

OVERVIEW & SETUP

Multimodal AI AI/ML framework  
pip install multimodal-ai  
# Verify installation  
python -c "import multimodal\_ai;  
print('OK')"  
# GPU support check  
import torch;  
print(torch.cuda.is\_available()),  
Use virtual environment: python -m  
venv venv

DATA PREPARATION

import pandas as pd  
df = pd.read\_csv('data.csv')  
X = df.drop('target', axis=1).values  
y = df['target'].values  
from sklearn.model\_selection import  
train\_test\_split  
X\_train, X\_test, y\_train, y\_test =  
train\_test\_split(X, y, test\_size=0.2)  
Normalize/standardize features for  
best results

MODEL DEFINITION

# Simple neural network:  
import torch.nn as nn  
class Model(nn.Module):  
def \_\_init\_\_(self):  
super().\_\_init\_\_()  
self.layers = nn.Sequential(  
nn.Linear(input\_size, 128),  
nn.ReLU(),  
nn.Linear(128, num\_classes)  
)  
def forward(self, x): return  
self.layers(x)

EVALUATION

from sklearn.metrics import  
accuracy\_score, classification\_report  
model.eval()  
with torch.no\_grad():  
preds = model(X\_test)  
accuracy = accuracy\_score(y\_test,  
preds.argmax(1))  
print(classification\_report(y\_test,  
preds.argmax(1)))  
Track precision, recall, F1; use  
confusion matrix

INFERENCE & SERVING

# Save model  
torch.save(model.state\_dict(),  
'model.pt')  
# Load model  
model.load\_state\_dict(torch.load('model.pt'))  
model.eval()  
# Inference  
with torch.no\_grad():  
prediction = model(input\_tensor)

HYPERPARAMETERS

Learning rate: start with 1e-3, tune  
with scheduler  
Batch size: 32-256 depending on GPU  
memory  
Epochs: use early stopping to  
prevent overfitting  
scheduler =  
torch.optim.lr\_scheduler.ReduceLROnPlateau(optimizer)  
scheduler.step(val\_loss)  
Grid search or Optuna for systematic  
tuning

DATASETS & LOADERS

from torch.utils.data import Dataset,  
DataLoader  
class MyDataset(Dataset):  
def \_\_len\_\_(self): return len(self.data)  
def \_\_getitem\_\_(self, idx): return  
self.data[idx]  
loader = DataLoader(dataset,  
batch\_size=32, shuffle=True,  
num\_workers=4)  
Pin memory for faster GPU data  
transfer

OPTIMIZATION TECHNIQUES

Mixed precision:  
torch.cuda.amp.autocast()",  
Gradient clipping:  
clip\_grad\_norm\_(model.parameters(),  
1.0)",  
Batch normalization:  
nn.BatchNorm1d()",  
Dropout: nn.Dropout(0.5) for  
regularization",  
# Learning rate warmup  
from transformers import  
get\_linear\_schedule\_with\_warmup

FRAMEWORKS & TOOLS

|                    |                |
|--------------------|----------------|
| TensorFlow / Keras | production ML  |
| PyTorch            | research & NLP |
| Scikit-learn       | classical ML   |
| Hugging Face       | pretrained     |
| transformers       |                |
| MLflow / W&B       | experiment     |
| tracking           |                |
| ONNX               | model          |
| export/portability |                |

DEPLOYMENT

# Export to ONNX  
torch.onnx.export(model, dummy\_input,  
'model.onnx')  
# FastAPI serving example  
@app.post('/predict')  
def predict(data: InputData):  
tensor = preprocess(data)  
return model(tensor).tolist()  
Use TorchServe, BentoML, or Triton  
for production

COMMON GOTCHAS

Gradient explosion: use gradient  
clipping",  
Memory leak: always zero\_grad()  
before backward()",  
Train/eval mode: model.train() vs  
model.eval()",  
GPU-CPU mismatch: ensure all tensors  
on same device",  
Data leakage: split before  
normalization, not after",