

使用可上雲的檔案生成更多問題解答對，並用生成的資料微調模型

Step1: 生成資料

Step2: 調整資料格式

Step3: 微調模型

Step1: 資料生成

- 使用gemini-1.5-flash生成共6413筆資料
- 參考文檔:
 - 中央銀行_申報系統表單與細項清單2023
 - 本國銀行財報類檔案目錄暨各申報截止日
 - Mantis_AI問題分析測試案例(所有試算表)

產生的6413筆資料形式皆為以下格式:

“instruction”:生成的問題,“input”:(全為空白),“output”:生成的解答

```
{"instruction": "CDIC008 債券名稱可否簡寫不帶全名?", "input": "", "output": "債券名稱為自動帶出行內提供的債票名稱欄位資訊。"}  
{"instruction": "CDIC008 到期日可否直接填值不帶公式?", "input": "", "output": "製表說明註明無到期日證券不需填寫年期及到期日欄，並以「--」表示。"}  
{"instruction": "001報表沒有抓到證券金融公司發行的票債券，請問如何解決?", "input": "", "output": "請與前端資料源系統確認發行機構類別是否正確。"}
```

Step2:調整資料格式

原本的格式：

```
{
  "instruction": "CDIC008 債券名稱可否簡寫不帶全名？",
  "input": "",
  "output": "債券名稱為自動帶出行內提供的債票名稱欄位資訊。"
},
{
  "instruction": "CDIC008 到期日可否直接填值不帶公式？",
  "input": "",
  "output": "製表說明註明無到期日證券不需填寫年期及到期日欄，並以「--」表示。"
},
{
  "instruction": "001報表沒有抓到證券金融公司發行的票債券，請問如何解決？",
  "input": "",
  "output": "請與前端資料源系統確認發行機構類別是否正確。"
}
```

轉換後的格式：

```
{
  "converted": "Human: CDIC008 債券名稱可否簡寫不帶全名? \n/s>>Assistant: 債券名稱為自動帶出行內提供的債票名稱欄位資訊。 \n/s>"}
{
  "converted": "Human: CDIC008 到期日可否直接填值不帶公式? \n/s>>Assistant: 製表說明註明無到期日證券不需填寫年期及到期日欄，並以「-」表示。 \n/s>"}
{
  "converted": "Human: 001報表沒有抓到證券金融公司發行的票債券，請問如何解決? \n/s>>Assistant: 請與前端資料源系統確認發行機構類別是否正確。 \n/s>"}
{
  "converted": "Human: CD_Rpt_Typ2BuzCd這張設定檔是在什麼情境下使用? \n/s>>Assistant: 單表有金融機構廣義、狹義定義衝突時，或報表有特別定義金融機構範圍時會使用。 \n/s>"}
}
```

將生成的6413筆資料轉換為以下的格式:

“converted”:<s>Human: “原本資料的instruction”

Assistant: “原本資料的output”

```
#SFT Trainer
trainer = SFTTrainer(
    model = model,
    train_dataset = dataset,
    peft_config = peft_config,
    dataset_text_field = "converted",
    max_seq_length = max_seq_length,
    args = training_arguments,
    tokenizer = tokenizer,
    packing = packing,
)
```

Step3: 微調模型

使用 **SFTTrainer + QLoRA** 來訓練adapter，最後再與base model的權重合併

使用的base model為Llama3-Chinese-8B-Instruct

- SFTTrainer
 - 用於執行Supervised Fine-Tuning(SFT)的訓練任務
 - 主要功能是簡化微調的過程，特別是對已經預訓練的模型進行監督式的微調
 - 與PEFT(Parameter Efficient Fine-Tuning)方法兼容
 - LoRA
 - QLoRA

Step3: 微調模型

使用QLoRA微調模型

- 能將預訓練模型量化為4-bit節省計算消耗

以下為QLoRA的參數設定：

```
[7] lora_r = 64 #lora attention dimension/ rank  
    lora_alpha = 16 #lora scaling parameter  
    lora_dropout = 0.1 #lora dropout probability
```

```
[8] use_4bit = True  
    bnb_4bit_compute_dtype = "float16"  
    bnb_4bit_quant_type = "nf4"  
    use_nested_quant = False
```

Step3: 微調模型

(1) 使用SFTTrainer + QLoRA來訓練adapter

(2) 與base model的權重合併

```
base_model = AutoModelForCausalLM.from_pretrained(  
    model_name,  
    low_cpu_mem_usage=True,  
    return_dict=True,  
    torch_dtype=torch.float16,  
    device_map=device_map,  
)  
model = PeftModel.from_pretrained(base_model, new_model)  
model = model.merge_and_unload()  
  
model.save_pretrained(merge_model)
```