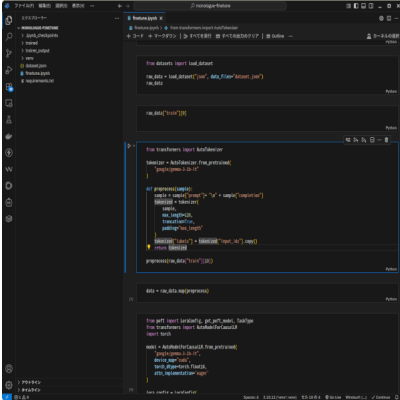




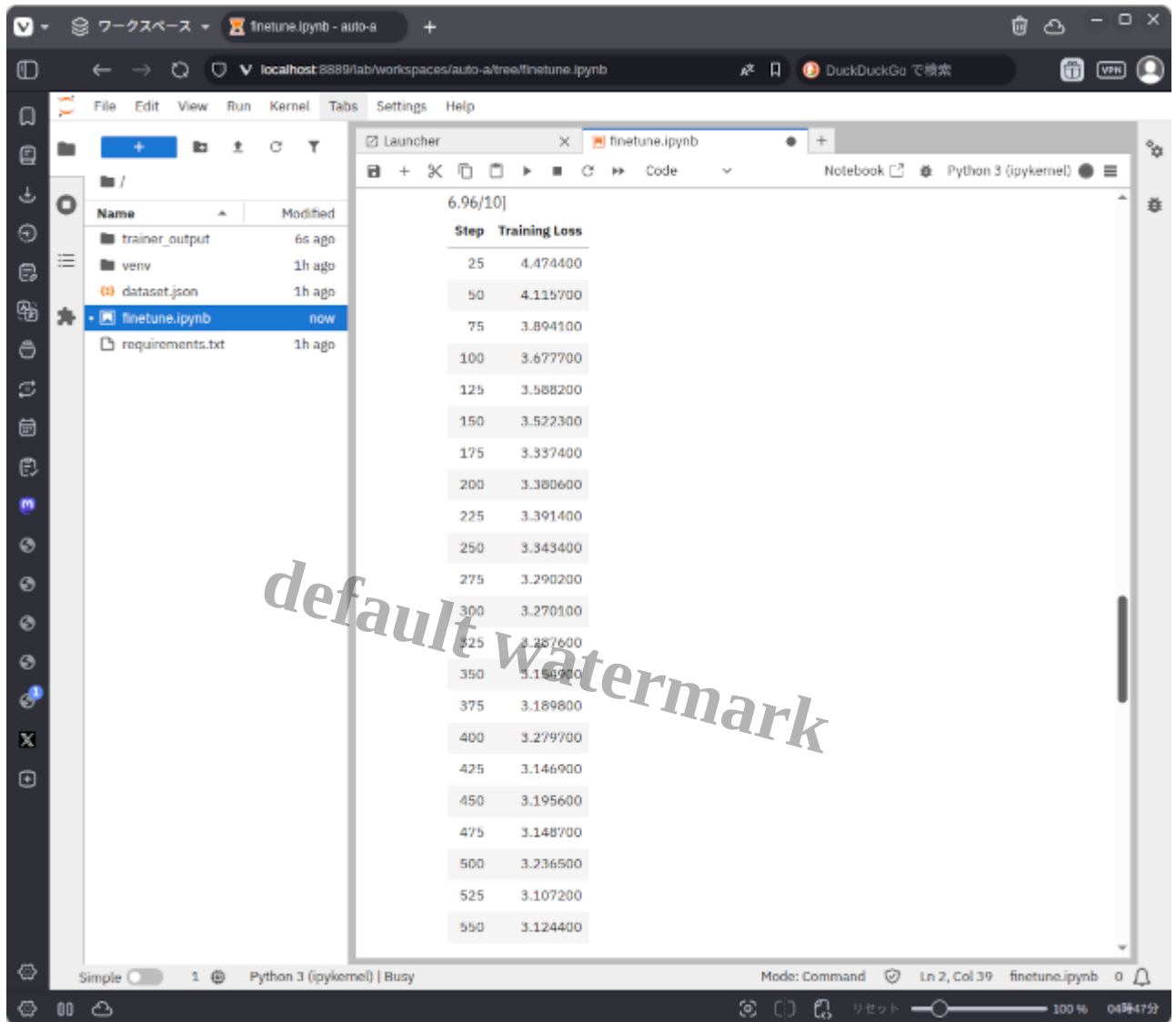
ç?¥è?è??ç©•ã?»±æ??ã•®ç ?ç©¶ã??ã•ã•®i¼? LLMã??ã?jã?ã?³ã?
ã?¥ã?¼ã??ã?³ã?°ã•®è©|ã•¿

Description

å ?æ?¥ã?•ã•®ã??ã?°ã??ã?¬ã?ã?¼ã?³ã?°ã?ã?æ??ç?¿ã??æ©?æ¢°å¿çç?ç?¬ã•
®ã??ã?¼ã?¿ã?»ã??ã??ã•«å±?æ•?ã•?ã??ã??ã?ã?°ã?©ã?ã??æ?ã•ã•¼ã•?ã•?ã??ã?»å??ã¬ã•ã•
®ç¶?ã•ã•§ã?•ã•ã•®ã??ã?¼ã?¿ã?»ã??ã??ã??ã½¿ã•£ã•?æ?¢ã?ã•®å-
¿ç¿æ¿?è?°è?ã?¢ã??ã?«i¼?Pretreined Language Modeli¼?ã??Loraã?•ã?¥ã?¼ã??ã?³ã?°ã•
?ã??ã??ã?ã?ã?©ã?ã??æ?ã•?ã?ã?¿ã?¿ã•¼ã•?ã•?ã??i¼?ã¿ã•®ã?¹ã¬¬ã?ã?¼ã?³ã?¬ã?§ã??ã??i¼?

ä½¿çç?¬ã?ã?ã?³ã?³ã??ã?¥ã?¼ã?¿ã?¼ã¬Core i5 13500 CPUã¬NVIDIA RTX A4000 (16GB
VRAM)ã??æ•?¼?ã•?ã?ã?ã?ã•®ã??ã¬¼è±jã¬ã?ã?è?°è?ã?¢ã??ã?«ã¬Googleã¬Gemma 3 1Bã•
®ã?ã?³ã?¹ã??ã?©ã¬¬ã?ã?§ã?³ã?ã?¥ã?¼ã??ã?³ã?°æ¿ã•®ã??ã•®ã•§ã?ã??

çµæ??ã¬å±±æ??ã??å¿ç¿ã?ã?¼ã?æ¿ã??ã?å®?çç?ç?ã•ã??ã•®ã¬«ã¬¬ã??ã?¼ã•?ã??ã•§ã?ã•
?ã??



ã•?ã•®ç`?â!ã•®â?!ç•?ã•§ã•?ã•`ã•?ã•³ã•³ã•??ã•¥ã•¹¼ã•?¿ã•¹¼ã•®â?!ç•?é??â!ã•«â¡§ã•-ä•ä•æ°?ã•-ã•ã•
?ã•?ã•®ã•®ã••VRAMã•-ã•?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•¹¼ã•?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?¼ã?¿ã?¼ã?_ã½?ã°|ã??ã?æ?_ã?ã?è©|ã?ã?¾ã?ã?ã?ã?ã?è?_ã?çµ•
æ??ã?_ã¾?ã??ã??ã¾ã?ã??ã?§ã?ã?ã??

ã??ã?¼ã?¿ã?»ã??ã??ã?®è³ªã?·æ?ªã?·?¼?ã??ã?ã?°æ??¿·¿ã·ã·®ã·¾ã·¾ã·¯æ??ã¾ã·ã·ã·ªã·?¿¼?ã·
 ¨ã·ã·ã·ã·¨ã·ã??ã·ã??ã·¾ã·?ã??ã??

ǎ•?è??ǎ•¼ǎ•šǎ•«ǎ•?ǎ•?ǎ•-ǎ•¼ǎ•?ǎ•?ǎ•°ǎ•?ǎ•¼ǎ•?ǎ•-æ-ǎ•ǎ•ǎ•?ǎ•?ǎ•ǎ•?ǎ•ǎ•šǎ•?ǎ•?

/* ã??ã?¼ã?¿ã?»ã??ã??ã??è^aã•¿:è^{¾¼}ã?? */

```
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?°å¬¼è±¼ä?•@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"]
)

/* ă??ă?-ă?¼ă??ă?³ă?°ă@?èj? */
trainer.train()

/* ă??ă?-ă?¼ă??ă?³ă?°æ, ?ă•@ă?¢ă??ă?«ă•"ă??ă?¼ă?¬ă??ă?¤ă?¶ă?¼ă•@ă¿•ă? */

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

ă!¿¿?ă?ă?¹ă¬1.0ă»¥ă, ?ă?¿?@æ" ?ă•$ă?ă?ă?•3.0ă•¼ă•$ă?ă?ă, ?ă??ă?¼ă•?ă??ă•$ă?ă?ă??
ă»¥ă, ?ă•@ă?³ă?¼ă??ă?ă?ă?ă?jă?¤ă?³ă?•ă?¥ă?¼ă??ă?³ă?°ă¼?ă•@ă?¢ă??ă?«ă??ă½¿ă•£ă•
?æ?"è«?ă??èj?ă?ă??ă•@ă•$ă?ă??

from transformers import pipeline

ask_llm = pipeline(
    model="./trained",
    tokenizer="./trained",
    device="cuda"
)

print(ask_llm("ă??ă?ă?³ă??ă??") [0] ["generated_text"])

```

Category

1. æ?¥ã? ã•®æ'»å??

Tags

1. AI
2. Linux
3. LLM
4. Lora
5. ã??ã?;ã?♠ã?³ã?•ã?¥ã?¼ã??ã?³ã?°
6. ä_?é??å??å_?
7. å♠§è!_œ``jè``èª?ã?¢ã??ã?«
8. å±±±æ¢`ç??

Date Created

2025å¹¸æ??3æ?¥

Author

kazuotsubaki