Total sales		\$2,278.21	\$2,157.86

		Revenue	
		December 29, 2017	December 29, 2018
		(in thousands)	(in thousands)
U.S.		\$426.991	\$411.934
International:			
China		276.30	254.83
United Kingdom		133.611	133.611
Germany		99.153	105.735
Japan		95.676	114.900
Italy		76.280	69.599
France		84.945	81.795
Other international		587.478	583.189
Total international		1,330.955	1,292.545
Total sales		\$2,157.586	\$2,105.188

组别	平均分	中位数	方差	合格率	优秀率
甲组	6.8	a	3.76	90%	30%
乙组	b	7.5	1.96	80%	20%

组别	平均分	中位数	方差	合格率	优秀率
甲组	6.8	α	3.76	90%	30%
乙组	b	7.5	1.96	80%	20%

digits of precision	16	32	64	128	256	512	1024	2048
n for $\theta = 1^\circ$	3	6	12	22	41	76	≥ 100	
n for $\theta = 1^\circ$ alt.	3	6	12	24	43	79	≥ 100	
n for $\theta = 45^\circ$ alt.	3	5	10	18	32	58	≥ 100	
n for $\theta = 76^\circ$ alt.	2	5	9	16	29	53	95	≥ 100

digits of precision	16	32	64	128	256	512	1024	2048
n for $\theta = 1^\circ$:	3	6	12	22	41	76	≥ 100	
n for $\theta = 1^\circ$ alt.	3	6	12	24	43	79	≥ 100	
n for $\theta = 45^\circ$ alt.	3	5	10	18	32	58	≥ 100	
n for $\theta = 76^\circ$ alt.	2	5	9	16	29	53	95	≥ 100

篇目	相关句例
《庄子》二则	北冥有鱼 庄子与惠子游于濠梁之上
《礼记》二则	虽有佳肴 大道之行也

篇目	哲思与情
《庄子》二则	北冥有鱼 庄子与惠子游于濠梁之上
《礼记》二则	虽有佳肴 大道之行

	Years Ended December 31,			
	2008	2007	2006	2005
Revenue				
Cost of goods and services	(4,838.81)	(4,497.78)	(4,175.52)	
Gross profit	\$7,558.88	\$7,517.29	\$6,824.83	
Expenses from continuing operations	(2,204.27)	(2,414.92)	(2,294.20)	
Operating earnings	(629.33)	(1,895.49)	(963.81)	
Interest expense, net	(96.07)	8.78	1.00	
Income tax expense	(172.26)	(144.88)	(144.88)	
Total nonrecurring expense	\$3,091.01	\$3,136.15	\$3,136.15	
Earnings before provision for income taxes and discontinued operations	(93.39)	(912.36)	(81.63)	
Provision for income taxes	(251.26)	(242.67)	(291.87)	
Earnings from continuing operations	(494.78)	(686.58)	(373.50)	
Income from discontinued operations, net	(2,022.42)	(4,420.29)	(2,022.42)	
Net earnings	(2,517.26)	(5,106.88)	(2,395.92)	
Basic earnings (loss) per common share				
Earnings from continuing operations	\$ 3.49	\$ 3.33	\$ 2.96	
Income from discontinued operations, net	(8.55)	(18.00)	(8.01)	
Net earnings	(5.06)	(14.67)	(5.05)	
Weighted average shares outstanding	200,377	200,377	200,377	
Diluted earnings (loss) per common share				
Earnings from continuing operations	\$ 3.67	\$ 3.33	\$ 2.96	
Income from discontinued operations, net	(8.55)	(18.00)	(8.01)	
Net earnings	(4.88)	(14.67)	(5.05)	
Weighted average shares outstanding	200,377	200,377	200,377	
Dividends paid per common share	\$ 0.90	\$ 0.77	\$ 0.75	
	Years Ended December 31,			
	2008	2007	2006	2005
	(In thousands per share figures)			
Revenue				
Cost of goods and services	(4,838.81)	(4,497.78)	(4,175.52)	
Gross profit	\$7,558.88	\$7,517.29	\$6,824.83	
Selling and administrative expenses	(1,050.10)	(1,054.97)	(962.67)	
Interest expense, net	(96.07)	8.78	1.00	
Other expense (income)	(125.26)	14.14	(19.99)	
Income tax expense	(172.26)	(144.88)	(144.88)	
Earnings before provision for income taxes and discontinued operations	(93.39)	(912.36)	(81.63)	
Provision for income taxes	(251.26)	(242.67)	(291.87)	
Earnings from continuing operations	(494.78)	(686.58)	(373.50)	
Income from discontinued operations	(2,022.42)	(4,420.29)	(2,022.42)	
Net earnings	(2,517.26)	(5,106.88)	(2,395.92)	
Basic earnings (loss) per common share				
Earnings from continuing operations	\$ 3.49	\$ 3.33	\$ 2.96	
Income from discontinued operations, net	(8.55)	(18.00)	(8.01)	
Net earnings	(5.06)	(14.67)	(5.05)	
Weighted average shares outstanding	200,377	200,377	200,377	
Diluted earnings (loss) per common share				
Earnings from continuing operations	\$ 3.67	\$ 3.33	\$ 2.96	
Income from discontinued operations, net	(8.55)	(18.00)	(8.01)	
Net earnings	(4.88)	(14.67)	(5.05)	
Weighted average shares outstanding	200,377	200,377	200,377	
Dividends paid per common share	\$ 0.90	\$ 0.77	\$ 0.75	

	For the Years Ended December 31,			
	in \$ millions			
	2008	2007	2006	2005
REVENUE				
Industrial Products	\$2,490.90	\$2,487.26	\$2,333.36	\$2,136.94
Plant Management	2,080.10	2,055.04	1,984.99	1,849.10
Fluid Management	1,714.04	1,602.08	1,529.09	1,379.00
Electronic Technologies	1,361.91	1,390.10	1,411.26	1,261.26
Inters. & construction	(11.14)	(14.79)	(11.07)	(10.77)
Total consolidated revenue	\$7,566.88	\$7,319.70	\$6,643.61	\$5,636.29
EARNINGS FROM CONTINUING OPERATIONS				
Segment Earnings				
Industrial Products	\$ 290.34	\$ 312.48	\$ 284.21	\$ 242.12
Plant Management	278.55	291.27	281.33	247.17
Fluid Management	385.17	348.87	332.37	287.37
Electronic Technologies	290.33	287.33	291.27	214.64
Corporate operations	1,657.25	1,699.28	1,666.00	1,548.87
Net interest expense	(96.03)	(97.19)	(89.35)	(77.00)
Income tax expense	(115.89)	(97.13)	(88.30)	(84.85)
Losses and discontinued operations	(224.67)	(242.67)	(207.37)	(207.37)
Goodwill impairment	(600.78)	—	—	—
Earnings from continuing operations — total consolidated	\$ 606.58	\$ 669.76	\$ 666.26	\$ 583.08
OPERATING MARGINS (pre-tax)				
Segment Margins				
Industrial Products	12.29	13.09	12	12
Plant Management	13.99	14.26	14	14
Fluid Management	22.96	21.61	21	21
Electronic Technologies	21.59	20.99	21	17
Total Segment	15.35	14.99	15	15
Company from continuing operations	8.12	9.16	10	10
INCOME TAXES				
Segment Taxes				
Industrial Products	\$ 109.18	\$ 114.48	\$ 94.24	\$ 80.24
Plant Management	105.53	109.32	104.72	94.60
Fluid Management	177.68	165.08	150.00	130.00
Electronic Technologies	136.98	136.38	131.00	100.00
Inters. & construction	(11.14)	(14.79)	(11.07)	(10.77)
Total consolidated income tax	\$ 327.13	\$ 309.97	\$ 281.03	\$ 215.61
EARNINGS FROM CONTINUING OPERATIONS				
Segment Earnings				
Industrial Products	\$ 290.34	\$ 312.48	\$ 284.21	\$ 242.12
Plant Management	278.55	291.27	281.33	247.17
Fluid Management	385.17	348.87	332.37	287.37
Electronic Technologies	290.33	287.33	291.27	214.64
Corporate operations	1,657.25	1,699.28	1,666.00	1,548.87
Net interest expense	(96.03)	(97.19)	(89.35)	(77.00)
Income tax expense	(115.89)	(97.13)	(88.30)	(84.85)
Losses and discontinued operations	(224.67)	(242.67)	(207.37)	(207.37)
Goodwill impairment	(600.78)	—	—	—
Earnings from continuing operations — total consolidated	\$ 606.58	\$ 669.76	\$ 666.26	\$ 583.08
OPERATING MARGINS (pre-tax)				
Segment Margins				
Industrial Products	12.29	13.09	12	12
Plant Management	13.99	14.26	14	14
Fluid Management	22.96	21.61	21	21
Electronic Technologies	21.59	20.99	21	17
Total Segment	15.35	14.99	15	15
Company from continuing operations	8.12	9.16	10	10

(shillings in thousands)	Year Ended December 31			% Change
	2013	2012	2011	
Refinancing and Financial				
Refinancing and Financial	1,546,044	1,373,209	1,248,058	20.3 %
Other Interest	1,246,144	1,465,022	1,248,058	1.2 %
Eliminations	(12,247)	(2,052,022)	(2,425,018)	(10.0) %
Field Losses and				
Eliminations	861,703	872,642	672,621	18.4 %
Eliminations	(12,247)	(1,445,444)	(1,110,000)	11.0 %
Segment costs				
Operating margin	57,559.8	58,165.2	41,946.9	14.7 %
Operating margin	12.2	14.7	14.0	1.0 %
Segment EBITDA				
Operating margin	57,559.8	57,825.1	55,882.6	18.7 %
Operating margin	12.2	14.7	14.0	1.0 %
Other revenues				
Operating margin	51,215.4	47,825.1	37,787.0	40.1 %
Operating margin	11.2	10.2	8.7	25.0 %
Backlog				
Refinancing and Financial	3,294,522	3,283,108	2,622,336	12.3 %
Operating margin	57,559.8	58,165.2	41,946.9	14.7 %
Eliminations	(12,247)	(2,052,022)	(2,425,018)	(10.0) %
Eliminations	861,703	872,642	672,621	18.4 %
Backlog				
Operating margin	57,559.8	57,825.1	55,882.6	18.7 %
Operating margin	12.2	14.7	14.0	1.0 %
Refinancing and Financial				
Operating margin	57,559.8	57,825.1	55,882.6	18.7 %
Operating margin	12.2	14.7	14.0	1.0 %
Components of revenue growth				
Organic growth	1.4	1.6	1.4	0.2 %
Acquisition	9.1	9.1	8.1	1.0 %
Foreign currency translation	(0.5)	(0.5)	(0.5)	0.0 %
Operating margin	12.2	14.7	14.0	1.0 %
Years Ended December 31				
(shillings in thousands)	2013	2012	2011	% Change
Refinancing and Financial				
Refinancing and Financial	1,546,044	1,373,209	1,248,058	20.3 %
Other Interest	1,246,144	1,465,022	1,248,058	1.2 %
Eliminations	(12,247)	(2,052,022)	(2,425,018)	(10.0) %
Field Losses and				
Eliminations	861,703	872,642	672,621	18.4 %
Eliminations	(12,247)	(1,445,444)	(1,110,000)	11.0 %
Segment earnings				
Operating margin	57,559.8	58,165.2	41,946.9	14.7 %
Operating margin	12.2	14.7	14.0	1.0 %
Segment EBITDA				
Operating margin	57,559.8	57,825.1	55,882.6	18.7 %
Operating margin	12.2	14.7	14.0	1.0 %
Other revenues				
Operating margin	51,215.4	47,825.1	37,787.0	40.1 %
Operating margin	11.2	10.2	8.7	25.0 %
Backlog				
Refinancing and Financial	3,294,522	3,283,108	2,622,336	12.3 %
Operating margin	57,559.8	58,165.2	41,946.9	14.7 %
Eliminations	(12,247)	(2,052,022)	(2,425,018)	(10.0) %
Eliminations	861,703	872,642	672,621	18.4 %
Backlog				
Operating margin	57,559.8	57,825.1	55,882.6	18.7 %
Operating margin	12.2	14.7	14.0	1.0 %
Refinancing and Financial				
Operating margin	57,559.8	57,825.1	55,882.6	18.7 %
Operating margin	12.2	14.7	14.0	1.0 %
Components of revenue growth				
Organic growth	1.4	1.6	1.4	0.2 %
Acquisition	9.1	9.1	8.1	1.0 %
Foreign currency translation	(0.5)	(0.5)	(0.5)	0.0 %
Operating margin	12.2	14.7	14.0	1.0 %

Figure 4: Some samples from the proposed multi-modal image-based table recognition dataset, including some table images with various styles and the corresponding HTML representations.

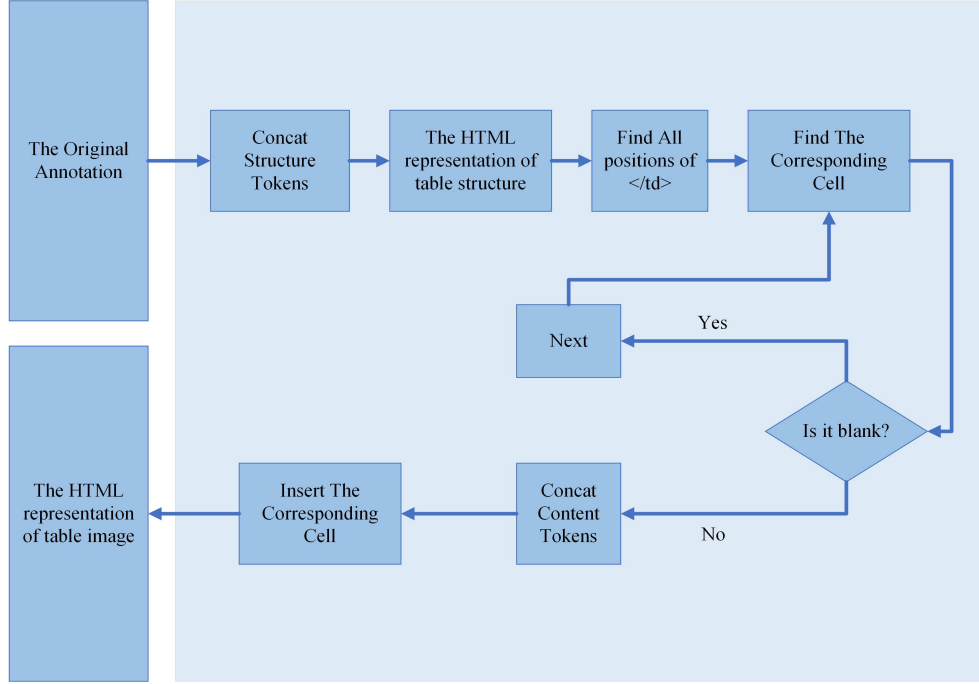


Figure 5: The method to generate the corresponding HTML representation of table images according to the original annotation of PubTabNet.

Models	Train	Test	TEDS	TEDS-Struct	Acc	EDSC	Efficiency
Qwen2-VL-7B-Instruct	PubTabNet	PubTabNet	79.12%	90.94%	54.21%	68.94%	9.53
	PubTabNet	FinTabNet	68.18%	85.40%	32.66%	61.36%	8.21
	PubTabNet	SciTSR	84.04%	94.82%	69.57%	79.37%	10.13
	PubTabNet	TableOCR	43.29%	68.29%	40.44%	64.91%	5.22
LLaVA1.5-7B-Instruct	PubTabNet	PubTabNet	59.31%	74.59%	11.76%	32.66%	8.35
	PubTabNet	FinTabNet	32.10%	62.36%	2.68%	20.64%	4.52
	PubTabNet	SciTSR	45.45%	75.05%	16.47%	25.47%	6.40
	PubTabNet	TableOCR	35.77%	54.73%	5.06%	28.48%	5.04
LLaVA1.5-13B-Instruct	PubTabNet	PubTabNet	60.21%	75.86%	12.86%	33.71%	4.49
	PubTabNet	FinTabNet	31.92%	62.55%	3.00%	20.64%	2.38
	PubTabNet	SciTSR	45.38%	75.12%	16.30%	25.56%	3.39
	PubTabNet	TableOCR	34.41%	57.24%	7.25%	29.77%	2.57
MiniCPM-V-V2.5	PubTabNet	PubTabNet	71.91%	88.03%	39.14%	55.24%	8.46
	PubTabNet	FinTabNet	52.18%	77.92%	19.96%	38.83%	6.14
	PubTabNet	SciTSR	67.95%	88.96%	48.83%	56.85%	7.99
	PubTabNet	TableOCR	45.90%	66.34%	28.50%	40.82%	5.40
TableVLM-4B	PubTabNet	PubTabNet	81.86%	92.65%	59.61%	73.63%	19.49
	PubTabNet	FinTabNet	62.29%	79.13%	26.67%	55.35%	14.83
	PubTabNet	SciTSR	80.88%	90.80%	67.33%	75.40%	19.26
	PubTabNet	TableOCR	45.64%	71.57%	45.88%	41.95%	10.87
TableVLM-14B	PubTabNet	PubTabNet	83.90%	95.51%	71.10%	76.16%	6.04
	PubTabNet	FinTabNet	64.17%	83.94%	36.81%	56.61%	4.62
	PubTabNet	SciTSR	83.24%	95.06%	72.93%	77.81%	5.99
	PubTabNet	TableOCR	53.33%	89.97%	77.41%	48.20%	3.84

Table 7: The comparison of the generalization capabilities of TableVLM and existing multi-modal large language models after fine-tuning on PubTabNet.

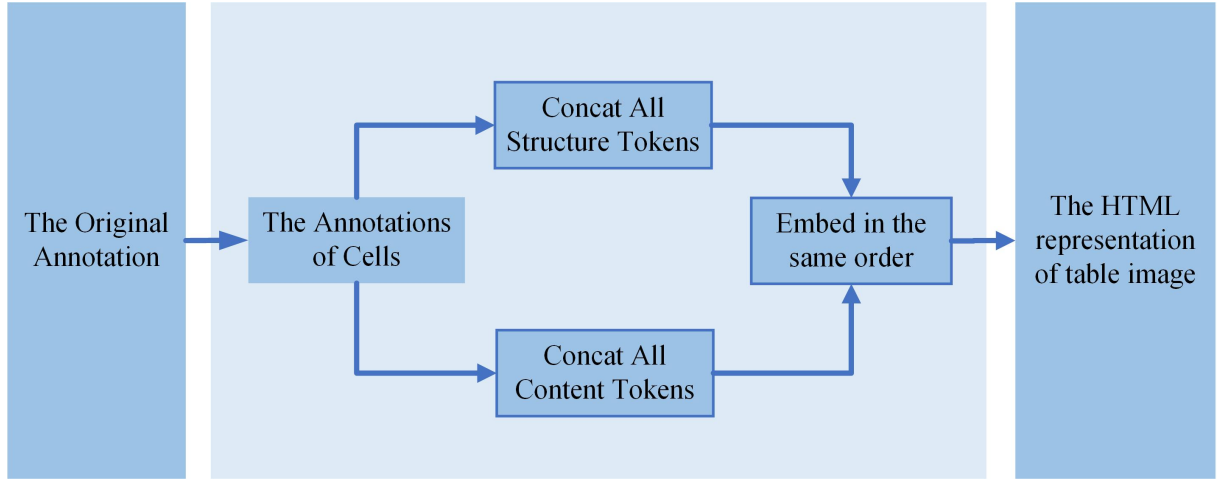


Figure 6: The method to generate the corresponding HTML representation of table images according to the original annotation of FinTabNet.

Models	Train	Test	TEDS	TEDS-Struct	Acc	EDSC	Efficiency
Qwen2-VL-7B-Instruct	FinTabNet	PubTabNet	64.60%	83.67%	36.46%	53.99%	7.78
	FinTabNet	FinTabNet	90.93%	95.55%	75.76%	86.78%	10.96
	FinTabNet	SciTSR	89.47%	92.71%	67.83%	86.28%	10.78
	FinTabNet	TableOCR	72.67%	81.24%	57.31%	75.73%	8.76
LLaVA1.5-7B-Instruct	FinTabNet	PubTabNet	35.74%	59.87%	1.90%	22.75%	5.03
	FinTabNet	FinTabNet	35.07%	71.35%	7.79%	25.75%	4.94
	FinTabNet	SciTSR	25.85%	62.82%	4.27%	19.45%	3.64
	FinTabNet	TableOCR	14.72%	48.09%	1.94%	20.52%	2.07
LLaVA1.5-13B-Instruct	FinTabNet	PubTabNet	36.67%	61.78%	2.12%	22.56%	2.74
	FinTabNet	FinTabNet	37.08%	73.20%	10.08%	26.78%	2.77
	FinTabNet	SciTSR	27.57%	64.38%	5.03%	20.02%	2.06
	FinTabNet	TableOCR	17.36%	52.48%	3.06%	20.96%	1.30
MiniCPM-V-V2.5	FinTabNet	PubTabNet	57.68%	78.73%	22.41%	43.89%	6.79
	FinTabNet	FinTabNet	83.97%	92.31%	59.95%	77.18%	9.88
	FinTabNet	SciTSR	70.14%	82.07%	37.80%	62.93%	8.25
	FinTabNet	TableOCR	51.14%	69.29%	28.28%	51.06%	6.02
TableVLM-4B	FinTabNet	PubTabNet	70.41%	88.48%	44.21%	61.05%	16.76
	FinTabNet	FinTabNet	93.50%	96.35%	78.07%	89.94%	22.26
	FinTabNet	SciTSR	89.89%	93.04%	66.37%	85.22%	21.40
	FinTabNet	TableOCR	56.06%	76.97%	47.69%	56.96%	13.35
TableVLM-14B	FinTabNet	PubTabNet	73.92%	91.44%	53.09%	66.05%	5.32
	FinTabNet	FinTabNet	95.82%	97.82%	84.94%	93.53%	6.89
	FinTabNet	SciTSR	92.47%	95.01%	71.93%	88.75%	6.65
	FinTabNet	TableOCR	74.78%	87.77%	69.41%	73.37%	5.38

Table 8: The comparison of the generalization capabilities of TableVLM and existing multi-modal large language models after fine-tuning on FinTabNet.

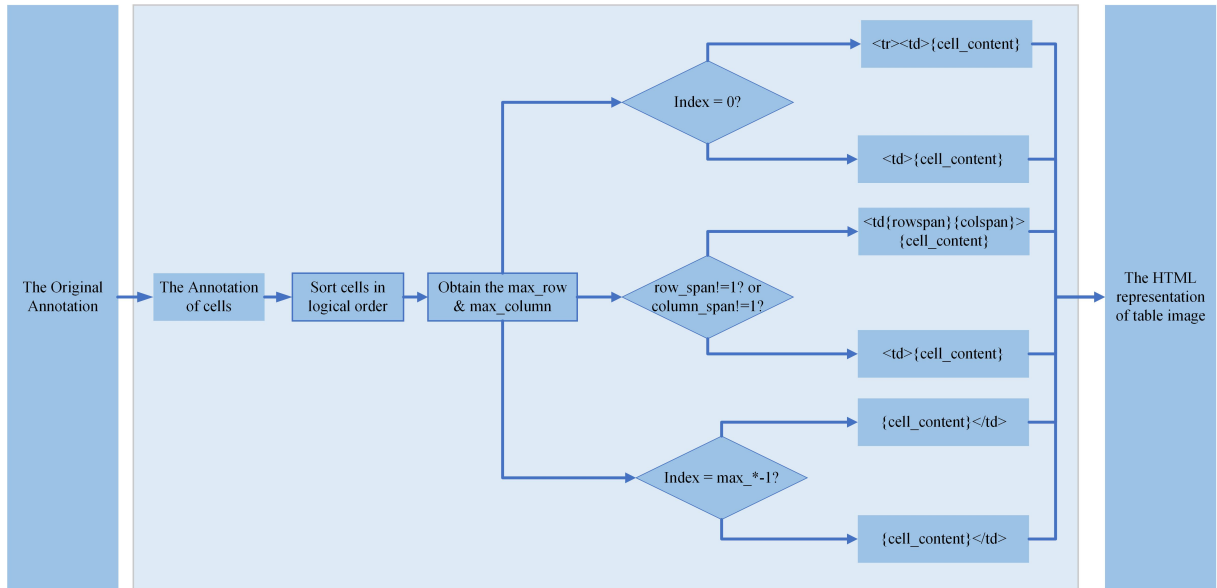


Figure 7: The method to generate the corresponding HTML representation of table images according to the original annotation of SciTSR.

Models	Train	Test	TEDS	TEDS-Struct	Acc	EDSC	Efficiency
Qwen2-VL-7B-Instruct	SciTSR	PubTabNet	63.30%	81.92%	22.16%	48.26%	7.63
	SciTSR	FinTabNet	70.47%	81.26%	21.93%	54.45%	8.49
	SciTSR	SciTSR	96.60%	97.36%	82.57%	95.26%	11.64
	SciTSR	TableOCR	73.38%	81.68%	52.72%	76.78%	8.84
LLaVA1.5-7B-Instruct	SciTSR	PubTabNet	35.93%	60.08%	3.35%	24.65%	5.06
	SciTSR	FinTabNet	17.21%	55.22%	1.42%	19.83%	2.42
	SciTSR	SciTSR	41.19%	73.61%	16.10%	27.64%	5.80
	SciTSR	TableOCR	18.43%	53.43%	5.31%	19.55%	2.60
LLaVA1.5-13B-Instruct	SciTSR	PubTabNet	36.62%	61.90%	3.41%	24.14%	2.73
	SciTSR	FinTabNet	18.01%	55.52%	1.66%	19.19%	1.34
	SciTSR	SciTSR	42.91%	74.79%	18.17%	28.92%	3.20
	SciTSR	TableOCR	18.35%	52.32%	5.59%	19.05%	1.37
MiniCPM-V-V2.5	SciTSR	PubTabNet	60.43%	79.07%	21.37%	46.46%	7.11
	SciTSR	FinTabNet	55.87%	70.63%	11.56%	45.83%	6.57
	SciTSR	SciTSR	88.94%	93.77%	61.53%	85.14%	10.46
	SciTSR	TableOCR	63.60%	72.72%	37.75%	69.79%	7.48
TableVLM-4B	SciTSR	PubTabNet	66.16%	82.91%	29.70%	56.33%	15.75
	SciTSR	FinTabNet	59.28%	68.66%	17.31%	53.47%	14.11
	SciTSR	SciTSR	96.55%	97.38%	81.03%	95.15%	22.99
	SciTSR	TableOCR	53.94%	73.23%	43.28%	59.56%	12.84
TableVLM-14B	SciTSR	PubTabNet	70.96%	88.18%	39.81%	61.59%	5.10
	SciTSR	FinTabNet	72.07%	78.36%	28.28%	60.32%	5.54
	SciTSR	SciTSR	97.89%	98.33%	87.67%	97.02%	7.04
	SciTSR	TableOCR	81.06%	88.04%	70.59%	82.68%	5.83

Table 9: The comparison of the generalization capabilities of TableVLM and existing multi-modal large language models after fine-tuning on SciTSR.

	Years Ended December 31,		
	2010	2011	2013
Revenues:			
Energy	\$ 1,493,040	\$ 2,617,129	\$ 1,835,833
Engineering	3,642,413	2,396,966	2,717,098
Platts	1,181,279	1,605,296	1,287,835
Refining and Food Equipment	3,731,480	3,621,149	3,981,480
Inter-segment transactions	(985)	(2,231)	1
Unaudited total revenue	\$ 6,075,117	\$ 10,239,579	\$ 9,813,227
Earnings from continuing operations:			
Segment revenues:			
Energy	\$ 173,190	\$ 61,813	\$ 459,649
Engineering Systems	37,061	38,993	36,767
Platts	262,117	215,439	224,523
Refining and Food Equipment	221,299	228,784	260,336
Inter-segment transactions	(2,335)	(1,719)	1
Corporate expense - "cfe"	(1)	(155,760)	(13,204)
Goodwill impairment	122,257	177,176	177,176
Earnings before provision for income taxes and discontinued operations	\$ 602,819	\$ 1,062,747	\$ 905,266
Income tax expense	(104,262)	(168,243)	(152,827)
Earnings from continuing operations	\$ 498,557	\$ 894,504	\$ 752,439
Operating results:			
Energy	18.7%	22.9%	24.7%
Engineering Systems	16.1%	16.2%	16.0%
Platts	18.1%	17.6%	17.6%
Refining and Food Equipment	12.8%	12.4%	14.2%
Inter-segment transactions	0.0%	0.0%	0.0%
Earnings from continuing operations	8.0%	10.0%	11.1%
Dispositions and settlements			
Energy	\$ 141,719	\$ 113,956	\$ 90,875
Engineering Systems	39,514	41,416	39,051
Platts	56,678	60,961	56,678
Refining and Food Equipment	66,278	60,761	32,217
Inter-segment transactions	0	0	0
Goodwill impairment	\$ 137,000	\$ 166,716	\$ 250,831
Capital expenditures:			
Energy	\$ 33,662	\$ 26,749	\$ 26,744
Engineering Systems	37,509	26,268	26,268
Platts	44,842	32,149	32,149
Refining and Food Equipment	33,511	33,010	27,377
Goodwill impairment	\$ 154,251	\$ 166,673	\$ 149,484

[illegible]

	Years Ended December 31,			
	2013	2012	2011	2010
Revenues:				
Energy	\$ 1,226,654	\$ 2,172,684	\$ 1,980,199	\$ 1,865,199
Engineering Systems	2,796,330	3,619,244	3,961,241	3,161,861
Printing & Identification	77,217	99,631	130,516	127,141
Information Technology	1,611,787	1,966,238	1,966,238	1,966,238
Other segment divisions	1,411,287	1,411,287	1,411,287	1,411,287
Unallocated revenue	11,392	192	192	192
Total revenues	\$ 5,044,667	\$ 9,270,786	\$ 9,451,273	\$ 8,532,818
Expenses from continuing operations:				
Segment revenues				
Energy	\$ 552,800	\$ 376,806	\$ 470,637	\$ 470,637
Engineering Systems	373,088	507,062	440,186	366,186
Information Technology	152,418	331,139	331,139	331,139
Communications Technologies	237,609	238,609	238,362	238,362
Other segment divisions	1,318,972	1,294,772	1,294,772	1,294,772
Corporate expense - other (1)	108,681	136,009	136,009	136,009
Unallocated expense	120,215	120,215	120,215	120,215
Expenses before provision for income taxes and discontinued operations	\$ 2,572,812	\$ 1,814,613	\$ 1,895,295	\$ 1,652,861
Provision for income taxes	271,467	464,425	464,425	464,425
Expenses from continuing operations	\$ 1,605,850	\$ 1,313,114	\$ 1,313,114	\$ 1,188,436
Operating revenues:				
Energy	24.1%	24.8%	24.9%	25.9%
Engineering Systems	12.2%	14.1%	14.6%	11.9%
Information Technology	14.9%	13.6%	13.6%	13.6%
Communications Technologies	14.7%	14.4%	14.4%	16.6%
Other segment divisions	17.4%	17.2%	17.2%	17.2%
Unallocated revenue	0.2%	0.2%	0.2%	0.2%
Expenses from continuing operations	11.1%	10.3%	10.3%	9.7%
Dispositions and amortization:				
Energy				
Engineering Systems	\$ 107,344	\$ 491,871	\$ 77,819	\$ 77,819
Information Technology	131,354	158,548	158,548	158,548
Communications Technologies	36,792	31,602	31,602	31,602
Other segment divisions	148,475	132,619	130,109	130,109
Unallocated expense	3,651	3,651	3,651	3,651
Dispositions and amortization	\$ 427,616	\$ 818,191	\$ 391,739	\$ 391,739
Comprehended total	\$ 812,466	\$ 1,131,305	\$ 1,204,874	\$ 1,204,874
Capital expenditures:				
Energy	\$ 47,654	\$ 70,334	\$ 74,951	\$ 74,951
Engineering Systems	58,637	60,618	60,618	60,618
Information Technology	8,864	6,251	6,251	6,251
Communications Technologies	99,124	152,245	151,141	151,141
Other segment divisions	2,568	2,568	2,568	2,568
Corporate	1,000	1,000	1,000	1,000
Capital expenditures	\$ 168,837	\$ 253,526	\$ 253,526	\$ 253,526

[illegible]

Table 10
 $\alpha = 0.05$, $\beta = 0.01$, and $\gamma = 0.9$ Year End December 31, 1994, 1995, 1996, 1997, 1998, 1999, 2000, 2001, 2002, 2003, 2004, 2005, 2006, 2007, 2008, 2009, 2010, 2011, 2012, 2013, 2014, 2015, 2016, 2017, 2018, 2019, 2020, 2021, 2022, 2023, 2024, 2025, 2026, 2027, 2028, 2029, 2030, 2031, 2032, 2033, 2034, 2035, 2036, 2037, 2038, 2039, 2040, 2041, 2042, 2043, 2044, 2045, 2046, 2047, 2048, 2049, 2050, 2051, 2052, 2053, 2054, 2055, 2056, 2057, 2058, 2059, 2060, 2061, 2062, 2063, 2064, 2065, 2066, 2067, 2068, 2069, 2070, 2071, 2072, 2073, 2074, 2075, 2076, 2077, 2078, 2079, 2080, 2081, 2082, 2083, 2084, 2085, 2086, 2087, 2088, 2089, 2090, 2091, 2092, 2093, 2094, 2095, 2096, 2097, 2098, 2099, 2100, 2101, 2102, 2103, 2104, 2105, 2106, 2107, 2108, 2109, 2110, 2111, 2112, 2113, 2114, 2115, 2116, 2117, 2118, 2119, 2120, 2121, 2122, 2123, 2124, 2125, 2126, 2127, 2128, 2129, 2130, 2131, 2132, 2133, 2134, 2135, 2136, 2137, 2138, 2139, 2140, 2141, 2142, 2143, 2144, 2145, 2146, 2147, 2148, 2149, 2150, 2151, 2152, 2153, 2154, 2155, 2156, 2157, 2158, 2159, 2160, 2161, 2162, 2163, 2164, 2165, 2166, 2167, 2168, 2169, 2170, 2171, 2172, 2173, 2174, 2175, 2176, 2177, 2178, 2179, 2180, 2181, 2182, 2183, 2184, 2185, 2186, 2187, 2188, 2189, 2190, 2191, 2192, 2193, 2194, 2195, 2196, 2197, 2198, 2199, 2200, 2201, 2202, 2203, 2204, 2205, 2206, 2207, 2208, 2209, 2210, 2211, 2212, 2213, 2214, 2215, 2216, 2217, 2218, 2219, 2220, 2221, 2222, 2223, 2224, 2225, 2226, 2227, 2228, 2229, 2230, 2231, 2232, 2233, 2234, 2235, 2236, 2237, 2238, 2239, 2240, 2241, 2242, 2243, 2244, 2245, 2246, 2247, 2248, 2249, 2250, 2251, 2252, 2253, 2254, 2255, 2256, 2257, 2258, 2259, 2260, 2261, 2262, 2263, 2264, 2265, 2266, 2267, 2268, 2269, 2270, 2271, 2272, 2273, 2274, 2275, 2276, 2277, 2278, 2279, 2280, 2281, 2282, 2283, 2284, 2285, 2286, 2287, 2288, 2289, 2290, 2291, 2292, 2293, 2294, 2295, 2296, 2297, 2298, 2299, 2300, 2301, 2302, 2303, 2304, 2305, 2306, 2307, 2308, 2309, 2310, 2311, 2312, 2313, 2314, 2315, 2316, 2317, 2318, 2319, 2320, 2321, 2322, 2323, 2324, 2325, 2326, 2327, 2328, 2329, 2330, 2331, 2332, 2333, 2334, 2335, 2336, 2337, 2338, 2339, 2340, 2341, 2342, 2343, 2344, 2345, 2346, 2347, 2348, 2349, 2350, 2351, 2352, 2353, 2354, 2355, 2356, 2357, 2358, 2359, 2360, 2361, 2362, 2363, 2364, 2365, 2366, 2367, 2368, 2369, 2370, 2371, 2372, 2373, 2374, 2375, 2376, 2377, 2378, 2379, 2380, 2381, 2382, 2383, 2384, 2385, 2386, 2387, 2388, 2389, 2390, 2391, 2392, 2393, 2394, 2395, 2396, 2397, 2398, 2399, 2400, 2401, 2402, 2403, 2404, 2405, 2406, 2407, 2408, 2409, 2410, 2411, 2412, 2413, 2414, 2415, 2416, 2417, 2418, 2419, 2420, 2421, 2422, 2423, 2424, 2425, 2426, 2427, 2428, 2429, 2430, 2431, 2432, 2433, 2434, 2435, 2436, 2437, 2438, 2439, 2440, 2441, 2442, 2443, 2444, 2445, 2446, 2447, 2448, 2449, 2450, 2451, 2452, 2453, 2454, 2455, 2456, 2457, 2458, 2459, 2460, 2461, 2462, 2463, 2464, 2465, 2466, 2467, 2468, 2469, 2470, 2471, 2472, 2473, 2474, 2475, 2476, 2477, 2478, 2479, 2480, 2481, 2482, 2483, 2484, 2485, 2486, 2487, 2488, 2489, 2490, 2491, 2492, 2493, 2494, 2495, 2496, 2497, 2498, 2499, 2500, 2501, 2502, 2503, 2504, 2505, 2506, 2507, 2508, 2509, 2510, 2511, 2512, 2513, 2514, 2515, 2516, 2517, 2518, 2519, 2520, 2521, 2522, 2523, 2524, 2525, 2526, 2527, 2528, 2529, 2530, 2531, 2532, 2533, 2534, 2535, 2536, 2537, 2538, 2539, 2540, 2541, 2542, 2543, 2544, 2545, 2546, 2547, 2548, 2549, 2550, 2551, 2552, 2553, 2554, 2555, 2556, 2557, 2558, 2559, 2560, 2561, 2562, 2563, 2564, 2565, 2566, 2567, 2568, 2569, 2570, 2571, 2572, 2573, 2574, 2575, 2576, 2577, 2578, 2579, 2580, 2581, 2582, 2583, 2584, 2585, 2586, 2587, 2588, 2589, 2590, 2591, 2592, 2593, 2594, 2595, 2596, 2597, 2598, 2599, 2600, 2601, 2602, 2603, 2604, 2605, 2606, 2607, 2608, 2609, 2610, 2611, 2612, 2613, 2614, 2615, 2616, 2617, 2618, 2619, 2620, 2621, 2622, 2623, 2624, 2625, 2626, 2627, 2628, 2629, 2630, 2631, 2632, 2633, 2634, 2635, 2636, 2637, 2638, 2639, 2640, 2641, 2642, 2643, 2644, 2645, 2646, 2647, 2648, 2649, 2650, 2651, 2652, 2653, 2654, 2655, 2656, 2657, 2658, 2659, 2660, 2661, 2662, 2663, 2664, 2665, 2666, 26

[illegible]

a) Meaningless Output b) Incomplete Output c) Hallucination

Figure 8: Some failed recognition samples correspond to the three main reasons mentioned above, including meaningless output, incomplete output, language preference, etc.

Models	Train	Test	TEDS	TEDS-Struct	Acc	EDSC	Efficiency
Qwen2-VL-7B-Instruct	TableOCR	PubTabNet	61.37%	74.75%	23.05%	42.39%	7.39
	TableOCR	FinTabNet	59.51%	68.64%	14.57%	43.89%	7.17
	TableOCR	SciTSR	88.97%	92.45%	64.47%	83.66%	10.72
	TableOCR	TableOCR	93.67%	96.26%	86.03%	93.25%	11.29
LLaVA1.5-7B-Instruct	TableOCR	PubTabNet	40.58%	44.75%	1.23%	23.48%	5.72
	TableOCR	FinTabNet	13.07%	37.30%	0.82%	17.07%	1.84
	TableOCR	SciTSR	24.86%	56.02%	4.53%	17.61%	3.50
	TableOCR	TableOCR	39.21%	72.47%	20.13%	31.17%	5.52
LLaVA1.5-13B-Instruct	TableOCR	PubTabNet	40.63%	47.60%	1.39%	24.32%	3.03
	TableOCR	FinTabNet	14.22%	38.67%	0.76%	17.85%	1.06
	TableOCR	SciTSR	26.57%	57.66%	6.00%	19.16%	1.98
	TableOCR	TableOCR	42.72%	76.04%	24.59%	33.21%	3.19
MiniCPM-V-V2.5	TableOCR	PubTabNet	56.79%	67.90%	16.51%	38.70%	6.68
	TableOCR	FinTabNet	42.01%	55.77%	7.84%	33.75%	4.94
	TableOCR	SciTSR	69.75%	79.35%	37.50%	62.24%	8.21
	TableOCR	TableOCR	87.34%	92.66%	72.69%	86.09%	10.28
TableVLM-4B	TableOCR	PubTabNet	61.94%	70.82%	26.08%	48.96%	14.75
	TableOCR	FinTabNet	41.42%	48.33%	7.93%	34.92%	9.86
	TableOCR	SciTSR	67.00%	73.19%	41.33%	61.53%	15.95
	TableOCR	TableOCR	88.22%	95.55%	82.31%	86.97%	21.00
TableVLM-14B	TableOCR	PubTabNet	74.28%	90.60%	51.60%	66.72%	5.34
	TableOCR	FinTabNet	63.76%	70.75%	20.02%	52.89%	4.59
	TableOCR	SciTSR	83.53%	87.28%	59.33%	78.83%	6.01
	TableOCR	TableOCR	97.70%	99.24%	95.56%	97.44%	7.03

Table 10: The comparison of the generalization capabilities of TableVLM and existing multi-modal large language models after fine-tuning on TableOCR.