

- Grammars & derivation

- Parsing $\begin{cases} \text{Robust Parsing} \\ \text{Shallow parsing} \\ \text{Probabilistic parsing} \end{cases}$

Grammar $\rightarrow$

$$G = (V, \Sigma, S, P), \quad P = \{ S \rightarrow \underline{a S b}, \\ S \rightarrow \epsilon \}$$

$$\text{Is } a a b b \in L(G)? \quad \Sigma = \{a, b\}, \quad V = \{S\}$$

$$\downarrow \quad S \Rightarrow \overset{(1)}{a} S b \Rightarrow \overset{(2)}{a} a S b b \Rightarrow \overset{(3)}{a} a b b \in L(G).$$

Grammar is for generating the sentences.
or used for deriving the sentence.

But, $a b a b \notin L(G).$

Syntax = Formal or structure.

Whether the sentence is correct or not syntax, we do it using grammar. This process is called syntax analysis.

In the above (1), (2), (3) are configurations of derivation process.

The PDA if used to recognize, then it moves through

sequence of configurations; if input is accepted, then the last configuration is accepting configuration.

The other purpose of Grammar is Semantic analysis, i.e. what is meaning associated with that sentence.

Grammar

Regular Grammar	— type-3 —	FA
Context-free Grammar	— —2 —	PDA
Context-sensitive "	—1 —	LBA
Unrestricted Grammar,	—0 —	TM

$A \in V$

Type-3

$$\begin{cases} A \rightarrow QA, \\ A \rightarrow \epsilon \end{cases} \quad \therefore W = aq, W = q, W = aqa.$$

Type-2

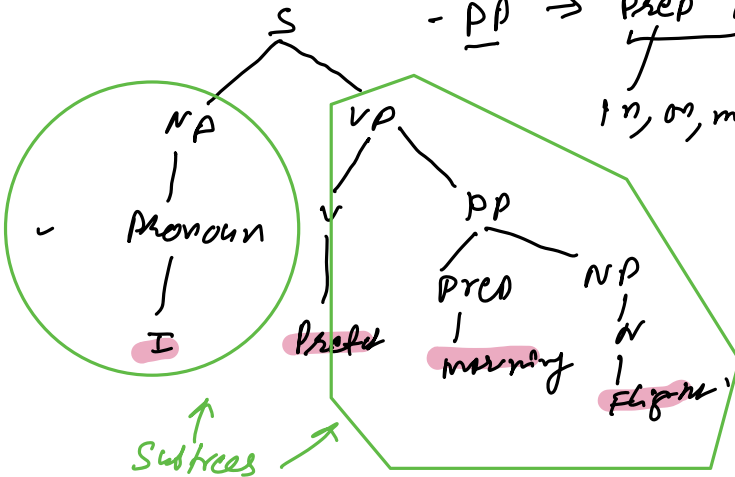
$$A \rightarrow (VU\epsilon)^*$$

Type-1

$$\underbrace{(VU\epsilon)^*}_{\alpha} \rightarrow \underbrace{(VU\epsilon)^*}_{\beta}, \quad \text{with at least one } \underline{var} \text{ in } \underline{LHS}, |\alpha| \leq |\beta|$$

type-0: $\alpha \rightarrow \beta$, with no restriction like $|\alpha| \leq |\beta|$.

But in NLP, only CFG is used, because the type 0 and 1 grammars are very difficult to implement. But, in fact, Nat. Lang. is Context-sensitive (figure - 1).

[illegible]

- If wrong Rule has been used, then we **backtrack**.
- If **nothing works**, then fail.

A Good parsing must :-

- Construct parse tree, its subtrees, and further subtrees.

↑ "Chunking", i.e. segmentation of text

- it must do disambiguation

When parsing is correctly done, like above, it is called

"Robust Parsing".

Challenge → As the sentence becomes larger & larger, it is highly time consuming to parse it. Hence, it cannot be used in real time speech and in machine translation.

In that case we do not go for complete parsing, but partial, we break the sentence and parse its units.

"The book is on table", it will remain there if you
↑ 3-4 words sentence do not collect str

This is called "Shallow parsing", this uses large Corpus of text where millions of small size parsed sentences are already available. Thus, a parser need to search those already parsed sentences. (Corpora)

It is annotated.

Present trend is shallow & probabilistic parsing.

Corpus: NLTK. (3-4 GB)

annotated: Book [n] is [Aux] on [prep] table [n]
Example.

— x — x —

