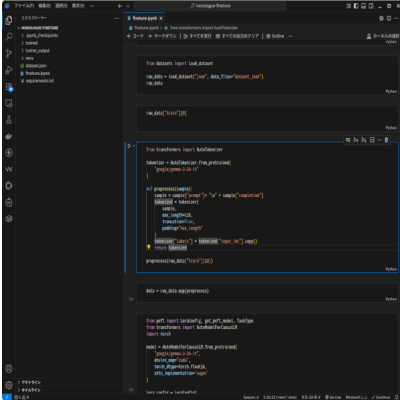




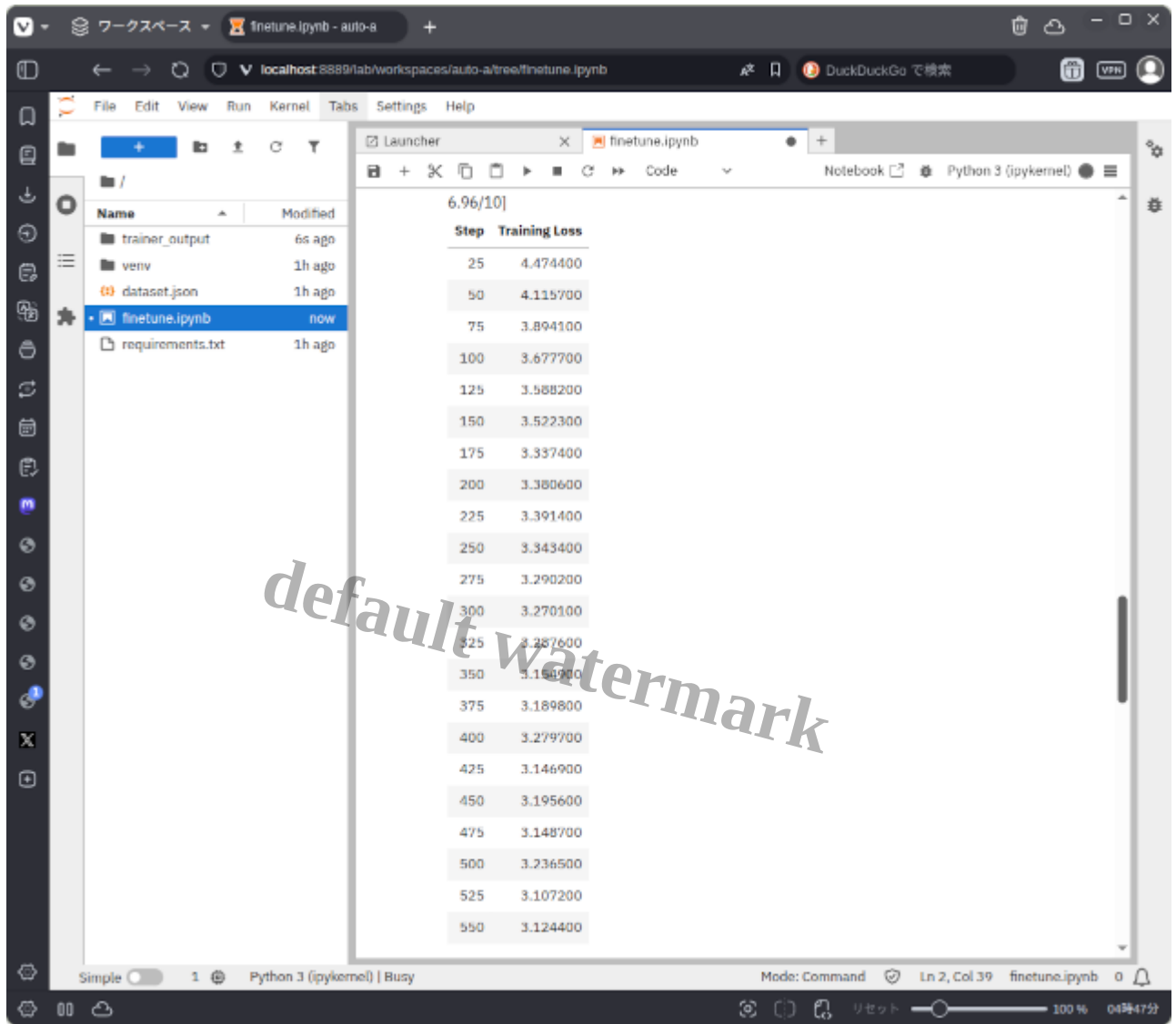
ç?¥è?è??ç©•ã?»å ±æ??ã•®ç ?ç©¶ã??ã•ã•®i¼? LLMã??ã?|ã?▯ã?³ã?•  
ã?¥ã?¼ã??ã?³ã?°ã•®è©|ã•¿

## Description

ǎ?æ?¥ǎ?•ǎ?ǎ•®ǎ??ǎ?ǎ?°ǎ??ǎ?¬ǎ?ǎ?¼ǎ?ªǎ?ǎ?°ǎ?ǎ;æ??ç¨¿ǎ??æ©?æ¢°ǎ|ç¿?ç?“ǎ•  
®ǎ??ǎ?¼ǎ?¿ǎ?»ǎ??ǎ??ǎ«ǎ□?æ•?ǎ•?ǎ??ǎ??ǎ?ǎ?°ǎ?©ǎ?ǎ??æ?ǎ•ǎ•¾ǎ?ǎ•?ǎ??ǎ»?ǎ??ǎ¬ǎ•ǎ•  
®ç¶||ǎ•ǎ•§ǎ?ǎ•ǎ•®ǎ??ǎ?¼ǎ?¿ǎ?»ǎ??ǎ??ǎ??ǎ½¿ǎ•£ǎ•?æ?¢ǎ?ǎ•®ǎ-  
|ç¿?æ,?è¨?èª?ǎ?¢ǎ??ǎ?«i¼?Pretreined Language Modeli¼?ǎ??Loraǎ?ǎ?¥ǎ?¼ǎ??ǎ?³ǎ?°ǎ•  
?ǎ??ǎ??ǎ?ǎ?°ǎ?©ǎ?ǎ??æ?ǎ•?ǎ;ǎ•¿ǎ•¾ǎ?ǎ•?ǎ??i¼?ǎ,?ǎ•®ǎ?¹ǎ¬ǎªǎ?¼ǎ?³ǎ?ǎ?§ǎ??ǎ??i¼?

ä½ç?ä•?ä•?ä³ä³ä??ä?¥ä?¼ä?¿ä?¼ä•Core i5 13500 CPUä•NVIDIA RTX A4000 (16GB VRAN)ä??æ•è¼ä•?ä•?ä?ä•®ä??ä½±jã•ä•?ä•?è•èªä?¢ä??ä?«ä•Googleä•®Gemma 3 1Bä•®ä?ªä³ä?¹ä??ä?©ä?¬ä?ä?§ä³ä³ä?¥ä?¼ä??ä³ä?°æ•?ä•®ä??ä•®ä•Sä•?ä??

çµ•æ??ã•¯â±æ??ã??â|ç¿?ã?ã?¹ã?æ,?ã??ã•?â®?ç?`ç??ã•ã??ã•®ã•«ã•¯ã•ã??ã•¼ã•?ã??ã•§ã•?ã•  
?ã??



ã•?ã•®ç`?å¡|ã•®å?|ç•?ã•§ã•?ã•`ã•?ã•³ã•³ã•??ã•?¥ã•¹¼ã•?¿ã•¹¼ã•®å?|ç•?é???å¡|ã•«å¡§ã•-ã•ä, •æ°?ã•-ã•ã•  
 ?ã•?ã•®ã•®ã•?VRAMã•-ã•?£ã•-æ-²ã•?ã•?ã•§ã•?ã•?ã•?ã•®ã•³ã•³ã•??ã•?¥ã•¹¼ã•?¿ã•¹¼ã•®æ§?æ?ã•  
 §ã•\_1Bã•®ã•?£ã•?ã•?«ã•?ã•?ã•?¼ã•?-ã•³é•.128ã•ã•?ã•?ã•?é?•ç•?ã•§ã•?14GBã•»ã•©ã•  
 ®VRAMã•?ã•½¿ã•£ã•!ã•?ã•¾ã•?ã•?ã•?

ã?ã?ã??ã?¼ã??ã?©ã?|ã?¼ã?¿ã?¼ã•\_ã½?â°|ã??â¤?æ?´ã?ã•|è©|ã•?ã¾ã•?ã•?ã•?ã?è?-ã?çµ•  
æ??ã•\_ã¾?ã??ã??ã¾ã•?ã??ã•§ã•?ã•?ã??

ã??ã?¼ã?¿ã?»ã??ã??ã?®è³ã?æ?ªã?·?¼?ã??ã?ã?°æ??ç·¿ã·ã·®ã·¾ã·¾ã·¯æ??ã¾ã·?ã·ã·ã·?¼?ã·  
·ã·?ã·?ã·?ã··ã·?ã??ã·?ã??ã·¾ã·?ã??ã??

ǎ?è??ǎ¼ǎ•šă«ǎ?ǎ??ǎ?-ǎ¹¼ǎ??ǎ³ǎ°ǎ³ǎ¹¼ǎ??ǎ•æ-ja®ǎ??ǎ?ǎ•ǎ??ǎ®ǎ•šă?ǎ??

/\* ã??ã?¼ã?¿ã?»ã??ã??ã??èªã•¿è¾¼ã?? \*/

```
from datasets import load_dataset
```

```
raw_data = load_dataset("json", data_files="dataset.json")
```

```
/* ä??ä?¼ä?¬ä??ä?¼ä?¶ä?¼ä?•@æ°?å?? */  
  
from transformers import AutoTokenizer  
  
tokenizer = AutoTokenizer.from_pretrained(  
    "google/gemma-3-1b-it"  
)  
  
/* ä??ä?¼ä?¿ä?»ä??ä??ä?•@å?•å?|ç?• */  
  
def preprocess(sample):  
    sample = sample["prompt"]+ "\n" + sample["completion"]  
    tokenized = tokenizer(  
        sample,  
        max_length=128,  
        truncation=True,  
        padding="max_length"  
    )  
    tokenized["labels"] = tokenized["input_ids"].copy()  
    return tokenized  
  
data = raw_data.map(preprocess)  
  
/* ä?•ä?¥ä?¼ä??ä?³ä?°ä¬¼è±¡ä?•@ä?¢ä??ä?«ä?•@æ°?å?? */  
  
from peft import LoraConfig, get_peft_model, TaskType  
from transformers import AutoModelForCausalLM  
import torch  
  
model = AutoModelForCausalLM.from_pretrained(  
    "google/gemma-3-1b-it",  
    device_map="cuda",  
    torch_dtype=torch.float16,  
    attn_implementation='eager'  
)  
  
/* LORAä?•@è¬å@? */  
lora_config = LoraConfig(  
    task_type=TaskType.CAUSAL_LM,  
    target_modules=["q_proj", "k_proj", "v_proj"]  
)  
  
model = get_peft_model(model, lora_config)  
  
/* ä??ä?¬ä?¼ä??ä?¼ä?•@æ°?å?? */  
  
from transformers import TrainingArguments, Trainer  
  
train_args = TrainingArguments(  
    num_train_epochs=10,  
    learning_rate=3e-4,  
    logging_steps=25,  
    fp16=True  
)
```

```
trainer = Trainer(
    args=train_args,
    model=model,
    train_dataset=data["train"]
)

/* ???-?%???^??°@è;? */
trainer.train()

/* ??-?%???^??°æ_ ?•@?¢???«ã· ``??%?-??-?□?¶?¼ã•@äç•å? */

trainer.save_model("./trained")
tokenizer.save_pretrained("./trained")

!ç¿?aa¹â¬1.0ä»¥ä,?ç?®æ¨?ã§ã?ã?ã?•3.0ã³¼ã§ã?ã?æ_,?ã?ã¾ã?ã??ã§ã?ã?ã??
ä»¥ä,?ã®ã³ã?¼ã??ã?ã??ã?jä²ã³ã?ã?¥ã¼ã??ã³ã?°å¾ã®ã¢ã??ã«ã??½ã£ã•
?æ?“è«ã??èj?ã?ã??ã®ã§ã?ã?
```

## Category

1. æ?¥ã? ã•®æ'»å??

## Tags

1. AI
2. Linux
3. LLM
4. Lora
5. ã??ã?jã?□ã?³ã?•ã?¥ã?¼ã??ã?³ã?°
6. ä,?é??å??å,?
7. å□Sè!•æ"jè"èª?ã?¢ã??ã?«
8. å±±æ¢`ç??

**Date Created**

2025å¹¸æ??3æ?¥

## Author

kazu0-tsubaki