

Artificial Intelligence :-

"Man" made + "Thinking Power".

It is branch of Computer Science by which we can create intelligent machine, which can behave like a human, think like humans & able to make decision.

→ AI do not need to Pre Program machine to do some work.

→ You have to create a machine with Programmed algorithm. which can work with own intelligence.

NLP :-

machine / Computer (AI) know only binary code (0101...), they don't know the human language that achieved by using NLP

code
↓
binary code
0 - low voltage
1 - high voltage
logic gate.

→ why we study natural language processing?

→ NLP stands for natural language processing.

→ It is a subfield of AI.

→ It deals with the interaction b/w humans & computers in natural language.

→ It uses computational techniques to process & analyse the data.

Applications:-

→ Question answering.

ex:- Alexa, Siri &

→ Spam detection.

ex:- If you get a mail, that mail is dangerous. Spam.

→ Machine translation.

ex:- Convert to one language to another.

→ Speech Recognition.

ex:- Google assistant, If you have Youtube, that video has no subtitles it will click CC it shows text.

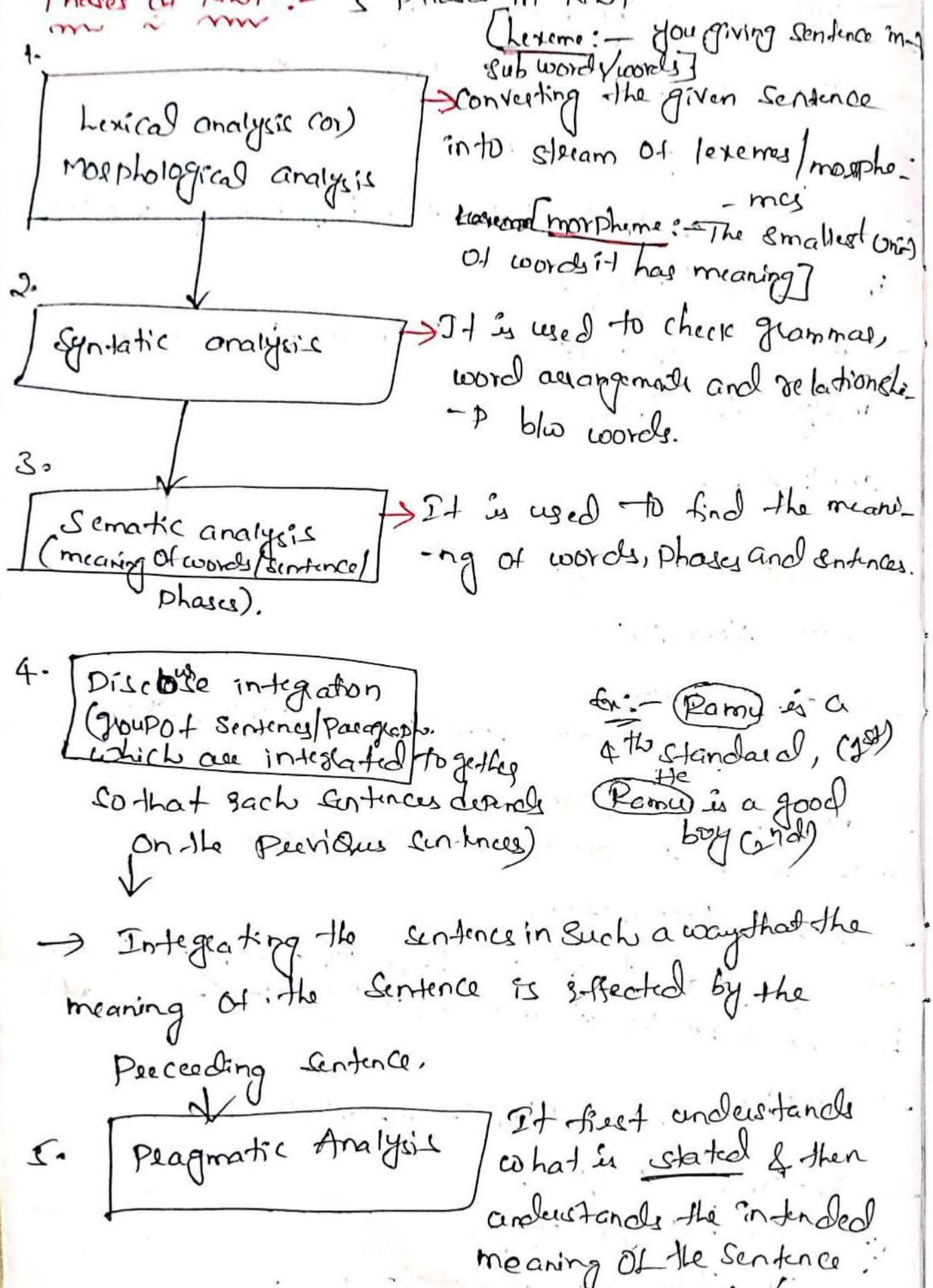
→ chatbot  chat box in websites.

→ Sentiment analysis.

ex:- Sad, happy, crying.

ex:- I am enjoying right now,

Phases of NLP:- 5 phases in NLP



Ex:- Ramu is a 4th standard, (1st)
 He Ramu is a good boy (and)

① words and their components:-

words:- words are nothing but a meaningful unit of a sentence, by using the meaningful words only we construct the meaningful sentence.

→ we have 4 components.

1. tokens
2. lexemes
3. morphemes
4. typography.

tokens:-

It refers to a sequence of characters (or) units that are treated as a single entity.

→ Tokens are words that are created by dividing the text into smaller units.

→ Process:- Tokenization. (By using tokenization we obtained the tokens)

- (1) word tokens
- (2) character tokens.
- (3) subword tokens.

Ex:- I | love | reading | news Papers.

↳ words tokens

n|e|l|a|l|a|l|n|g

↳ character tokens.

news | PaPee |

↳

subword tokens

2. Lexemes:-

→ They are the base (or) canonical form of words

ex:- The base word of running, ran is 'run' (base word)
inflection/variation

ex2:- largest, larger is large (base word).

→ fundamental unit of meaning associated with a word, regardless of the inflection (or) variations.

Process:- Lemmatization.

3. morphemes:-

→ They are the smallest meaningful unit.

→ Many words are formed by combining more than one morpheme.

Types:-

1. Free morpheme :- ex:- agree

2. Bound morpheme :- ex:- disagreement

→ ex:- disagreement.

Process:- morphological process

4. Typology:-

→ Typology refers to the classification or categorization of language based on their structural or grammatical features.

→ Aim:- To identify common pattern & variation across language to understand their properties or relationships.

1. Isolating or Analytical language:-

→ Have no or relatively few words that are comprised of more than one morpheme.

Ex:- Chinese, Thai

(In the kind of languages ~~there~~ there are very less words which are formed by one or more morphemes).

2. Synthetic languages:-

It is opposite of 1st language.
Combines more morphemes in one word.

3. Agglutinative language:-

morphemes associated with only a single function at a time. (Complex words to explain a simple topic).

Ex:- Korean, Japanese, Tamil, Finnish.

4. Fusional languages:-

Ex:- bat ← animal
cricket bat.

feature (meaning) — Per-morpheme relation higher than one.

Ex:- Arabic, Sanskrit, German etc.

① Issues and challenges:-

There are 3 issues and challenges.

1. Irregularity
2. Ambiguity
3. Productivity

Irregularity:-

The phenomenon where certain words (or) words and forms does not follow regular patterns (or) rules in terms of their morphology (or) syntax.

→ It is a challenge for algorithms which follow particular patterns

(i) Irregular verbs & nouns:-

which does not follow standard pattern (or) inflection.

<u>Ex:-</u>	<u>Past</u>	<u>Present</u>	<u>Future</u>	
	choose	chose	chosen	(choose)
	bite	bit	bitten	(bite)
	<u>went</u>			

} regular words.

(go) Irregular words

(ii) Exceptional Inflection:-

Comparative and superlative adjectives.

Ex:- big, bigger, biggest

dark, darker, darkest

good, better, best } Irregular words.

② Ambiguity:- (The word which is having multiple meanings, same spelling).

word forms that can be understood in multiple ways out of context.

→ word forms that look same but have distinct function or meaning (homonyms).

→ Ambiguity arises in morphological processing and language processing.

(a) word sense ambiguity:- Meaning depending on the context in which they are used.

Ex:- bat

animal
cricket bat.

bat

money
give banks

(b) Parts of speech ambiguity:- Different parts of speech on their usage.

Ex:- I run (Verb)

He went for a run (noun).

(c) Structural ambiguity:- multiple valid syntactic structure.

→ Perseum

(d) Referential Ambiguity:-

Referring a person by he/she
girl:- she/her, boy:- he/his.

(3) Productivity:-

Ability to generate new words (or) word forms using Productive rules. (It generates some new words by taking some unknown words).

Ex:- According to Webster's,

-googol means 1 followed by 100 zero's
↓

From this word, an unknown word [google] is generated.

↳ By Productivity rules.

new words are generated based on that word

i.e., googling, googlies, googology etc.

→ Names of People, Organisation, Location have unique structure where the Productivity issues arises.

①

Morphological Models:-

The morphological models used to analyze the structure and also formation of the words.

→ It also helps us to identify the morphemes in a particular sentence.

5 morphological models.

1. Dictionary lookup.
2. finite-state morphology.
3. Unification-based morphology.
4. Functional morphology.
5. Morphology Induction

(a) Dictionary lookup:- This used to analyze the structure of a particular word.

↓
word → base form (or) canonical form
↓
eg:- beautiful → beauty
search in dictionary
↓

returning these information.

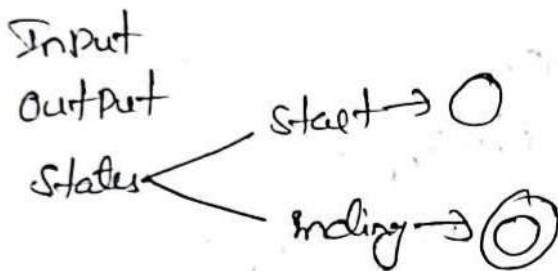
→ If you don't find the word in dictionary (or) you see more topics in that particular word.

It will combine the another algorithm and show the result word.

(b) Finite-state morphology:-

- Based on finite-automation and formal language theory.
- used for generation & recognition tasks.

formal language theory:-

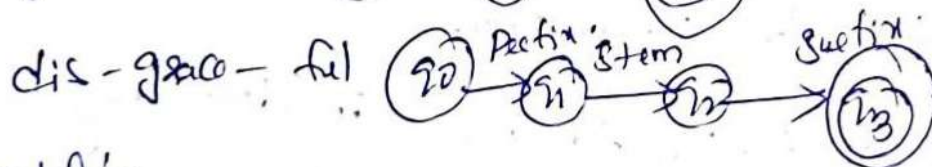
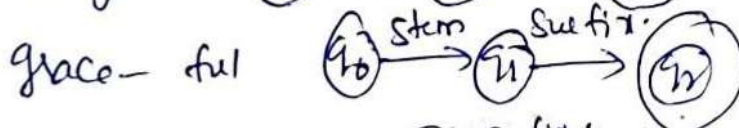
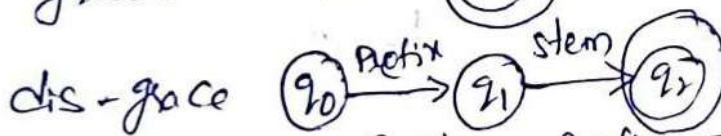
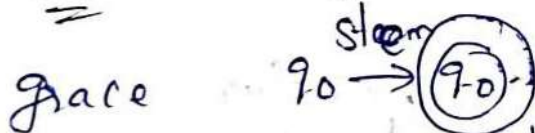


Transitions:- (→) ↘

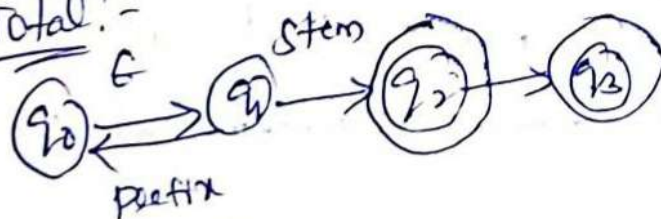
Finite-state transitions:-

This process is represented by FST's

ex! - word - grace



Total:-



(c) Unification-based morphology:-

One word ^{which} is formed by combining these different words. That formed word should be meaningful.

$$\boxed{\text{word} = \text{word 1} + \text{word 2} + \text{word 3.}}$$

↙
meaningful.

eg:- $\text{wow} = \text{world} + \text{wide} + \text{web}$.

$\text{DIY} = \text{Do it + yourself}$

$\text{HDTV} = \text{High + definition + Television}$

(d) Functional-based morphology:-

→ used to analyze the role of a morpheme in a word. (which is helpful in forming another word).

→ and also analyze their contribution in overall grammatical structure of a sentence.

eg:- walk, walking,
 ↓
 word

(e) morphological Induction:-

→ It is used to discover the pattern of a particular word (or) structures without helping in a any human being. The system only automatically detects:

→ By ~~using~~ ^{Deepening} classical morphological

Finding the structure of documents:-

→ Extracting the structure of documents helps Natural language Processing tasks like Parsing, machine translation etc.

How?

→ By chunking (Parsing) the input text (or) Speech to blocks → Segmentation.

→ Different approaches for different languages.

For example:- (No white spaces).

Chinese documents → 1. Segment character sequence into words

Morphologically rich language. → 1. Preprocessing step (determines the tokens).
2. Apply algorithms.

→ In order to chunk the text, Segmentation is used.

Segmentation:- (Given text into some sentences).

→ Decides whether to mark a boundary b/w 2 tokens as sentence boundary (or) topic

2 types.

(i) Sentence-boundary detection.

(ii) TOPIC-boundary detection.

→ Sentence-boundary detection:-

→ Also called Sentence Segmentation.

→ The Process of Segmenting sequence of words into units.

Sentence can be identified by:-

In written English:-

→ Beginning of the sentence (UPPER case).

→ end of sentence (?, ., !)

But sometimes.

→ capitalizations are ~~either~~ used for nouns and abbreviations.

→ Punctuation marks are used inside the sentences.

Example:-

I talked to Dr. Smith and my house is on mountain Dr.

(i) Optical Character Recognition:- Segments the I/p text into sentence. These systems confuses with Periods and commas.

→ It result in meaningful sentence.

Speech Character Recognition

(i) Automatic Speech Recognition: Segments the speech into sentence. These systems confuse with Punctuation marks (?, !, ,)

→ Code switching Problem arises when segmentation done for other languages.

Example: - Spanish uses inverted question marks & exclaiming marks (¿ and ¡)

→ Conventional rules-based patterns are used to identify potential ends of sentence.

(ii) Topic boundary detection: -

→ It is also called as discourse segmentation/text segmentation

→ The process of dividing the speech into homogenous blocks is called topic segmentation.

Applications: -

- 1) Text Extraction and Retrieval.
- 2) Text Summarization.

Identification: -

(i) Text Segmentation Clues: -

- by following the headlines.
- By Paragraph breaks.

(ii) Speech Segmentation Clues: -

- Pause duration
- Speaker changes.

Methods:-

1) Generative Sequence Classification Method:-

Observation :- words, punctuation. (., !)

Labels :- sentence boundary, topic boundary.

→ These methods have models which are not only capable of Predicting the most likely label but also generating the sequence itself based on learned probabilities.

One example of this method is Hidden Markov:-

How HMM ~~works~~ works?

- 1) Learn from data about the observations and its corresponding hidden states (POS).
- 2) Predict the labels for sequence generation is done.
- 3) Classify the new sequence.

Example:-

Sentence 1:- 'I / PRON Love / VERB Coding / NOUN.'

Sentence 2:- "The / DET quick / ADJ brown / ADJ Fox / NOUN Jumps / VERB".

Sentence 3:- 'She / PRON sells / VERB Sausages / NOUN'.

→ Based on these 3 sentences sequence generation and classification is done.

2) Discriminative Local Classification method:-

Local features:- word identities, Prefix, Suffixes & nearby Pos (part of speeches).

→ These approaches focus on making decision based on local information, typically without considering the entire sequence of data.

→ These methods directly model the relationship b/w input features and output labels.

Examples of these method be maximum entropy marker model (or) support vector machine.

Applications:-

- 1) Parts-of-Speech Tagging. (It Predict the ^{pos} ~~Product~~ At the Particular Sequence)
- 2) Speech Recognition. (ex:- Youtube)
- 3) Named Entity Recognition. (Person name, location name etc).

Complexity of the approaches:-

* Whether comes to complexity / one approach may be better than other.

* The complexity can be measured by rating their training and Prediction.

* Also measure their Performance.

Performance of the Approaches:-

$$\text{Error \%} = \frac{\text{no. of errors}}{\text{no. of examples}}.$$

$$\text{Precision} = \frac{\text{total no. of correct positive predictions}}{\text{total no. of positive predictions.}}$$

$$\text{Recall} = \frac{\text{total no. of all correct positive prediction}}{\text{total no. of positive instances.}}$$

$$f1 \text{ score} = \frac{2 * \text{Precision} * \text{Recall}}{\text{Precision} + \text{Recall.}}$$