

## Introduction

In HW9, I aim to predict sentiment of text using “Financial Sentiment Analysis” dataset. For easier implementation, I employ GRU in Pytorch, torch.nn.GRU.

## Data Preparation

For the data preprocessing, I follow the homework instruction’s tokenization and word embedding. More importantly, I run the experiments respectively to check which one is better among ‘word level’ and ‘subword level tokenization’.

## Training

### Data Loading

I split the train-test into 80:20. Also, I made a one-hot vector of 3 dimensions corresponding to each emotion.

```
1
2 # Split data. Using Word-level embedding or subword-level embedding is user's choice.
3
4 # X_train, X_test, y_train, y_test = train_test_split(word_embeddings, sentiments, test_size=0.2,
5               random_state=42)
6 X_train, X_test, y_train, y_test = train_test_split(subword_embeddings, sentiments, test_size=0.2)
7
8 # Convert sentiments to one-hot vectors
9 le = LabelEncoder()
10 y_train_encoded = le.fit_transform(y_train)
11 y_train_one_hot = torch.nn.functional.one_hot(torch.tensor(y_train_encoded), num_classes=3)
12
13 y_test_encoded = le.transform(y_test)
14 y_test_one_hot = torch.nn.functional.one_hot(torch.tensor(y_test_encoded), num_classes=3)
15
16 class SentimentDataset(Dataset):
17     def __init__(self, embeddings, sentiments):
18         self.embeddings = embeddings
19         self.sentiments = sentiments
20
21     def __len__(self):
22         return len(self.embeddings)
23
24     def __getitem__(self, idx):
25         embedding = self.embeddings[idx].squeeze(0)
26         return embedding, self.sentiments[idx]
27
28 train_dataset = SentimentDataset(X_train, y_train_one_hot)
29 test_dataset = SentimentDataset(X_test, y_test_one_hot)
30
31 train_loader = DataLoader(train_dataset, batch_size=128, shuffle=True)
32 test_loader = DataLoader(test_dataset, batch_size=128, shuffle=False)
```

## Network

My architecture uses a GRU (Gated Recurrent Unit) layer to process the input text sequence, capturing sequential dependencies. The GRU can be unidirectional or bidirectional, depending on the `bidirectional` parameter. The last hidden state of the GRU is then passed through a fully connected layer with ReLU activation to obtain the output logits. Finally, a softmax function is applied to obtain the predicted probabilities for each sentiment class. This architecture is suitable for sentiment analysis tasks where the goal is to classify the sentiment of a given text into predefined categories.

```
1 class MyGRU(nn.Module):
2     def __init__(self, input_dim, hidden_dim, output_dim, num_layers, bidirectional):
3         super(MyGRU, self).__init__()
4         self.gru = nn.GRU(input_dim, hidden_dim, num_layers, batch_first=True, bidirectional=
        bidirectional)
5         self.fc = nn.Linear(hidden_dim * (2 if bidirectional else 1), output_dim)
6         self.relu = nn.ReLU()
7         self.softmax = nn.Softmax(dim=1)
8
9     def forward(self, x):
10        out, hidden = self.gru(x)
11
12        out = self.fc(self.relu(out[:,-1]))
13        return self.softmax(out)
```

## Training

I'm optimizing unidirectional model and bidirectional model simultaneously, using different optimizer. Each model consists of 256 hidden dimension with two layers. Instead of NLL loss with LogSoftmax, I used cross-entropy loss with softmax, which is equivalent each other. I run the training 100 epochs without validation, and use the last model for the evaluation.

```
1
2
3 def train(model, train_loader, optimizer, criterion, device):
4     model.train()
5     total_loss = 0
6     for inputs, labels in train_loader:
7         inputs, labels = inputs.to(device), labels.to(device)
8         optimizer.zero_grad()
9         outputs = model(inputs.float())
10
11         loss = criterion(outputs, torch.argmax(labels, 1))
12
13         loss.backward()
14         optimizer.step()
15         total_loss += loss.item()
16     return total_loss / len(train_loader)
17 device = torch.device("cuda" if torch.cuda.is_available() else "cpu")
18
19 unidirectional_model = MyGRU(input_dim=768, hidden_dim=256, output_dim=3, num_layers=2,
        bidirectional=False).to(device)
20 bidirectional_model = MyGRU(input_dim=768, hidden_dim=256, output_dim=3, num_layers=2,
        bidirectional=True).to(device)
21
22 # optimizer_uni = optim.SGD(unidirectional_model.parameters(), lr=0.0001, momentum=0.9)
23 # optimizer_bi = optim.SGD(bidirectional_model.parameters(), lr=0.0001, momentum=0.9)
24 optimizer_uni = optim.Adam(unidirectional_model.parameters(), lr=0.001)
25 optimizer_bi = optim.Adam(bidirectional_model.parameters(), lr=0.001)
26 criterion = nn.CrossEntropyLoss()
27
28 num_epochs = 100
29
```

```

30 uni_losses = []
31 bi_losses = []
32
33 for epoch in range(num_epochs):
34     uni_loss = train(unidirectional_model, train_loader, optimizer_uni, criterion, device)
35     uni_losses.append(uni_loss)
36
37     bi_loss = train(bidirectional_model, train_loader, optimizer_bi, criterion, device)
38     bi_losses.append(bi_loss)
39
40     print(f'Epoch {epoch+1}/{num_epochs}, Unidirectional Loss: {uni_loss:.4f}, Bidirectional Loss:
        {bi_loss:.4f}')

```

## Evaluation

```

1
2
3 def evaluate(model, test_loader, criterion, device):
4     model.eval()
5     total_loss = 0
6     correct = 0
7     total = 0
8     all_predicted = []
9     all_labels = []
10    with torch.no_grad():
11        for inputs, labels in test_loader:
12            inputs, labels = inputs.to(device), labels.to(device)
13            outputs = model(inputs.float())
14            loss = criterion(outputs, torch.argmax(labels, 1))
15            total_loss += loss.item()
16            _, predicted = torch.max(outputs.data, 1)
17            _, labels = torch.max(labels.data, 1)
18            total += labels.size(0)
19            correct += (predicted == labels).sum().item()
20            all_predicted.extend(predicted.cpu().numpy())
21            all_labels.extend(labels.cpu().numpy())
22    return total_loss / len(test_loader), correct / total, confusion_matrix(all_labels,
        all_predicted)
23
24 uni_test_loss, uni_accuracy, uni_conf_matrix = evaluate(unidirectional_model, test_loader,
        criterion, device)
25 bi_test_loss, bi_accuracy, bi_conf_matrix = evaluate(bidirectional_model, test_loader, criterion,
        device)

```

## Result

In the case of training loss with subword-level embedding, bi-directional GRU shows faster convergence and slightly better accuracy compared to unidirectional GRU as shown in Figure 1 and 2. I guess bi-directional GRU can capture information from both past and future contexts within a sequence, thereby is beneficial for understanding the context. This can lead to faster convergence and improved accuracy.

As shown in Figure 3 and 4, when I change the experiments to the word-level embedding, it doesn't converge well, and shows poor test performance. It implies that the subword-level embedding is suitable for current network.

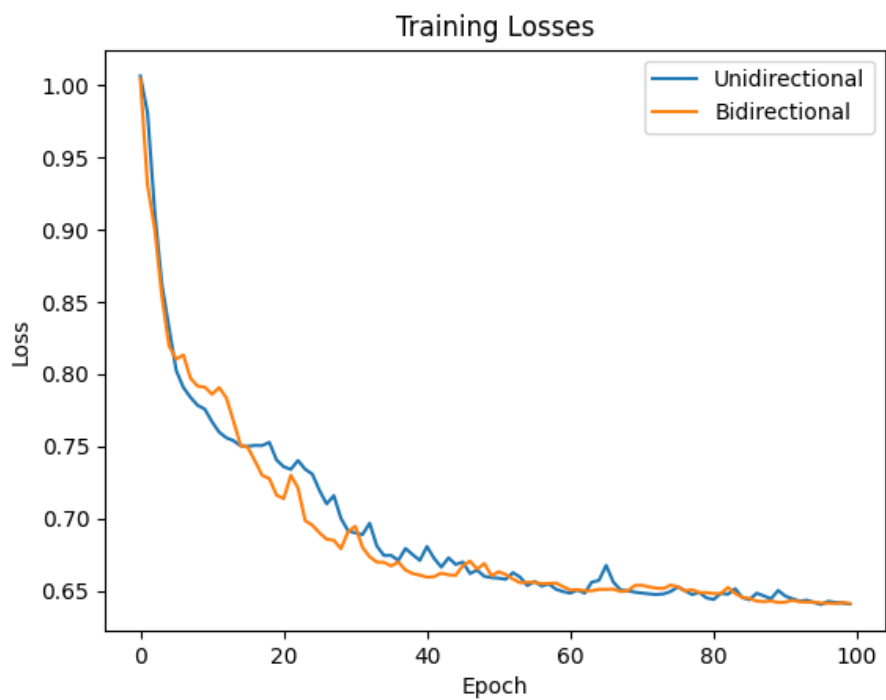


Figure 1: GRU training loss with subword-level embedding.

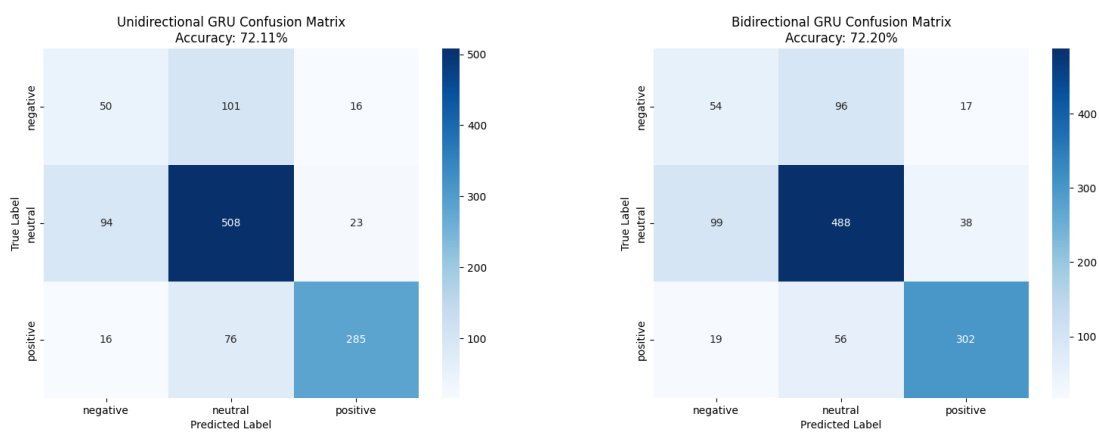


Figure 2: GRU confusion matrix and accuracy with subword-level embedding.

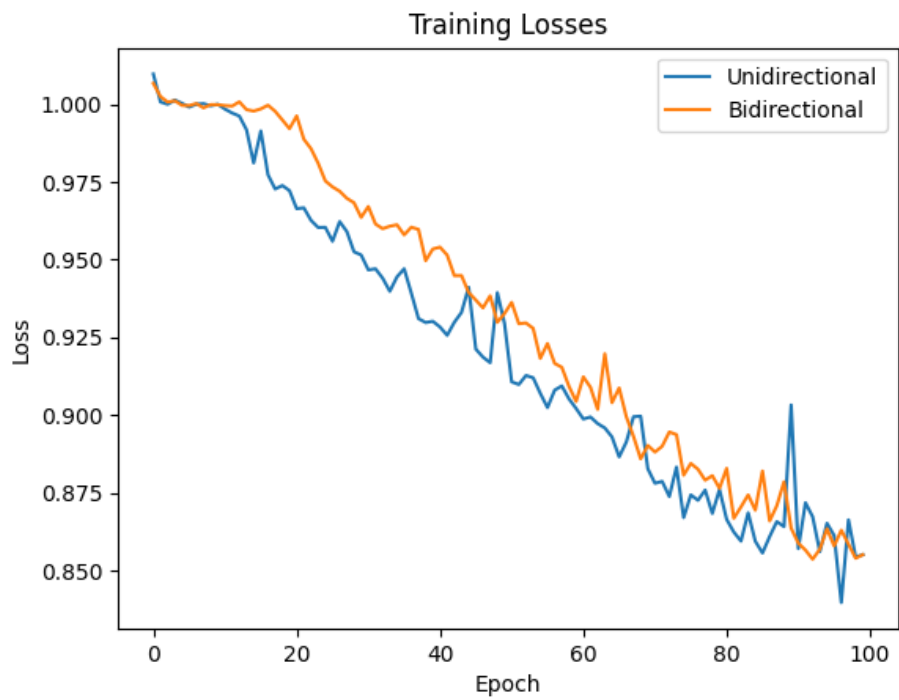


Figure 3: GRU training loss with word-level embedding.

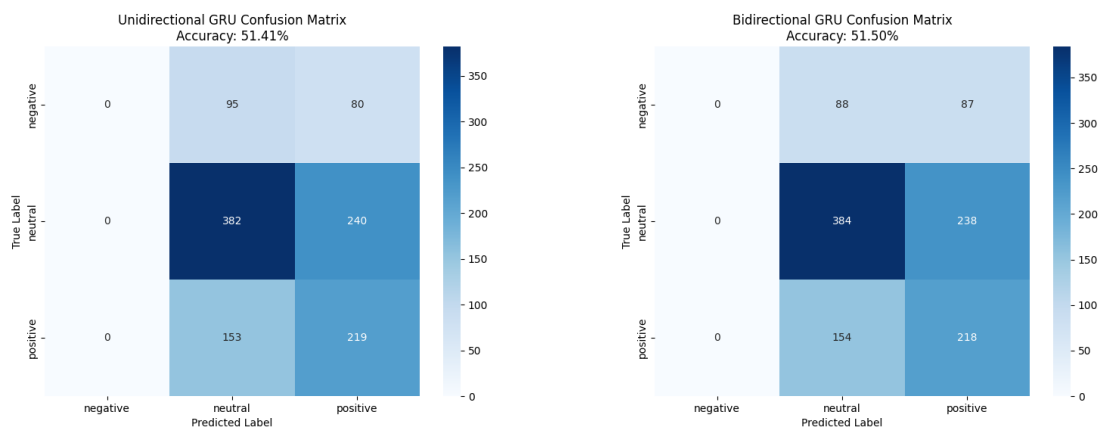


Figure 4: GRU confusion matrix and accuracy with word-level embedding.