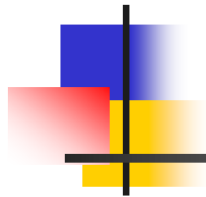




FORMAL LANGUAGES AND AUTOMATA THEORY

UNIT 1

Introduction to Automata Theory





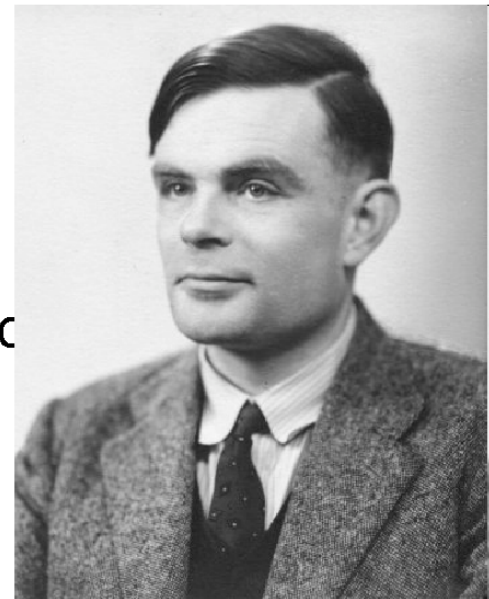
What is Automata Theory?

- *Study of abstract computing devices, or “machines”*
- **Automaton = an abstract computing device**
 - Note: A “device” need not even be a physical hardware!
- **A fundamental question in computer science:**
 - Find out what different models of machines can do and cannot do
 - The *theory of computation*
- Computability vs. Complexity

(A pioneer of automata theory)

Alan Turing (1912-1954)

- Father of Modern Computer Science
- English mathematician
- Studied abstract machines called **Turing machines** even before computers existed
- Heard of the Turing test?



Languages & Grammars

An **alphabet** is a set of symbols:

$\{0,1\}$

Or “**words**”

↓
Sentences are strings of symbols:

0,1,00,01,10,1,...

A **language** is a set of sentences:

$L = \{000,0100,0010, \dots\}$

A **grammar** is a finite list of rules defining a language.

$S \longrightarrow 0A$

$B \longrightarrow 1B$

$A \longrightarrow 1A$

$B \longrightarrow 0F$

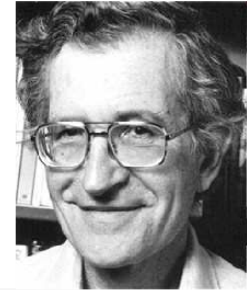
$A \longrightarrow 0B$

$F \longrightarrow \epsilon$

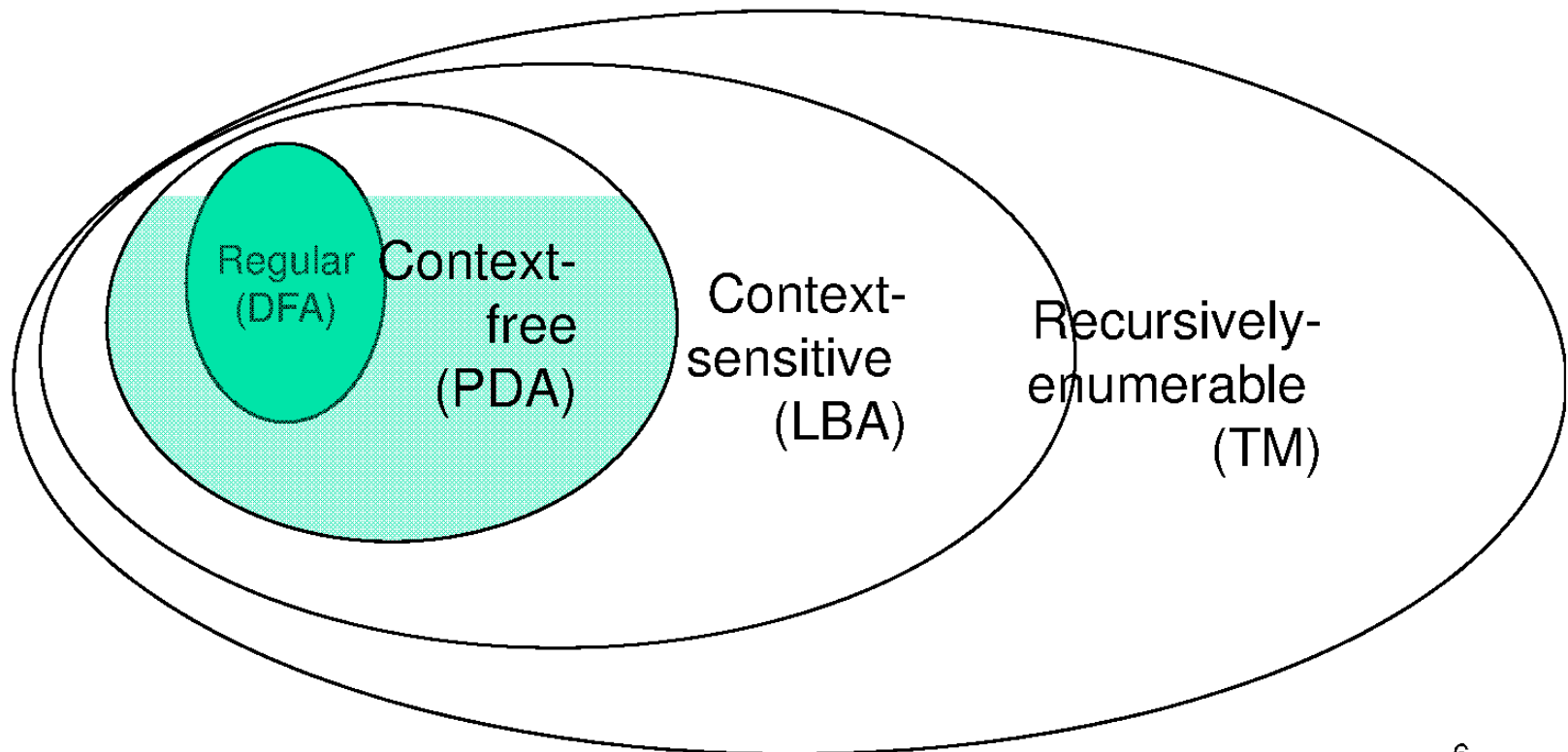
- Languages: “A language is a collection of sentences of finite length all constructed from a finite alphabet of symbols”
- Grammars: “A grammar can be regarded as a device that enumerates the sentences of a language” - nothing more, nothing less
- N. Chomsky, *Information and Control*, Vol 2, 1959



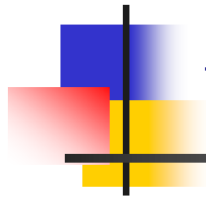
The Chomsky Hierarchy



- A containment hierarchy of classes of formal languages



The Central Concepts of Automata Theory





Alphabet

An alphabet is a finite, non-empty set of symbols

- We use the symbol Σ (sigma) to denote an alphabet
- Examples:
 - Binary: $\Sigma = \{0,1\}$
 - All lower case letters: $\Sigma = \{a,b,c,\dots z\}$
 - Alphanumeric: $\Sigma = \{a-z, A-Z, 0-9\}$
 - DNA molecule letters: $\Sigma = \{a,c,g,t\}$
 - ...



Strings

A string or word is a finite sequence of symbols chosen from Σ

- **Empty string is ε (or “epsilon”)**
- Length of a string w , denoted by “ $|w|$ ”, is equal to the *number of (non- ε) characters in the string*
 - E.g., $x = 010100$ $|x| = 6$
 - $x = 01\varepsilon 0\varepsilon 1\varepsilon 00\varepsilon$ $|x| = ?$
- xy = concatenation of two strings x and y



Powers of an alphabet

Let Σ be an alphabet.

- Σ^k = the set of all strings of length k
- $\Sigma^* = \Sigma^0 \cup \Sigma^1 \cup \Sigma^2 \cup \dots$
- $\Sigma^+ = \Sigma^1 \cup \Sigma^2 \cup \Sigma^3 \cup \dots$



Languages

L is said to be a language over alphabet Σ , only if $L \subseteq \Sigma^$*

→ this is because Σ^* is the set of all strings (of all possible length including 0) over the given alphabet Σ

Examples:

1. Let L be *the* language of all strings consisting of n 0's followed by n 1's:

$$L = \{\epsilon, 01, 0011, 000111, \dots\}$$

2. Let L be *the* language of all strings of with equal number of 0's and 1's:

$$L = \{\epsilon, 01, 10, 0011, 1100, 0101, 1010, 1001, \dots\}$$

Definition: $\emptyset$ denotes the Empty language

- Let $L = \{\epsilon\}$; Is $L = \emptyset$?

NO



The Membership Problem

Given a string $w \in \Sigma^$ and a language L over Σ , decide whether or not $w \in L$.*

Example:

Let $w = 100011$

Q) Is $w \in$ the language of strings with equal number of 0s and 1s?

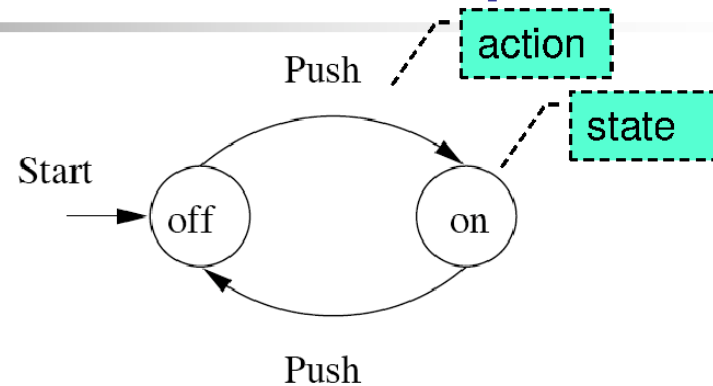


Finite Automata

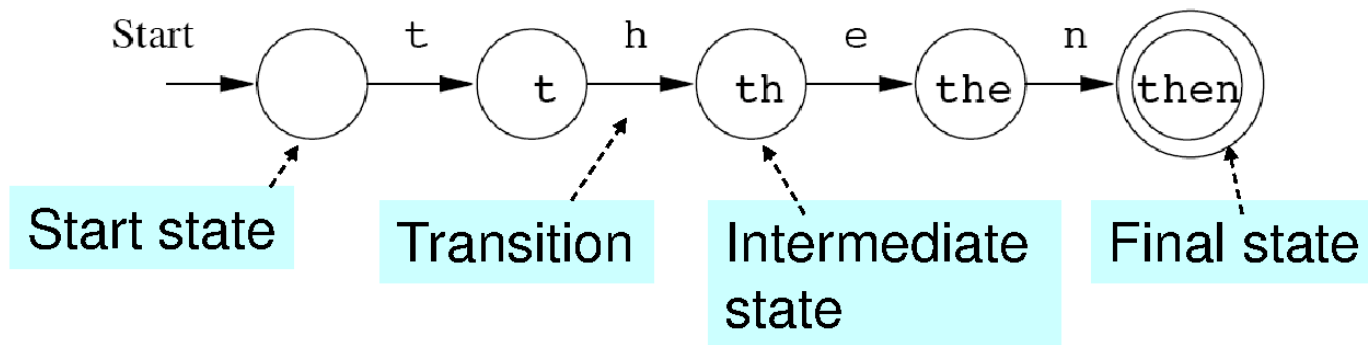
- Some Applications
 - Software for designing and checking the behavior of digital circuits
 - Lexical analyzer of a typical compiler
 - Software for scanning large bodies of text (e.g., web pages) for pattern finding
 - Software for verifying systems of all types that have a finite number of states (e.g., stock market transaction, communication/network protocol)

Finite Automata : Examples

- On/Off switch



- Modeling recognition of the word “*then*”



Structural expressions

- Grammars
- Regular expressions
 - E.g., unix style to capture city names such as “Palo Alto CA”:

■ `[A-Z][a-z]*([ ][A-Z][a-z]*)*[ ][A-Z][A-Z]`

Start with a letter

A string of other letters (possibly empty)

Other space delimited words (part of city name)

Should end w/ 2-letter state code



Summary

- Automata theory & a historical perspective
- Chomsky hierarchy
- Finite automata
- Alphabets, strings/words/sentences, languages
- Membership problem



Finite Automata



Finite Automaton (FA)

- Informally, a state diagram that comprehensively captures all possible states and transitions that a machine can take while responding to a stream or sequence of input symbols
- Recognizer for “Regular Languages”
- **Deterministic Finite Automata (DFA)**
 - The machine can exist in only one state at any given time
- **Non-deterministic Finite Automata (NFA)**
 - The machine can exist in multiple states at the same time



Deterministic Finite Automata

- Definition

- A Deterministic Finite Automaton (DFA) consists of:
 - $Q \Rightarrow$ a finite set of states
 - $\Sigma \Rightarrow$ a finite set of input symbols (alphabet)
 - $q_0 \Rightarrow$ a start state
 - $F \Rightarrow$ set of final states
 - $\delta \Rightarrow$ a transition function, which is a mapping between $Q \times \Sigma \Rightarrow Q$
- A DFA is defined by the 5-tuple:
 - $\{Q, \Sigma, q_0, F, \delta\}$



What does a DFA do on reading an input string?

- Input: a word w in Σ^*
- Question: Is w acceptable by the DFA?
- Steps:
 - Start at the “start state” q_0
 - For every input symbol in the sequence w do
 - Compute the next state from the current state, given the current input symbol in w and the transition function
 - If after all symbols in w are consumed, the current state is one of the final states (F) then *accept* w ;
 - Otherwise, *reject* w .



Regular Languages

- Let $L(A)$ be a language *recognized* by a DFA A .
 - Then $L(A)$ is called a “*Regular Language*”.
- Locate regular languages in the Chomsky Hierarchy



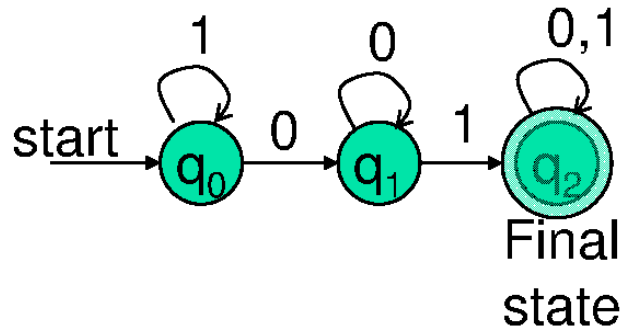
Example #1

- Build a DFA for the following language:
 - $L = \{w \mid w \text{ is a binary string that contains } 01 \text{ as a substring}\}$
- Steps for building a DFA to recognize L:
 - $\Sigma = \{0,1\}$
 - Decide on the states: Q
 - Designate start state and final state(s)
 - δ : Decide on the transitions:
- Final states == same as “accepting states”
- Other states == same as “non-accepting states”

Regular expression: $(0+1)^*01(0+1)^*$

DFA for strings containing 01

- What makes this DFA deterministic?



- What if the language allows empty strings?

- $Q = \{q_0, q_1, q_2\}$
- $\Sigma = \{0, 1\}$
- start state = q_0
- $F = \{q_2\}$
- Transition table

		symbols	
δ		0	1
states	q_0	q_1	q_0
	q_1	q_1	q_2
	q_2	q_2	q_2



Example #2

Clamping Logic:

- A clamping circuit waits for a "1" input, and turns on forever. However, to avoid clamping on spurious noise, we'll design a DFA that waits for *two consecutive 1s* in a row before clamping on.
- Build a DFA for the following language:
$$L = \{ w \mid w \text{ is a bit string which contains the substring } 11 \}$$
- State Design:
 - q_0 : start state (initially off), also means the most recent input was not a 1
 - q_1 : has never seen 11 but the most recent input was a 1
 - q_2 : has seen 11 at least once



Example #3

- Build a DFA for the following language:
 $L = \{ w \mid w \text{ is a binary string that has even number of 1s and even number of 0s} \}$
- ?



Extension of transitions (δ) to Paths ($\hat{\delta}$)

- $\hat{\delta}(q, w)$ = *destination state* from state q on input string w
- $\hat{\delta}(q, wa) = \delta(\hat{\delta}(q, w), a)$
 - Work out example #3 using the input sequence $w=10010$, $a=1$:
 - $\hat{\delta}(q_0, wa) = ?$



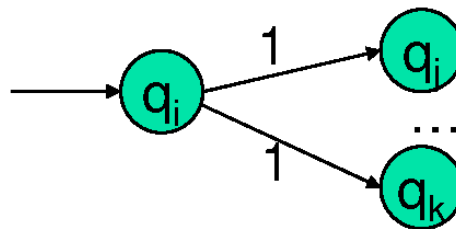
Language of a DFA

A DFA A accepts string w if there is a path from q_0 to an accepting (or final) state that is labeled by w

- *i.e.*, $L(A) = \{ w \mid \hat{\delta}(q_0, w) \in F \}$
- *i.e.*, $L(A) =$ *all strings that lead to a final state from q_0*

Non-deterministic Finite Automata (NFA)

- A Non-deterministic Finite Automaton (NFA)
 - is of course “non-deterministic”
 - Implying that the machine can exist in more than one state at the same time
 - Transitions could be non-deterministic



• Each transition function therefore maps to a set of states



Non-deterministic Finite Automata (NFA)

- A Non-deterministic Finite Automaton (NFA) consists of:
 - $Q \Rightarrow$ a finite set of states
 - $\Sigma \Rightarrow$ a finite set of input symbols (alphabet)
 - $q_0 \Rightarrow$ a start state
 - $F \Rightarrow$ set of final states
 - $\delta \Rightarrow$ a transition function, which is a mapping between $Q \times \Sigma \Rightarrow$ subset of Q
- An NFA is also defined by the 5-tuple:
 - $\{Q, \Sigma, q_0, F, \delta\}$



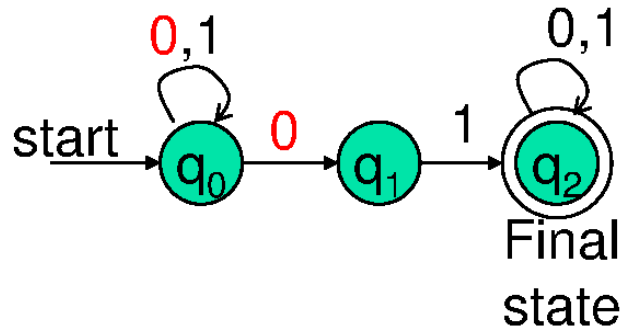
How to use an NFA?

- Input: a word w in Σ^*
- Question: Is w acceptable by the NFA?
- Steps:
 - Start at the “start state” q_0
 - For every input symbol in the sequence w do
 - Determine **all possible next states from all current states**, given the current input symbol in w and the transition function
 - If after all symbols in w are consumed and if at least **one of** the current states is a final state then *accept* w ;
 - Otherwise, *reject* w .

Regular expression: $(0+1)^*01(0+1)^*$

NFA for strings containing 01

Why is this non-deterministic?



What will happen if at state q_1 an input of 0 is received?

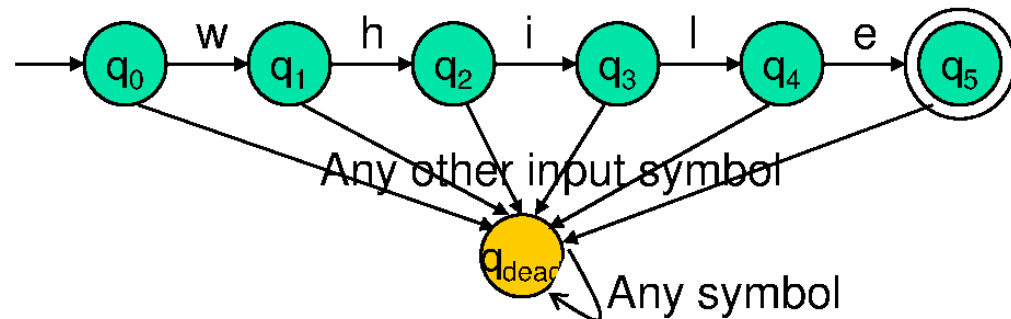
- $Q = \{q_0, q_1, q_2\}$
- $\Sigma = \{0, 1\}$
- start state = q_0
- $F = \{q_2\}$
- Transition table

symbols		
δ	0	1
states q_0	$\{q_0, q_1\}$	$\{q_0\}$
q_1	$\emptyset$	$\{q_2\}$
$*q_2$	$\{q_2\}$	$\{q_2\}$

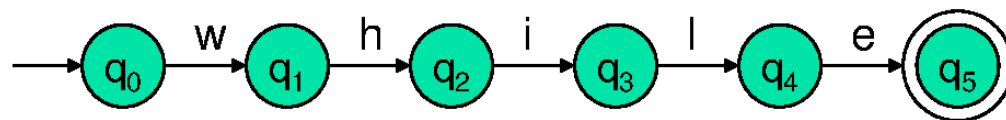
Note: Explicitly specifying dead states is just a matter of design convenience (one that is generally followed in NFAs), and this feature does not make a machine deterministic or non-deterministic.

What is a “dead state”?

- A DFA for recognizing the key word “while”



- An NFA for the same purpose:



Transitions into a dead state are implicit



Example #2

- Build an NFA for the following language:
 $L = \{ w \mid w \text{ ends in } 01 \}$
- ?
- Other examples
 - Keyword recognizer (e.g., if, then, else, while, for, include, etc.)
 - Strings where the first symbol is present somewhere later on at least once



Extension of δ to NFA Paths

- Basis: $\hat{\delta}(q, \varepsilon) = \{q\}$
- Induction:
 - Let $\hat{\delta}(q_0, w) = \{p_1, p_2, \dots, p_k\}$
 - $\delta(p_i, a) = S_i$ for $i=1, 2, \dots, k$
 - Then, $\hat{\delta}(q_0, wa) = S_1 \cup S_2 \cup \dots \cup S_k$



Language of an NFA

- An NFA accepts w if *there exists at least one* path from the start state to an accepting (or final) state that is labeled by w
- $L(N) = \{ w \mid \hat{\delta}(q_0, w) \cap F \neq \Phi \}$



Advantages & Caveats for NFA

- Great for modeling regular expressions
 - String processing - e.g., grep, lexical analyzer
- Could a non-deterministic state machine be implemented in practice?
 - A parallel computer could exist in multiple “states” at the same time
 - Probabilistic models could be viewed as extensions of non-deterministic state machines (e.g., toss of a coin, a roll of dice)

But, DFAs and NFAs are equivalent in their power to capture languages !!

Differences: DFA vs. NFA

■ DFA

1. All transitions are deterministic
 - Each transition leads to exactly one state
2. For each state, transition on all possible symbols (alphabet) should be defined
3. Accepts input if the last state is in F
4. Sometimes harder to construct because of the number of states
5. Practical implementation is feasible

■ NFA

1. Some transitions could be non-deterministic
 - A transition could lead to a subset of states
2. Not all symbol transitions need to be defined explicitly (if undefined will go to a dead state – this is just a design convenience, not to be confused with “non-determinism”)
3. Accepts input if *one of* the last states is in F
4. Generally easier than a DFA to construct
5. Practical implementation has to be deterministic (convert to DFA) or in the form of parallelism



Equivalence of DFA & NFA

- Theorem:

Should be
true for
any L

- → A language L is accepted by a DFA if and only if it is accepted by an NFA.

- Proof:

1. If part:

- Prove by showing every NFA can be converted to an equivalent DFA (in the next few slides...)

2. Only-if part is trivial:

- Every DFA is a special case of an NFA where each state has exactly one transition for every input symbol. Therefore, if L is accepted by a DFA, it is accepted by a corresponding NFA. □



Proof for the if-part

- If-part: A language L is accepted by a DFA if it is accepted by an NFA
 - rephrasing...
 - Given any NFA N , we can construct a DFA D such that $L(N)=L(D)$
-
- How to convert an NFA into a DFA?
 - Observation: In an NFA, each transition maps to a *subset* of states
 - Idea: Represent:
each “subset of NFA_states” → a single “DFA_state”

Subset construction



NFA to DFA by subset construction

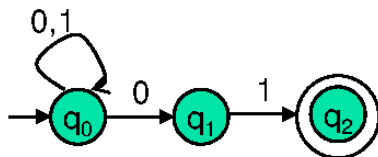
- Let $N = \{Q_N, \Sigma, \delta_N, q_0, F_N\}$
- Goal: Build $D = \{Q_D, \Sigma, \delta_D, \{q_0\}, F_D\}$ s.t.
 $L(D) = L(N)$
- Construction:
 1. Q_D = all subsets of Q_N (i.e., power set)
 2. F_D = set of subsets S of Q_N s.t. $S \cap F_N \neq \emptyset$
 3. δ_D : for each subset S of Q_N and for each input symbol a in Σ :
 - $\delta_D(S, a) = \bigcup_{p \in S} \delta_N(p, a)$

Idea: To avoid enumerating all of power set, do “lazy creation of states”

NFA to DFA construction: Example

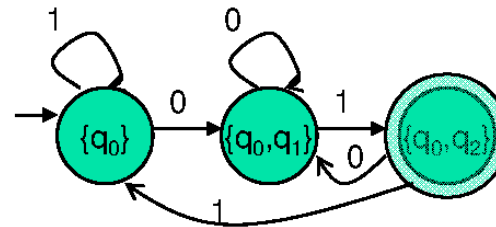
- $L = \{w \mid w \text{ ends in } 01\}$

NFA:



δ_N	0	1
$\rightarrow q_0$	$\{q_0, q_1\}$	$\{q_0\}$
q_1	$\emptyset$	$\{q_2\}$
$*q_2$	$\emptyset$	$\emptyset$

DFA:



δ_D	
$\emptyset$	
$\rightarrow \{q_0\}$	
$\{q_1\}$	
$*\{q_2\}$	
$\{q_0, q_1\}$	
$*\{q_0, q_2\}$	
$\{q_1, q_2\}$	
$\{q_0, q_1, q_2\}$	

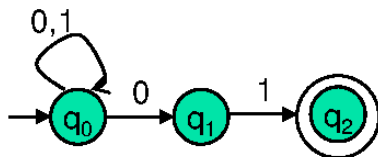
δ_D	0	1
$\rightarrow \{q_0\}$	$\{q_0, q_1\}$	$\{q_0\}$
$\{q_0, q_1\}$	$\{q_0, q_1\}$	$\{q_0, q_2\}$
$*\{q_0, q_2\}$	$\{q_0, q_1\}$	$\{q_0\}$

0. Enumerate all possible subsets
1. Determine transitions
2. Retain only those states reachable from $\{q_0\}$

NFA to DFA: Repeating the example using *LAZY CREATION*

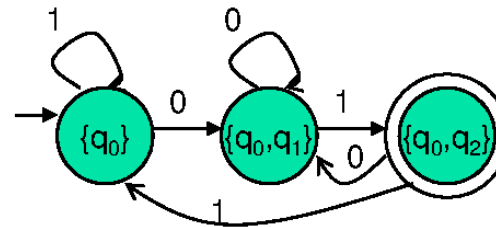
- $L = \{w \mid w \text{ ends in } 01\}$

NFA:



δ_N	0	1
q_0	$\{q_0, q_1\}$	$\{q_0\}$
q_1	$\emptyset$	$\{q_2\}$
$*q_2$	$\emptyset$	$\emptyset$

DFA:



δ_D	0	1
$\{q_0\}$	$\{q_0, q_1\}$	$\{q_0\}$

Main Idea:

Introduce states as you go
(on a need basis)



Correctness of subset construction

Theorem: *If D is the DFA constructed from NFA N by subset construction, then $L(D)=L(N)$*

■ Proof:

- Show that $\hat{\delta}_D(\{q_0\}, w) \equiv \hat{\delta}_N(q_0, w)$, for all w
- Using induction on w 's length:
 - Let $w = xa$
 - $\hat{\delta}_D(\{q_0\}, xa) \equiv \hat{\delta}_D(\hat{\delta}_N(q_0, x), a) \equiv \hat{\delta}_N(q_0, w)$



A bad case where $\#states(DFA) \gg \#states(NFA)$

- $L = \{w \mid w \text{ is a binary string s.t., the } k^{\text{th}} \text{ symbol from its end is a } 1\}$
 - NFA has $k+1$ states
 - But an equivalent DFA needs to have at least 2^k states

(Pigeon hole principle)

- m holes and $>m$ pigeons
 - $\Rightarrow$ at least one hole has to contain two or more pigeons



Applications

- Text indexing
 - inverted indexing
 - For each unique word in the database, store all locations that contain it using an NFA or a DFA
- Find pattern P in text T
 - Example: Google querying
- Extensions of this idea:
 - PATRICIA tree, suffix tree



A few subtle properties of DFAs and NFAs

- The machine never really terminates.
 - It is always waiting for the next input symbol or making transitions.
- The machine decides when to consume the next symbol from the input and when to ignore it.
 - (but the machine can never skip a symbol)
- => A transition can happen even *without* really consuming an input symbol (think of consuming ϵ as a free token)
- A single transition cannot consume more than one symbol.



FA with ϵ -Transitions

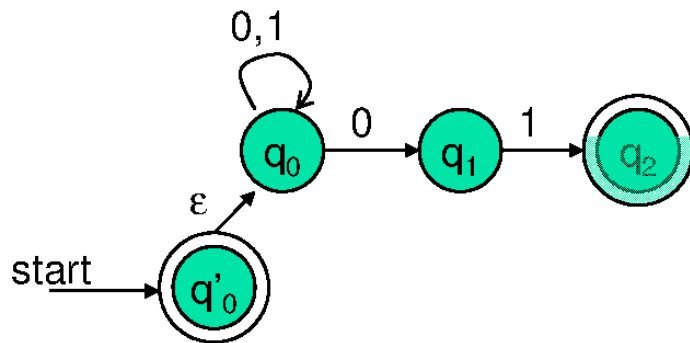
- We can allow explicit ϵ -transitions in finite automata
 - i.e., a transition from one state to another state without consuming any additional input symbol
 - Makes it easier sometimes to construct NFAs

Definition: ϵ -NFAs are those NFAs with at least one explicit ϵ -transition defined.

- ϵ -NFAs have one more column in their transition table

Example of an ϵ -NFA

$L = \{w \mid w \text{ is empty, or if non-empty will end in } 01\}$



δ_E	0	1	ϵ
* q'_0	$\emptyset$	$\emptyset$	$\{q'_0, q_0\}$
q_0	$\{q_0, q_1\}$	$\{q_0\}$	$\{q_0\}$
q_1	$\emptyset$	$\{q_2\}$	$\{q_1\}$
* q_2	$\emptyset$	$\emptyset$	$\{q_2\}$

ECLOSE(q'_0)

ECLOSE(q_0)

ECLOSE(q_1)

ECLOSE(q_2)

- ϵ -closure of a state q , **ECLOSE(q)**, is the set of all states (including itself) that can be *reached* from q by repeatedly making an arbitrary number of ϵ -transitions.

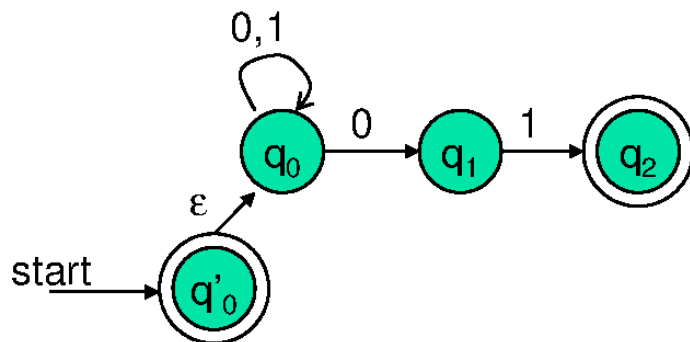
To simulate any transition:

Step 1) Go to all immediate destination states.

Step 2) From there go to all their ϵ -closure states as well.

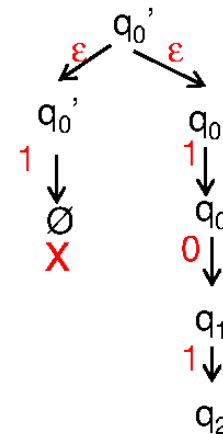
Example of an ϵ -NFA

$L = \{w \mid w \text{ is empty, or if non-empty will end in } 01\}$



δ_E	0	1	ϵ
$*q'_0$	$\emptyset$	$\emptyset$	$\{q'_0, q_0\}$ ← ECLOSE(q'_0)
q_0	$\{q_0, q_1\}$	$\{q_0\}$	$\{q_0\}$ ← ECLOSE(q_0)
q_1	$\emptyset$	$\{q_2\}$	$\{q_1\}$
$*q_2$	$\emptyset$	$\emptyset$	$\{q_2\}$

Simulate for $w=101$:

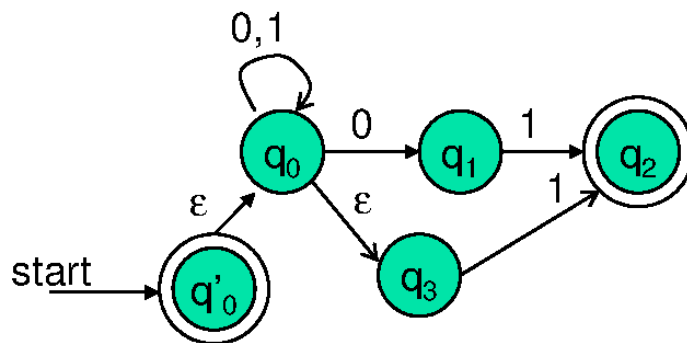


To simulate any transition:

Step 1) Go to all immediate destination states.

Step 2) From there go to all their ϵ -closure states as well.

Example of another ϵ -NFA



Simulate for $w=101$:

?

δ_E	0	1	ϵ
$\rightarrow *q'_0$	$\emptyset$	$\emptyset$	$\{q'_0, q_0, q_3\}$
q_0	$\{q_0, q_1\}$	$\{q_0\}$	$\{q_0, q_3\}$
q_1	$\emptyset$	$\{q_2\}$	$\{q_1\}$
$*q_2$	$\emptyset$	$\emptyset$	$\{q_2\}$
q_3	$\emptyset$	$\{q_2\}$	$\{q_3\}$



Equivalency of DFA, NFA, ϵ -NFA

- Theorem: A language L is accepted by some ϵ -NFA if and only if L is accepted by some DFA
- Implication:
 - $\text{DFA} \equiv \text{NFA} \equiv \epsilon\text{-NFA}$
 - (all accept Regular Languages)



Eliminating ϵ -transitions

Let $E = \{Q_E, \Sigma, \delta_E, q_0, F_E\}$ be an ϵ -NFA

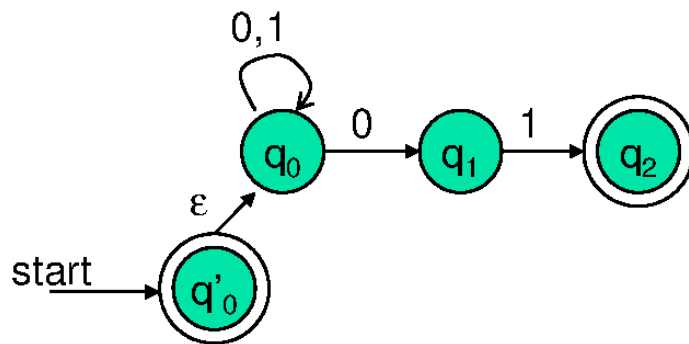
Goal: To build DFA $D = \{Q_D, \Sigma, \delta_D, \{q_D\}, F_D\}$ s.t. $L(D) = L(E)$

Construction:

1. $Q_D =$ all reachable subsets of Q_E factoring in ϵ -closures
2. $q_D = \text{ECLOSE}(q_0)$
3. $F_D =$ subsets S in Q_D s.t. $S \cap F_E \neq \emptyset$
4. δ_D : for each subset S of Q_E and for each input symbol $a \in \Sigma$:
 - Let $R = \bigcup_{p \in S} \delta_E(p, a)$ // go to destination states
 - $\delta_D(S, a) = \bigcup_{r \in R} \text{ECLOSE}(r)$ // from there, take a union of all their ϵ -closures

Example: ϵ -NFA $\rightarrow$ DFA

$L = \{w \mid w \text{ is empty, or if non-empty will end in } 01\}$

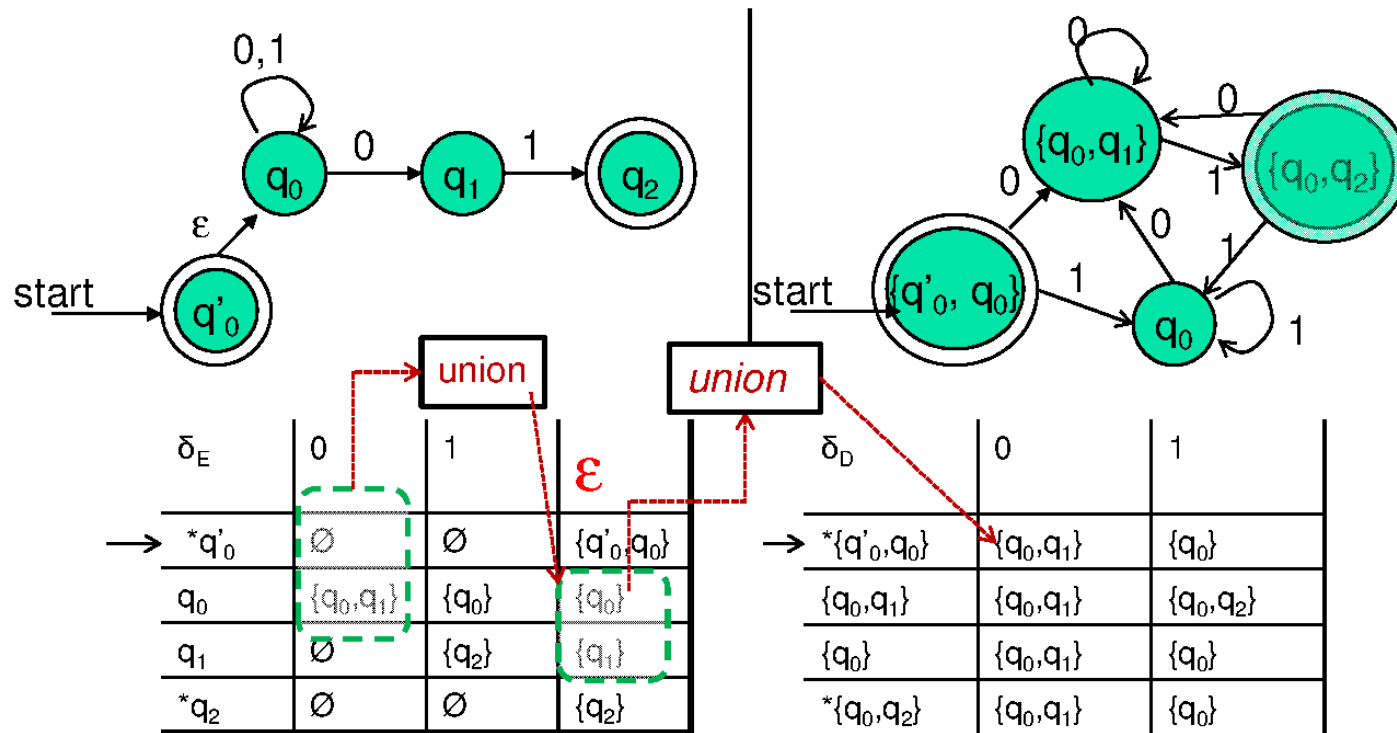


δ_E	0	1	ϵ
$\rightarrow *q'_0$	$\emptyset$	$\emptyset$	$\{q'_0, q_0\}$
q_0	$\{q_0, q_1\}$	$\{q_0\}$	$\{q_0\}$
q_1	$\emptyset$	$\{q_2\}$	$\{q_1\}$
$*q_2$	$\emptyset$	$\emptyset$	$\{q_2\}$

δ_D	0	1
$\rightarrow * \{q'_0, q_0\}$		
...		

Example: ϵ -NFA $\rightarrow$ DFA

$L = \{w \mid w \text{ is empty, or if non-empty will end in } 01\}$

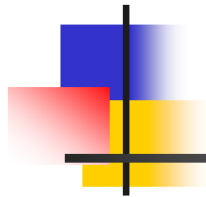




Summary

- DFA
 - Definition
 - Transition diagrams & tables
- Regular language
- NFA
 - Definition
 - Transition diagrams & tables
- DFA vs. NFA
- NFA to DFA conversion using subset construction
- Equivalency of DFA & NFA
- Removal of redundant states and including dead states
- ϵ -transitions in NFA
- Pigeon hole principles
- Text searching applications

Equivalence & Minimization of DFAs





Applications of interest

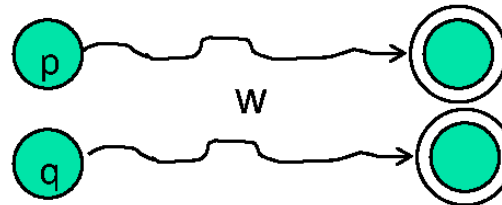
- Comparing two DFAs:
 - $L(\text{DFA}_1) == L(\text{DFA}_2)$?
- How to minimize a DFA?
 1. Remove unreachable states
 2. Identify & condense equivalent states into one

When to call two states in a DFA “equivalent”?

Past doesn't matter - only future does!

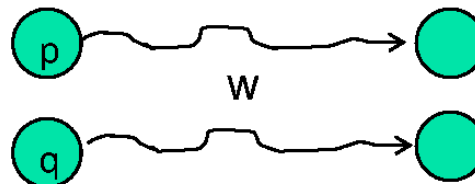
Two states p and q are said to be *equivalent* iff:

- i) Any string w accepted by starting at p is also accepted by starting at q ;



AND

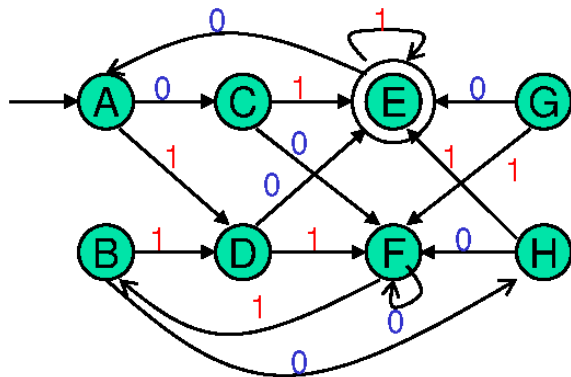
- i) Any string w rejected by starting at p is also rejected by starting at q .



→ $p \equiv q$

Computing equivalent states in a DFA

Table Filling Algorithm



Pass #0

1. Mark accepting states $\neq$ non-accepting states

Pass #1

1. Compare every pair of states
2. Distinguish by one symbol transition
3. Mark = or $\neq$ or blank(tbd)

Pass #2

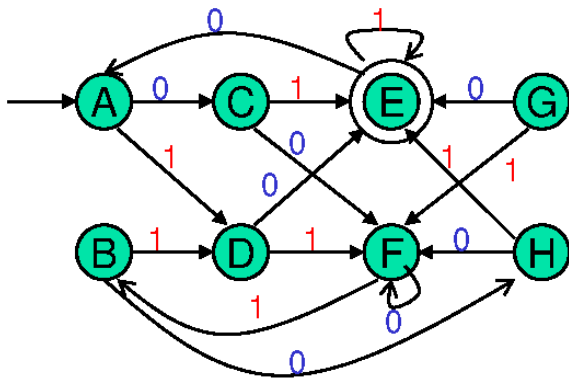
1. Compare every pair of states
2. Distinguish by up to two symbol transitions (until different or same or tbd)

....

(keep repeating until table complete)

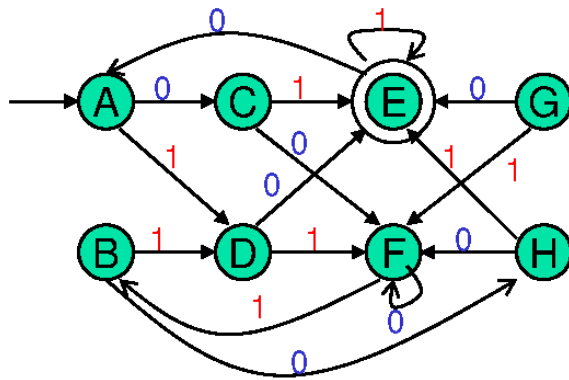
A	=							
B	=	=						
C	x	x	=					
D	x	x	x	=				
E	x	x	x	x	=			
F	x	x	x	x	x	=		
G	x	x	x	=	x	x	=	
H	x	x	=	x	x	x	x	=
	A	B	C	D	E	F	G	H

Table Filling Algorithm - step by step



A	=							
B		=						
C			=					
D				=				
E					=			
F						=		
G							=	
H								=
	A	B	C	D	E	F	G	H

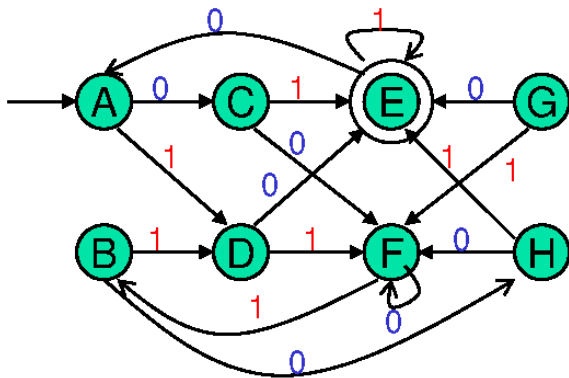
Table Filling Algorithm - step by step



1. Mark X between accepting vs. non-accepting state

A	=							
B		=						
C			=					
D				=				
E	X	X	X	X	=			
F					X	=		
G					X		=	
H					X			=
	A	B	C	D	E	F	G	H

Table Filling Algorithm - step by step

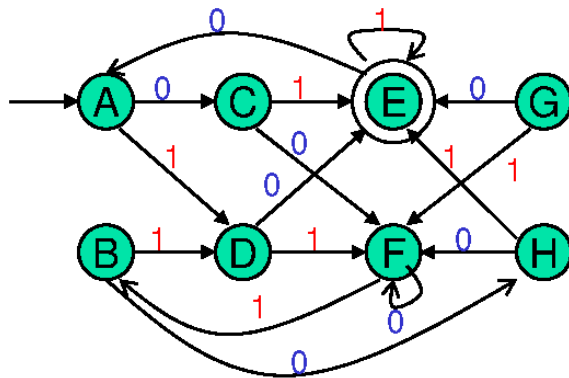


1. Mark X between accepting vs. non-accepting state
2. Look 1- hop away for distinguishing states or strings

A	=							
B		=						
C	X		=					
D	X			=				
E	X	X	X	X	=			
F					X	=		
G	X				X		=	
H	X				X			=
	A	B	C	D	E	F	G	H

↑

Table Filling Algorithm - step by step

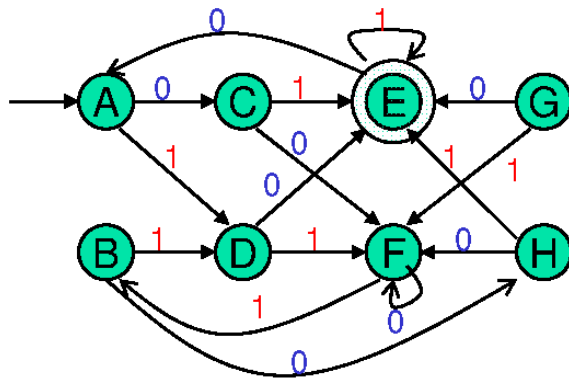


1. Mark X between accepting vs. non-accepting state
2. Look 1- hop away for distinguishing states or strings

A	=							
B		=						
C	X	X	=					
D	X	X		=				
E	X	X	X	X	=			
F					X	=		
G	X	X			X		=	
H	X	X			X			=
	A	B	C	D	E	F	G	H

↑

Table Filling Algorithm - step by step

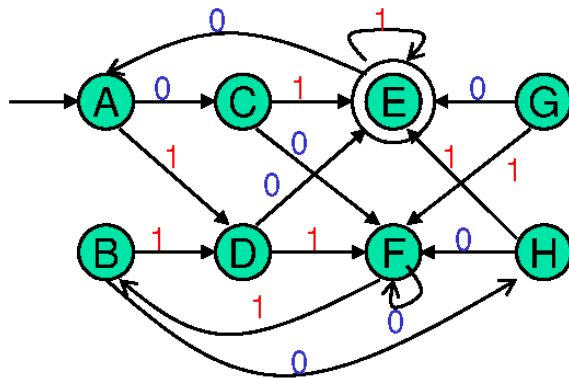


1. Mark X between accepting vs. non-accepting state
2. Look 1- hop away for distinguishing states or strings

A	=							
B		=						
C	X	X	=					
D	X	X	X	=				
E	X	X	X	X	=			
F			X		X	=		
G	X	X	X		X		=	
H	X	X	=		X			=
	A	B	C	D	E	F	G	H

↑

Table Filling Algorithm - step by step

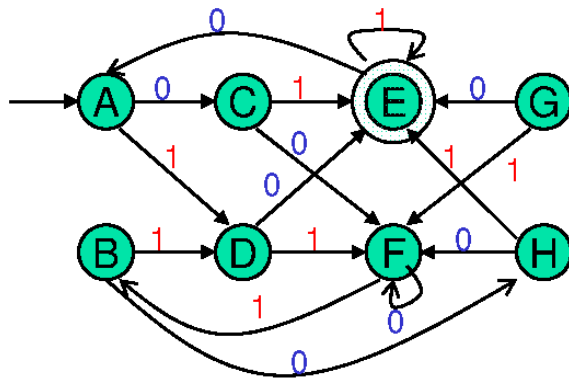


1. Mark X between accepting vs. non-accepting state
2. Look 1- hop away for distinguishing states or strings

A	=							
B		=						
C	X	X	=					
D	X	X	X	=				
E	X	X	X	X	=			
F			X	X	X	=		
G	X	X	X	=	X		=	
H	X	X	=	X	X			=
	A	B	C	D	E	F	G	H

↑

Table Filling Algorithm - step by step

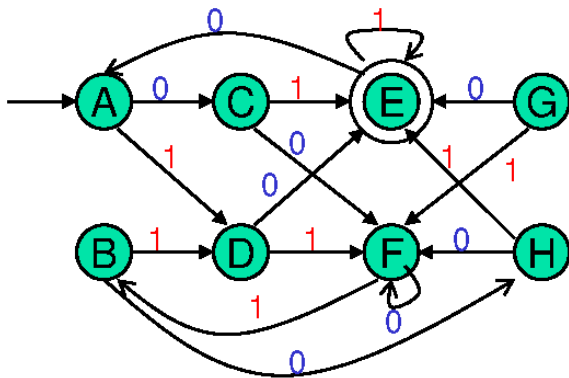


1. Mark X between accepting vs. non-accepting state
2. Look 1- hop away for distinguishing states or strings

A	=							
B		=						
C	X	X	=					
D	X	X	X	=				
E	X	X	X	X	=			
F			X	X	X	=		
G	X	X	X	=	X	X	=	
H	X	X	=	X	X	X		=
	A	B	C	D	E	F	G	H

↑

Table Filling Algorithm - step by step

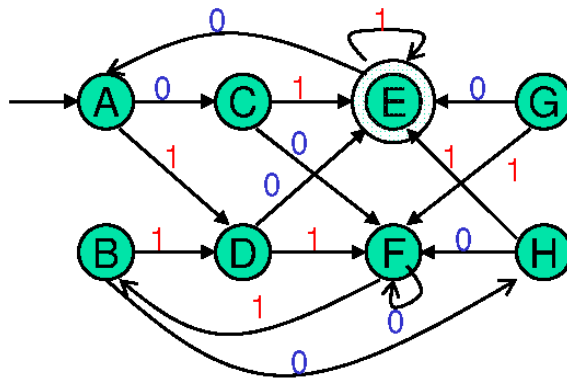


1. Mark X between accepting vs. non-accepting state
2. Look 1- hop away for distinguishing states or strings

A	=							
B		=						
C	X	X	=					
D	X	X	X	=				
E	X	X	X	X	=			
F			X	X	X	=		
G	X	X	X	=	X	X	=	
H	X	X	=	X	X	X	X	=
	A	B	C	D	E	F	G	H



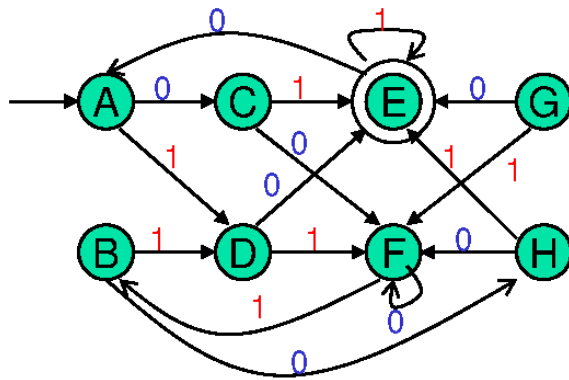
Table Filling Algorithm - step by step



1. Mark X between accepting vs. non-accepting state
2. Look 1- hop away for distinguishing states or strings
3. Look 2-hops away for distinguishing states or strings

A	=							
B	=	=						
C	X	X	=					
D	X	X	X	=				
E	X	X	X	X	=			
F	X	X	X	X	X	=		
G	X	X	X	=	X	X	=	
H	X	X	=	X	X	X	X	=
	A	B	C	D	E	F	G	H

Table Filling Algorithm - step by step



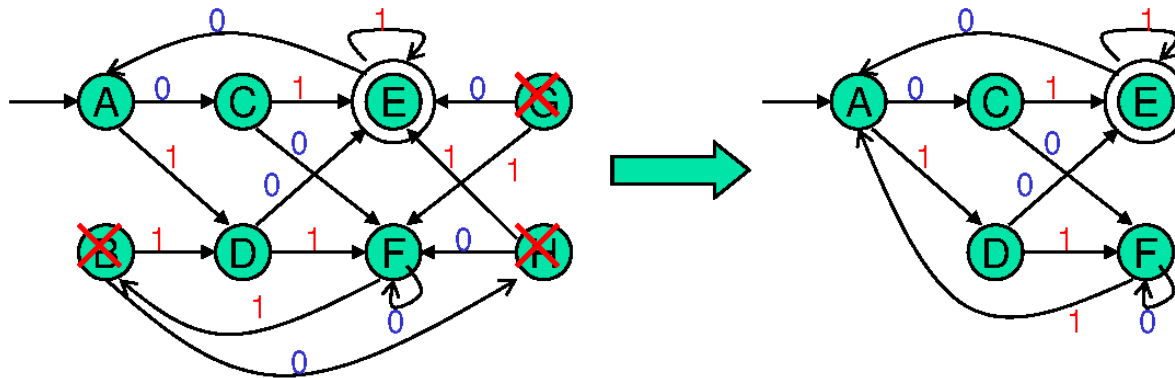
1. Mark X between accepting vs. non-accepting state
2. Look 1- hop away for distinguishing states or strings
3. Look 2-hops away for distinguishing states or strings

A	=							
B		=						
C	X	X	=					
D	X	X	X	=				
E	X	X	X	X	=			
F	X	X	X	X	X	=		
G	X	X	X	=	X	X	=	
H	X	X	=	X	X	X	X	=
	A	B	C	D	E	F	G	H

Equivalences:

- A=B
- C=H
- D=G

Table Filling Algorithm - step by step

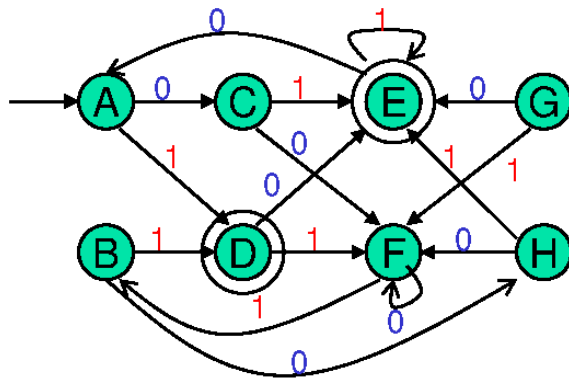


Retrain only one copy for
each equivalence set of states

Equivalences:

- A=B
- C=H
- D=G

Table Filling Algorithm – special case



A	=							
B		=						
C			=					
D				=				
E					?	=		
F							=	
G								=
H								=
	A	B	C	D	E	F	G	H

Q) What happens if the input DFA has more than one final state?
Can all final states initially be treated as equivalent to one another?

Putting it all together ...



How to minimize a DFA?

- Goal: Minimize the number of states in a DFA
- Algorithm:
 - 1. Eliminate states unreachable from the start state
 - 2. Identify and remove equivalent states
 - 3. Output the resultant DFA

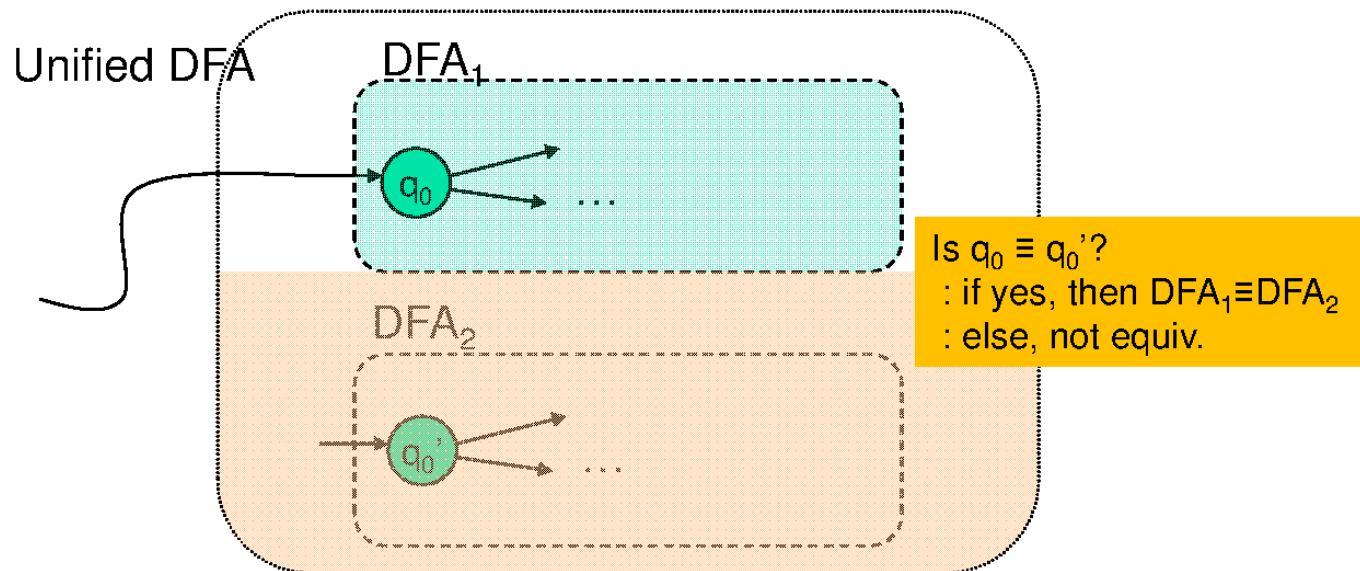
Depth-first traversal from the start state



Table filling algorithm



Are Two DFAs Equivalent?



1. Make a new dummy DFA by just putting together both DFAs
2. Run table-filling algorithm on the unified DFA
3. IF the start states of both DFAs are found to be equivalent,
 THEN: $\text{DFA}_1 \equiv \text{DFA}_2$
 ELSE: different



Summary

- Simplification of DFAs
 - How to remove unreachable states?
 - How to identify and collapse equivalent states?
 - How to minimize a DFA?
 - How to tell whether two DFAs are equivalent?



FORMAL LANGUAGES AND AUTOMATA THEORY

UNIT 2



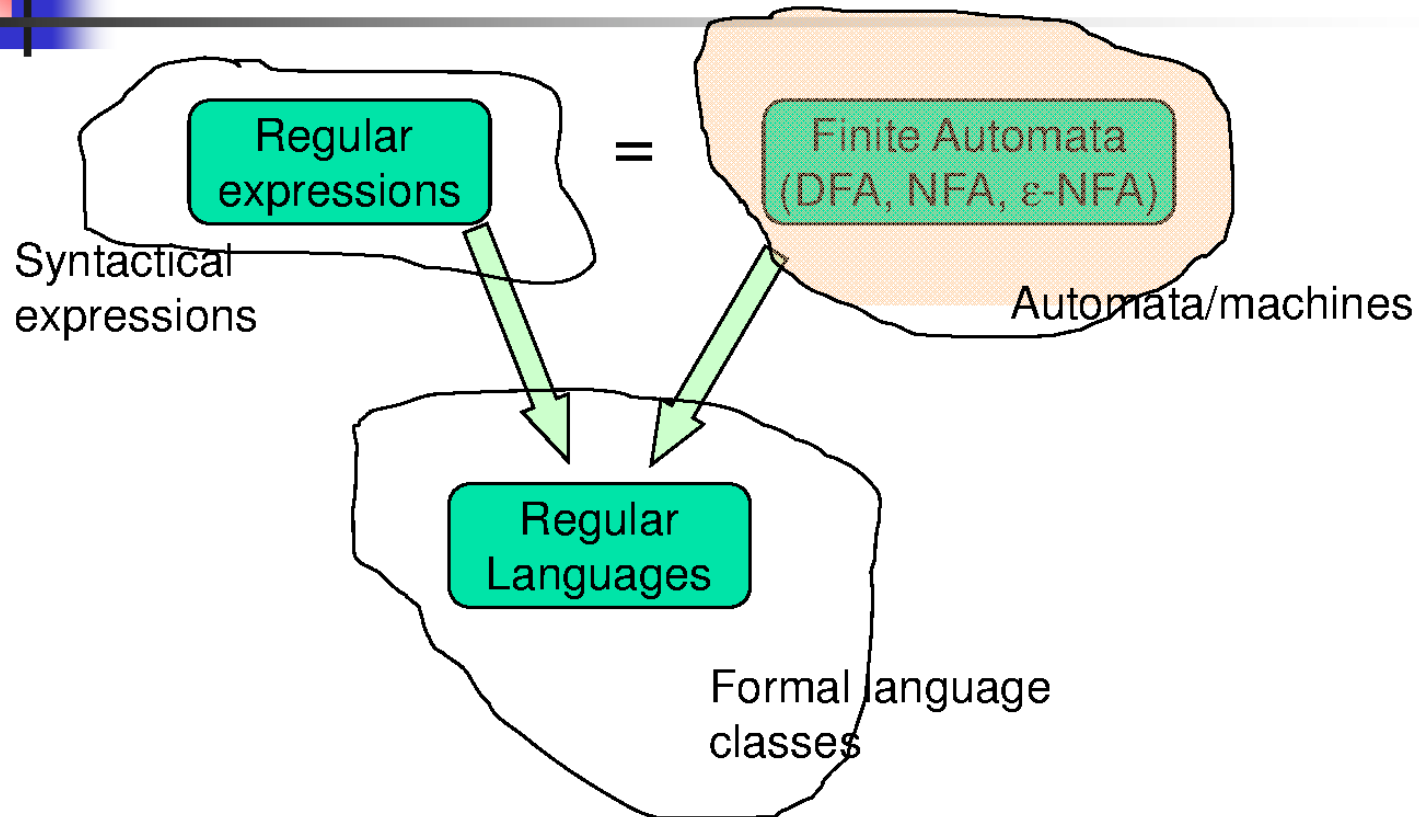
Regular Expressions



Regular Expressions vs. Finite Automata

- Offers a declarative way to express the pattern of any string we want to accept
 - E.g., $01^* + 10^*$
- Automata $\Rightarrow$ more machine-like
 - < input: string , output: [accept/reject] >
- Regular expressions $\Rightarrow$ more program syntax-like
- Unix environments heavily use regular expressions
 - E.g., bash shell, grep, vi & other editors, sed
- Perl scripting – good for string processing
- Lexical analyzers such as Lex or Flex

Regular Expressions





Language Operators

- Union of two languages:
 - $L \cup M$ = all strings that are either in L or M
 - Note: A union of two languages produces a third language
- Concatenation of two languages:
 - $L . M$ = all strings that are of the form xy
s.t., $x \in L$ and $y \in M$
 - The *dot* operator is usually omitted
 - i.e., LM is same as $L.M$

"i" here refers to how many strings to concatenate from the parent language L to produce strings in the language L^i

Kleene Closure (the * operator)

- Kleene Closure of a given language L:
 - $L^0 = \{\epsilon\}$
 - $L^1 = \{w \mid \text{for some } w \in L\}$
 - $L^2 = \{w_1w_2 \mid w_1 \in L, w_2 \in L \text{ (duplicates allowed)}\}$
 - $L^i = \{w_1w_2\dots w_i \mid \text{all } w\text{'s chosen are } \in L \text{ (duplicates allowed)}\}$
 - (Note: the choice of each w_i is independent)
 - $L^* = \bigcup_{i \geq 0} L^i$ (arbitrary number of concatenations)

Example:

- Let $L = \{1, 00\}$
 - $L^0 = \{\epsilon\}$
 - $L^1 = \{1, 00\}$
 - $L^2 = \{11, 100, 001, 0000\}$
 - $L^3 = \{111, 1100, 1001, 10000, 000000, 00001, 00100, 0011\}$
 - $L^* = L^0 \cup L^1 \cup L^2 \cup \dots$



Kleene Closure (special notes)

- L^* is an infinite set iff $|L| \geq 1$ and $L \neq \{\epsilon\}$
- If $L = \{\epsilon\}$, then $L^* = \{\epsilon\}$
- If $L = \Phi$, then $L^* = \{\epsilon\}$

Σ^* denotes the set of all words over an alphabet Σ

- Therefore, an abbreviated way of saying there is an arbitrary language L over an alphabet Σ is:
 - $L \subseteq \Sigma^*$



Building Regular Expressions

- Let E be a regular expression and the language represented by E is $L(E)$
- Then:
 - $(E) = E$
 - $L(E + F) = L(E) \cup L(F)$
 - $L(E F) = L(E) L(F)$
 - $L(E^*) = (L(E))^*$

Example: how to use these regular expression properties and language operators?

- $L = \{ w \mid w \text{ is a binary string which does not contain two consecutive 0s or two consecutive 1s anywhere} \}$
 - E.g., $w = 01010101$ is in L , while $w = 10010$ is not in L
- Goal: Build a regular expression for L
- Four cases for w :
 - Case A: w starts with 0 and $|w|$ is even
 - Case B: w starts with 1 and $|w|$ is even
 - Case C: w starts with 0 and $|w|$ is odd
 - Case D: w starts with 1 and $|w|$ is odd
- Regular expression for the four cases:
 - Case A: $(01)^*$
 - Case B: $(10)^*$
 - Case C: $0(10)^*$
 - Case D: $1(01)^*$
- Since L is the union of all 4 cases:
 - Reg Exp for $L = (01)^* + (10)^* + 0(10)^* + 1(01)^*$
- If we introduce ϵ then the regular expression can be simplified to:
 - Reg Exp for $L = (\epsilon + 1)(01)^*(\epsilon + 0)$



Precedence of Operators

- Highest to lowest

- * operator (star)
- . (concatenation)
- + operator

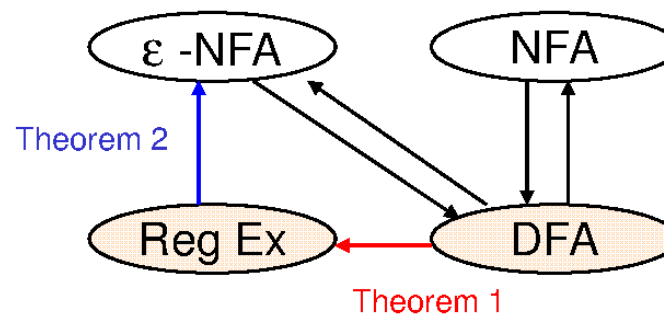
- Example:

- $01^* + 1 = (0 \cdot ((1)^*)) + 1$

Finite Automata (FA) & Regular Expressions (Reg Ex)

- To show that they are interchangeable, consider the following theorems:
 - *Theorem 1: For every DFA A there exists a regular expression R such that $L(R)=L(A)$*
 - *Theorem 2: For every regular expression R there exists an ϵ -NFA E such that $L(E)=L(R)$*

Proofs
in the book



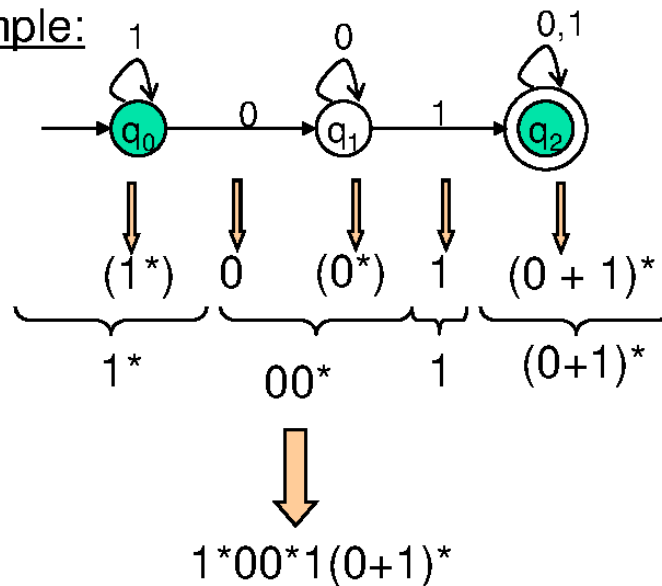
Kleene Theorem



DFA to RE construction

Informally, trace all distinct paths (traversing cycles only once) from the start state to *each of the* final states and enumerate all the expressions along the way

Example:



Q) What is the language?



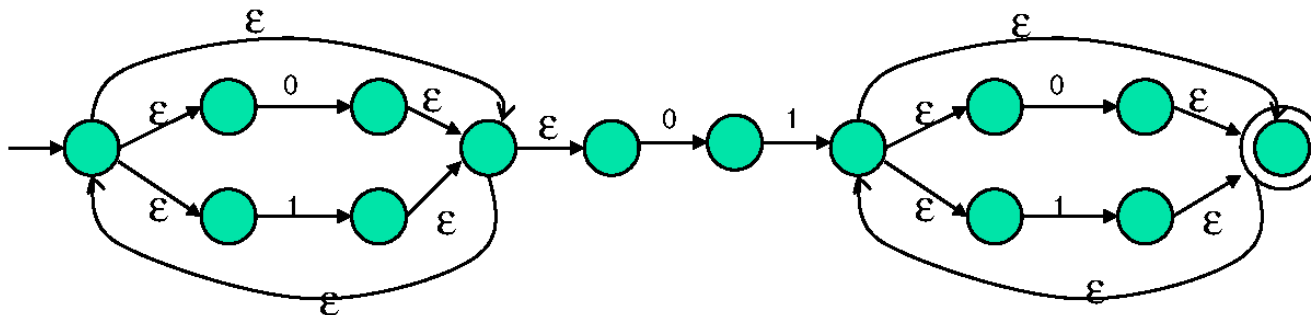
RE to ϵ -NFA construction

Example: $(0+1)^*01(0+1)^*$

$(0+1)^*$

01

$(0+1)^*$





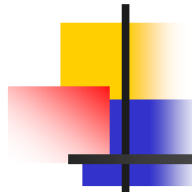
Algebraic Laws of Regular Expressions

- Commutative:
 - $E + F = F + E$
- Associative:
 - $(E + F) + G = E + (F + G)$
 - $(EF)G = E(FG)$
- Identity:
 - $E + \Phi = E$
 - $\varepsilon E = E \varepsilon = E$
- Annihilator:
 - $\Phi E = E \Phi = \Phi$



Algebraic Laws...

- Distributive:
 - $E(F+G) = EF + EG$
 - $(F+G)E = FE+GE$
- Idempotent: $E + E = E$
- Involving Kleene closures:
 - $(E^*)^* = E^*$
 - $\Phi^* = \epsilon$
 - $\epsilon^* = \epsilon$
 - $E^+ = EE^*$
 - $E? = \epsilon + E$



True or False?

Let R and S be two regular expressions. Then:

1. $((R^*)^*)^* = R^*$?

2. $(R+S)^* = R^* + S^*$?

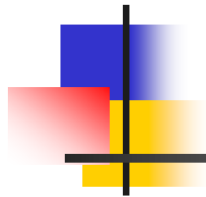
3. $(RS + R)^* RS = (RR^*S)^*$?



Summary

- Regular expressions
- Equivalence to finite automata
- DFA to regular expression conversion
- Regular expression to ϵ -NFA conversion
- Algebraic laws of regular expressions
- Unix regular expressions and Lexical Analyzer

Properties of Regular Languages





Topics

- 1) How to prove whether a given language is regular or not?
- 2) Closure properties of regular languages



Some languages are *not* regular

When is a language is regular?

if we are able to construct one of the following: DFA *or* NFA *or* ϵ -NFA *or* regular expression

When is it not?

If we can show that no FA can be built for a language



How to prove languages are *not* regular?

What if we cannot come up with any FA?

- A) Can it be language that is not regular?
- B) Or is it that we tried wrong approaches?

How do we *decisively* prove that a language is not regular?

"The hardest thing of all is to find a black cat in a dark room, especially if there is no cat!" -Confucius



Example of a non-regular language

Let $L = \{w \mid w \text{ is of the form } 0^n 1^n, \text{ for all } n \geq 0\}$

- Hypothesis: L is not regular
- Intuitive rationale: How do you keep track of a running count in an FA?
- A more formal rationale:
 - By contradiction, if L is regular then there should exist a DFA for L .
 - Let k = number of states in that DFA.
 - Consider the special word $w = 0^k 1^k \Rightarrow w \in L$
 - DFA is in some state p_i , after consuming the first i symbols in w



Rationale...

- Let $\{p_0, p_1, \dots, p_k\}$ be the sequence of states that the DFA should have visited after consuming the first k symbols in w which is 0^k
- But there are only k states in the DFA!
- $\implies$ at least one state should repeat somewhere along the path (by  +  Principle)
- $\implies$ Let the repeating state be $p_i = p_j$ for $i < j$
- $\implies$ We can fool the DFA by inputting $0^{(k-(j-i))}1^k$ and still get it to accept (note: $k-(j-i)$ is at most $k-1$).
- $\implies$ DFA accepts strings w with unequal number of 0s and 1s, implying that the DFA is wrong!



The Pumping Lemma for Regular Languages



A technique that is used to show
that a given language is not
regular



Pumping Lemma for Regular Languages

Let L be a regular language

Then there exists some constant N such that for every string $w \in L$ s.t. $|w| \geq N$, there exists a way to break w into three parts, $w = xyz$, such that:

1. $y \neq \varepsilon$
2. $|xy| \leq N$
3. For all $k \geq 0$, all strings of the form $xy^kz \in L$

This clause should hold for all regular languages.

Definition: N is called the “Pumping Lemma Constant”



Pumping Lemma: Proof

- L is regular $\Rightarrow$ it should have a DFA.
 - Set $N :=$ number of states in the DFA
- Any string $w \in L$, s.t. $|w| \geq N$, should have the form: $w = a_1 a_2 \dots a_m$, where $m \geq N$
- Let the states traversed after reading the first N symbols be: $\{p_0, p_1, \dots, p_N\}$
 - $\Rightarrow$ There are $N+1$ p-states, while there are only N DFA states
 - $\Rightarrow$ at least one state has to repeat
i.e, $p_i = p_j$ where $0 \leq i < j \leq N$ (by PHP)

Pumping Lemma: Proof...

- => We should be able to break $w = \mathbf{x} \mathbf{y} \mathbf{z}$ as follows:

- $x = a_1 a_2 \dots a_i$; $y = a_{i+1} a_{i+2} \dots a_j$; $z = a_{j+1} a_{j+2} \dots a_m$
- x 's path will be $p_0 \dots p_i$
- y 's path will be $p_i p_{i+1} \dots p_j$ (but $p_i = p_j$ implying a loop)
- z 's path will be $p_j p_{j+1} \dots p_m$

- Now consider another string $w_k = \mathbf{x} \mathbf{y}^k \mathbf{z}$, where $k \geq 0$

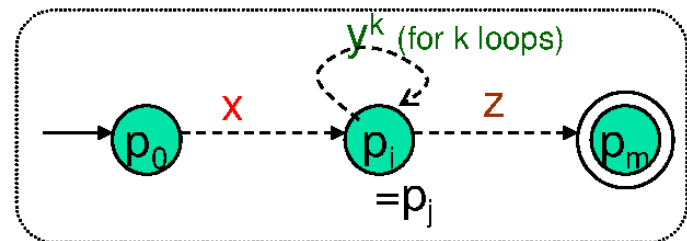
- Case $k=0$

- DFA will reach the accept state p_m

- Case $k > 0$

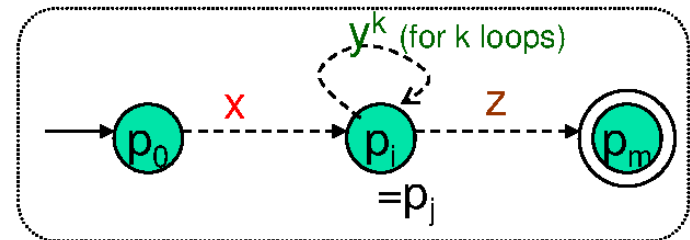
- DFA will loop for y^k , and finally reach the accept state p_m for z

- In either case, $w_k \in L$ This proves part (3) of the lemma



Pumping Lemma: Proof...

- For part (1):
 - Since $i < j$, $y \neq \varepsilon$



- For part (2):
 - By PHP, the repetition of states has to occur within the first N symbols in w
 - $\implies |xy| \leq N$

□



The Purpose of the Pumping Lemma for RL

- To prove that some languages *cannot* be regular.



How to use the pumping lemma?

Think of playing a 2 person game

- Role 1: **You** claim that the language cannot be regular
- Role 2: An **adversary** who claims the language is regular
- You show that the adversary's statement will lead to a contradiction that implies pumping lemma *cannot* hold for the language.
- You win!!



How to use the pumping lemma? (The Steps)

1. (you) L is not regular.
 2. (adv.) Claims that L is regular and gives you a value for N as its P/L constant
 3. (you) Using N , choose a string $w \in L$ s.t.,
 1. $|w| \geq N$,
 2. Using w as the template, construct other words w_k of the form xy^kz and show that at least one such $w_k \notin L$

$\Rightarrow$ this implies you have successfully broken the pumping lemma for the language, and hence that the adversary is wrong.
- (Note: In this process, you may have to try many values of k , starting with $k=0$, and then 2, 3, .. so on, until $w_k \notin L$)

Note: This N can be anything (need not necessarily be the #states in the DFA.
It's the adversary's choice.)

Example of using the Pumping Lemma to prove that a language is not regular

Let $L_{eq} = \{w \mid w \text{ is a binary string with equal number of 1s and 0s}\}$

■ Your Claim: L_{eq} is not regular

■ Proof:

➤ By contradiction, let L_{eq} be regular

→ adv.

➤ P/L constant should exist

→ adv.

➤ Let $N =$ that P/L constant

➤ Consider input $w = 0^N 1^N$

→ you

(your choice for the template string)

➤ By pumping lemma, we should be able to break $w = xyz$, such that:

→ you

1) $y \neq \epsilon$

2) $|xy| \leq N$

3) For all $k \geq 0$, the string xy^kz is also in L

Template string $w = 0^N 1^N = \underset{\leftarrow}{00} \cdot \ddots \cdot \ddots \cdot \underset{\rightarrow}{011} \cdot \ddots \cdot \ddots \cdot \underset{\rightarrow}{1}$

Proof...

- Because $|xy| \leq N$, xy should contain only 0s
 - (This and because $y \neq \varepsilon$, implies $y = 0^+$)
- Therefore x can contain *at most* $N-1$ 0s
- Also, all the N 1s must be inside z
- By (3), any string of the form $xy^kz \in L_{eq}$ for all $k \geq 0$
- Case $k=0$: xz has at most $N-1$ 0s but has N 1s
- Therefore, $xy^0z \notin L_{eq}$
- This violates the P/L (a contradiction) ↯

→ you



Setting $k=0$ is referred to as “**pumping down**”

Setting $k>1$ is referred to as “**pumping up**”

Another way of proving this will be to show that if the #0s is arbitrarily pumped up (e.g., $k=2$), then the #0s will become exceed the #1s



Exercise 2

Prove $L = \{0^n 1 0^n \mid n \geq 1\}$ is not regular

Note: This n is not to be confused with the pumping lemma constant N . That *can* be different.

In other words, the above question is same as proving:

- $L = \{0^m 1 0^m \mid m \geq 1\}$ is not regular



Example 3: Pumping Lemma

Claim: $L = \{ 0^i \mid i \text{ is a perfect square} \}$ is not regular

■ Proof:

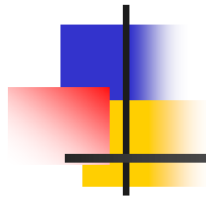
- By contradiction, let L be regular.
- P/L should apply
- Let $N = P/L$ constant
- Choose $w = 0^{N^2}$
- By pumping lemma, $w = xyz$ satisfying all three rules
- By rules (1) & (2), y has between 1 and N 0s
- By rule (3), any string of the form xy^kz is also in L for all $k \geq 0$
- Case $k=0$:
 - $\#zeros(xy^0z) = \#zeros(xyz) - \#zeros(y)$
 - $N^2 - N \leq \#zeros(xy^0z) \leq N^2 - 1$
 - $(N-1)^2 < N^2 - N \leq \#zeros(xy^0z) \leq N^2 - 1 < N^2$
 - $xy^0z \notin L$
 - But the above will complete the proof ONLY IF $N > 1$.
 - ... (proof contd.. Next slide)



Example 3: Pumping Lemma

- (proof contd...)
 - If the adversary pick $N=1$, then $(N-1)^2 \leq N^2 - N$, and therefore the $\#zeros(xy^0z)$ could end up being a perfect square!
 - This means that pumping down (i.e., setting $k=0$) is not giving us the proof!
 - So lets try pumping up next...
- Case $k=2$:
 - $\#zeros(xy^2z) = \#zeros(xyz) + \#zeros(y)$
 - $N^2 + 1 \leq \#zeros(xy^2z) \leq N^2 + N$
 - $N^2 < N^2 + 1 \leq \#zeros(xy^2z) \leq N^2 + N < (N+1)^2$
 - $xy^2z \notin L$ 
- (Notice that the above should hold for all possible N values of $N > 0$. Therefore, this completes the proof.)

Closure properties of Regular Languages



Closure properties for Regular Languages (RL)

This is different from Kleene closure

- Closure property:
 - If a set of regular languages are combined using an operator, then the resulting language is also regular
- Regular languages are closed under:
 - Union, intersection, complement, difference
 - Reversal
 - Kleene closure
 - Concatenation
 - Homomorphism
 - Inverse homomorphism

Now, let's prove all of this!



RLs are closed under union

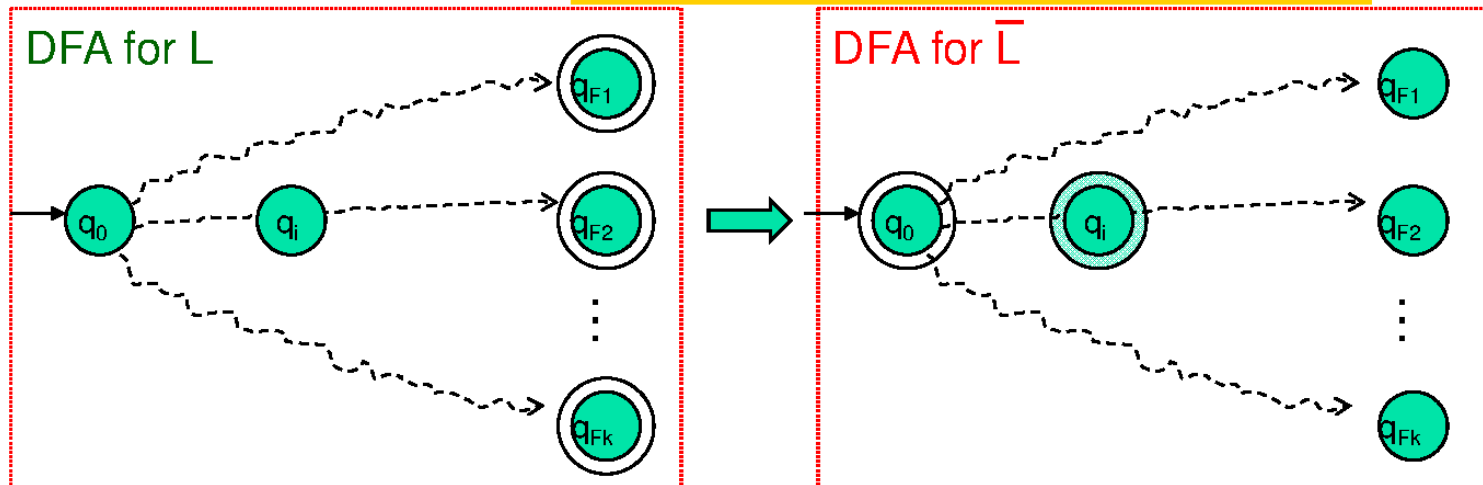
- IF L and M are two RLs THEN:
 - they both have two corresponding regular expressions, R and S respectively
 - $(L \cup M)$ can be represented using the regular expression $R+S$
 - Therefore, $(L \cup M)$ is also regular □

How can this be proved using FAs?

RLs are closed under complementation

- If L is an RL over Σ , then $\bar{L} = \Sigma^* - L$
- To show $\bar{L}$ is also regular, make the following construction

Convert every final state into non-final, and every non-final state into a final state



Assumes q_0 is a non-final state. If not, do the opposite.



RLs are closed under intersection

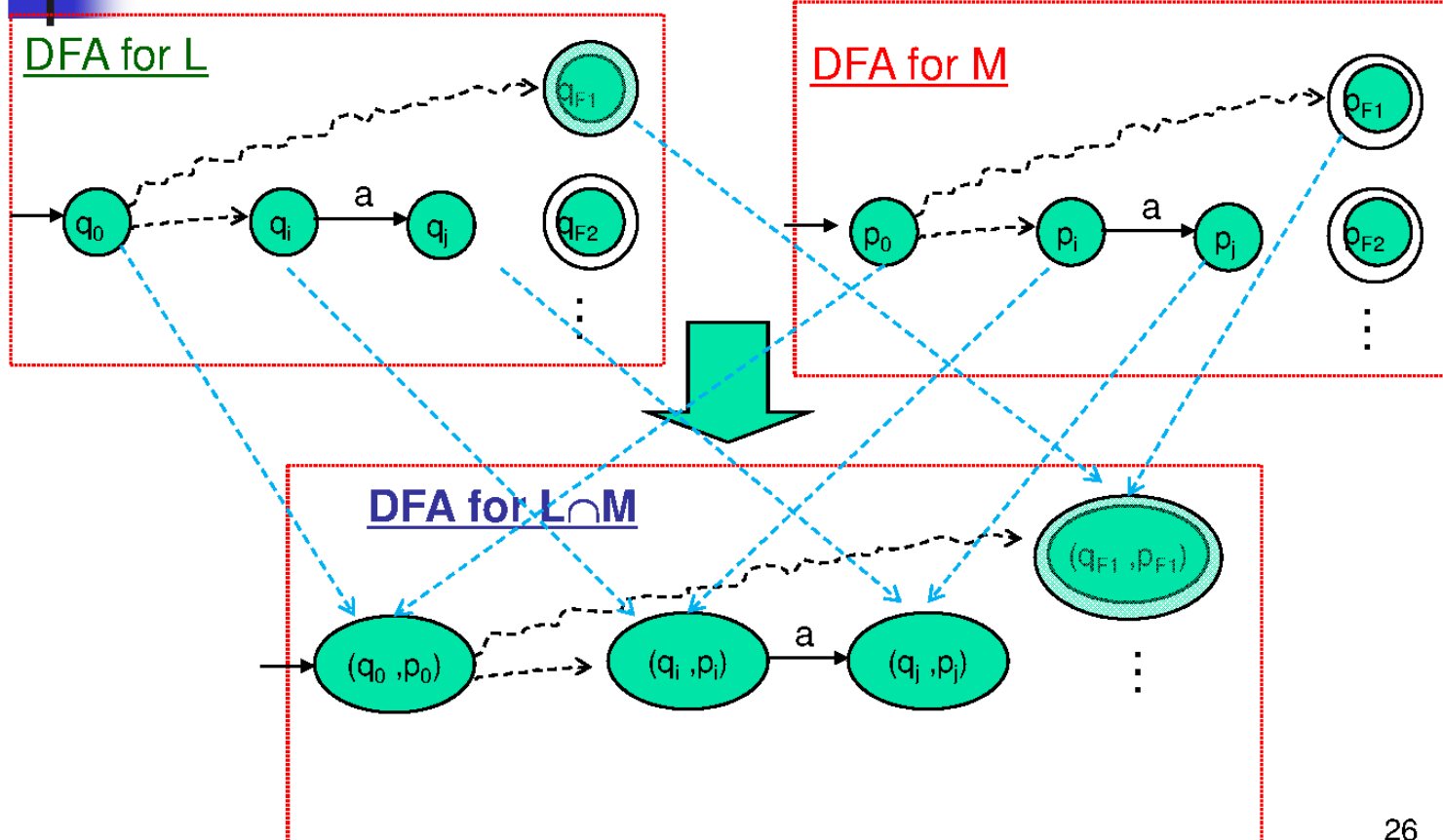
- A quick, indirect way to prove:
 - By DeMorgan's law:
 - $L \cap M = \overline{(\bar{L} \cup \bar{M})}$
 - Since we know RLs are closed under union and complementation, they are also closed under intersection
- A more direct way would be construct a finite automaton for $L \cap M$



DFA construction for $L \cap M$

- $A_L = \text{DFA for } L = \{Q_L, \Sigma, q_L, F_L, \delta_L\}$
- $A_M = \text{DFA for } M = \{Q_M, \Sigma, q_M, F_M, \delta_M\}$
- Build $A_{L \cap M} = \{Q_L \times Q_M, \Sigma, (q_L, q_M), F_L \times F_M, \delta\}$ such that:
 - $\delta((p, q), a) = (\delta_L(p, a), \delta_M(q, a))$, where p in Q_L , and q in Q_M
- This construction ensures that a string w will be accepted if and only if w reaches an accepting state in both input DFAs.

DFA construction for $L \cap M$



RLs are closed under set difference

- We observe:

- $L - M = L \cap \bar{M}$

Closed under intersection

Closed under
complementation

- Therefore, $L - M$ is also regular



RLs are closed under reversal

Reversal of a string w is denoted by w^R

- E.g., $w=00111$, $w^R=11100$

Reversal of a language:

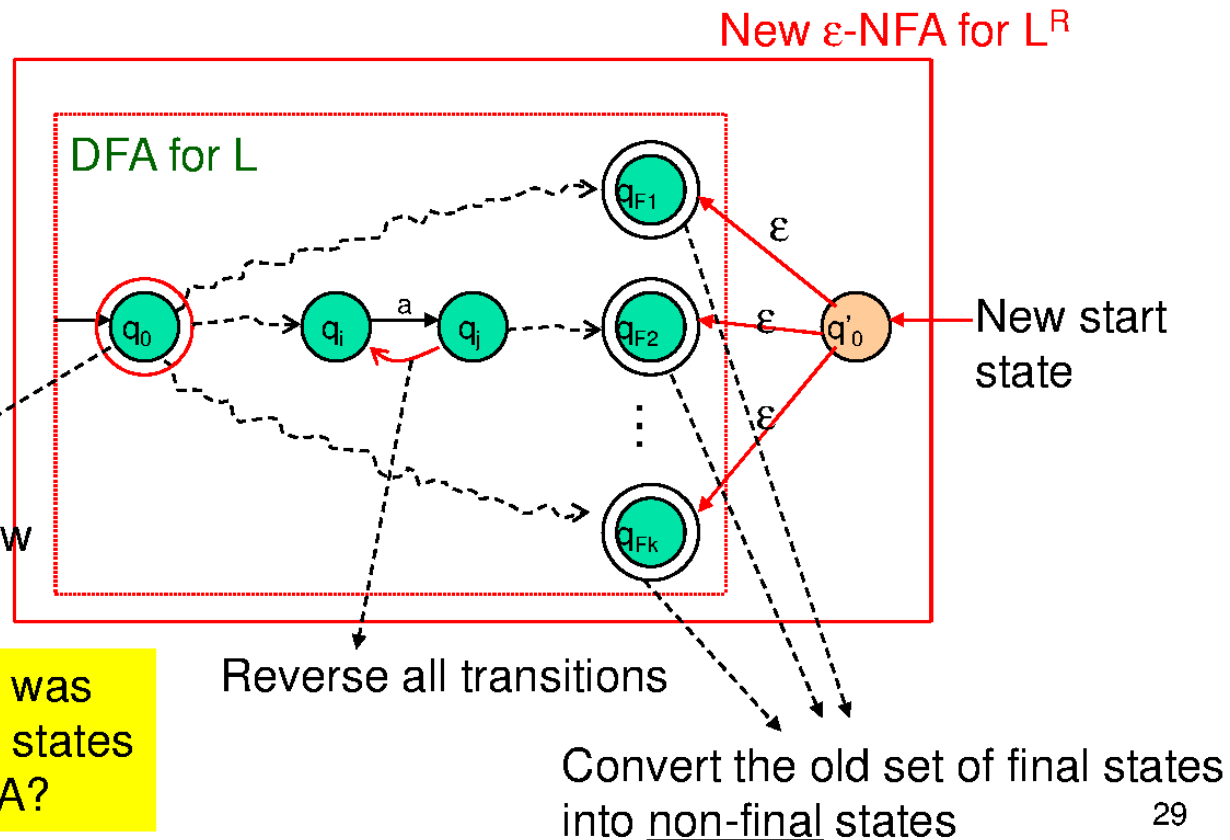
- L^R = The language generated by reversing all strings in L

Theorem: If L is regular then L^R is also regular

ϵ -NFA Construction for L^R

Make the
old start state
as the only new
final state

What to do if q_0 was
one of the final states
in the input DFA?





If L is regular, L^R is regular (proof using regular expressions)

- Let E be a regular expression for L
- Given E , how to build E^R ?
- Basis: If $E = \varepsilon$, $\emptyset$, or a , then $E^R = E$
- Induction: Every part of E (refer to the part as “ F ”) can be in only *one* of the three following forms:
 1. $F = F_1 + F_2$
 - $F^R = F_1^R + F_2^R$
 2. $F = F_1 F_2$
 - $F^R = F_2^R F_1^R$
 3. $F = (F_1)^*$
 - $(F^R)^* = (F_1^R)^*$



Homomorphisms

- Substitute each symbol in Σ (main alphabet) by a corresponding string in T (another alphabet)
 - $h: \Sigma \rightarrow T^*$
- Example:
 - Let $\Sigma = \{0, 1\}$ and $T = \{a, b\}$
 - Let a homomorphic function h on Σ be:
 - $h(0) = ab, h(1) = \epsilon$
 - If $w = 10110$, then $h(w) = \epsilon ab \epsilon \epsilon ab = abab$
- In general,
 - $h(w) = h(a_1) h(a_2) \dots h(a_n)$



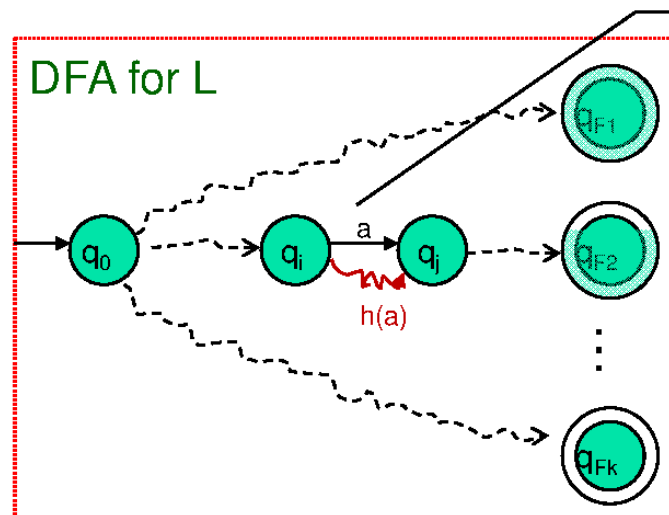
RLs are closed under homomorphisms

- Theorem: If L is regular, then so is $h(L)$
 - Proof: If E is a RE for L , then show $L(h(E)) = h(L(E))$
 - Basis: If $E = \epsilon, \emptyset$, or a , then the claim holds.
 - Induction: There are three forms of E :
 1. $E = E_1 + E_2$
 - $L(h(E)) = L(h(E_1) + h(E_2)) = L(h(E_1)) \cup L(h(E_2)) \text{ ----- (1)}$
 - $h(L(E)) = h(L(E_1) + L(E_2)) = h(L(E_1)) \cup h(L(E_2)) \text{ ----- (2)}$
 - By inductive hypothesis, $L(h(E_1)) = h(L(E_1))$ and $L(h(E_2)) = h(L(E_2))$
 - Therefore, $L(h(E)) = h(L(E))$
 2. $E = E_1 E_2$
 3. $E = (E_1)^*$
- } Similar argument

Think of a DFA based construction

Given a DFA for L , how to convert it into an FA for $h(L)$?

FA Construction for $h(L)$

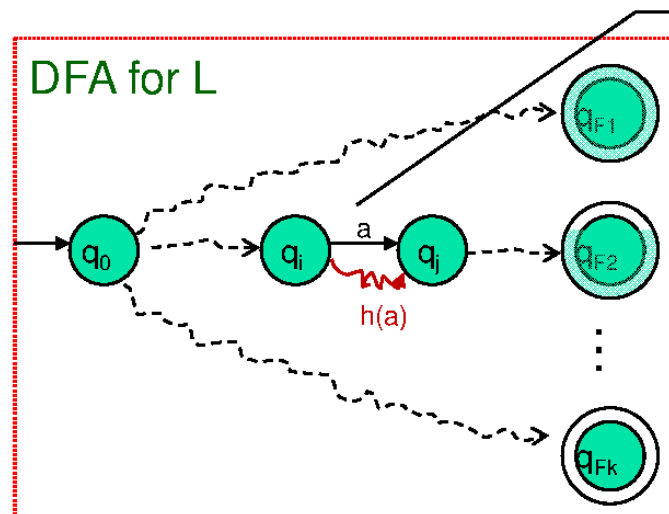


Replace every edge "a" by a path labeled $h(a)$ in the new DFA

- Build a new FA that simulates $h(a)$ for every symbol a transition in the above DFA
- The resulting FA (may or may not be a will be for $h(L)$)

Given a DFA for L , how to convert it into an FA for $h(L)$?

FA Construction for $h(L)$



Replace every edge "a" by a path labeled $h(a)$ in the new DFA

- Build a new FA that simulates $h(a)$ for every symbol a transition in the above DFA
- The resulting FA may or may not be a DFA, but will be a FA for $h(L)$

Given a DFA for M , how to convert it into an FA for $h^{-1}(M)$?

The set of strings in Σ^* whose homomorphic translation results in the strings of M

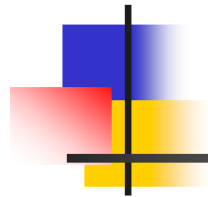
Inverse homomorphism

- Let $h: \Sigma \rightarrow T^*$
- Let M be a language over alphabet T
- $h^{-1}(M) = \{w \mid w \in \Sigma^* \text{ s.t.}, h(w) \in M\}$

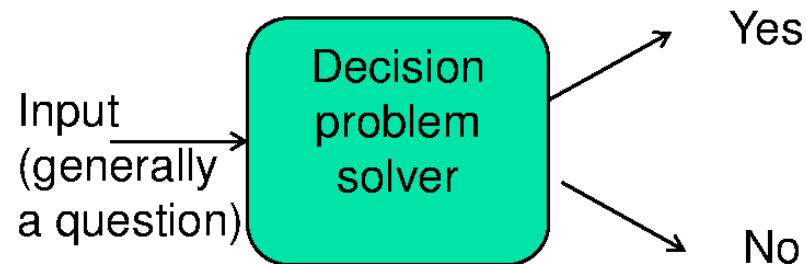
Claim: If M is regular, then so is $h^{-1}(M)$

- Proof:
 - Let A be a DFA for M
 - Construct another DFA A' which encodes $h^{-1}(M)$
 - A' is an exact replica of A , except that its transition functions are s.t. for any input symbol a in Σ , A' will simulate $h(a)$ in A .
 - $\delta(p, a) = \hat{\delta}(p, h(a))$

Decision properties of regular languages



Any “decision problem” looks like this:





Membership question

- Decision Problem: Given L , is w in L ?
- Possible answers: Yes or No
- Approach:
 1. Build a DFA for L
 2. Input w to the DFA
 3. If the DFA ends in an accepting state, then yes; otherwise no.



Emptiness test

- Decision Problem: Is $L = \emptyset$?
- Approach:
 1. Build a DFA for L
 2. From the start state, run a *reachability* test, which returns:
 1. success: if there is at least one final state that is reachable from the start state
 2. failure: otherwise
 3. $L = \emptyset$ if and only if the reachability test fails

How to implement the reachability test?



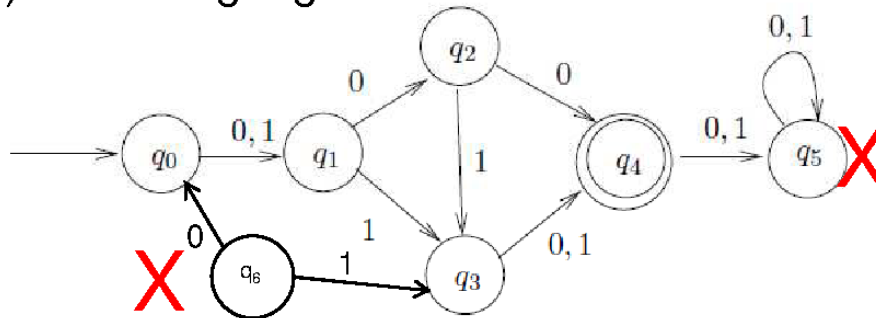
Finiteness

- Decision Problem: Is L finite or infinite?
- Approach:
 1. Build DFA for L
 2. Remove all states unreachable from the start state
 3. Remove all states that cannot lead to any accepting state.
 4. After removal, check for cycles in the resulting FA
 5. L is finite if there are no cycles; otherwise it is infinite
- Another approach
 - Build a regular expression and look for Kleene closure

How to implement steps 2 and 3?

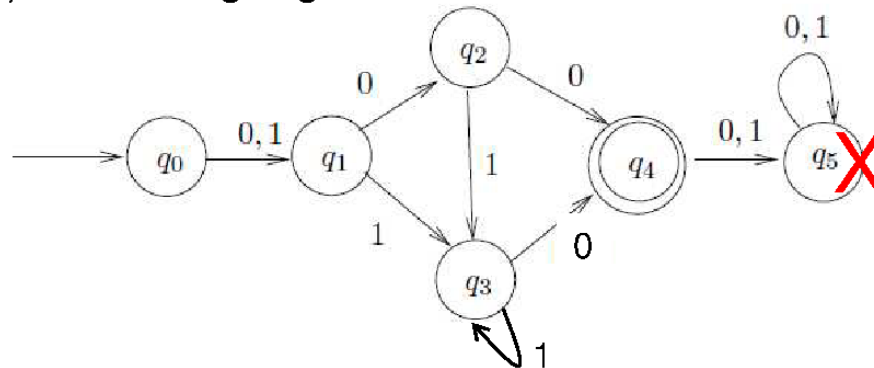
Finiteness test - examples

Ex 1) Is the language of this DFA finite or infinite?



FINITE

Ex 2) Is the language of this DFA finite or infinite?



INFINITE



Summary

- How to prove languages are not regular?
 - Pumping lemma & its applications
- Closure properties of regular languages

Context-Free Languages & Grammars (CFLs & CFGs)



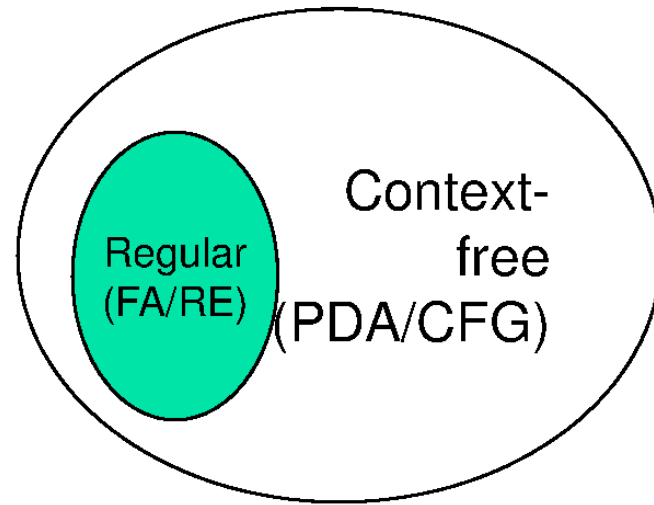
Not all languages are regular

- So what happens to the languages which are not regular?
- Can we still come up with a language recognizer?
 - i.e., something that will accept (or reject) strings that belong (or do not belong) to the language?



Context-Free Languages

- A language class larger than the class of regular languages
- Supports natural, recursive notation called “context-free grammar”
- Applications:
 - Parse trees, compilers
 - XML





An Example

- A palindrome is a word that reads identical from both ends

- E.g., madam, redivider, malayalam, 010010010

- Let $L = \{ w \mid w \text{ is a binary palindrome} \}$
- Is L regular?

- No.

- Proof:

- Let $w = 0^N 1 0^N$ (assuming N to be the p/l constant)
 - By Pumping lemma, w can be rewritten as xyz , such that xy^kz is also L (for any $k \geq 0$)
 - But $|xy| \leq N$ and $y \neq \epsilon$
 - $\implies y = 0^+$
 - $\implies xy^kz$ will NOT be in L for $k \neq 0$
 - $\implies$ Contradiction

But the language of palindromes...

is a CFL, because it supports recursive substitution (in the form of a CFG)

- This is because we can construct a “grammar” like this:

- Productions
1. $A \Rightarrow \epsilon$
 2. $A \Rightarrow 0$
 3. $A \Rightarrow 1$
 4. $A \Rightarrow 0A0$
 5. $A \Rightarrow 1A1$

Terminal

Variable or non-terminal

Same as:

$A \Rightarrow 0A0 \mid 1A1 \mid 0 \mid 1 \mid \epsilon$

How does this grammar work?



How does the CFG for palindromes work?

An input string belongs to the language (i.e., accepted) iff it can be generated by the CFG

- Example: $w=01110$
- G can generate w as follows:

1. $A \Rightarrow 0A0$
2. $\Rightarrow 01A10$
3. $\Rightarrow 01110$

G :

$A \Rightarrow 0A0 \mid 1A1 \mid 0 \mid 1 \mid \epsilon$

Generating a string from a grammar:

1. Pick and choose a sequence of productions that would allow us to generate the string.
2. At every step, substitute one variable with one of its productions.



Context-Free Grammar: Definition

- A context-free grammar $G=(V,T,P,S)$, where:
 - V : set of variables or non-terminals
 - T : set of terminals (= alphabet $\cup \{\epsilon\}$)
 - P : set of *productions*, each of which is of the form
 $V \Rightarrow \alpha_1 \mid \alpha_2 \mid \dots$
 - Where each α_i is an arbitrary string of variables and terminals
 - $S \Rightarrow$ start variable

CFG for the language of binary palindromes:
 $G=({A},{0,1},P,A)$
 $P: A \Rightarrow 0 A 0 \mid 1 A 1 \mid 0 \mid 1 \mid \epsilon$



More examples

- Parenthesis matching in code
- Syntax checking
- In scenarios where there is a general need for:
 - Matching a symbol with another symbol, or
 - Matching a count of one symbol with that of another symbol, or
 - Recursively substituting one symbol with a string of other symbols



Example #2

- Language of balanced paranthesis
e.g., $()((((())))(())\dots$
- CFG?

G:
 $S \Rightarrow (S) \mid SS \mid \varepsilon$

How would you “interpret” the string “((((()))())” using this grammar?



Example #3

- A grammar for $L = \{0^m 1^n \mid m \geq n\}$
- CFG?

G:
 $S \Rightarrow 0S1 \mid A$
 $A \Rightarrow 0A \mid \varepsilon$

How would you interpret the string “00000111”
using this grammar?



Example #4

A program containing **if-then(-else)** statements

if *Condition* **then** *Statement* **else** *Statement*

(Or)

if *Condition* **then** *Statement*

CFG?



More examples

- $L_1 = \{0^n \mid n \geq 0\}$
- $L_2 = \{0^n \mid n \geq 1\}$
- $L_3 = \{0^i 1^j 2^k \mid i=j \text{ or } j=k, \text{ where } i, j, k \geq 0\}$
- $L_4 = \{0^i 1^j 2^k \mid i=j \text{ or } i=k, \text{ where } i, j, k \geq 1\}$



Applications of CFLs & CFGs

- Compilers use parsers for syntactic checking
- Parsers can be expressed as CFGs
 1. Balancing parenthesis:
 - $B \Rightarrow BB \mid (B) \mid \text{Statement}$
 - $\text{Statement} \Rightarrow \dots$
 2. If-then-else:
 - $S \Rightarrow SS \mid \text{if Condition then Statement else Statement} \mid \text{if Condition then Statement} \mid \text{Statement}$
 - $\text{Condition} \Rightarrow \dots$
 - $\text{Statement} \Rightarrow \dots$
 3. C parenthesis matching $\{ \dots \}$
 4. Pascal *begin-end* matching
 5. YACC (Yet Another Compiler-Compiler)



More applications

- Markup languages

- Nested Tag Matching

- HTML

- `<html> ...<p> ... <a href=...> ... </a> </p> ... </html>`

- XML

- `<PC> ... <MODEL> ... </MODEL> .. <RAM> ...
</RAM> ... </PC>`



Tag-Markup Languages

Roll $\Rightarrow$ **<ROLL>** Class Students **</ROLL>**

Class $\Rightarrow$ **<CLASS>** Text **</CLASS>**

Text $\Rightarrow$ Char Text | Char

Char $\Rightarrow$ **a | b | ... | z | A | B | .. | Z**

Students $\Rightarrow$ Student Students | ϵ

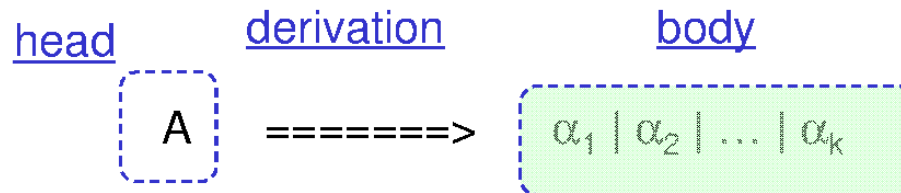
Student $\Rightarrow$ **<STUD>** Text **</STUD>**

Here, the left hand side of each production denotes one non-terminals (e.g., "Roll", "Class", etc.)

Those symbols on the right hand side for which no productions (i.e., substitutions) are defined are terminals (e.g., 'a', 'b', '|', '<', '>', "ROLL", etc.)



Structure of a production



The above is same as:

1. $A \Longrightarrow \alpha_1$
2. $A \Longrightarrow \alpha_2$
3. $A \Longrightarrow \alpha_3$
- ...
- K. $A \Longrightarrow \alpha_k$

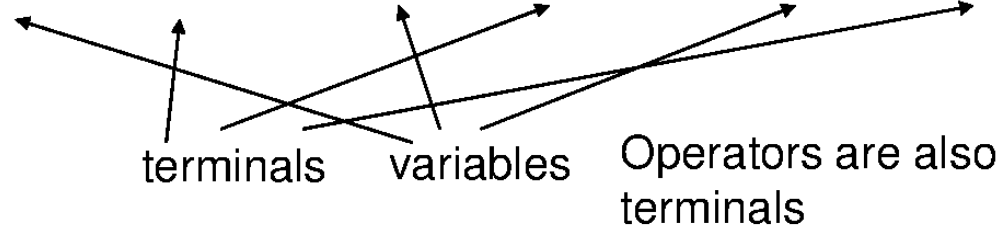


CFG conventions

- Terminal symbols $\leq a, b, c \dots$
- Non-terminal symbols $\leq A, B, C, \dots$
- Terminal or non-terminal symbols $\leq X, Y, Z$
- Terminal strings $\leq w, x, y, z$
- Arbitrary strings of terminals and non-terminals $\leq \alpha, \beta, \gamma, \dots$

Syntactic Expressions in Programming Languages

*result = a*b + score + 10 * distance + c*



Regular languages have only terminals

- Reg expression = $[a-z][a-z0-1]^*$
- If we allow only letters a & b, and 0 & 1 for constants (for simplification)
 - Regular expression = $(a+b)(a+b0+1)^*$

String membership

How to say if a string belong to the language defined by a CFG?

1. Derivation
 - Head to body
2. Recursive inference
 - Body to head

Both are equivalent forms

Example:

- $w = 01110$
- Is w a palindrome?

G:

$A \Rightarrow 0A0 \mid 1A1 \mid 0 \mid 1 \mid \epsilon$

$A \Rightarrow 0A0$
 $\Rightarrow 01A10$
 $\Rightarrow 01110$



Simple Expressions...

- We can write a CFG for accepting simple expressions
- $G = (V, T, P, S)$
 - $V = \{E, F\}$
 - $T = \{0, 1, a, b, +, *, (,)\}$
 - $S = \{E\}$
 - P :
 - $E \Rightarrow E + E \mid E * E \mid (E) \mid F$
 - $F \Rightarrow aF \mid bF \mid 0F \mid 1F \mid a \mid b \mid 0 \mid 1$



Generalization of derivation

- Derivation is *head* $\Rightarrow$ *body*
- $A \Rightarrow X$ (A derives X in a single step)
- $A \Rightarrow_G^* X$ (A derives X in a multiple steps)
- Transitivity:
If $A \Rightarrow_G^* B$, and $B \Rightarrow_G^* C$, THEN $A \Rightarrow_G^* C$



Context-Free Language

- The language of a CFG, $G=(V,T,P,S)$, denoted by $L(G)$, is the set of terminal strings that have a derivation from the start variable S .
 - $L(G) = \{ w \text{ in } T^* \mid S \Rightarrow^*_G w \}$

Left-most & Right-most Derivation Styles

G:

$E \Rightarrow E + E \mid E * E \mid (E) \mid F$
 $F \Rightarrow aF \mid bF \mid 0F \mid 1F \mid \varepsilon$

Derive the string $a^*(ab+10)$ from G:

$E \Rightarrow_G a^*(ab+10)$

Left-most
derivation:

Always
substitute
leftmost
variable

$\bullet E$
 $\bullet \Rightarrow E * E$
 $\bullet \Rightarrow F * E$
 $\bullet \Rightarrow a * E$
 $\bullet \Rightarrow a * (E)$
 $\bullet \Rightarrow a * (E + E)$
 $\bullet \Rightarrow a * (F + E)$
 $\bullet \Rightarrow a * (aF + E)$
 $\bullet \Rightarrow a * (abF + E)$
 $\bullet \Rightarrow a * (ab + E)$
 $\bullet \Rightarrow a * (ab + F)$
 $\bullet \Rightarrow a * (ab + 1F)$
 $\bullet \Rightarrow a * (ab + 10F)$
 $\bullet \Rightarrow a * (ab + 10)$

Right-most
derivation:

Always
substitute
rightmost
variable

$\bullet E$
 $\bullet \Rightarrow E * E$
 $\bullet \Rightarrow E * (E)$
 $\bullet \Rightarrow E * (E + E)$
 $\bullet \Rightarrow E * (E + F)$
 $\bullet \Rightarrow E * (E + 1F)$
 $\bullet \Rightarrow E * (E + 10F)$
 $\bullet \Rightarrow E * (E + 10)$
 $\bullet \Rightarrow E * (F + 10)$
 $\bullet \Rightarrow E * (aF + 10)$
 $\bullet \Rightarrow E * (abF + 10)$
 $\bullet \Rightarrow E * (ab + 10)$
 $\bullet \Rightarrow F * (ab + 10)$
 $\bullet \Rightarrow aF * (ab + 10)$
 $\bullet \Rightarrow a * (ab + 10)$



Leftmost vs. Rightmost derivations

Q1) For every leftmost derivation, there is a rightmost derivation, and vice versa. True or False?

True - will use parse trees to prove this

Q2) Does every word generated by a CFG have a leftmost and a rightmost derivation?

Yes – easy to prove (reverse direction)

Q3) Could there be words which have more than one leftmost (or rightmost) derivation?

Yes – depending on the grammar



How to prove that your CFGs are correct?

(using induction)



CFG & CFL

$$\begin{array}{l} \underline{G_{\text{pal}}}: \\ A \Rightarrow 0A0 \mid 1A1 \mid 0 \mid 1 \mid \epsilon \end{array}$$

- Theorem: A string w in $(0+1)^*$ is in $L(G_{\text{pal}})$, if and only if, w is a palindrome.

- Proof:
 - Use induction
 - on string length for the IF part
 - On length of derivation for the ONLY IF part

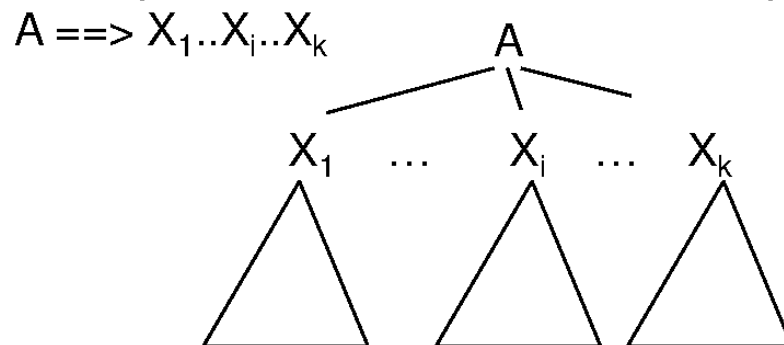


Parse trees

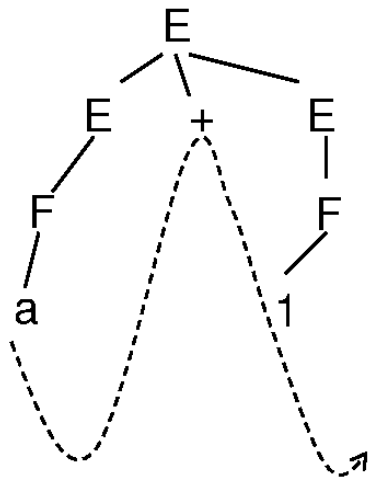
Parse Trees

- Each CFG can be represented using a *parse tree*:
 - Each internal node is labeled by a variable in V
 - Each leaf is terminal symbol
 - For a production, $A \Rightarrow X_1 X_2 \dots X_k$, then any internal node labeled A has k children which are labeled from $X_1, X_2, \dots, X_k$ from left to right

Parse tree for production and all other subsequent productions:



Examples



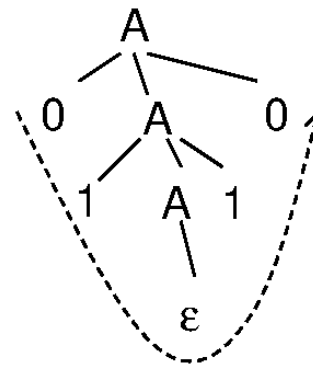
Parse tree for $a + 1$

G:

$E \Rightarrow E + E \mid E * E \mid (E) \mid F$

$F \Rightarrow aF \mid bF \mid 0F \mid 1F \mid 0 \mid 1 \mid a \mid b$

Recursive inference



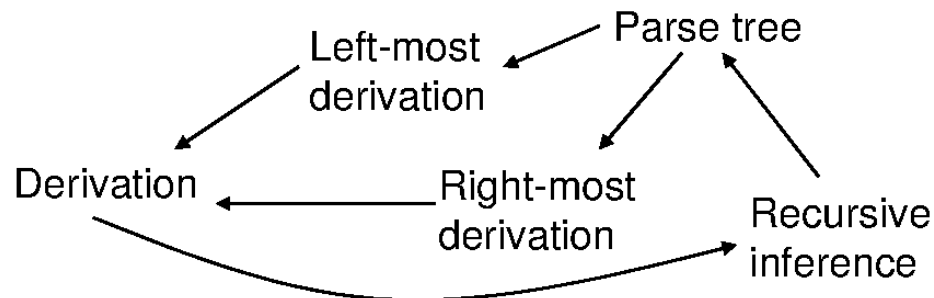
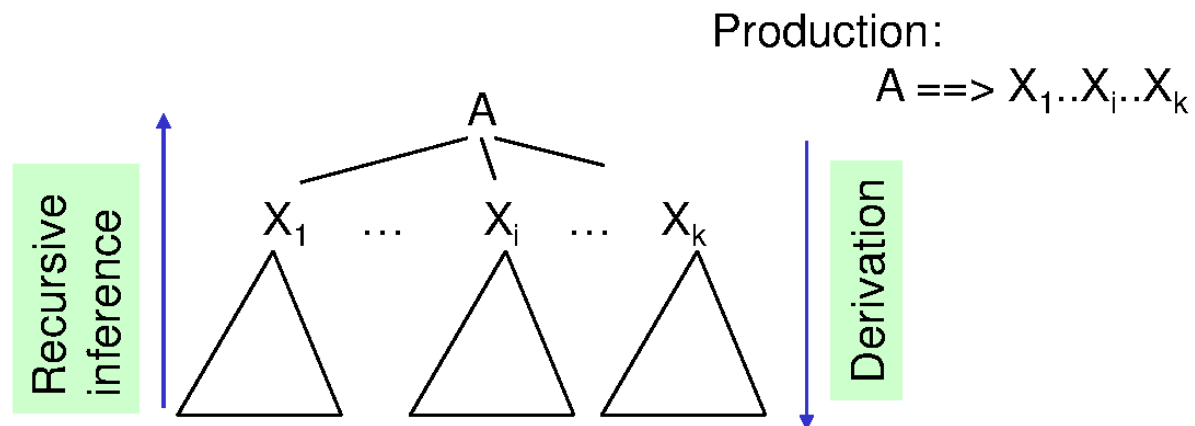
Parse tree for 0110

Derivation

G:

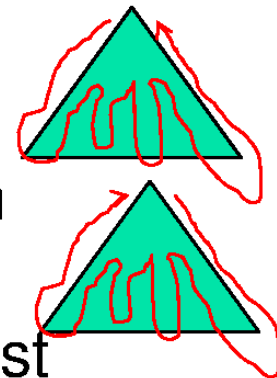
$A \Rightarrow 0A0 \mid 1A1 \mid 0 \mid 1 \mid \epsilon$

Parse Trees, Derivations, and Recursive Inferences

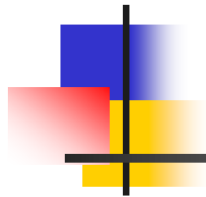


Interchangeability of different CFG representations

- Parse tree $\Rightarrow$ left-most derivation
 - DFS left to right
- Parse tree $\Rightarrow$ right-most derivation
 - DFS right to left
- $\Rightarrow$ left-most derivation $\equiv$ right-most derivation
- Derivation $\Rightarrow$ Recursive inference
 - Reverse the order of productions
- Recursive inference $\Rightarrow$ Parse trees
 - bottom-up traversal of parse tree



Connection between CFLs and RLs



What kind of grammars result for regular languages?

CFLs & Regular Languages

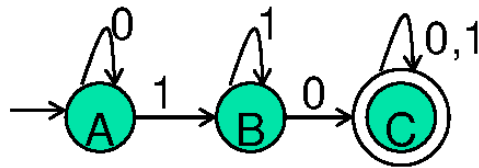
- A CFG is said to be *right-linear* if all the productions are one of the following two forms: $A \Rightarrow wB$ (or) $A \Rightarrow w$

Where:

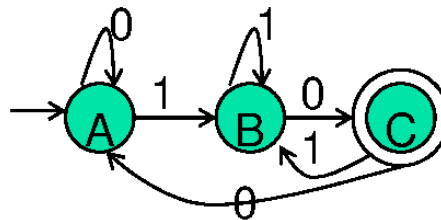
- A & B are variables,
- w is a string of terminals

- Theorem 1: Every right-linear CFG generates a regular language
- Theorem 2: Every regular language has a right-linear grammar
- Theorem 3: Left-linear CFGs also represent RLs

Some Examples



Right linear CFG?



Right linear CFG?

➤ $A \Rightarrow 01B \mid C$
 $B \Rightarrow 11B \mid 0C \mid 1A$
 $C \Rightarrow 1A \mid 0 \mid 1$

Finite Automaton?



Summary

- Context-free grammars
- Context-free languages
- Productions, derivations, recursive inference, parse trees
- Left-most & right-most derivations



FORMAL LANGUAGES AND AUTOMATA THEORY

UNIT 3



Ambiguity in CFGs and CFLs



Ambiguity in CFGs

- A CFG is said to be *ambiguous* if there exists a string which has more than one left-most derivation

Example:

$S \Rightarrow AS \mid \varepsilon$

$A \Rightarrow A1 \mid 0A1 \mid 01$

Input string: 00111

Can be derived in two ways

LM derivation #1:

$S \Rightarrow AS$
 $\Rightarrow 0A1S$
 $\Rightarrow 00A11S$
 $\Rightarrow 00111S$
 $\Rightarrow 00111$

LM derivation #2:

$S \Rightarrow AS$
 $\Rightarrow A1S$
 $\Rightarrow 0A11S$
 $\Rightarrow 00111S$
 $\Rightarrow 00111$



Why does ambiguity matter?

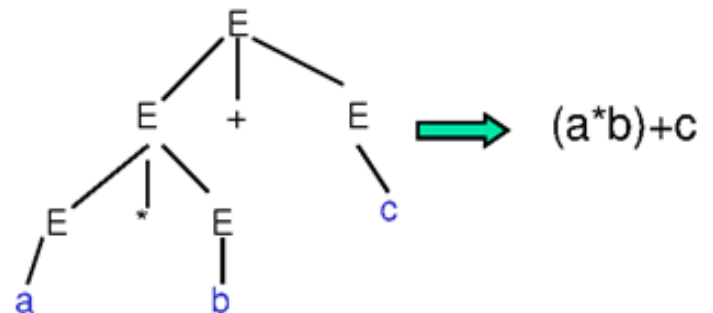
$E \Rightarrow E + E \mid E * E \mid (E) \mid a \mid b \mid c \mid 0 \mid 1$

Values are
different !!!

string = $a * b + c$

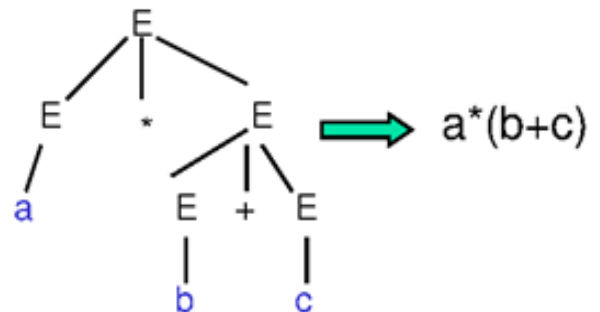
• LM derivation #1:

• $E \Rightarrow E + E \Rightarrow E * E + E \Rightarrow a * b + c$



• LM derivation #2

• $E \Rightarrow E * E \Rightarrow a * E \Rightarrow a * E + E \Rightarrow a * b + c$



The calculated value depends on which of the two parse trees is actually used.



Removing Ambiguity in Expression Evaluations

- It MAY be possible to remove ambiguity for some CFLs
 - E.g., in a CFG for expression evaluation by imposing rules & restrictions such as precedence
 - This would imply rewrite of the grammar

Modified unambiguous version:

- Precedence: $()$, $*$, $+$

$$\begin{aligned} E &\Rightarrow E + T \mid T \\ T &\Rightarrow T * F \mid F \\ F &\Rightarrow I \mid (E) \\ I &\Rightarrow a \mid b \mid c \mid 0 \mid 1 \end{aligned}$$

Ambiguous version:

$$E \Rightarrow E + E \mid E * E \mid (E) \mid a \mid b \mid c \mid 0 \mid 1$$

How will this avoid ambiguity?



Inherently Ambiguous CFLs

- However, for some languages, it may not be possible to remove ambiguity
- A CFL is said to be *inherently ambiguous* if every CFG that describes it is ambiguous

Example:

- $L = \{ a^n b^n c^m d^m \mid n, m \geq 1 \} \cup \{ a^n b^m c^m d^n \mid n, m \geq 1 \}$
- L is inherently ambiguous
- Why?

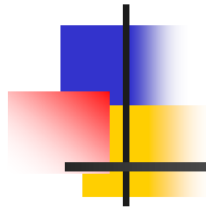
Input string: $a^n b^n c^n d^n$



Summary

- Ambiguous grammars
- Removing ambiguity

Properties of Context-free Languages





Topics

- 1) Simplifying CFGs, Normal forms
- 2) Pumping lemma for CFLs
- 3) Closure and decision properties of CFLs



How to “simplify” CFGs?



Three ways to simplify/clean a CFG

(clean)

1. Eliminate *useless symbols*

(simplify)

2. Eliminate ϵ -productions

$A \not\Rightarrow \epsilon$

3. Eliminate *unit productions*

$A \not\Rightarrow B$



Eliminating useless symbols

Grammar cleanup



Eliminating *useless symbols*

A symbol X is reachable if there exists:

- $S \rightarrow^* \alpha X \beta$

A symbol X is generating if there exists:

- $X \rightarrow^* w,$
 - for some $w \in T^*$

For a symbol X to be “useful”, it has to be both reachable *and* generating

- $S \rightarrow^* \alpha X \beta \rightarrow^* w',$ for some $w' \in T^*$

reachable generating



Algorithm to detect useless symbols

1. First, eliminate all symbols that are *not* generating
2. Next, eliminate all symbols that are *not* reachable

Is the order of these steps important,
or can we switch?



Example: Useless symbols

- $S \rightarrow AB \mid a$
- $A \rightarrow b$

1. A, S are generating
2. B is *not generating* (and therefore B is useless)
3. $\Rightarrow$ Eliminating B ... (i.e., remove all productions that involve B)
 1. $S \rightarrow a$
 2. $A \rightarrow b$
4. Now, A is *not reachable* and therefore is useless

5. Simplified G:

1. $S \rightarrow a$

What would happen if you reverse the order:
i.e., test reachability before generating?

Will fail to remove:
 $A \rightarrow b$


$$X \rightarrow^* w$$

Algorithm to find all generating symbols

- Given: $G=(V,T,P,S)$
- Basis:
 - Every symbol in T is obviously generating.
- Induction:
 - Suppose for a production $A \rightarrow \alpha$, where α is generating
 - Then, A is also generating


$$S \rightarrow^* \alpha X \beta$$

Algorithm to find all reachable symbols

- Given: $G=(V,T,P,S)$
- Basis:
 - S is obviously reachable (from itself)
- Induction:
 - Suppose for a production $A \rightarrow \alpha_1 \alpha_2 \dots \alpha_k$, where A is reachable
 - Then, all symbols on the right hand side, $\{\alpha_1, \alpha_2, \dots, \alpha_k\}$ are also reachable.



Eliminating ϵ -productions

$A \Rightarrow \epsilon$

X

What's the point of removing ϵ -productions?

$$A \rightarrow \epsilon$$

Eliminating ϵ -productions

Caveat: It is *not* possible to eliminate ϵ -productions for languages which include ϵ in their word set

So we will target the grammar for the *rest of the language*

Theorem: If $G=(V,T,P,S)$ is a CFG for a language L , then $L-\{\epsilon\}$ has a CFG without ϵ -productions

Definition: A is “nullable” if $A \rightarrow^* \epsilon$

- If A is nullable, then any production of the form “ $B \rightarrow CAD$ ” can be simulated by:
 - $B \rightarrow CD \mid CAD$
 - This can allow us to remove ϵ transitions for A



Algorithm to detect all nullable variables

- Basis:

- If $A \rightarrow \varepsilon$ is a production in G , then A is nullable
(note: A can still have other productions)

- Induction:

- If there is a production $B \rightarrow C_1 C_2 \dots C_k$, where *every* C_i is nullable, then B is also nullable



Eliminating ϵ -productions

Given: $G=(V,T,P,S)$

Algorithm:

1. Detect all nullable variables in G
2. Then construct $G_1=(V,T,P_1,S)$ as follows:
 - i. For each production of the form: $A \rightarrow X_1 X_2 \dots X_k$, where $k \geq 1$, suppose m out of the k X_i 's are nullable symbols
 - ii. Then G_1 will have 2^m versions for this production
 - i. i.e, all combinations where each X_i is either present or absent
 - iii. Alternatively, if a production is of the form: $A \rightarrow \epsilon$, then remove it

Example: Eliminating ϵ -productions

- Let L be the language represented by the following CFG G :

- i. $S \rightarrow AB$
- ii. $A \rightarrow aAA \mid \epsilon$
- iii. $B \rightarrow bBB \mid \epsilon$

Simplified
grammar

Goal: To construct G_1 , which is the grammar for $L - \{\epsilon\}$

- Nullable symbols: $\{A, B\}$
- G_1 can be constructed from G as follows:
 - $B \rightarrow b \mid bB \mid bBB$
 - $\implies B \rightarrow b \mid bB \mid bBB$
 - Similarly, $A \rightarrow a \mid aA \mid aAA$
 - Similarly, $S \rightarrow A \mid B \mid AB$

- Note: $L(G) = L(G_1) \cup \{\epsilon\}$

G_1 :

- $S \rightarrow A \mid B \mid AB$
- $A \rightarrow a \mid aA \mid aAA$
- $B \rightarrow b \mid bB \mid bBB$

+

$S \rightarrow \epsilon$

Eliminating unit productions

$A \Rightarrow B$ ← B has to be a variable

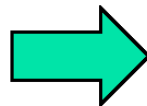
X

What's the point of removing unit transitions ?

Will save #substitutions

E.g.,

$A \Rightarrow B \mid \dots$
$B \Rightarrow C \mid \dots$
$C \Rightarrow D \mid \dots$
$D \Rightarrow xxx \mid yyy \mid zzz$



$A \Rightarrow xxx \mid yyy \mid zzz \mid \dots$
$B \Rightarrow xxx \mid yyy \mid zzz \mid \dots$
$C \Rightarrow xxx \mid yyy \mid zzz \mid \dots$
$D \Rightarrow xxx \mid yyy \mid zzz$

before

after

$$A \rightarrow B$$

Eliminating unit productions

- Unit production is one which is of the form $A \rightarrow B$, where both A & B are variables
- E.g.,
 1. $E \rightarrow T \mid E+T$
 2. $T \rightarrow F \mid T^*F$
 3. $F \rightarrow I \mid (E)$
 4. $I \rightarrow a \mid b \mid Ia \mid Ib \mid I0 \mid I1$
- How to eliminate unit productions?
 - Replace $E \rightarrow T$ with $E \rightarrow F \mid T^*F$
 - Then, upon recursive application wherever there is a unit production:
 - $E \rightarrow F \mid T^*F \mid E+T$ (substituting for T)
 - $E \rightarrow I \mid (E) \mid T^*F \mid E+T$ (substituting for F)
 - $E \rightarrow a \mid b \mid Ia \mid Ib \mid I0 \mid I1 \mid (E) \mid T^*F \mid E+T$ (substituting for I)
 - Now, E has no unit productions
 - Similarly, eliminate for the remainder of the unit productions



The Unit Pair Algorithm: to remove unit productions

- Suppose $A \rightarrow B_1 \rightarrow B_2 \rightarrow \dots \rightarrow B_n \rightarrow \alpha$
- Action: Replace all intermediate productions to produce α directly
 - i.e., $A \rightarrow \alpha; B_1 \rightarrow \alpha; \dots B_n \rightarrow \alpha;$

Definition: (A,B) to be a “*unit pair*” if $A \rightarrow^* B$

- We can find all unit pairs inductively:
 - Basis: Every pair (A,A) is a unit pair (by definition). Similarly, if $A \rightarrow B$ is a production, then (A,B) is a unit pair.
 - Induction: If (A,B) and (B,C) are unit pairs, and $A \rightarrow C$ is also a unit pair.



The Unit Pair Algorithm: to remove unit productions

Input: $G=(V,T,P,S)$

Goal: to build $G_1=(V,T,P_1,S)$ devoid of unit productions

Algorithm:

1. Find all unit pairs in G
2. For each unit pair (A,B) in G :
 1. Add to P_1 a new production $A \rightarrow \alpha$, for every $B \rightarrow \alpha$ which is a *non-unit* production
 2. If a resulting production is already there in P , then there is no need to add it.

Example: eliminating unit productions

G:

1. $E \rightarrow T \mid E+T$
2. $T \rightarrow F \mid T^*F$
3. $F \rightarrow () \mid (E)$
4. $I \rightarrow a \mid b \mid la \mid lb \mid l0 \mid l1$

Unit pairs	Only non-unit productions to be added to P_1
(E,E)	$E \rightarrow E+T$
(E,T)	$E \rightarrow T^*F$
(E,F)	$E \rightarrow (E)$
(E,I)	$E \rightarrow a \mid b \mid la \mid lb \mid l0 \mid l1$
(T,T)	$T \rightarrow T^*F$
(T,F)	$T \rightarrow (E)$
(T,I)	$T \rightarrow a \mid b \mid la \mid lb \mid l0 \mid l1$
(F,F)	$F \rightarrow (E)$
(F,I)	$F \rightarrow a \mid b \mid la \mid lb \mid l0 \mid l1$
(I,I)	$I \rightarrow a \mid b \mid la \mid lb \mid l0 \mid l1$

G_1 :

1. $E \rightarrow E+T \mid T^*F \mid (E) \mid a \mid b \mid la \mid lb \mid l0 \mid l1$
2. $T \rightarrow T^*F \mid (E) \mid a \mid b \mid la \mid lb \mid l0 \mid l1$
3. $F \rightarrow (E) \mid a \mid b \mid la \mid lb \mid l0 \mid l1$
4. $I \rightarrow a \mid b \mid la \mid lb \mid l0 \mid l1$



Putting all this together...

- Theorem: If G is a CFG for a language that contains at least one string other than ϵ , then there is another CFG G_1 , such that $L(G_1) = L(G) - \epsilon$, and G_1 has:
 - no ϵ -productions
 - no unit productions
 - no useless symbols
- Algorithm:
 - Step 1) eliminate ϵ -productions
 - Step 2) eliminate unit productions
 - Step 3) eliminate useless symbols

Again,
the order is
important!

Why?



Normal Forms



Why normal forms?

- If all productions of the grammar could be expressed in the same form(s), then:
 - a. It becomes easy to design algorithms that use the grammar
 - b. It becomes easy to show proofs and properties



Chomsky Normal Form (CNF)

Let G be a CFG for some $L - \{\epsilon\}$

Definition:

G is said to be in **Chomsky Normal Form** if all its productions are in one of the following two forms:

- i. $A \rightarrow BC$ where A, B, C are variables, or
- ii. $A \rightarrow a$ where a is a terminal

- G has no useless symbols
- G has no unit productions
- G has no ϵ -productions



CNF checklist

Is this grammar in CNF?

G₁:

1. $E \rightarrow E+T \mid T^*F \mid (E) \mid Ia \mid Ib \mid I0 \mid I1$
2. $T \rightarrow T^*F \mid (E) \mid Ia \mid Ib \mid I0 \mid I1$
3. $F \rightarrow (E) \mid Ia \mid Ib \mid I0 \mid I1$
4. $I \rightarrow a \mid b \mid Ia \mid Ib \mid I0 \mid I1$

Checklist:

- G has no ϵ -productions ✓
- G has no unit productions ✓
- G has no useless symbols ✓
- But...
 - the normal form for productions is violated



So, the grammar is not in CNF

How to convert a G into CNF?

- Assumption: G has no ϵ -productions, unit productions or useless symbols
- 1) For every terminal a that appears in the body of a production:
 - i. create a unique variable, say X_a , with a production $X_a \rightarrow a$, and
 - ii. replace all other instances of a in G by X_a
- 2) Now, all productions will be in one of the following two forms:
 - $A \rightarrow B_1 B_2 \dots B_k$ ($k \geq 3$) or $A \rightarrow a$
- 3) Replace each production of the form $A \rightarrow B_1 B_2 B_3 \dots B_k$ by:

$\begin{array}{c} \xleftrightarrow{\quad} \\ B_2 \quad C_2 \\ \xleftrightarrow{\quad} \\ B_1 \quad C_1 \end{array}$
 and so on...

 - $A \rightarrow B_1 C_1 \quad C_1 \rightarrow B_2 C_2 \quad \dots \quad C_{k-3} \rightarrow B_{k-2} C_{k-2} \quad C_{k-2} \rightarrow B_{k-1} B_k$

Example #1

G:

$S \Rightarrow AS \mid BABC$

$A \Rightarrow A1 \mid 0A1 \mid 01$

$B \Rightarrow 0B \mid 0$

$C \Rightarrow 1C \mid 1$



G in CNF:

$X_0 \Rightarrow 0$

$X_1 \Rightarrow 1$

$S \Rightarrow AS \mid BY_1$

$Y_1 \Rightarrow AY_2$

$Y_2 \Rightarrow BC$

$A \Rightarrow AX_1 \mid X_0Y_3 \mid X_0X_1$

$Y_3 \Rightarrow AX_1$

$B \Rightarrow X_0B \mid 0$

$C \Rightarrow X_1C \mid 1$

All productions are of the form: $A \Rightarrow BC$ or $A \Rightarrow a$

Example #2

G:

1. $E \rightarrow E+T \mid T^*F \mid (E) \mid I_a \mid I_b \mid I_0 \mid I_1$
2. $T \rightarrow T^*F \mid (E) \mid I_a \mid I_b \mid I_0 \mid I_1$
3. $F \rightarrow (E) \mid I_a \mid I_b \mid I_0 \mid I_1$
4. $I \rightarrow a \mid b \mid I_a \mid I_b \mid I_0 \mid I_1$

Step (1)

1. $E \rightarrow EX_+T \mid TX_+F \mid X_+(EX_+) \mid IX_a \mid IX_b \mid IX_0 \mid IX_1$
2. $T \rightarrow TX_+F \mid X_+(EX_+) \mid IX_a \mid IX_b \mid IX_0 \mid IX_1$
3. $F \rightarrow X_+(EX_+) \mid IX_a \mid IX_b \mid IX_0 \mid IX_1$
4. $I \rightarrow X_a \mid X_b \mid IX_a \mid IX_b \mid IX_0 \mid IX_1$
5. $X_+ \rightarrow +$
6. $X_* \rightarrow *$
7. $X_+ \rightarrow +$
8. $X_+ \rightarrow ($
9.

Step (2)

1. $E \rightarrow EC_1 \mid TC_2 \mid X_+C_3 \mid IX_a \mid IX_b \mid IX_0 \mid IX_1$
2. $C_1 \rightarrow X_+T$
3. $C_2 \rightarrow X_*F$
4. $C_3 \rightarrow EX_+$
5. $T \rightarrow \dots\dots\dots$
6.

Languages with ϵ

- For languages that include ϵ ,
 - Write down the rest of grammar in CNF
 - Then add production “ $S \Rightarrow \epsilon$ ” at the end

E.g., consider:

G:

$S \Rightarrow AS \mid BABC$

$A \Rightarrow A1 \mid 0A1 \mid 01 \mid \epsilon$

$B \Rightarrow 0B \mid 0 \mid \epsilon$

$C \Rightarrow 1C \mid 1 \mid \epsilon$



G in CNF:

$X_0 \Rightarrow 0$

$X_1 \Rightarrow 1$

$S \Rightarrow AS \mid BY_1 \mid \epsilon$

$Y_1 \Rightarrow AY_2$

$Y_2 \Rightarrow BC$

$A \Rightarrow AX_1 \mid X_0Y_3 \mid X_0X_1$

$Y_3 \Rightarrow AX_1$

$B \Rightarrow X_0B \mid 0$

$C \Rightarrow X_1C \mid 1$



Other Normal Forms

- Griebach Normal Form (GNF)
 - All productions of the form

$$A \Rightarrow a \alpha$$



Return of the Pumping Lemma !!

Think of languages that cannot be CFL

== think of languages for which a stack will not be enough

e.g., the language of strings of the form ww



Why pumping lemma?

- A result that will be useful in proving languages that *are not* CFLs
 - (just like we did for regular languages)
- But before we prove the pumping lemma for CFLs
 - Let us first prove an important property about parse trees

Observe that any parse tree generated by a CNF will be a binary tree, where all internal nodes have exactly two children (except those nodes connected to the leaves).

The “parse tree theorem”

Given:

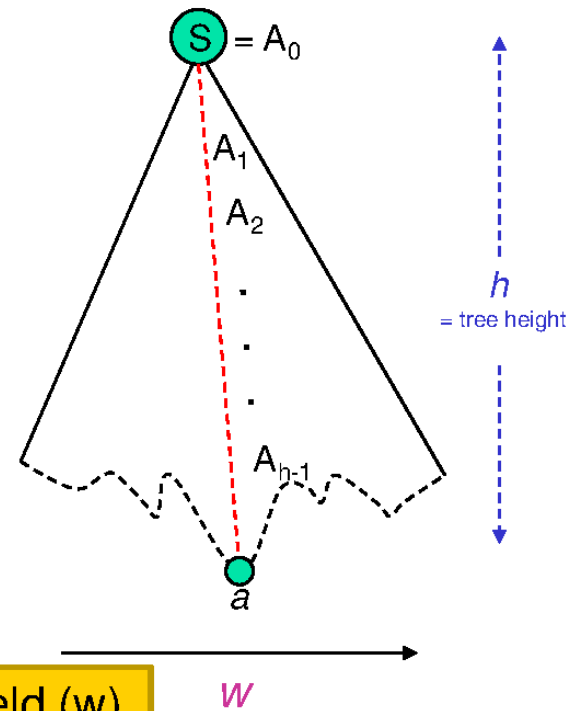
- Suppose we have a parse tree for a string w , according to a CNF grammar, $G=(V,T,P,S)$
- Let h be the height of the parse tree

Implies:

- $|w| \leq 2^{h-1}$

In other words, a CNF parse tree's string yield (w) can no longer be 2^{h-1}

Parse tree for w



To show: $|w| \leq 2^{h-1}$

Proof...The size of parse trees

Proof: (using induction on h)

Basis: $h = 1$

- Derivation will have to be " $S \rightarrow a$ "
- $|w| = 1 = 2^{1-1}$.

Ind. Hyp: $h = k-1$

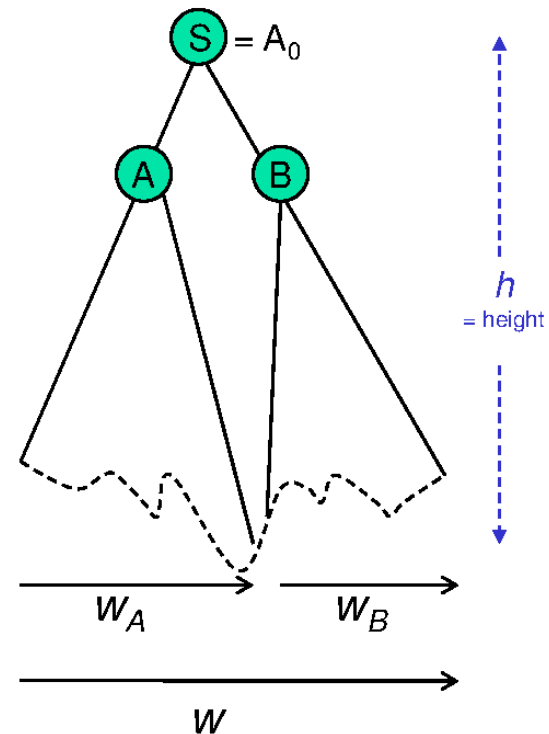
- $|w| \leq 2^{k-2}$

Ind. Step: $h = k$

S will have exactly two children:
 $S \rightarrow AB$

- Heights of A & B subtrees are at most $h-1$
- $w = w_A w_B$, where $|w_A| \leq 2^{k-2}$ and $|w_B| \leq 2^{k-2}$
- $|w| \leq 2^{k-1}$

Parse tree for w





Implication of the Parse Tree Theorem (assuming CNF)

Fact:

- If the height of a parse tree is h , then
 - $\implies |w| \leq 2^{h-1}$

Implication:

- If $|w| \geq 2^h$, then
 - Its parse tree's height is *at least* $h+1$



The Pumping Lemma for CFLs

Let L be a CFL.

Then there exists a constant N , s.t.,

- if $z \in L$ s.t. $|z| \geq N$, then we can write $z = uvwx^ky$, such that:

1. $|vwx| \leq N$
2. $vx \neq \epsilon$
3. For all $k \geq 0$: $uv^kwx^ky \in L$

Note: we are pumping in two places (v & x)



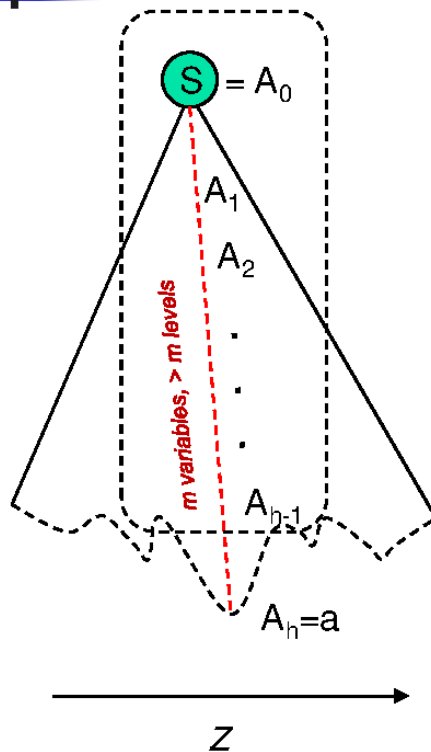
Proof: Pumping Lemma for CFL

- If $L = \Phi$ or contains only ϵ , then the lemma is trivially satisfied (as it cannot be violated)
- For any other L which is a CFL:
 - Let G be a CNF grammar for L
 - Let m = number of variables in G
 - Choose $N = 2^m$.
 - Pick any $z \in L$ s.t. $|z| \geq N$
 - ➔ the parse tree for z should have a height $\geq m+1$
(by the parse tree theorem)

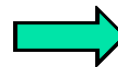
Parse tree for z

Meaning:
Repetition in the
last $m+1$ variables

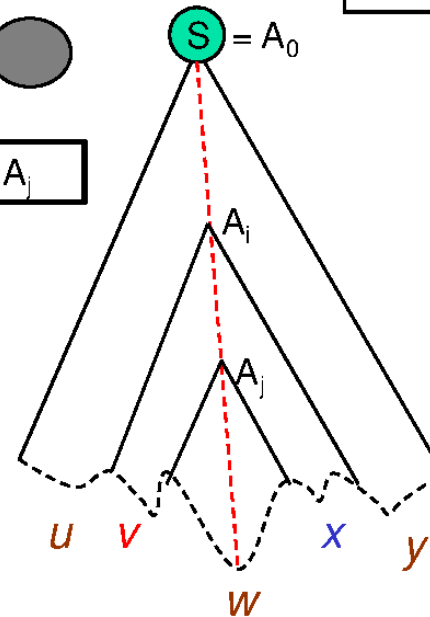
$$h-m \leq i < j \leq h$$



$$h \geq m+1$$



$$A_i = A_j$$

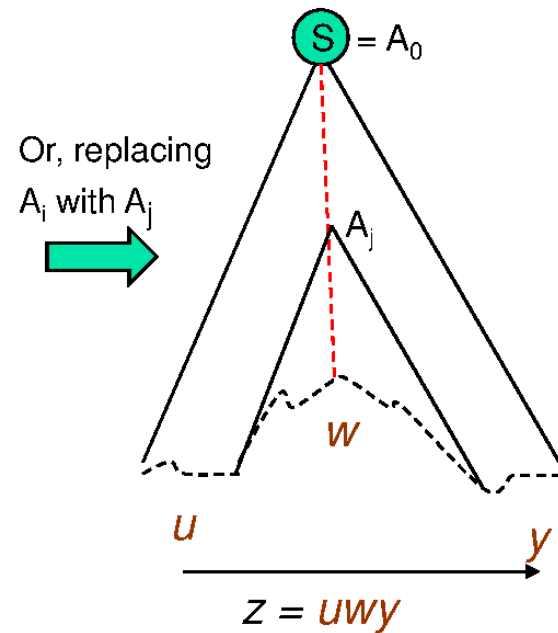
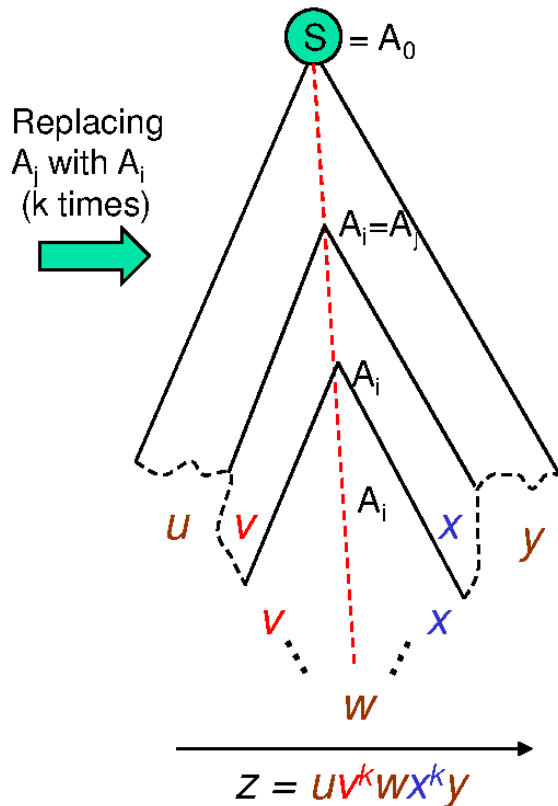


$$h \geq m+1$$

$$m+1$$

- Therefore, $vx \neq \epsilon$

Extending the parse tree...



$\Rightarrow$ For all $k \geq 0$: $uv^kwx^ky \in L$



Proof contd..

- Also, since A_i 's subtree no taller than $m+1$

$\Rightarrow$ the string generated under A_i 's subtree, which is vwX , cannot be longer than $2^m (=N)$

But, $2^m = N$

$\Rightarrow |vwX| \leq N$

This completes the proof for the pumping lemma.



Application of Pumping Lemma for CFLs

Example 1: $L = \{a^m b^m c^m \mid m > 0\}$

Claim: L is not a CFL

Proof:

- Let $N \leq P/L$ constant
- Pick $z = a^N b^N c^N$
- Apply pumping lemma to z and show that there exists at least one other string constructed from z (obtained by pumping up or down) that is $\notin L$



Proof contd...

- $z = uvwxy$
- As $z = a^N b^N c^N$ and $|vwx| \leq N$ and $vx \neq \epsilon$
 - $\Rightarrow v, x$ cannot contain all three symbols (a,b,c)
 - $\Rightarrow$ we can pump up or pump down to build another string which is $\notin L$



Example #2 for P/L application

- $L = \{ ww \mid w \text{ is in } \{0,1\}^* \}$
- Show that L is not a CFL
 - Try string $z = 0^N 0^N$
 - what happens?
 - Try string $z = 0^N 1^N 0^N 1^N$
 - what happens?



Example 3

- $L = \{ 0^{k^2} \mid k \text{ is any integer} \}$
- Prove L is not a CFL using Pumping Lemma



Example 4

- $L = \{a^i b^j c^k \mid i < j < k\}$
- Prove that L is not a CFL



CFL Closure Properties



Closure Property Results

- CFLs are closed under:
 - Union
 - Concatenation
 - Kleene closure operator
 - Substitution
 - Homomorphism, inverse homomorphism
 - reversal

-
- CFLs are *not* closed under:
 - Intersection
 - Difference
 - Complementation

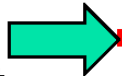
Note: Reg languages
are closed
under
these
operators

Strategy for Closure Property Proofs

- First prove “closure under **substitution**”
- Using the above result, prove other closure properties
- CFLs are closed under:

- Union ←
- Concatenation ←
- Kleene closure operator ←
- Substitution →
- Homomorphism, inverse homomorphism ←
- Reversal

Prove
this first



Note: $s(L)$ can use
a different alphabet

The *Substitution* operation

For each $a \in \Sigma$, then let $s(a)$ be a language
If $w = a_1 a_2 \dots a_n \in L$, then:

$$s(w) = \{ x_1 x_2 \dots \} \in s(L), \text{ s.t., } x_i \in s(a_i)$$

Example:

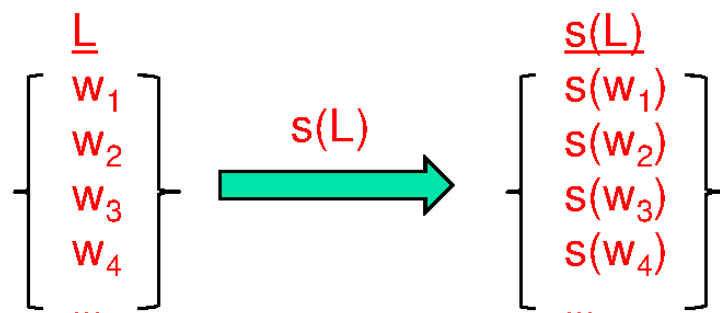
- Let $\Sigma = \{0, 1\}$
- Let: $s(0) = \{a^n b^n \mid n \geq 1\}$, $s(1) = \{aa, bb\}$
- If $w = 01$, $s(w) = s(0) \cdot s(1)$
 - E.g., $s(w)$ contains $a^1 b^1 aa$, $a^1 b^1 bb$,
 $a^2 b^2 aa$, $a^2 b^2 bb$,
... and so on.

CFLs are closed under Substitution

IF L is a CFL and a substitution defined on L , $s(L)$, is s.t., $s(a)$ is a CFL for every symbol a , THEN:

- $s(L)$ is also a CFL

What is $s(L)$?

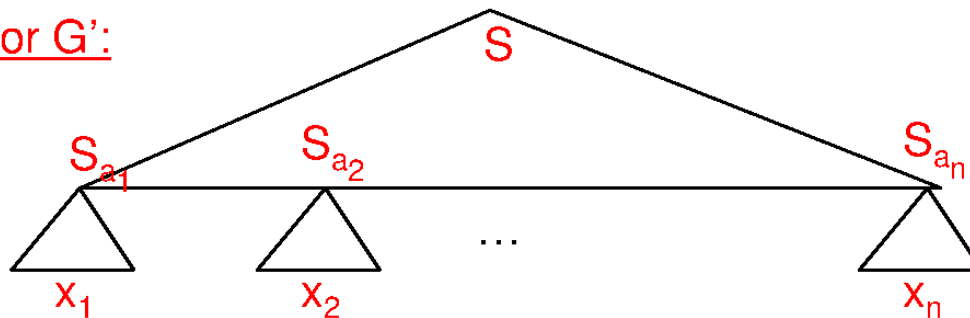


Note: each $s(w)$ is itself a set of strings

CFLs are closed under *Substitution*

- $G=(V,T,P,S)$: CFG for L
- Because every $s(a)$ is a CFL, there is a CFG for each $s(a)$
 - Let $G_a = (V_a, T_a, P_a, S_a)$
- Construct $G'=(V',T',P',S)$ for $s(L)$
- P' consists of:
 - The productions of P , but with every occurrence of terminal “ a ” in their bodies replaced by S_a .
 - All productions in any P_a , for any $a \in \Sigma$

Parse tree for G' :



Substitution of a CFL: example

- Let L = language of binary palindromes s.t., substitutions for 0 and 1 are defined as follows:
 - $s(0) = \{a^n b^n \mid n \geq 1\}$, $s(1) = \{xx, yy\}$
- Prove that $s(L)$ is also a CFL.

CFG for L :

$S \Rightarrow 0S0 \mid 1S1 \mid \epsilon$

CFG for $s(0)$:

$S_0 \Rightarrow aS_0b \mid ab$

CFG for $s(1)$:

$S_1 \Rightarrow xx \mid yy$



Therefore, CFG for $s(L)$:

$S \Rightarrow S_0 S S_0 \mid S_1 S S_1 \mid \epsilon$

$S_0 \Rightarrow aS_0b \mid ab$

$S_1 \Rightarrow xx \mid yy$



CFLs are closed under *union*

Let L_1 and L_2 be CFLs

To show: $L_1 \cup L_2$ is also a CFL

Let us show by using the result of *Substitution*

- Make a new language:

- $L_{\text{new}} = \{a, b\}$ s.t., $s(a) = L_1$ and $s(b) = L_2$

- $\implies s(L_{\text{new}}) == \text{same as} == L_1 \cup L_2$



-
- A more direct, alternative proof

- Let S_1 and S_2 be the starting variables of the grammars for L_1 and L_2

- Then, $S_{\text{new}} \Rightarrow S_1 \mid S_2$



CFLs are closed under *concatenation*

- Let L_1 and L_2 be CFLs

Let us show by using the result of *Substitution*

- Make $L_{\text{new}} = \{ab\}$ s.t.,
 $s(a) = L_1$ and $s(b) = L_2$
 $\implies L_1 L_2 = s(L_{\text{new}})$
-

- A proof without using substitution?



CFLs are closed under *Kleene Closure*

- Let L be a CFL
- Let $L_{\text{new}} = \{a\}^*$ and $s(a) = L_1$
 - Then, $L^* = s(L_{\text{new}})$

CFLs are closed under *Reversal*

- Let L be a CFL, with grammar $G=(V,T,P,S)$
- For L^R , construct $G^R=(V,T,P^R,S)$ s.t.,
 - If $A \Rightarrow \alpha$ is in P , then:
 - $A \Rightarrow \alpha^R$ is in P^R
 - (that is, reverse every production)

CFLs are *not* closed under Intersection

- Existential proof:
 - $L_1 = \{0^n 1^n 2^i \mid n \geq 1, i \geq 1\}$
 - $L_2 = \{0^i 1^n 2^n \mid n \geq 1, i \geq 1\}$
- Both L_1 and L_2 are CFLs
 - Grammars?
- But $L_1 \cap L_2$ *cannot* be a CFL
 - Why?
- We have an example, where intersection is not closed.
- Therefore, CFLs are not closed under intersection

CFLs are not closed under complementation

- Follows from the fact that CFLs are not closed under intersection

- $$L_1 \cap L_2 = \overline{\overline{L_1} \cup \overline{L_2}}$$

Logic: if CFLs were to be closed under complementation

- the whole right hand side becomes a CFL (because CFL is closed for union)
- the left hand side (intersection) is also a CFL
- but we just showed CFLs are NOT closed under intersection!
- CFLs cannot be closed under complementation.



CFLs are not closed under difference

- Follows from the fact that CFLs are not closed under complementation
- Because, if CFLs are closed under difference, then:
 - $\bar{L} = \Sigma^* - L$
 - So $\bar{L}$ has to be a CFL too
 - Contradiction



Decision Properties

- Emptiness test
 - Generating test
 - Reachability test
- Membership test
 - PDA acceptance



“Undecidable” problems for CFL

- Is a given CFG G ambiguous?
- Is a given CFL inherently ambiguous?
- Is the intersection of two CFLs empty?
- Are two CFLs the same?
- Is a given $L(G)$ equal to Σ^* ?



Summary

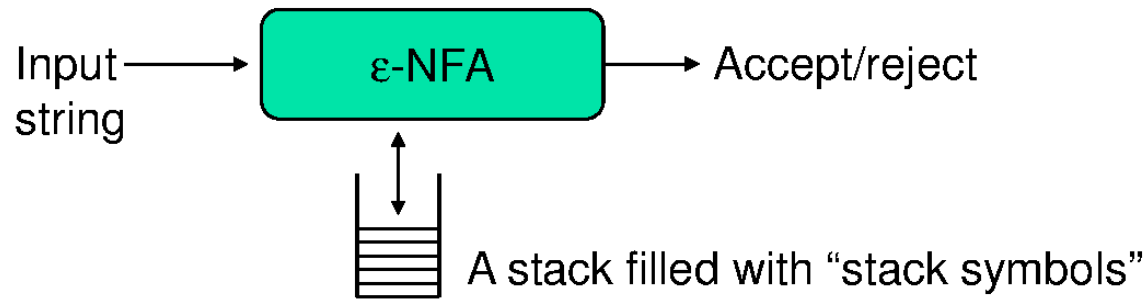
- Normal Forms
 - Chomsky Normal Form
 - Greibach Normal Form
 - Useful in proving P/L
- Pumping Lemma for CFLs
 - Main difference: $z = uv^iwx^iy$
- Closure properties
 - Closed under: union, concatenation, reversal, Kleen closure, homomorphism, substitution
 - Not closed under: intersection, complementation, difference



Pushdown Automata (PDA)

PDA - the automata for CFLs

- What is?
 - FA to Reg Lang, PDA is to CFL
- PDA == [ϵ -NFA + “a stack”]
- Why a stack?





Pushdown Automata - Definition

- A PDA $P := (Q, \Sigma, \Gamma, \delta, q_0, Z_0, F)$:
 - Q : states of the ε -NFA
 - Σ : input alphabet
 - Γ : stack symbols
 - δ : transition function
 - q_0 : start state
 - Z_0 : Initial stack top symbol
 - F : Final/accepting states

$$\delta : \overset{\text{old state}}{Q} \times \overset{\text{Stack top}}{\Gamma} \times \overset{\text{input symb.}}{\Sigma} \Rightarrow \overset{\text{new state(s)}}{Q} \times \overset{\text{new Stack top(s)}}{\Gamma}$$

δ : The Transition Function

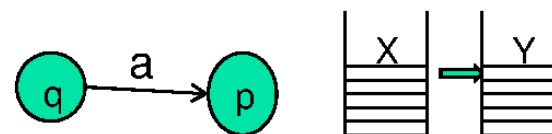
$$\delta(q, a, X) = \{(p, Y), \dots\}$$

2. state transition from q to p
 3. a is the next input symbol
 4. X is the current stack top symbol

Y is the replacement for X; it is in Γ^* (a string of stack symbols)

- i. Set $Y = \epsilon$ for: Pop(X)
- ii. If $Y = X$:
stack top is unchanged
- iii. If $Y = Z_1 Z_2 \dots Z_k$: X is popped and is replaced by Y in reverse order (i.e., Z_1 will be the

new stack top)



	Y = ?	Action
i)	$Y = \epsilon$	Pop(X)
ii)	$Y = X$	Pop(X) Push(X)
iii)	$Y = Z_1 Z_2 \dots Z_k$	Pop(X) Push(Z_k) Push(Z_{k-1}) ... Push(Z_2) Push(Z_1)



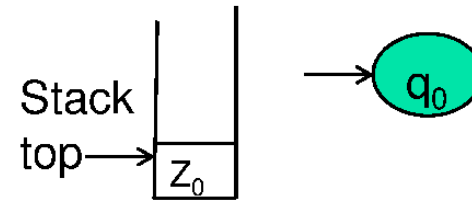
Example

Let $L_{ww^R} = \{ww^R \mid w \text{ is in } (0+1)^*\}$

- CFG for L_{ww^R} : $S \Rightarrow 0S0 \mid 1S1 \mid \epsilon$
- PDA for L_{ww^R} :
- $P := (Q, \Sigma, \Gamma, \delta, q_0, Z_0, F)$
 $= (\{q_0, q_1, q_2\}, \{0, 1\}, \{0, 1, Z_0\}, \delta, q_0, Z_0, \{q_2\})$

PDA for L_{ww^R}

Initial state of the PDA:



1. $\delta(q_0, 0, Z_0) = \{(q_0, 0Z_0)\}$
2. $\delta(q_0, 1, Z_0) = \{(q_0, 1Z_0)\}$
3. $\delta(q_0, 0, 0) = \{(q_0, 00)\}$
4. $\delta(q_0, 0, 1) = \{(q_0, 01)\}$
5. $\delta(q_0, 1, 0) = \{(q_0, 10)\}$
6. $\delta(q_0, 1, 1) = \{(q_0, 11)\}$
7. $\delta(q_0, \epsilon, 0) = \{(q_1, 0)\}$
8. $\delta(q_0, \epsilon, 1) = \{(q_1, 1)\}$
9. $\delta(q_0, \epsilon, Z_0) = \{(q_1, Z_0)\}$
10. $\delta(q_1, 0, 0) = \{(q_1, \epsilon)\}$
11. $\delta(q_1, 1, 1) = \{(q_1, \epsilon)\}$
12. $\delta(q_1, \epsilon, Z_0) = \{(q_2, Z_0)\}$

First symbol push on stack

Grow the stack by pushing new symbols on top of old (w-part)

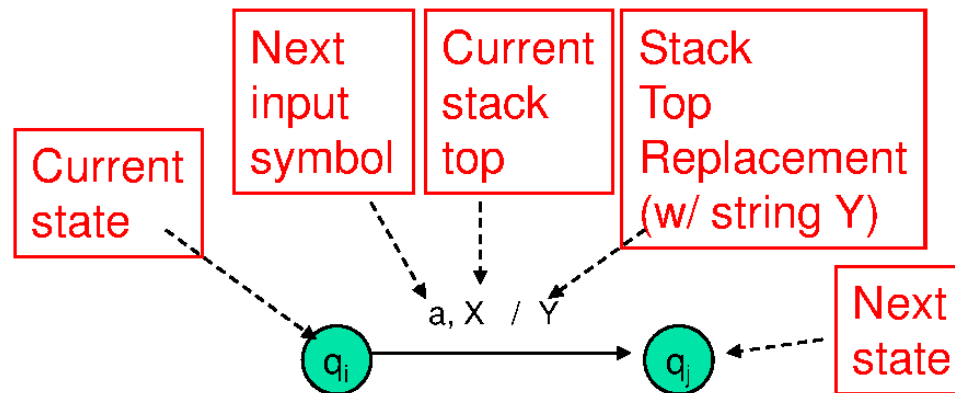
Switch to popping mode (boundary between w and w^R)

Shrink the stack by popping matching symbols (w^R -part)

Enter acceptance state

PDA as a state diagram

$$\delta(q_i, a, X) = \{(q_j, Y)\}$$



PDA for L_{wwr} : Transition Diagram

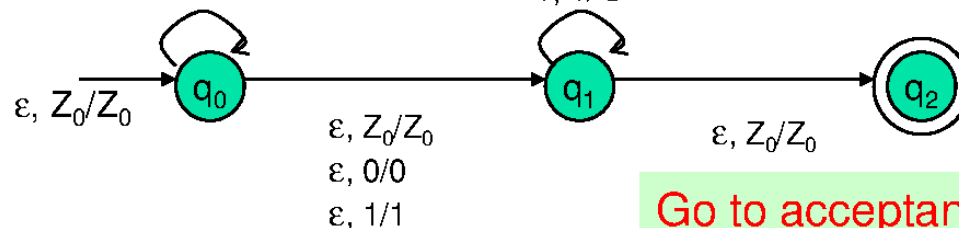
Grow stack

$0, Z_0/0Z_0$
 $1, Z_0/1Z_0$
 $0, 0/00$
 $0, 1/01$
 $1, 0/10$
 $1, 1/11$

Pop stack for matching symbols

$0, 0/\epsilon$
 $1, 1/\epsilon$

$\Sigma = \{0, 1\}$
 $\Gamma = \{Z_0, 0, 1\}$
 $Q = \{q_0, q_1, q_2\}$

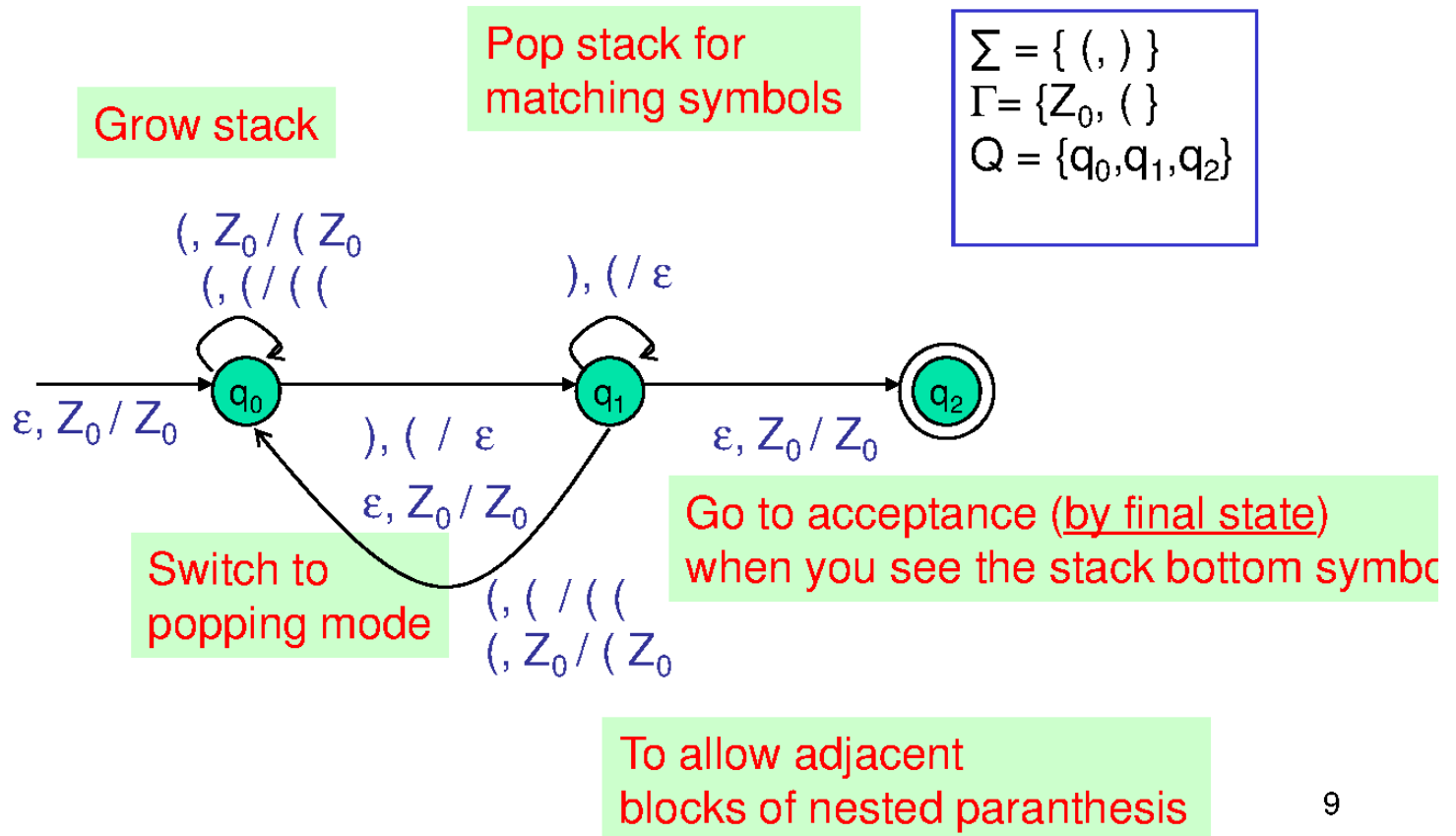


Go to acceptance

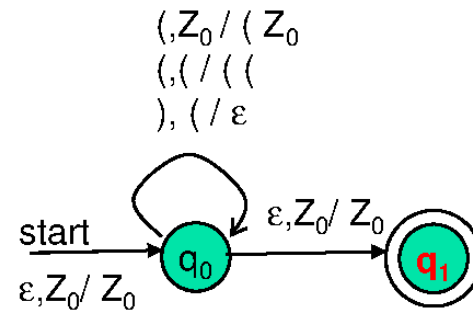
Switch to popping mode

This would be a non-deterministic PDA

Example 2: language of balanced paranthesis



Example 2: language of balanced parenthesis (another design)



$\Sigma = \{ (,) \}$
 $\Gamma = \{ Z_0, (\}$
 $Q = \{ q_0, q_1 \}$



PDA's Instantaneous Description (ID)

A PDA has a configuration at any given instance:

(q,w,y)

- q - current state
 - w - remainder of the input (i.e., unconsumed part)
 - y - current stack contents as a string from top to bottom of stack
-

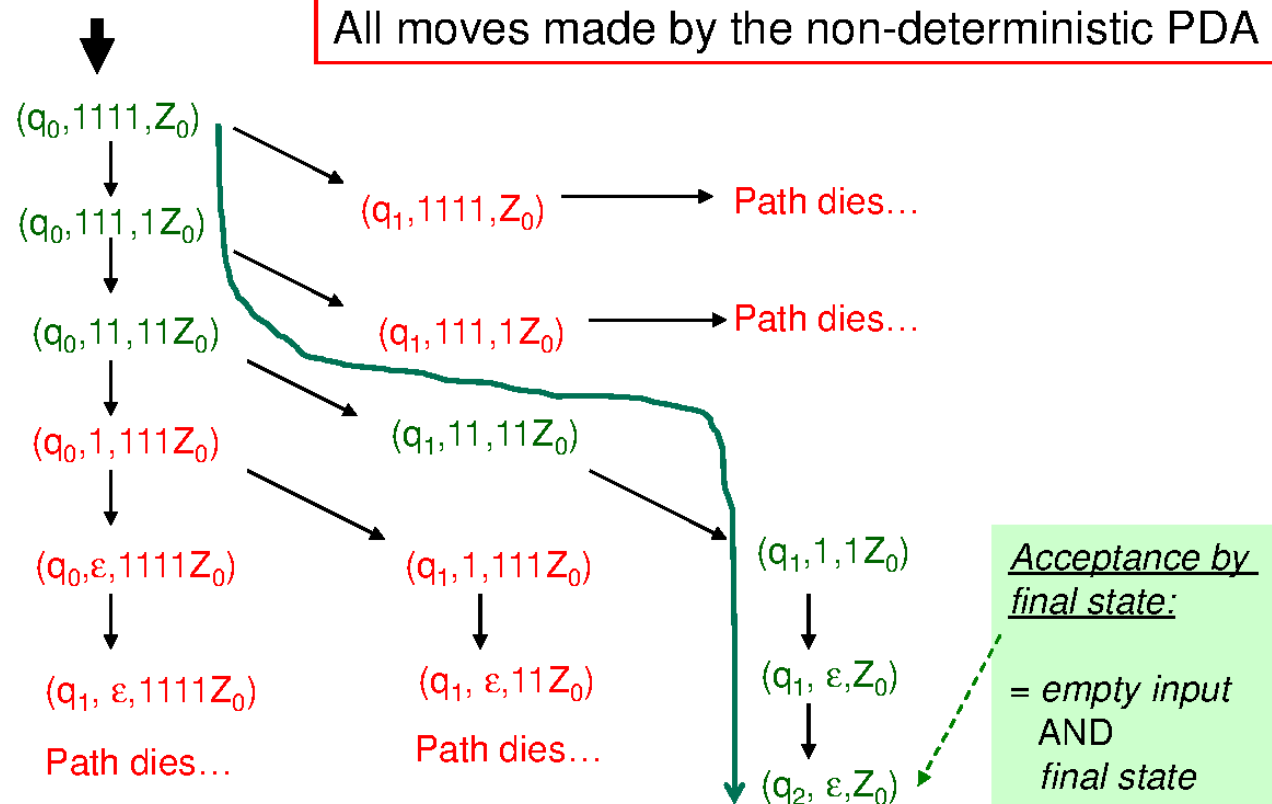
If $\delta(q, a, X) = \{(p, A)\}$ is a transition, then the following are also true:

- $(q, a, X) \vdash (p, \epsilon, A)$
 - $(q, aw, XB) \vdash (p, w, AB)$
-

$\vdash$ sign is called a “turnstile notation” and represents one move

$\vdash^*$ sign represents a sequence of moves

How does the PDA for L_{wwr} work on input “1111”?





Principles about IDs

- Theorem 1: If for a PDA,
 $(q, x, A) \vdash^{*} (p, y, B)$, then for any string $w \in \Sigma^{*}$ and $\gamma \in \Gamma^{*}$, it is also true that:
 - $(q, x w, A \gamma) \vdash^{*} (p, y w, B \gamma)$
- Theorem 2: If for a PDA,
 $(q, x w, A) \vdash^{*} (p, y w, B)$, then it is also true that:
 - $(q, x, A) \vdash^{*} (p, y, B)$

There are two types of PDAs that one can design:
those that accept by final state or by empty stack

Acceptance by...

■ PDAs that accept by **final state**:

- For a PDA P , the language accepted by P , denoted by $L(P)$ by *final state*, is:

- $\{w \mid (q_0, w, Z_0) \vdash^* (q, \epsilon, A)\}$, s.t., $q \in F$

Checklist:

- input exhausted?
- in a final state?

■ PDAs that accept by **empty stack**:

- For a PDA P , the language accepted by P , denoted by $N(P)$ by *empty stack*, is:

- $\{w \mid (q_0, w, Z_0) \vdash^* (q, \epsilon, \epsilon)\}$, for any $q \in Q$.

Q) Does a PDA that accepts by empty stack need any final state specified in the design?

Checklist:

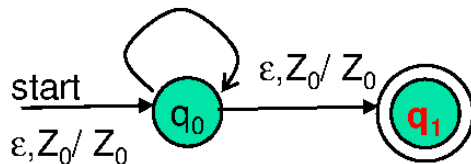
- input exhausted?
- is the stack empty?

Example: L of balanced parenthesis

PDA that accepts by final state

P_F :

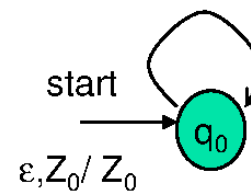
$(, Z_0 / (Z_0$
 $(, (/ (($
 $), (/ \epsilon$



An equivalent PDA that accepts by empty stack

P_N :

$(, Z_0 / (Z_0$
 $(, (/ (($
 $), (/ \epsilon$
 $\epsilon, Z_0 / \epsilon$



How will these two PDAs work on the input: $((())())()$



PDA for L_{ww^R} : Proof of correctness

- Theorem: The PDA for L_{ww^R} accepts a string x by final state **if and only if** x is of the form ww^R .
- Proof:
 - *(if-part)* If the string is of the form ww^R then there exists a sequence of IDs that leads to a final state:
 $(q_0, ww^R, Z_0) \vdash^* (q_0, w^R, wZ_0) \vdash^* (q_1, w^R, wZ_0) \vdash^* (q_1, \epsilon, Z_0) \vdash^* (\mathbf{q_2}, \epsilon, Z_0)$
 - *(only-if part)*
 - Proof by induction on $|x|$



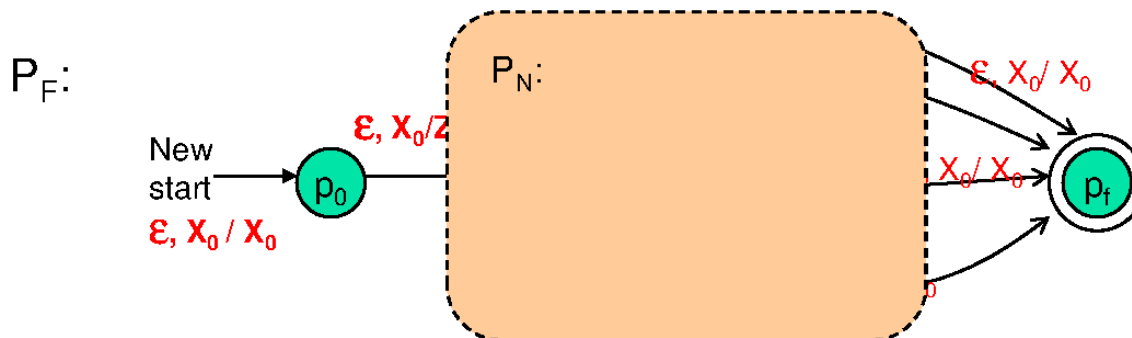
PDAs accepting by final state and empty stack are equivalent

- $P_F \leq$ PDA accepting by final state
 - $P_F = (Q_F, \Sigma, \Gamma, \delta_F, q_0, Z_0, F)$
- $P_N \leq$ PDA accepting by empty stack
 - $P_N = (Q_N, \Sigma, \Gamma, \delta_N, q_0, Z_0)$
- Theorem:
 - $(P_N \Rightarrow P_F)$ For every P_N , there exists a P_F s.t. $L(P_F) = L(P_N)$
 - $(P_F \Rightarrow P_N)$ For every P_F , there exists a P_N s.t. $L(P_F) = L(P_N)$

How to convert an empty stack PDA into a final state PDA?

$P_N \Rightarrow P_F$ construction

- Whenever P_N 's stack becomes empty, make P_F go to a final state without consuming any addition symbol
- To detect empty stack in P_N : P_F pushes a new stack symbol X_0 (not in Γ of P_N) initially before simulating P_N



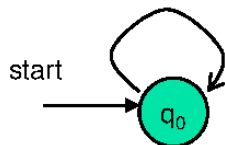
$$P_F = (Q_N \cup \{p_0, p_f\}, \Sigma, \Gamma \cup \{X_0\}, \delta_F, p_0, X_0, \{p_f\})$$

Example: Matching parenthesis “(” “)”

$P_N:$ $(\{q_0\}, \{(\cdot)\}, \{Z_0, Z_1\}, \delta_N, q_0, Z_0)$

$\delta_N:$
 $\delta_N(q_0, (, Z_0) = \{ (q_0, Z_1 Z_0) \}$
 $\delta_N(q_0, (, Z_1) = \{ (q_0, Z_1 Z_1) \}$
 $\delta_N(q_0,), Z_1) = \{ (q_0, \epsilon) \}$
 $\delta_N(q_0, \epsilon, Z_0) = \{ (q_0, \epsilon) \}$

$(, Z_0 / Z_1 Z_0$
 $(, Z_1 / Z_1 Z_1$
 $), Z_1 / \epsilon$
 $\epsilon, Z_0 / \epsilon$

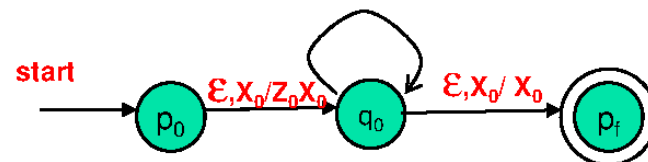


Accept by empty stack

$P_f:$ $(\{p_0, q_0, p_f\}, \{(\cdot)\}, \{X_0, Z_0, Z_1\}, \delta_f, p_0, X_0, p_f)$

$\delta_f:$
 $\delta_f(p_0, \epsilon, X_0) = \{ (q_0, Z_0) \}$
 $\delta_f(q_0, (, Z_0) = \{ (q_0, Z_1 Z_0) \}$
 $\delta_f(q_0, (, Z_1) = \{ (q_0, Z_1 Z_1) \}$
 $\delta_f(q_0,), Z_1) = \{ (q_0, \epsilon) \}$
 $\delta_f(q_0, \epsilon, Z_0) = \{ (q_0, \epsilon) \}$
 $\delta_f(p_0, \epsilon, X_0) = \{ (p_f, X_0) \}$

$(, Z_0 / Z_1 Z_0$
 $(, Z_1 / Z_1 Z_1$
 $), Z_1 / \epsilon$
 $\epsilon, Z_0 / \epsilon$



Accept by final state

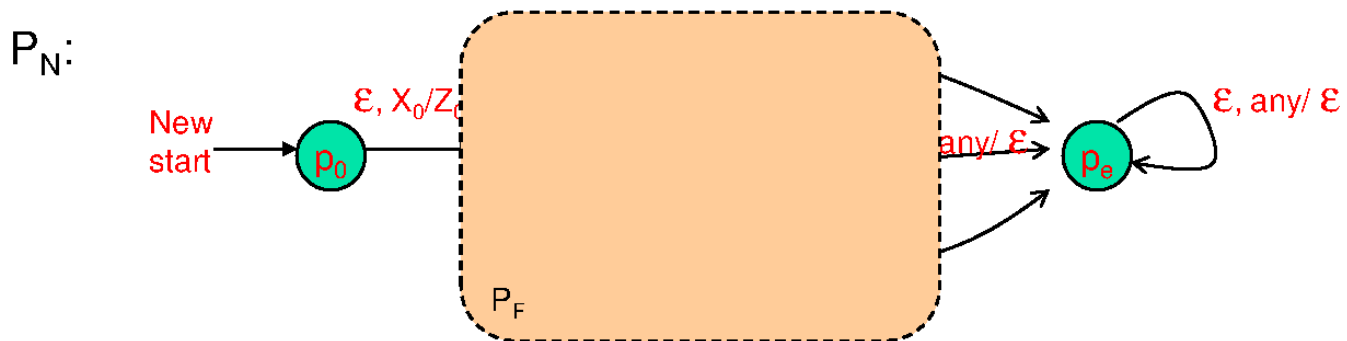
How to convert an final state PDA into an empty stack PDA?

$P_F \Rightarrow P_N$ construction

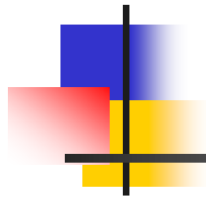
- Main idea:

- Whenever P_F reaches a final state, just make an ϵ -transition into a new end state, clear out the stack and accept
- Danger: What if P_F design is such that it clears the stack midway *without* entering a final state?
 → to address this, add a new start symbol X_0 (not in Γ of P_F)

$$P_N = (Q \cup \{p_0, p_e\}, \Sigma, \Gamma \cup \{X_0\}, \delta_N, p_0, X_0)$$

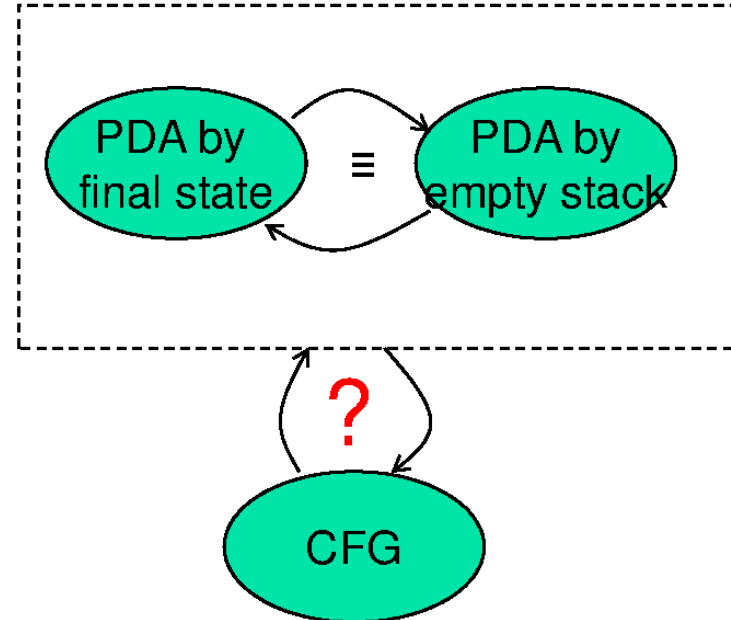


Equivalence of PDAs and CFGs





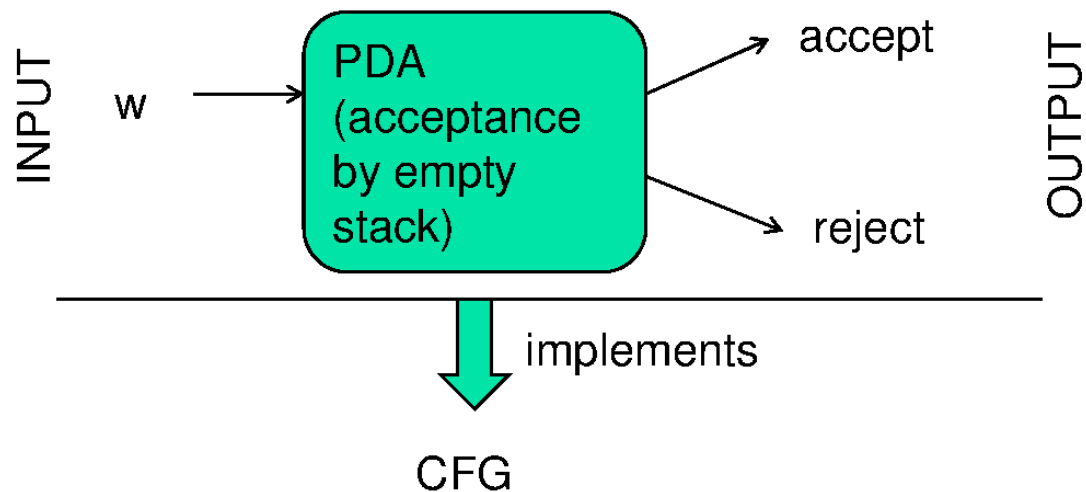
CFGs == PDAs ==> CFLs



This is same as: “implementing a CFG using a PDA”

Converting CFG to PDA

Main idea: The PDA simulates the leftmost derivation on a given w , and upon consuming it fully it either arrives at acceptance (by empty stack) or non-acceptance.



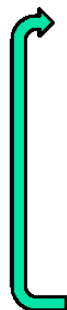
This is same as: “implementing a CFG using a PDA”



Converting a CFG into a PDA

Main idea: The PDA simulates the leftmost derivation on a given w , and upon consuming it fully it either arrives at acceptance (by empty stack) or non-acceptance.

Steps:

- 
1. Push the right hand side of the production onto the stack, with leftmost symbol at the stack top
 2. If stack top is the leftmost variable, then replace it by all its productions (each possible substitution will represent a distinct path taken by the non-deterministic PDA)
 3. If stack top has a terminal symbol, and if it matches with the next symbol in the input string, then pop it

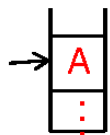
State is inconsequential (only one state is needed)

Formal construction of PDA from CFG

Note: Initial stack symbol (S) same as the start variable in the grammar

- Given: $G = (V, T, P, S)$
- Output: $P_N = (\{q\}, T, V \cup T, \delta, q, S)$
- δ :

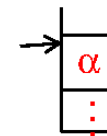
Before:



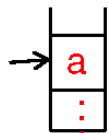
- For all $A \in V$, add the following transition(s) in the PDA:

- $\delta(q, \varepsilon, A) = \{ (q, \alpha) \mid "A \Rightarrow \alpha" \in P \}$

After:



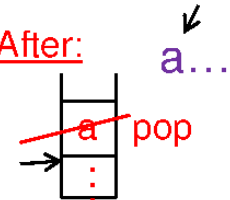
Before:



- For all $a \in T$, add the following transition(s) in the PDA:

- $\delta(q, a, a) = \{ (q, \varepsilon) \}$

After:



Example: CFG to PDA

- $G = (\{S,A\}, \{0,1\}, P, S)$
- P :
 - $S \Rightarrow AS \mid \varepsilon$
 - $A \Rightarrow 0A1 \mid A1 \mid 01$
- $PDA = (\{q\}, \{0,1\}, \{0,1,A,S\}, \delta, q, S)$
- δ :
 - $\delta(q, \varepsilon, S) = \{ (q, AS), (q, \varepsilon) \}$
 - $\delta(q, \varepsilon, A) = \{ (q, 0A1), (q, A1), (q, 01) \}$
 - $\delta(q, 0, 0) = \{ (q, \varepsilon) \}$
 - $\delta(q, 1, 1) = \{ (q, \varepsilon) \}$

$1,1/\varepsilon$
 $0,0/\varepsilon$
 $\varepsilon,A/01$
 $\varepsilon,A/A1$
 $\varepsilon,A/0A1$
 $\varepsilon,S/\varepsilon$
 $\varepsilon,S/AS$

How will this new PDA work?

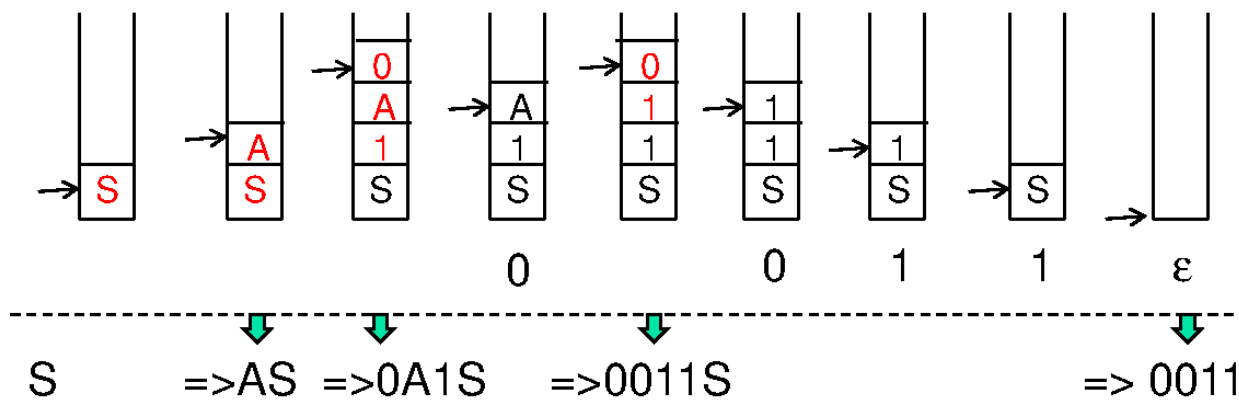
Lets simulate string 0011

Simulating string 0011 on the new PDA ...

PDA (δ):

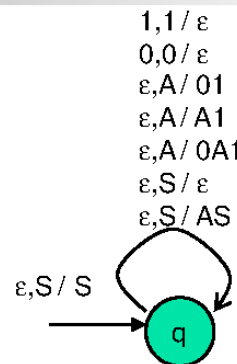
$\delta(q, \epsilon, S) = \{ (q, AS), (q, \epsilon) \}$
 $\delta(q, \epsilon, A) = \{ (q, 0A1), (q, A1), (q, 01) \}$
 $\delta(q, 0, 0) = \{ (q, \epsilon) \}$
 $\delta(q, 1, 1) = \{ (q, \epsilon) \}$

Stack moves (shows only the successful path):



Leftmost deriv.:

$S \Rightarrow AS$
 $\Rightarrow 0A1S$
 $\Rightarrow 0011S$
 $\Rightarrow 0011$



Accept by empty stack



Proof of correctness for $\text{CFG} \Rightarrow \text{PDA}$ construction

- Claim: A string is accepted by G iff it is accepted (by empty stack) by the PDA
- Proof:
 - *(only-if part)*
 - Prove by induction on the number of derivation steps
 - *(if part)*
 - If $(q, wx, S) \vdash^* (q, x, B)$ then $S \Rightarrow_{\text{Im}}^* wB$



Converting a PDA into a CFG

- Main idea: Reverse engineer the productions from transitions

If $\delta(q, a, Z) \Rightarrow (p, Y_1 Y_2 Y_3 \dots Y_k)$:

1. State is changed from q to p ;
2. Terminal a is consumed;
3. Stack top symbol Z is popped and replaced with a sequence of k variables.

- Action: Create a grammar variable called “[qZp]” which includes the following production:

- $[qZp] \Rightarrow a[pY_1q_1] [q_1Y_2q_2] [q_2Y_3q_3] \dots [q_{k-1}Y_kq_k]$

- Proof discussion (in the book)

Example: Bracket matching

- To avoid confusion, we will use $b = "("$ and $e = ")"$

$P_N: (\{q_0\}, \{b, e\}, \{Z_0, Z_1\}, \delta, q_0, Z_0)$

- $\delta(q_0, b, Z_0) = \{ (q_0, Z_1 Z_0) \}$
- $\delta(q_0, b, Z_1) = \{ (q_0, Z_1 Z_1) \}$
- $\delta(q_0, e, Z_1) = \{ (q_0, \epsilon) \}$
- $\delta(q_0, \epsilon, Z_0) = \{ (q_0, \epsilon) \}$

- $S \Rightarrow [q_0 Z_0 q_0]$
- $[q_0 Z_0 q_0] \Rightarrow b [q_0 Z_1 q_0] [q_0 Z_0 q_0]$
- $[q_0 Z_1 q_0] \Rightarrow b [q_0 Z_1 q_0] [q_0 Z_1 q_0]$
- $[q_0 Z_1 q_0] \Rightarrow e$
- $[q_0 Z_0 q_0] \Rightarrow \epsilon$

Let $A = [q_0 Z_0 q_0]$
Let $B = [q_0 Z_1 q_0]$

- $S \Rightarrow A$
- $A \Rightarrow b B A$
- $B \Rightarrow b B B$
- $B \Rightarrow e$
- $A \Rightarrow \epsilon$

Simplifying,

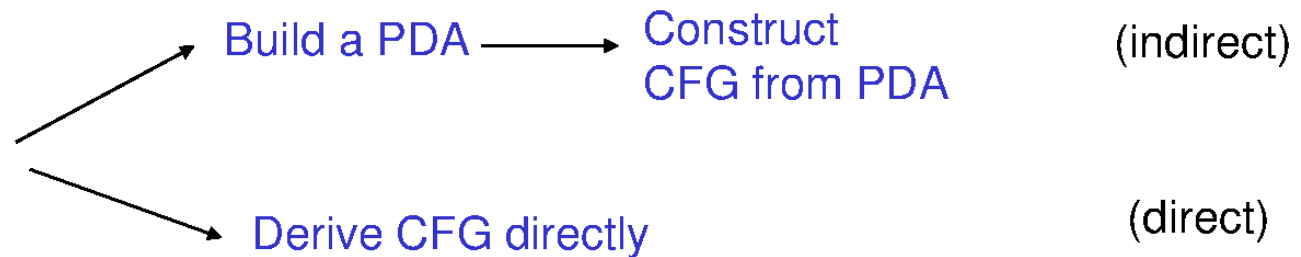
- $S \Rightarrow b B S \mid \epsilon$
- $B \Rightarrow b B B \mid e$

If you were to directly write a CFG:

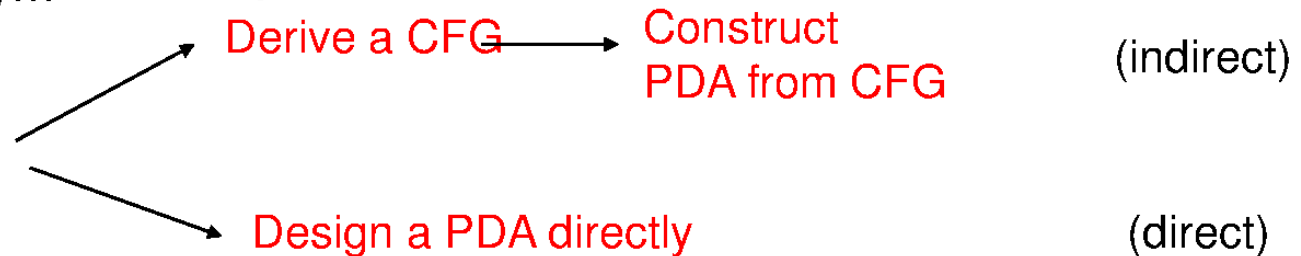
$S \Rightarrow b S e S \mid \epsilon$



Two ways to build a CFG



Similarly... Two ways to build a PDA





Deterministic PDAs

This PDA for L_{wwr} is non-deterministic

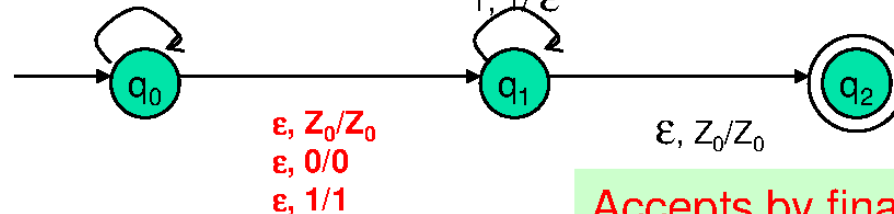
Grow stack

$0, Z_0/0Z_0$
 $1, Z_0/1Z_0$
 $0, 0/00$
 $0, 1/01$
 $1, 0/10$
 $1, 1/11$

Pop stack for
matching symbols

$0, 0/\epsilon$
 $1, 1/\epsilon$

Why does it have
to be non-
deterministic?



Switch to
popping mode

Accepts by final state

To remove
guessing,
impose the user
to insert c in the
middle

Example shows that: Nondeterministic PDAs $\neq$ D-PDAs

D-PDA for $L_{wcw^R} = \{wcw^R \mid c \text{ is some special symbol not in } w\}$

Grow stack

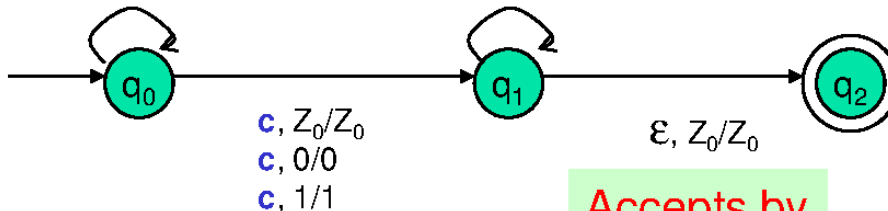
0, $Z_0/0Z_0$
 1, $Z_0/1Z_0$
 0, $0/00$
 0, $1/01$
 1, $0/10$
 1, $1/11$

Pop stack for
matching symbols

0, $0/\epsilon$
 1, $1/\epsilon$

Note:

- all transitions have become deterministic



Switch to
popping mode

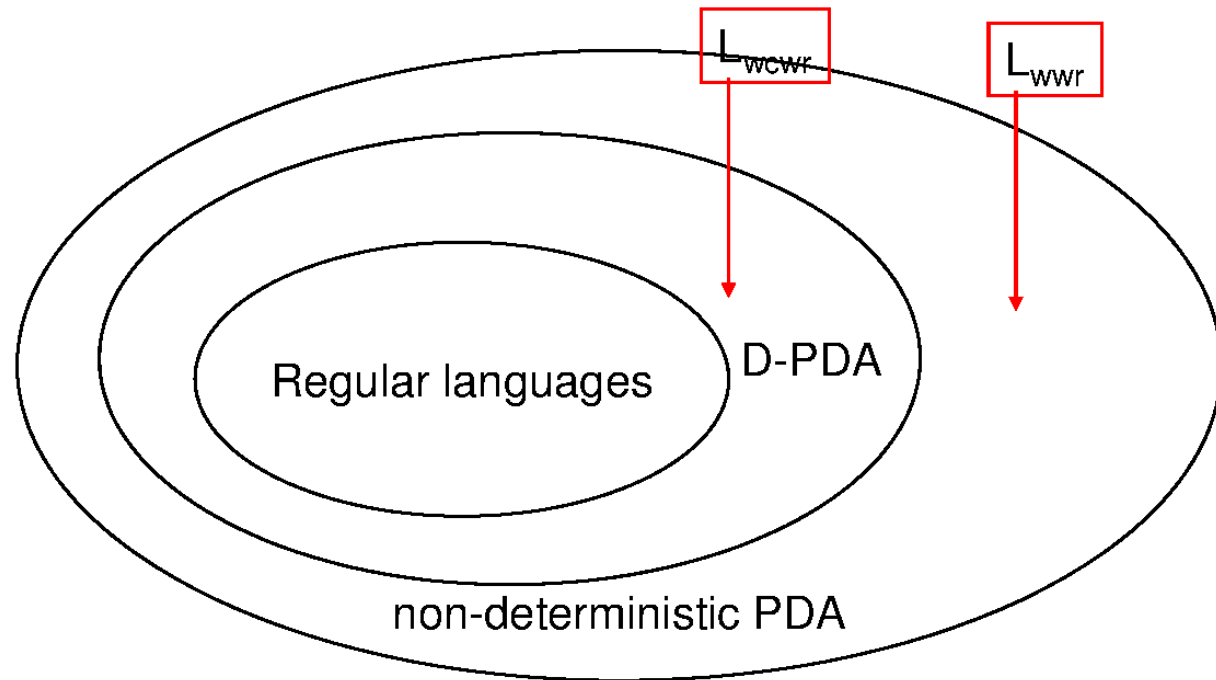
Accepts by
final state



Deterministic PDA: Definition

- A PDA is *deterministic* if and only if:
 1. $\delta(q, a, X)$ has *at most one* member for any $a \in \Sigma \cup \{\epsilon\}$
- ➔ If $\delta(q, a, X)$ is non-empty for some $a \in \Sigma$, then $\delta(q, \epsilon, X)$ must be empty.

PDA vs DPDA vs Regular languages





Summary

- PDAs for CFLs and CFGs
 - Non-deterministic
 - Deterministic
- PDA acceptance types
 1. By final state
 2. By empty stack
- PDA
 - IDs, Transition diagram
- Equivalence of CFG and PDA
 - CFG $\Rightarrow$ PDA construction
 - PDA $\Rightarrow$ CFG construction



FORMAL LANGUAGES AND AUTOMATA THEORY

UNIT 4



Turing Machines



Turing Machines are...

- Very powerful (abstract) machines that could simulate any modern day computer (although very, very slowly!)
- Why design such a machine?
 - If a problem cannot be “solved” even using a TM, then it implies that the problem is ***undecidable***
- Computability vs. Decidability

For every input,
answer YES or NO

A Turing Machine (TM)

■ $M = (Q, \Sigma, \Gamma, \delta, q_0, B, F)$

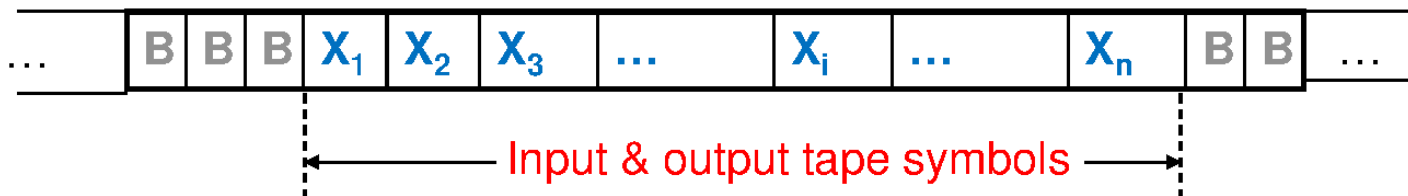
This is like the CPU & program counter

Finite control

Tape is the memory

Tape head

Infinite tape with tape symbols



B: blank symbol (special symbol reserved to indicate data boundary)



Transition function

You can also use:

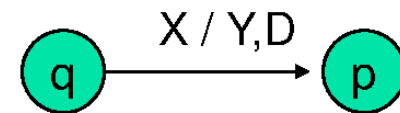
→ for R

← for L

- One move (denoted by |---) in a TM does the following:

- $\delta(q, X) = (p, Y, D)$

- q is the current state
- X is the current tape symbol pointed by tape head
- State changes from q to p
- After the move:
 - X is replaced with symbol Y
 - If D="L", the tape head moves "left" by one position.
Alternatively, if D="R" the tape head moves "right" by one position.



ID of a TM

- Instantaneous Description or ID :

- $X_1 X_2 \dots X_{i-1} q X_i X_{i+1} \dots X_n$

means:

- q is the current state
 - Tape head is pointing to X_i
 - $X_1 X_2 \dots X_{i-1} X_i X_{i+1} \dots X_n$ are the current tape symbols

- $\delta(q, X_i) = (p, Y, R)$ is same as:

$$X_1 \dots X_{i-1} q X_i \dots X_n \quad | \text{---} \quad X_1 \dots X_{i-1} Y p X_{i+1} \dots X_n$$

- $\delta(q, X_i) = (p, Y, L)$ is same as:

$$X_1 \dots X_{i-1} q X_i \dots X_n \quad | \text{---} \quad X_1 \dots p X_{i-1} Y X_{i+1} \dots X_n$$

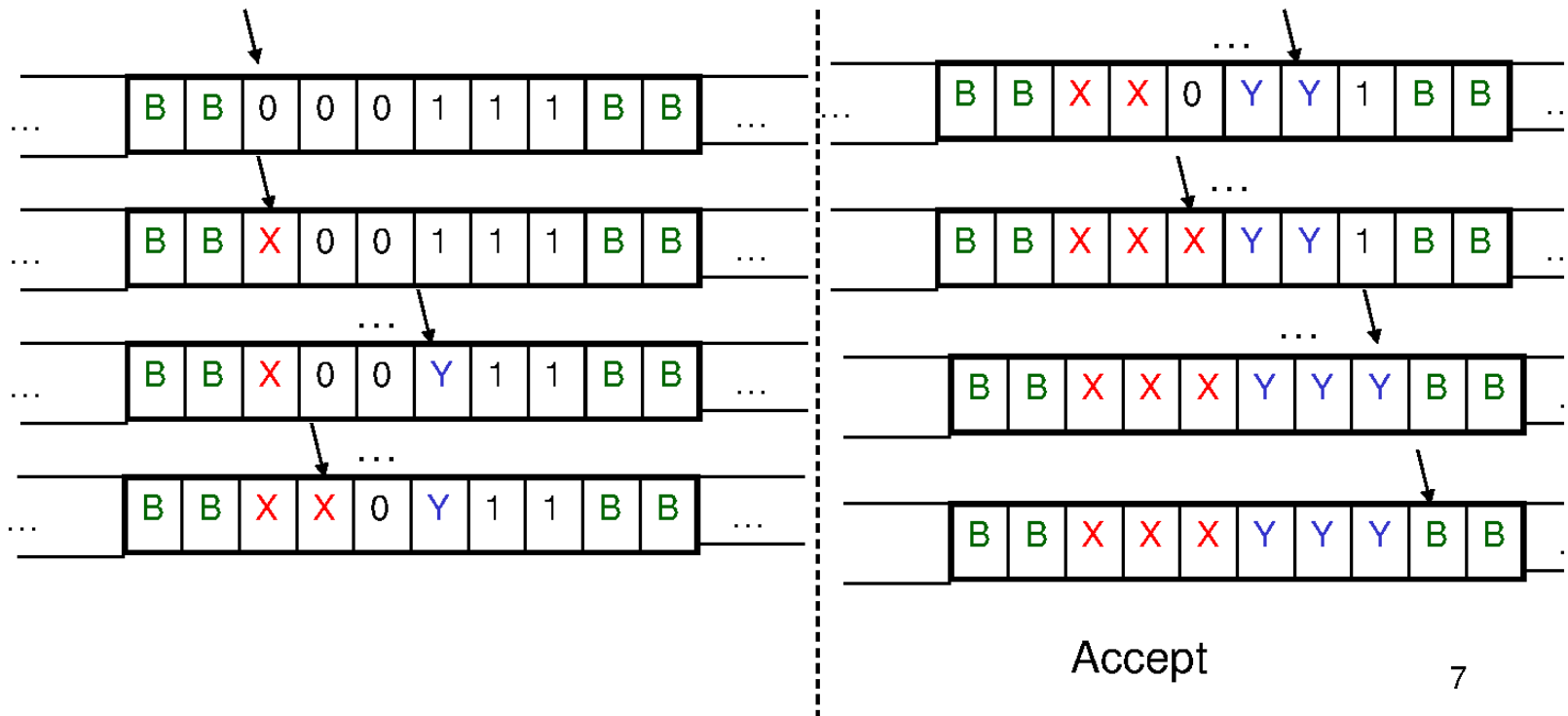


Way to check for Membership

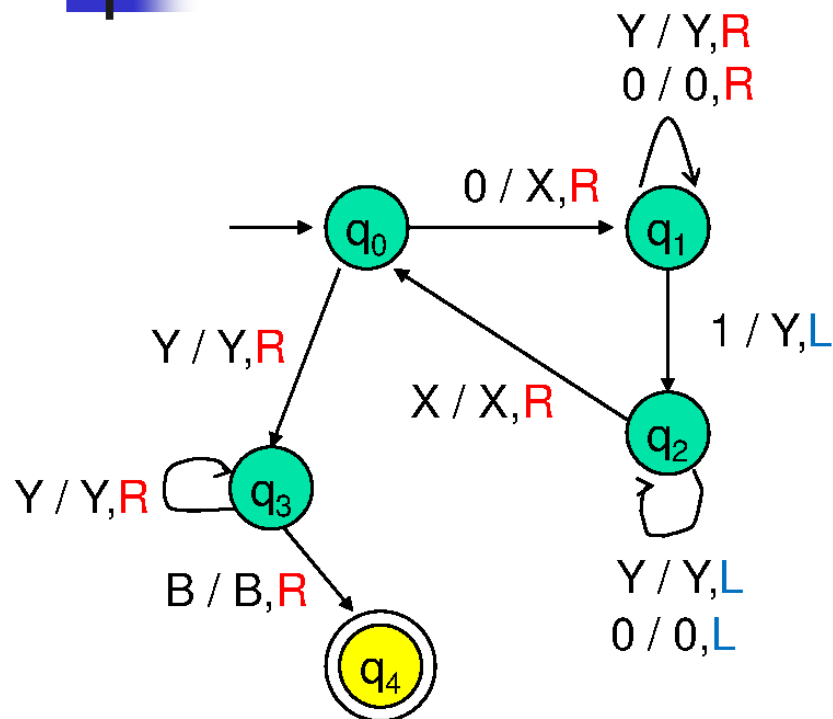
- Is a string w accepted by a TM?
- Initial condition:
 - The (whole) input string w is present in TM, preceded and followed by infinite blank symbols
- Final acceptance:
 - Accept w if TM enters final state and halts
 - If TM halts and not final state, then reject

Example: $L = \{0^n 1^n \mid n \geq 1\}$

■ Strategy: $w = 000111$



TM for $\{0^n 1^n \mid n \geq 1\}$



1. Mark next unread 0 with X and move right
2. Move to the right all the way to the first unread 1, and mark it with Y
3. Move back (to the left) all the way to the last marked X, and then move one position to the right
4. If the next position is 0, then goto step 1. Else move all the way to the right to ensure there are no excess 1s. If not move right to the next blank symbol and stop & accept.

*state diagram representation preferred



TM for $\{0^n 1^n \mid n \geq 1\}$

	Next Tape Symbol				
Curr. State	0	1	X	Y	B
→ q_0	(q_1, X, R)	-	-	(q_3, Y, R)	-
q_1	$(q_1, 0, R)$	(q_2, Y, L)	-	(q_1, Y, R)	-
q_2	$(q_2, 0, L)$	-	(q_0, X, R)	(q_2, Y, L)	-
q_3	-	-	-	(q_3, Y, R)	(q_4, B, R)
$*q_4$	-	--	-	-	-

Table representation of the state diagram



TMs for calculations

- TMs can also be used for calculating values
 - Like arithmetic computations
 - Eg., addition, subtraction, multiplication, etc.



Example 2: monus subtraction

$$“m \dot{-} n” = \max\{m-n, 0\}$$

$$0^m 1 0^n \rightarrow \begin{array}{l} \dots B 0^{m-n} B.. \text{ (if } m > n\text{)} \\ \dots BB\dots B.. \text{ (otherwise)} \end{array}$$

Give state diagram

1. For every 0 on the left (mark X), mark off a 0 on the right (mark Y)
2. Repeat process, until one of the following happens:
 1. // No more 0s remaining on the left of 1
Answer is 0, so flip all excess 0s on the right of 1 to Bs (and the 1 itself) and halt
 2. // No more 0s remaining on the right of 1
Answer is m-n, so simply halt after making 1 to B



Example 3: Multiplication

- $0^m 1 0^n 1$ (input), $0^{mn} 1$ (output)
- Pseudocode:
 1. Move tape head back & forth such that for every 0 seen in 0^m , write n 0s to the right of the last delimiting 1
 2. Once written, that zero is changed to B to get marked as finished
 3. After completing on all m 0s, make the remaining n 0s and 1s also as Bs

Give state diagram

Calculations vs. Languages

A “calculation” is one that takes an input and outputs a value (or values)

The “language” for a certain calculation is the set of strings of the form “<input, output>”, where the output corresponds to a valid calculated value for the input

A “language” is a set of strings that meet certain criteria

E.g., The language L_{add} for the addition operation

“<0#0,0>”

“<0#1,1>”

...

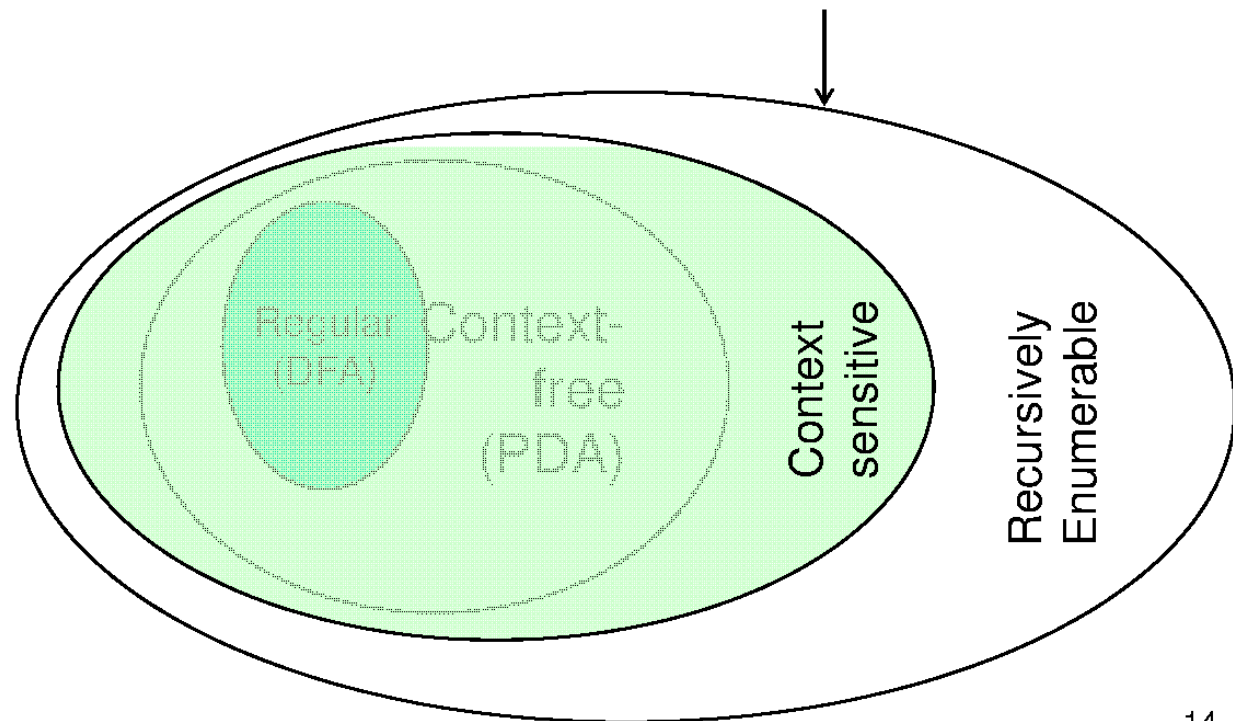
“<2#4,6>”

...

Membership question == verifying a solution
e.g., is “<15#12,27>” a member of L_{add} ?

Language of the Turing Machines

- *Recursive Enumerable (RE) language*



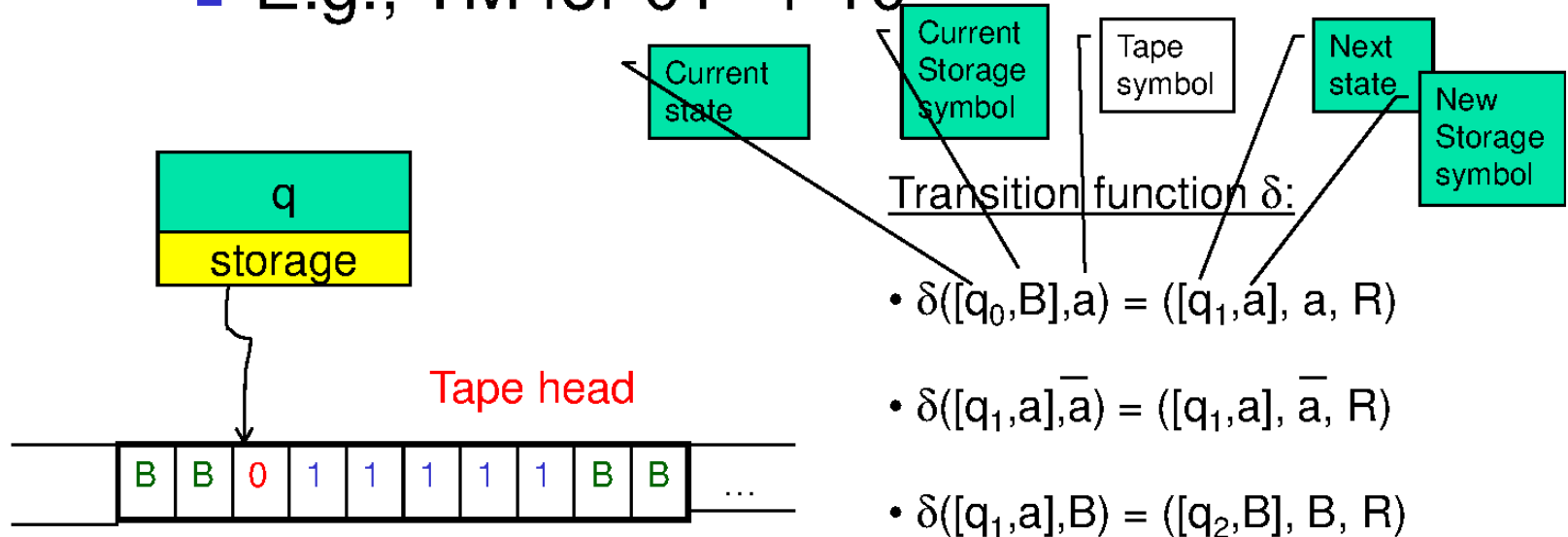


Variations of Turing Machines

TMs with *storage*

Generic description
Will work for both $a=0$ and $a=1$

- E.g., TM for $01^* + 10^*$



$[q, a]$: where q is current state,
 a is the symbol in storage

Are the standard TMs
equivalent to TMs with storage?

Yes



Standard TMs are equivalent to TMs with storage - Proof

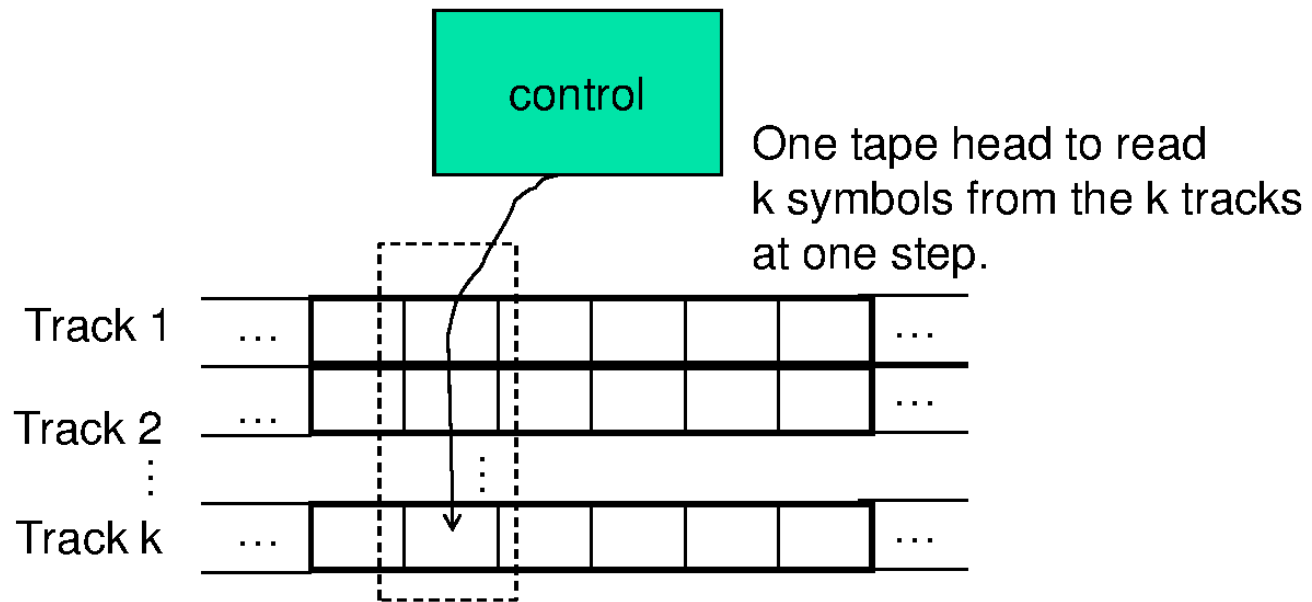
Claim: Every TM w/ storage can be simulated by a TM w/o storage as follows:

- For every [state, symbol] combination in the TM w/ storage:
 - Create a new state in the TM w/o storage
 - Define transitions induced by TM w/ storage

Since there are only finite number of states and symbols in the TM with storage, the number of states in the TM without storage will also be finite

Multi-track Turing Machines

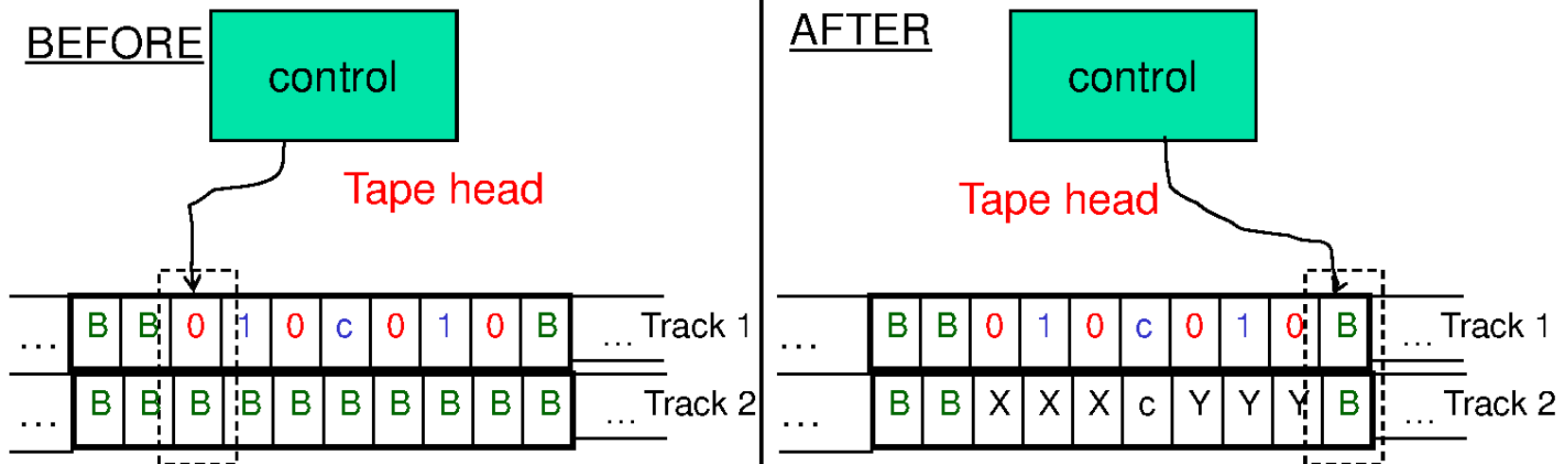
- TM with multiple tracks,
but just one unified tape head



Multi-Track TMs

- TM with multiple “tracks” but just one head

E.g., TM for $\{wcw \mid w \in \{0,1\}^*\}$
but w/o modifying original input string



Second track mainly used as a scratch space for marking



Multi-track TMs are equivalent to basic (single-track) TMs

- Let M be a single-track TM
 - $M = (Q, \Sigma, \Gamma, \delta, q_0, B, F)$
- Let M' be a multi-track TM (k tracks)
 - $M' = (Q', \Sigma', \Gamma', \delta', q'_0, B, F')$
 - $\delta'(q_i, \langle a_1, a_2, \dots, a_k \rangle) = (q_j, \langle b_1, b_2, \dots, b_k \rangle, L/R)$
- Claims:
 - *For every M , there is an M' s.t. $L(M) = L(M')$.*
 - (proof trivial here)



Multi-track TM $\Rightarrow$ TM (proof)

- For every M' , there is an M s.t. $L(M')=L(M)$.

- $M = (Q, \Sigma, \Gamma, \delta, q_0, [B, B, \dots], F)$
- Where:
 - $Q = Q'$
 - $\Sigma = \Sigma' \times \Sigma' \times \dots$ (k times for k-track)
 - $\Gamma = \Gamma' \times \Gamma' \times \dots$ (k times for k-track)
 - $q_0 = q'_0$
 - $F = F'$
 - $\delta(q_i, [a_1, a_2, \dots, a_k]) = \delta'(q_i, \langle a_1, a_2, \dots, a_k \rangle)$

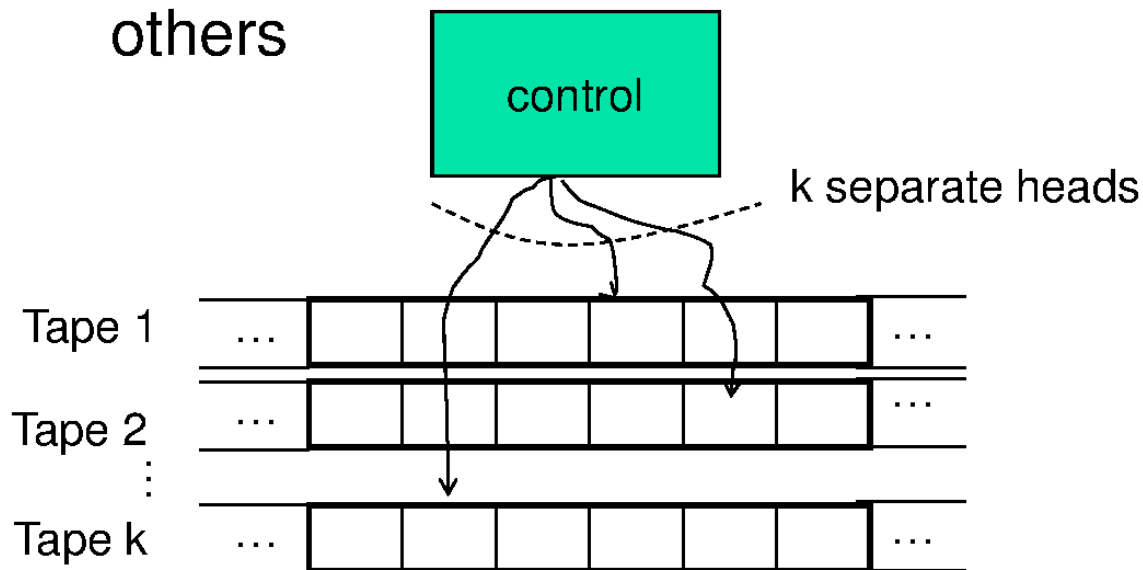
Main idea:

Create one composite symbol to represent every combination of k symbols

- Multi-track TMs are just a different way to represent single-track TMs, and is a matter of design convenience.

Multi-tape Turing Machines

- TM with multiple tapes, *each tape with a separate head*
 - Each head can move independently of the others





On how a Multi-tape TM would operate

- Initially:
 - The input is in tape #1 surrounded by blanks
 - All other tapes contain only blanks
 - The tape head for tape #1 points to the 1st symbol of the input
 - The heads for all other tapes point at an arbitrary cell (doesn't matter because they are all blanks anyway)
- A move:
 - Is a function (current state, the symbols pointed by all the heads)
 - After each move, each tape head can move independently (left or right) of one another

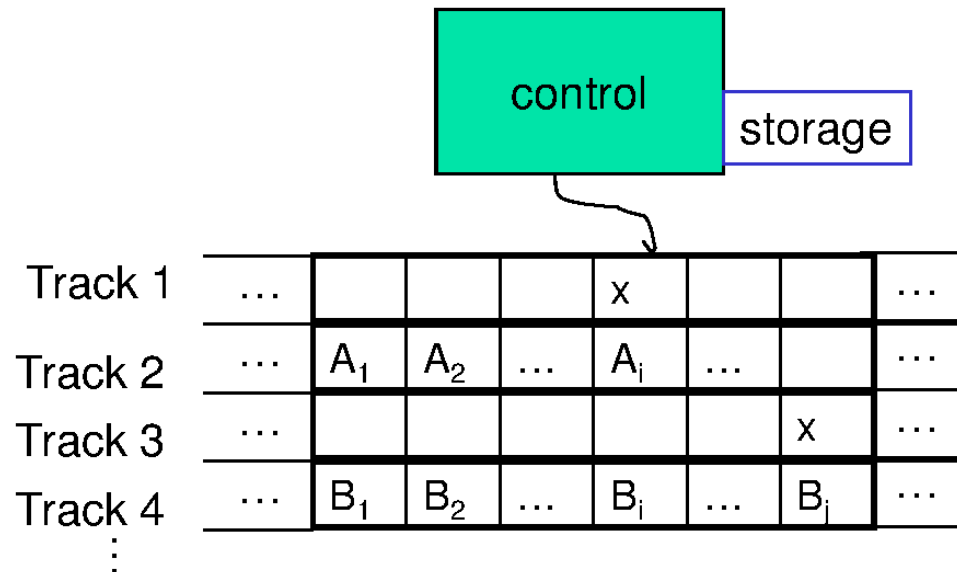


Multitape TMs $\equiv$ Basic TMs

- Theorem: Every language accepted by a k-**tape** TM is also accepted by a single-**tape** TM
- Proof by construction:
 - Construct a single-**tape** TM with $2k$ **tracks**, where each **tape** of the k-**tape** TM is simulated by 2 **tracks** of basic TM
 - k out of the $2k$ **tracks** simulate the k input **tapes**
 - The other k out of the $2k$ **tracks** keep track of the k tape **head** positions

Multitape TMs $\equiv$ Basic TMs ...

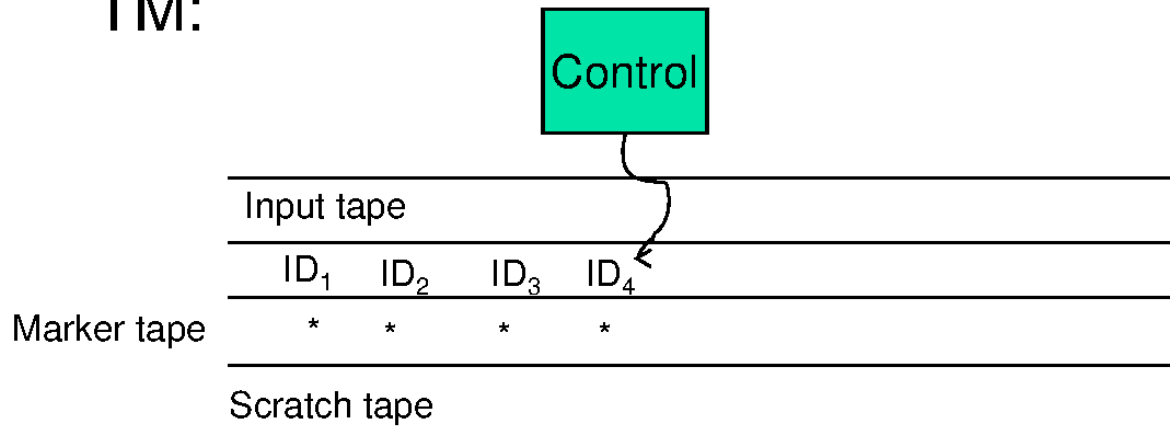
- To simulate one move of the k-tape TM:
 - Move from the leftmost marker to the rightmost marker (k markers) and in the process, gather all the input symbols into storage
 - Then, take the action same as done by the k-tape TM (rewrite tape symbols & move L/R using the markers)



Non-deterministic TMs $\equiv$ Deterministic TMs

Non-deterministic TMs

- A TM can have non-deterministic moves:
 - $\delta(q, X) = \{ (q_1, Y_1, D_1), (q_2, Y_2, D_2), \dots \}$
- Simulation using a multitape deterministic TM:





Summary

- TMs == Recursively Enumerable languages
- TMs can be used as both:
 - Language recognizers
 - Calculators/computers
- ***Basic TM is equivalent to all the below:***
 1. *TM + storage*
 2. *Multi-track TM*
 3. *Multi-tape TM*
 4. *Non-deterministic TM*
- TMs are like universal computing machines with unbounded storage



FORMAL LANGUAGES AND AUTOMATA THEORY

UNIT 5



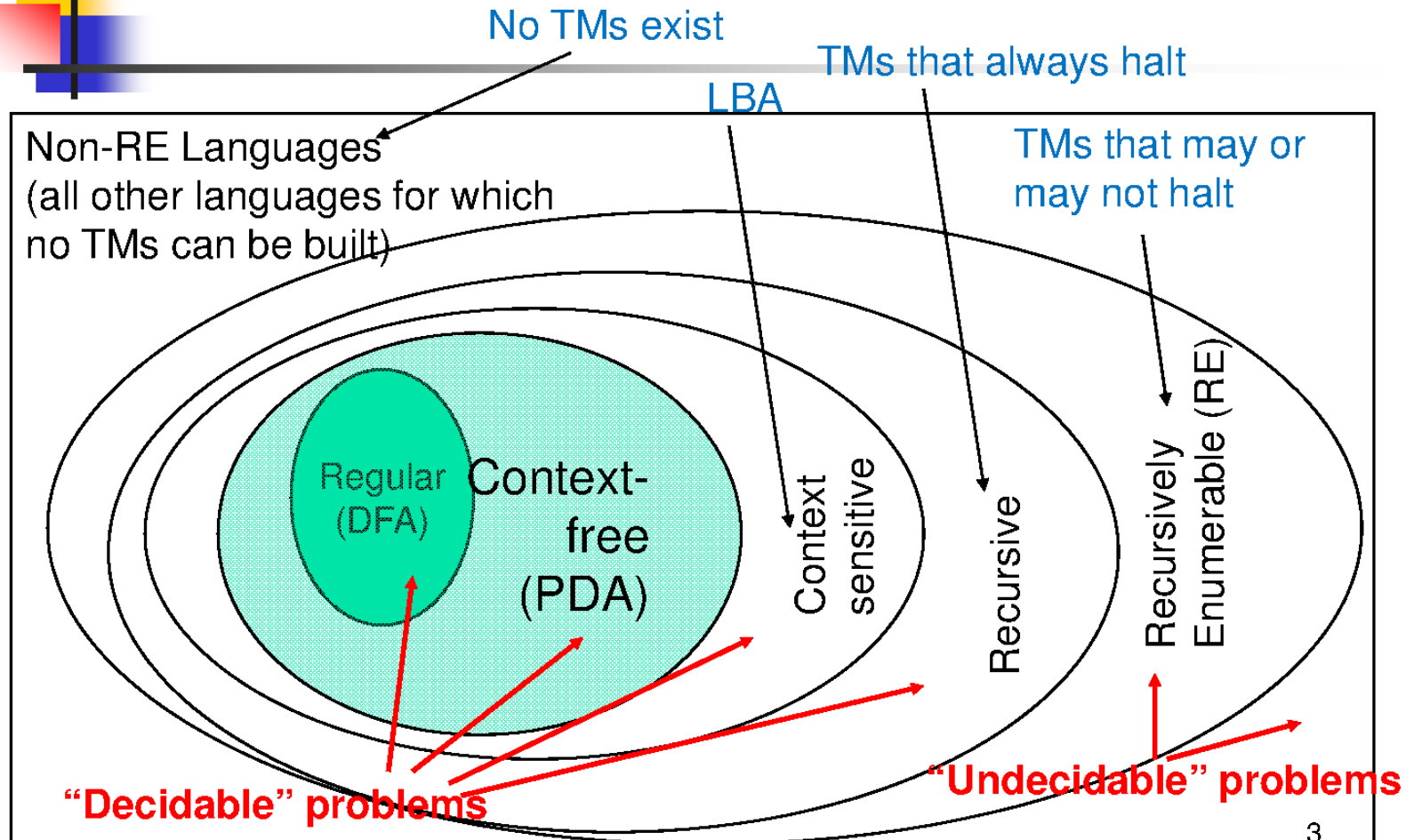
Undecidability



Decidability vs. Undecidability

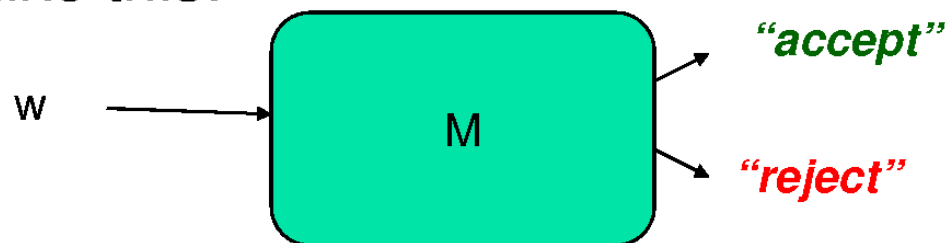
- There are two types of TMs (based on halting):
 - (*Recursive*)
 - TMs that *always* halt**, no matter accepting or non-accepting $\equiv$ **DECIDABLE PROBLEMS**
 - (*Recursively enumerable*)
 - TMs that *are guaranteed to halt only on acceptance***. If non-accepting, it may or may not halt (i.e., could loop forever).
- **Undecidability:**
 - Undecidable problems are those that are not recursive

Recursive, RE, Undecidable languages



Recursive Languages & Recursively Enumerable (RE) languages

- Any TM for a Recursive language is going to look like this:



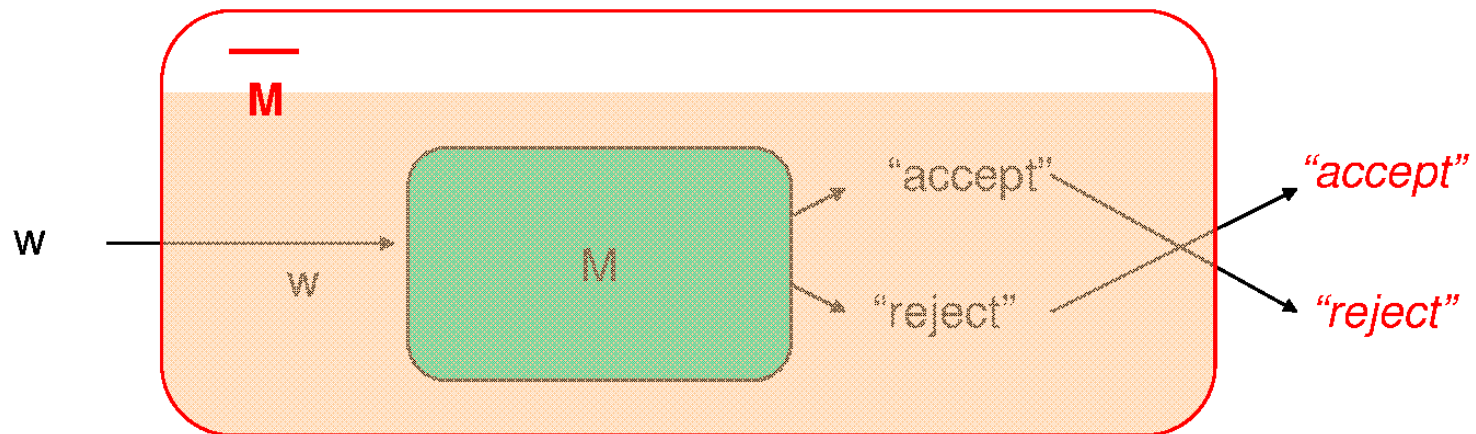
- Any TM for a Recursively Enumerable (RE) language is going to look like this:



- 
- Closure Properties of:
- the Recursive language class, and
 - the Recursively Enumerable language class

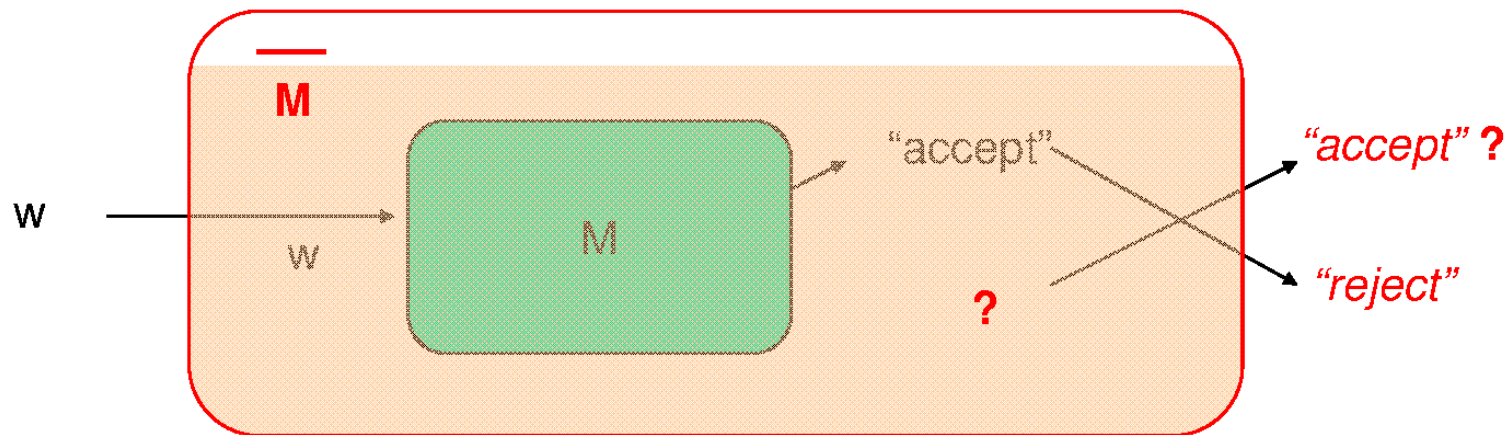
Recursive Languages are closed under complementation

- If L is Recursive, $\overline{L}$ is also Recursive



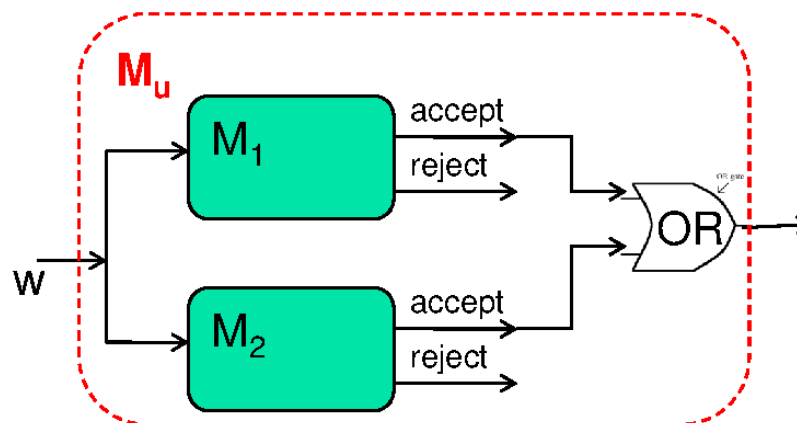
Are Recursively Enumerable Languages closed under complementation? (NO)

- If L is RE, $\overline{L}$ need not be RE



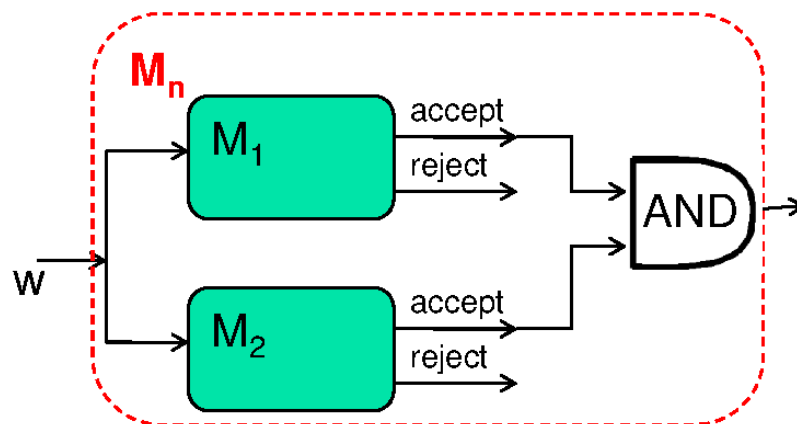
Recursive Langs are closed under Union

- Let $M_u = \text{TM for } L_1 \cup L_2$
- M_u construction:
 1. Make 2-tapes and copy input w on both tapes
 2. Simulate M_1 on tape 1
 3. Simulate M_2 on tape 2
 4. If either M_1 or M_2 accepts, then M_u accepts
 5. Otherwise, M_u rejects.



Recursive Langs are closed under Intersection

- Let $M_n = \text{TM for } L_1 \cap L_2$
- M_n construction:
 1. Make 2-tapes and copy input w on both tapes
 2. Simulate M_1 on tape 1
 3. Simulate M_2 on tape 2
 4. If either M_1 AND M_2 accepts, then M_n accepts
 5. Otherwise, M_n rejects.





Other Closure Property Results

- Recursive languages are also closed under:
 - Concatenation
 - Kleene closure (star operator)
 - Homomorphism, and inverse homomorphism
 - RE languages are closed under:
 - Union, intersection, concatenation, Kleene closure
-
- RE languages are *not* closed under:
 - complementation



“Languages” vs. “Problems”

A “language” is a set of strings

Any “problem” can be expressed as a set of all strings that are of the form:

- “<input, output>”

e.g., Problem (a+b) $\equiv$ Language of strings of the form { “a#b, a+b” }

$\implies$ Every problem also corresponds to a language!!

Think of the language for a “problem” $\equiv$ a *verifier* for the problem

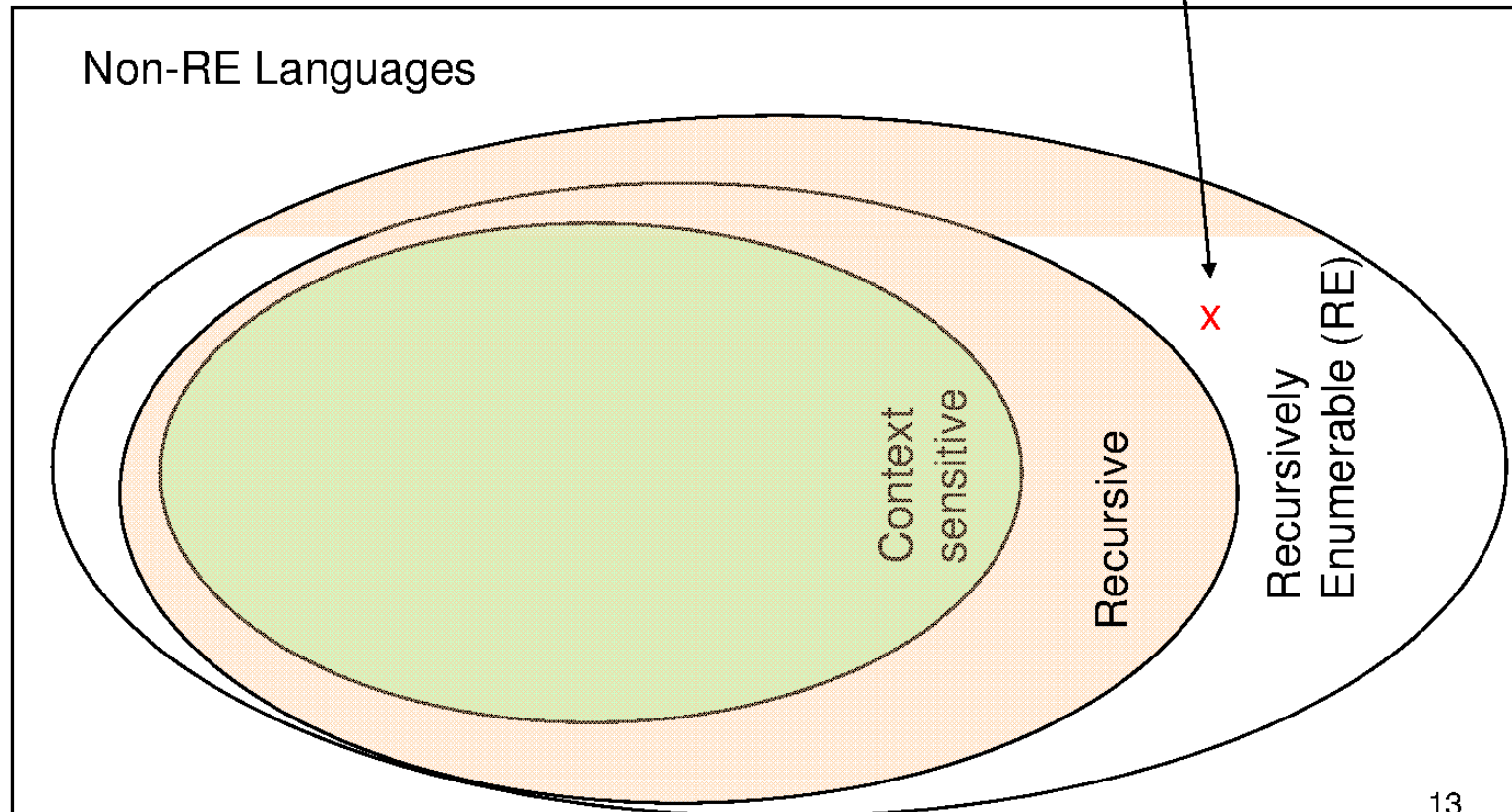


The Halting Problem

**An example of a recursive
enumerable problem that is
also undecidable**



The Halting Problem

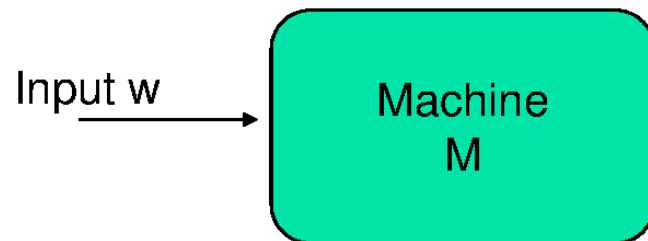




What is the Halting Problem?

Definition of the “halting problem”:

- *Does a given Turing Machine M halt on a given input w ?*



A Turing Machine simulator

The Universal Turing Machine

- Given: TM **M** & its input **w**
- Aim: Build another TM called “H”, that will output:
 - “*accept*” if M accepts w, and
 - “*reject*” otherwise

- An algorithm for H:
 - Simulate M on w

Implies: H is in RE

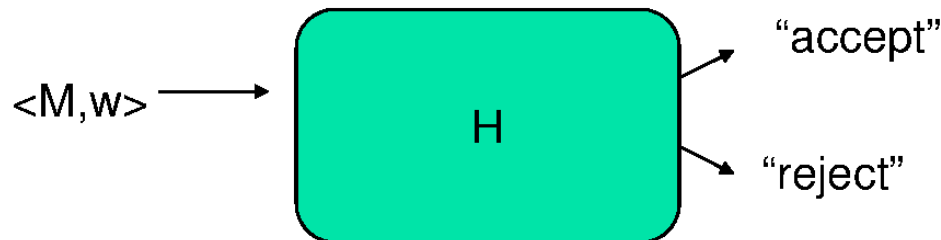
$$\text{■ } H(\langle M, w \rangle) = \begin{cases} \text{accept,} & \text{if } M \text{ accepts } w \\ \text{reject,} & \text{if } M \text{ does not accept } w \end{cases}$$

Question: If M does *not* halt on w, what will happen to H?



A Claim

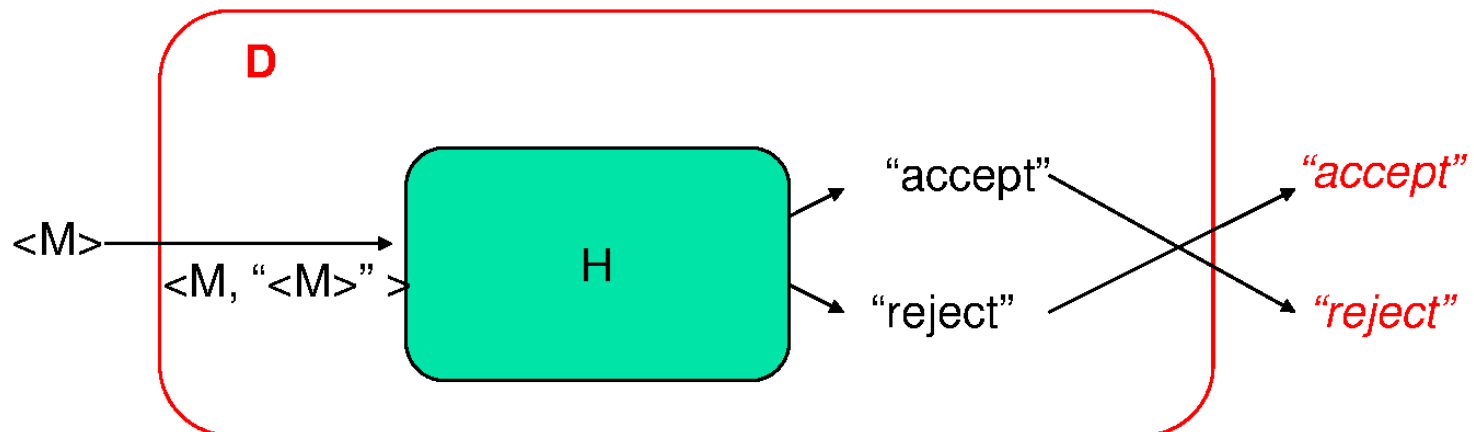
- Claim: No H that is always guaranteed to halt, can exist!
- Proof: (Alan Turing, 1936)
 - By contradiction, let us assume H exists



Therefore, if H exists $\rightarrow$ D also should exist.
But can such a D exist? (if not, then H also cannot exist)

HP Proof (step 1)

- Let us construct a new TM **D** using H as a subroutine:
 - On input $\langle M \rangle$:
 1. Run H on input $\langle M, \langle M \rangle \rangle$; //(i.e., run M on M itself)
 2. Output the *opposite* of what H outputs;





HP Proof (step 2)

- The notion of inputting “<M>” to M itself
 - A program can be input to itself (e.g., a compiler is a program that takes any program as input)

$$D(<M>) = \begin{cases} \text{accept,} & \text{if } M \text{ does not accept } <M> \\ \text{reject,} & \text{if } M \text{ accepts } <M> \end{cases}$$

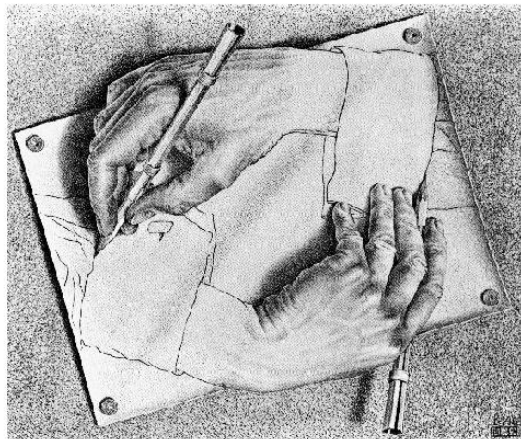
Now, what happens if D is input to itself?

$$D(<D>) = \begin{cases} \text{accept,} & \text{if } D \text{ does not accept } <D> \\ \text{reject,} & \text{if } D \text{ accepts } <D> \end{cases}$$

A contradiction!!! ==> Neither D nor H can exist.

Of Paradoxes & Strange Loops

E.g., Barber's paradox, Achilles & the Tortoise (Zeno's paradox)
MC Escher's paintings



A fun book for further reading:

"Godel, Escher, Bach: An Eternal Golden Braid"
by Douglas Hofstadter (Pulitzer winner, 1980)



The Diagonalization Language

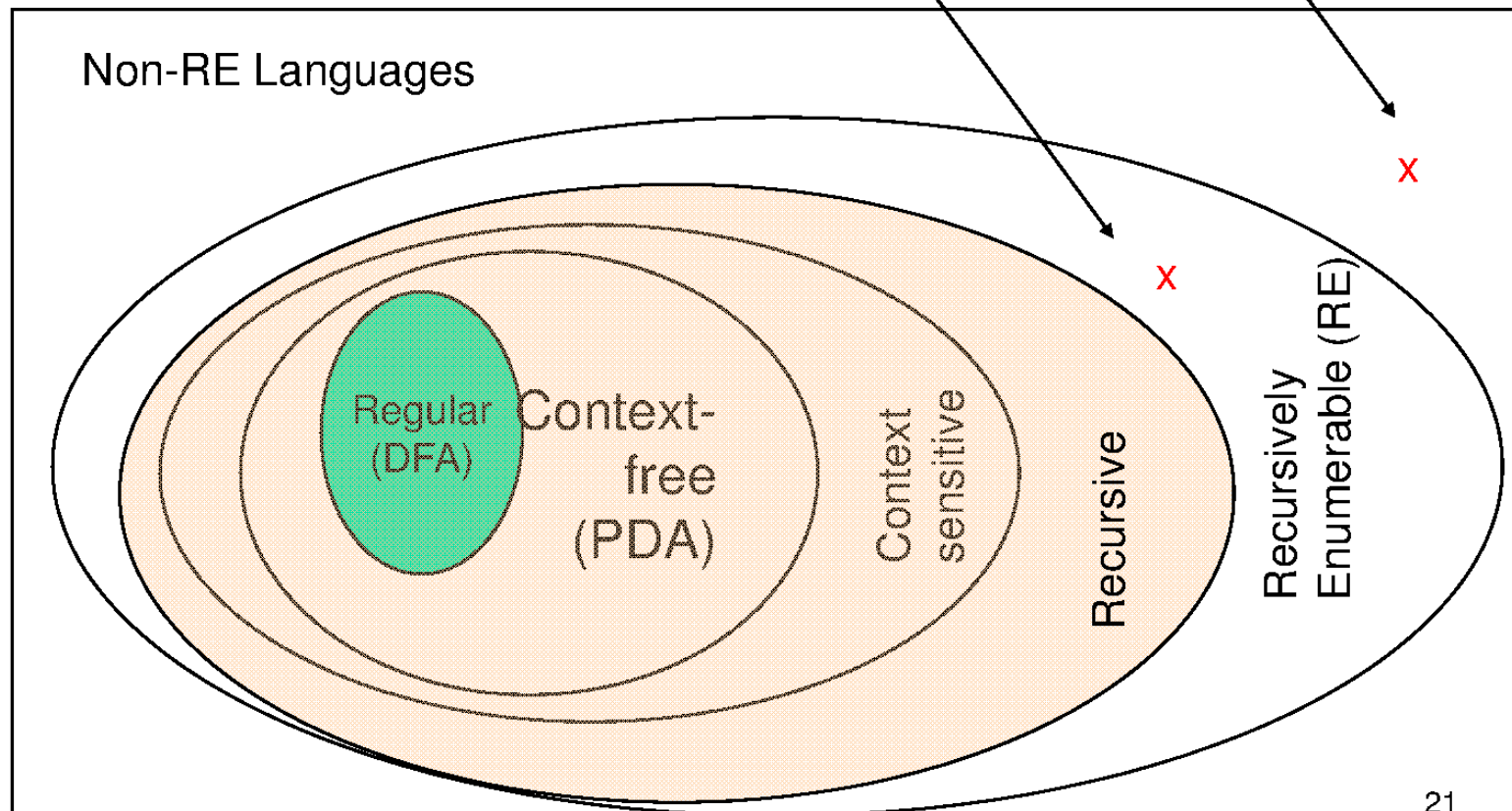
**Example of a language that is
not recursive enumerable**

(i.e, no TMs exist)



The Diagonalization language

The Halting Problem





A Language about TMs & acceptance

- Let L be the language of all strings $\langle M, w \rangle$ s.t.:
 1. M is a TM (coded in binary) with input alphabet also binary
 2. w is a binary string
 3. M accepts input w .



Enumerating all binary strings

- Let w be a binary string
- Then $1w \equiv i$, where i is some integer
 - E.g., If $w=\varepsilon$, then $i=1$;
 - If $w=0$, then $i=2$;
 - If $w=1$, then $i=3$; so on...
- If $1w \equiv i$, then call w as the i^{th} word or i^{th} binary string, denoted by w_i .
- **$\Rightarrow$ A canonical ordering of all binary strings:**
 - $\{\varepsilon, 0, 1, 00, 01, 10, 11, 000, 100, 101, 110, \dots\}$
 - $\{w_1, w_2, w_3, w_4, \dots, w_i, \dots\}$



Any TM M can also be binary-coded

- $M = \{ Q, \{0,1\}, \Gamma, \delta, q_0, B, F \}$
 - Map all states, tape symbols and transitions to integers ($\Rightarrow$ binary strings)
 - $\delta(q_i, X_j) = (q_k, X_l, D_m)$ will be represented as:
 - $\Rightarrow 0^i 1 0^j 1 0^k 1 0^l 1 0^m$
- Result: Each TM can be written down as a long binary string
- $\Rightarrow$ Canonical ordering of TMs:
 - $\{M_1, M_2, M_3, M_4, \dots, M_i, \dots\}$

The Diagonalization Language

- $L_d = \{ w_i \mid w_i \notin L(M_i) \}$
 - The language of all strings whose corresponding machine does *not* accept itself (i.e., its own code)

		(input word w)					
		j	1	2	3	4	...
(TMs)	i	1	0	1	0	1	...
	2	1	1	0	0	...	
	3	0	1	0	1	...	
	4	1	0	0	1	...	
	⋮						

• T

- Table: $T[i,j] = 1$, if M_i accepts w_j
 $= 0$, otherwise.

- Make a new language called
 $L_d = \{ w_i \mid T[i,i] = 0 \}$



L_d is not RE (i.e., has no TM)

- Proof (by contradiction):
- Let M be the TM for L_d
- $\implies M$ has to be equal to some M_k s.t.
$$L(M_k) = L_d$$
- $\implies$ Will w_k belong to $L(M_k)$ or not?
 1. If $w_k \in L(M_k) \implies T[k,k]=1 \implies w_k \notin L_d$
 2. If $w_k \notin L(M_k) \implies T[k,k]=0 \implies w_k \in L_d$
- A contradiction either way!!

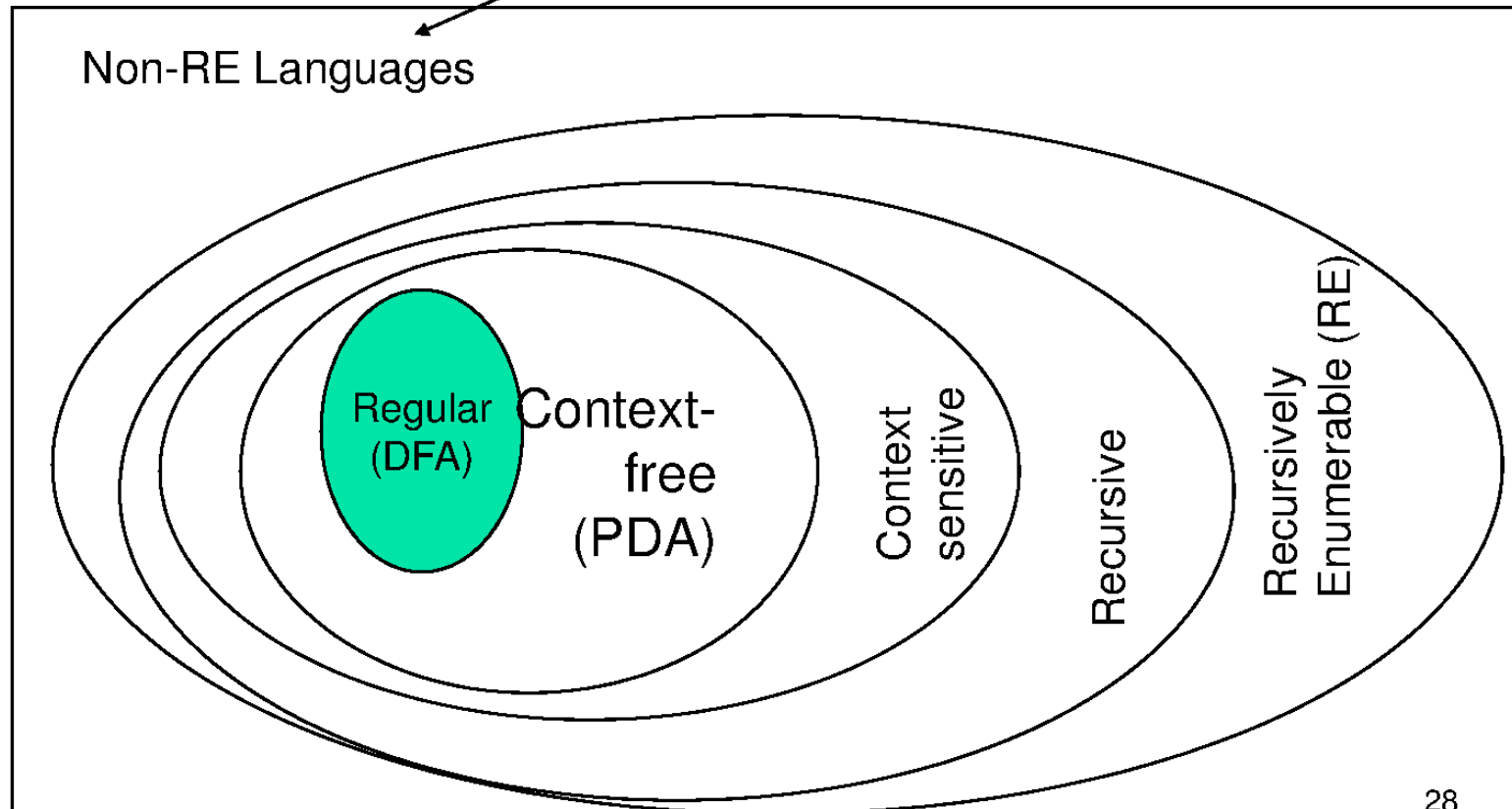
Why should there be
languages that do not have
TMs?



We thought TMs can solve
everything!!

Non-RE languages

How come there are languages here?
(e.g., diagonalization language)





One Explanation

There are more languages than TMs

- By pigeon hole principle:
 - $\implies$ some languages cannot have TMs
- But how do we show this?
- Need a way to “*count & compare*” two infinite sets (languages and TMs)



How to count elements in a set?

Let A be a set:

- If A is finite $\implies$ counting is trivial
- If A is infinite $\implies$ how do we count?
- And, how do we compare two infinite sets by their size?



Cantor's definition of set "size" for infinite sets (1873 A.D.)

Let $N = \{1, 2, 3, \dots\}$ (all natural numbers)

Let $E = \{2, 4, 6, \dots\}$ (all even numbers)

Q) Which is bigger?

■ A) Both sets are of the same size

■ "**Countably infinite**"

■ Proof: Show by one-to-one, onto set correspondence from

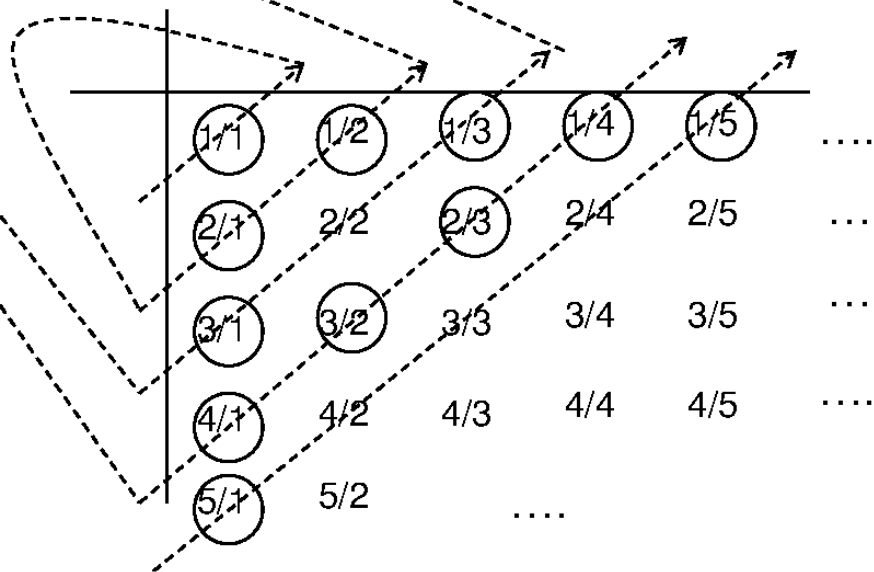
$N \implies E$

n	$f(n)$
1	2
2	4
3	6
$\vdots$	$\vdots$
$\vdots$	$\vdots$

i.e, for every element in N ,
there is a unique element in E ,
and vice versa.

Example #2

- Let Q be the set of all rational numbers
- $Q = \{ m/n \mid \text{for all } m, n \in \mathbb{N} \}$
- Claim: Q is also countably infinite; $\Rightarrow |Q| = |\mathbb{N}|$



Really, really big sets!
(even bigger than countably infinite sets)



Uncountable sets

Example:

- Let R be the set of all real numbers
- Claim: R is uncountable

n	$f(n)$
1	3 . <u>1</u> 4 1 5 9 ...
2	5 . 5 <u>5</u> 5 5 5 ...
3	0 . 1 2 <u>3</u> 4 5 ...
4	0 . 5 1 4 <u>3</u> 0 ...
.	
.	
.	

Build x s.t. x cannot possibly
occur in the table

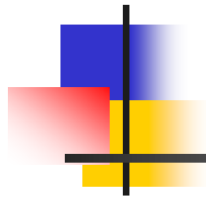
E.g. $x = 0 . 2 6 4 4 \dots$



Therefore, some languages cannot have TMs...

- The set of all TMs is countably infinite
- The set of all Languages is uncountable
- $\Rightarrow$ There should be some languages without TMs (by PHP)

The problem reduction technique, and reusing other constructions



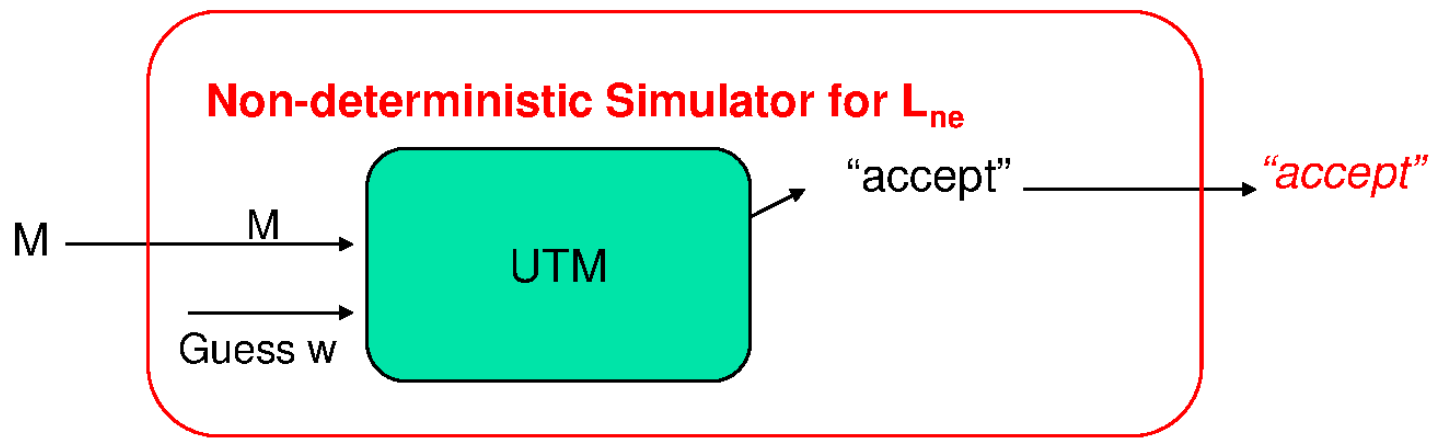


Languages that we know about

- *Language of a Universal TM (“UTM”)*
 - $L_u = \{ \langle M, w \rangle \mid M \text{ accepts } w \}$
 - Result: L_u is in RE but not recursive
- *Diagonalization language*
 - $L_d = \{ w_i \mid M_i \text{ does not accept } w_i \}$
 - Result: L_d is non-RE

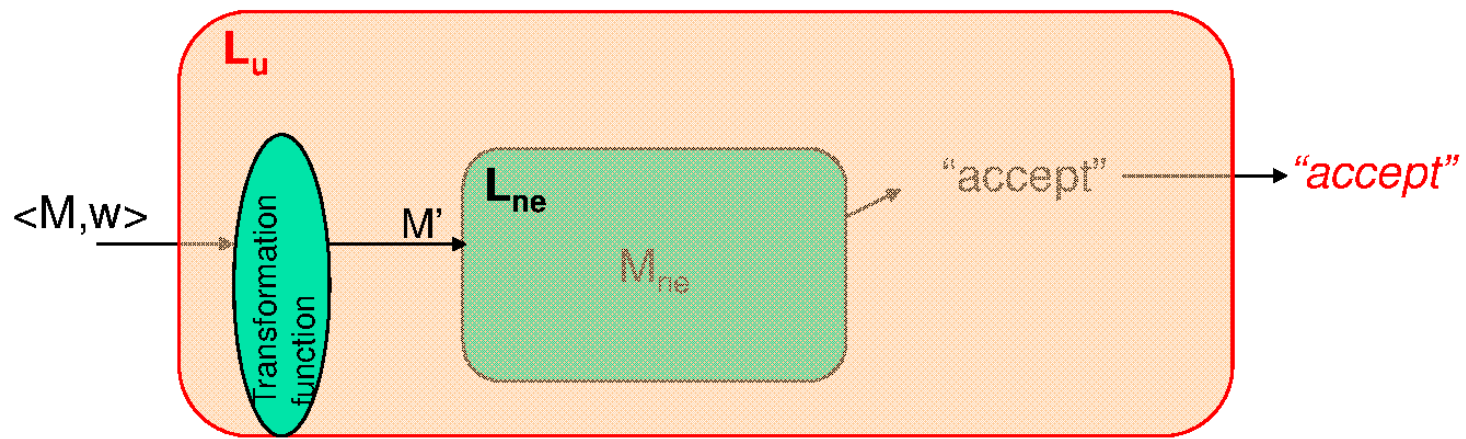
TMs that accept non-empty languages

- $L_{ne} = \{ M \mid L(M) \neq \emptyset \}$
- L_{ne} is RE
- Proof: (build a TM for L_{ne} using UTM)



TMs that accept non-empty languages

- L_{ne} is not recursive
- Proof: (“Reduce” L_u to L_{ne})
 - Idea: M accepts w if and only if $L(M') \neq \emptyset$





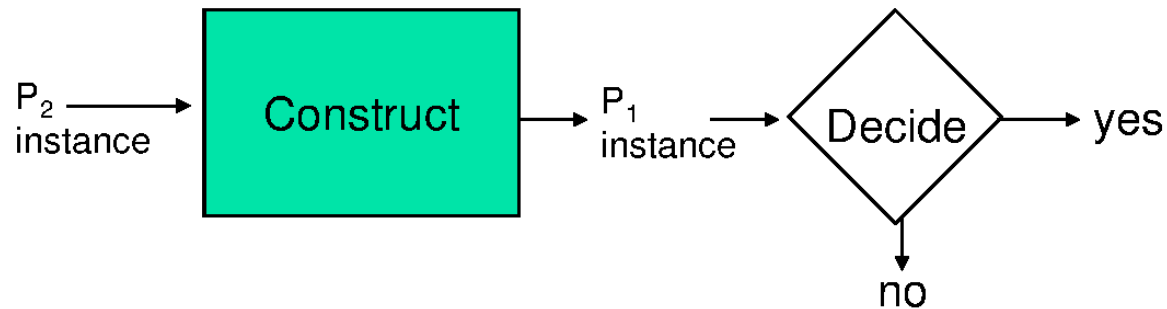
Reducability

- To prove: Problem P_1 is undecidable
- Given/known: Problem P_2 is undecidable
- Reduction idea:
 1. “Reduce” P_2 to P_1 :
 - Convert P_2 ’s input instance to P_1 ’s input instance s.t.
 - » P_2 decides only if P_1 decides
 2. Therefore, P_2 is decidable
 3. A contradiction
 4. Therefore, P_1 has to be undecidable

The Reduction Technique

Reduce P_2 to P_1 :

Note:
not same as
 $P_1 \implies P_2$



Conclusion: If we could solve P_1 , then we can solve P_2 as well



Summary

- Problems vs. languages
- Decidability
 - Recursive
- Undecidability
 - Recursively Enumerable
 - Not RE
 - Examples of languages
- The diagonalization technique
- Reducability