

* UNIT-5 — Uncertain knowledge:

→ Discourse Processing :

Discourse processing is a subfield of NLP that focuses on understanding the structure & meaning of text at beyond level of individual sentences.

→ The main aim is to extract meaningful information from the text which can be useful for summarization, information extraction, question answering

→ It involves 4 sub tasks,

i) coherence & cohesion analysis

ii) Discourse structure analysis

iii) Rhetorical parsing

iv) Discourse - level sentiment analysis.

i) coherence & cohesion analysis:

→ coherence refers to overall logical connection blw sentences in Text.

→ cohesion refers to grammatical connections blw sentences (such as: Pronouns, conjunctions ... etc)

eg: John went to store. He bought some milk

He in 2nd sentence refers to John in 1st sentence.

ii) Discourse structure analysis:

→ It involves in analyzing the hierarchical structure in the Text, identifying main topics, supporting details, & relations

eg: 1st we'll discuss Problem, Then we'll find solⁿ, By Discussion only we can find solⁿ.

Structure: Introduction: First, we'll discuss both
main point: we'll find solⁿ
supporting detail: By Discussion only
we can find solⁿ

iii) Rhetorical Passing:

↳ (To inform)

→ Rhetorical Passing involves in analyzing
structure of Text to understand
how different parts of a Text are
related.

eg: "she was clever because she studies
everyday".

because indicating its cause & effect
of sentence.

iv) Discourse -level sentiment Analysis:

Involves in analyzing the
sentiment of a Text.

eg: The battery life of mobile is great
(Positive) ↑

but camera quality is bad
(negative) ↓

* Key components of Discourse

Processing:

- 1) cohesion
- 2) Resolution Reference
- 3) Discourse cohesion & structure

→ cohesion ÷

→ The grammatical connections b/w the sentences

(or) Features that link different parts of text.

→ cohesion includes:

- a) Reference words: These are like Pronouns / articles which refers to entities mentioned in initial stage of text.

Eg: Lucky went to bookstore, He bought a book

he in 2nd sentence refers to lucky in 1st sentence

b) conjunctions: words that connect sentences

like: and, but, so

eg: I want to go outside, but it's raining
but is conjunction

c) Ellipsis: words which are understood
a necessary to repeat

eg: "John reading a book, and Mary is too"

ellipse— Mary also reading book

d) lexical cohesion:

creating a meaning using similar words

eg: "The restaurant had a tasty biryani."

The dish was liked by many people

Dish refers → Biryani

a) Reference Resolution:

It is the process of determining which words in the text refers to same entity.

a) Pronoun ÷ Identify ~~who~~ To which entity the pronoun is referring.

eg: After John finished work, he went outside.

he is a pronoun referring → John.

b) Noun Phrase ÷

Identify to which entity

the noun is referring

eg: The book is on table, I read it yesterday.

Here, book & it refers same object.

c) Anaphora ÷ Using pronoun to

refer previously mentioned entity

eg: I bought a new laptop.
It has high resolution.

3) Discourse cohesion & structure:

Involves in understanding how different parts of Text, related to form 1 sentence.

a) Discourse connectives: words which gives hints the relationships.

Eg: "The government should increase education funds because Education is key to economic growth" → {connective}

b) Discourse Markers:

words which helps to organise words.

Eg: "Firstly, we need to complete syllabus on the other hand we'll prepare for exam."

c) segmentation:

Dividing Texts into segments.

Eg: "The cat chased the mouse. The mouse ran away."

* PART - 2 *

Overview ÷

1) Introduction to Language Modelling

- overview

- Applications

2) Types of Language Models ÷

a) statistical Language model

→ NGrams

b) Neural Language Models

→ Feedforward neural networks

→ Recurrent neural network

→ Transformers (BERT, GPT)

3) Language Model Evaluation:

→ perplexity

→ Accuracy

4) Parameter Estimation:

→ Max likelihood

→ Bayesian Parameter Estimation.

→ large-scale language model

5) language model adaptation

6) language specific modelling

→ Word order

→ character - Based models.

7) multilingual & crosslingual modelling

Introduction :

Language modeling is a crucial task in NLP involves in predicting next word in a sequence.

Application:

Text Generation

Speech Recognition

Machine Translation.

eg: "The cat sat on the —"

language model will predict next word as 'mat'.

→ Types of Model:

* statistical lang model:

These model uses statistical methods to estimate probability of a word sequence.

* N-Gram :

N gram can be defined as contiguous sequence of n items from a given text.

→ The items can be letters, words, (or) pairs

→ N grams can be; ~~at~~ times, let's say

- unigram — (consider each word as independent)
- bigram — (considers prob. of previous word)
- Trigram — (considers prob of previous two words)
- ⋮
- ⋮
- ⋮

N Gram modeling :

Predicts the probability of given N-gram with sequence of previous words.

eg: The cat sat on the —

→ It contains 5 words:

"The", "cat", "sat", "on", "the"

→ To find the 6th word, find conditional probability of first 5 words

$$P\left(\frac{w_6}{w_1, w_2, w_3, \dots}\right)$$

→ It can be calculated using chain rule

Hence, we can find 6th word.

$$P(A/B) = \frac{P(A \cap B)}{P(B)}$$

Example:

S1 = "I am Happy"

S2 = "I am sad"

Predict the likelihood of S1 with S2.

Step-1: Let's Take Bigram (2 word/pair)

S1 → "I am", "am Happy"

S2 → "I am"; "am sad"

Step-2: count Frequency

"I am" → appeared 2 times

"am happy" → 1

step 3: calc probability

$$P\left(\frac{\text{am}}{g}\right) = \frac{\text{count}(g\text{am})}{\text{count}(g)} = \frac{2}{2} = 1$$

$$P\left(\frac{\text{happy}}{\text{am}}\right) = \frac{\text{count}(\text{happy})}{\text{count}(\text{am})} = \frac{1}{2} = 0.5$$

$$P\left(\frac{\text{sad}}{\text{am}}\right) = \frac{\text{count}(\text{sad})}{\text{count}(\text{am})} = \frac{1}{2} = 0.5$$

step - 4: Predict Prob

$$P("g\text{am happy}) = P\left(\frac{\text{am}}{g}\right) \times P\left(\frac{\text{happy}}{\text{am}}\right) = 1.0 \times 0.5 = 0.5$$

g\text{am happy has 0.5 probability.

* Neural Language models:

a) feed forward:

uses a simple feedforward neural network where each word is represented by embeddings.
↓ (vector representation)

Architecture:

Input: Fixed no. of previous words given

Process: word embedding given to Neural network

output: network predicts probability of next word.

eg: Input: "Sun rises in _____"

Process: each word represented as vector

Prediction: The neural network

outputs the probability of next word.

eg: east, west, North, south.

if east has higher probability

b) Recurrent Neural Network

designed to handle the sequences, it maintains hidden layers / state; updated box. Every word sequence.

eg: I love my —

step - 1) "I" → update hidden state

2) "love" → update hidden state

3) "my" → update hidden state

It uses final state (hidden state) to predict next word it might be
~~word~~, [dog, cat]

[c]

c) Transformer based (BERT, GPT);

BERT: (Bidirectional encoder Representation from Transformer)

It uses transformer to understand meaning of words from both

~~It~~ It mask some words
wantedly (mask)

Eg: "The [mask] barked at cat".

It predicts masked word as "dog"

* GPT: (Generative Pretrained Transformers)

It generates text by predicting
next word in a sequence.

→ It takes the input prompt
ex generates output.

Eg: "The weather today is—"

It takes the prompt and
predicts "The weather today is
warm".

* Language model evaluation *

1) Perplexity: used to measure how well a language model is predicting sequence of words.

low perplexity = better prediction of next word.

Eg: "The cat sat on the mat"

model A

$$P(\text{The}) = 0.4$$

$$P(\text{cat} | \text{The}) = 0.2$$

$$P(\text{sat} | \text{cat}) = 0.2$$

$$P\left(\frac{\text{on}}{\text{cat sat}}\right) = 0.35$$

$$P\left(\frac{\text{mat}}{\text{cat sat on}}\right) = 0.1$$

model B

$$P(\text{The}) = 0.5$$

$$P(\text{cat} | \text{The}) = 0.4$$

$$P(\text{sat} | \text{cat}) = 0.3$$

$$P\left(\frac{\text{on}}{\text{cat sat}}\right) = 0.1$$

$$P\left(\frac{\text{mat}}{\text{cat sat on}}\right) = 0.2$$

$$\text{Perplexity} = \text{PP}(w) = \left[\frac{1}{P(w)} \right]^{\frac{1}{N}}$$

model A

$$PP(w) = \frac{1}{0.4 \times 0.2 \times 0.2 \times 0.3 \times 0.1}$$

$$= 43.84$$

model B

$$PP(w) = \frac{1}{0.5 \times 0.4 \times 0.3 \times 0.1 \times 0.2}$$

$$= 16.67$$

↓
low perplexity
mean better prediction

⇒ Accuracy: measures % of correct prediction made by a model

$$\text{Accuracy} = \frac{\text{correct Prediction}}{\text{Total Prediction}}$$

eg: "I love my Dog"

Predictions, after "I":

after I → love ✓

after I → my ✓

after I → cat ✗

after I → (love, my) ✓

Accuracy

$$= \frac{2}{3} \times 100$$

$$= 66 \%$$

* Parameter Estimations *

1) max likelihood Estimation:
used to estimate Parameters
of a model by likelihood function.

$$\text{max likelihood} = \arg \max_{\theta} P\left(\frac{D}{\theta}\right)$$

$$\hat{\theta}_{MLE} = \arg \max_{\theta} P\left(\frac{D}{\theta}\right)$$

eg: Training Data

g	love	dogs ;	g like dogs
g	love	cats ;	g like cats
g	love	coding ;	g like coding

To find Probability of next word after

love

$$P(\text{love}) = \frac{\text{count}(\text{love})}{\text{total no. of words}}$$

$$= \frac{3}{6} = 0.5$$

50% of chance is love to
be appear in given context.

2) Bayesian Parameter Estimation

It compares prior belief of parameters & update these belief on data.

$$P\left(\frac{\theta}{D}\right) = \frac{P\left(\frac{\theta}{D}\right) \cdot P(\theta)}{P(D)}$$

eg: 1) Flipping coin

$\alpha \rightarrow \text{Head}$ } and
 $\beta \rightarrow \text{Tail}$ } prior belief is
 $\alpha = 2 ; \beta = 2$

2) Now, After Flipping 10 times and observes 7 head ; 3 tails

3) update : $\alpha = \alpha_{\text{prior}} + \text{no. of H} = (2 + 7)$
 $= 9$

$$\beta = \beta_{\text{prior}} + \text{no. of T} = (2 + 3)$$
$$= 5$$

Distribution (9, 5)

$$\text{Prob of head} = \frac{\alpha}{\alpha + \beta}$$

$$= \frac{9}{9 + 5} = 0.6$$

3) Large-scale Language model;

It consists of large amount of data where the Prediction occurs iteratively by assigning weights on previous work

eg: step-1: Assigning weights
step-2: model Predicts if correct
if wrong (calculate loss)
Re-assign weight - - - again
Predict.

eg: "The sun is shining"

step-1: initialize weights
(random weights)

step-2: model Predicts

"The sun is dull" X

calc loss: let's assume
loss = 0.8

Re-Assign weights,

Smoothing Technique; It is also used for Probability

Add-on smoothing;

Add something to avoid 'zero' if scenario never happen

eg: 'The bird barked' → Not possible
but, probability > 0 (small non zero)

* Language model Adaptation *

Process of adjusting pre-trained language model to perform better on a specific task.

→ models like GPT sometimes may not give proper answers
so, Adaptation makes them to provide better information.

Benefits of Adaptation;

- Improved performance
- Efficiency

Techniques of Adaptation:

1) Fine-tuning:

Adjust the pretrained model parameters by training on a smaller domains first, then increment further.

2) Prompt Engineering:

Providing a proper prompt / instructions to guide model's behavior.

Eg: "summarize about NLP" can give a full summary of NLP

3) Few-shot Learning:

Training on a small amount of labeled datasets.

4) Zero-shot Learning:

using pretrained model without any additional training data

* Language - specific modelling

It is a fundamental

Task in NLP involves in predicting
next word / sequence of words in

a Text by word order / character.

i) Word order:

Different languages have

different word orders.

→ English have fixed word order
German have flexible word order.

Eg: The Dog, The cat (English)
Der Hund Der Katze (German)

ii) Character - Based:

Breaking the words into
smaller units (characters / syllables)

Eg: "running" → "r" "u" "n" "n" "i" "n" "g"

* Multi Lingual & Cross Lingual

→ Multi Lingual:

Focuses on building models that generates text in multiple languages.

→ used for machine Translation,
Text summarization

→ e.g. Translating Text from English to Spanish.

→ Cross Lingual:

This involves in training model in one language to perform task in another language.

→ The trained data will be limited

& The language where it perform will be high level

low trained data → high language