



your roots to success...

23IT601: Automata theory and compiler design

UNIT-I

Topic: Introduction to finite automata

M. Rajitha

Assistant Professor



Your roots to success....

NARSIMHA REDDY ENGINEERING COLLEGE
UGC AUTONOMOUS INSTITUTION

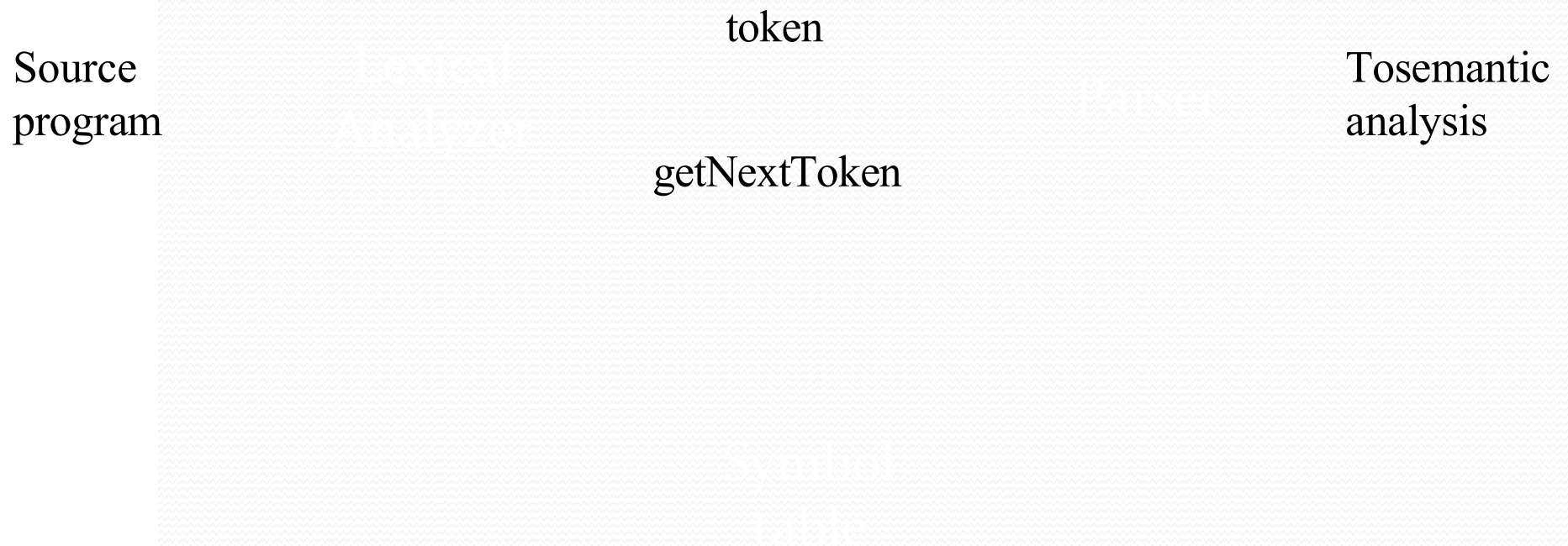
Maisammaguda (V), Kompally - 500100, Secunderabad, Telangana State, India

UGC - Autonomous Institute
Accredited by NBA & NAAC with 'A' Grade
Approved by AICTE
Permanently affiliated to JNTUH

Outline

- Role of lexical analyzer
- Specification of tokens
- Recognition of tokens
- Lexical analyzer generator
- Finite automata
- Design of lexical analyzer generator

The role of lexical analyzer



Why to separate Lexical analysis and parsing

1. Simplicity of design
2. Improving compiler efficiency
3. Enhancing compiler portability

Tokens, Patterns and Lexemes

- A token is a pair a token name and an optional token value
- A pattern is a description of the form that the lexemes of a token may take
- A lexeme is a sequence of characters in the source program that matches the pattern for a token

Example

Token	Informal description	Sample lexemes
if	Character s, f	if
else	Character e, l, s, e	else
comparison	< or > or <= or >= or == or !=	<=, !=
id	Letter followed by letter and digits	pi, score, D2
number	Any numeric constant	3.14159, 0, 6.02e23
literal	Anything but "surrounded by"	"core dumped"

```
printf("total=%d\n", score);
```

Attributesfortokens

- $E = M * C^{**}_2$
 - $\langle \text{id}, \text{pointertsymboltableentryforE} \rangle$
 - $\langle \text{assign-op} \rangle$
 - $\langle \text{id}, \text{pointertsymboltableentryforM} \rangle$
 - $\langle \text{mult-op} \rangle$
 - $\langle \text{id}, \text{pointertsymboltableentryforC} \rangle$
 - $\langle \text{exp-op} \rangle$
 - $\langle \text{number}, \text{integervalue}_2 \rangle$

Lexical errors

- Some errors are out of power of lexical analyzer to recognize:
 - $f(a=f(x))...$
- However it may be able to recognize errors like:
 - $d=2r$
- Such errors are recognized when no pattern for tokens matches a character sequence

Error recovery

- Panic mode: successive characters are ignored until we reach to a well formed token
- Delete one character from the remaining input
- Insert a missing character into the main input
- Replace a character by another character
- Transpose two adjacent characters



- E = M * C ** 2 eof



```
if(forwardisatendoffirstbuffer){ reload
    second buffer;
    forward=beginningofsecondbuffer;
```

}

```
Elseif {forward is at end of secondbuffer){
```

```
reload first buffer;\
```

```
forward=beginning offirst buffer;
```

}

```
else/*eof within a buffermarks the end of input*/
```

```
terminate lexical analysis;
```

```
break;
```

cases for the other characters;

Specification of tokens

- In theory of compilation regular expressions are used to formalize the specification of tokens
- Regular expressions are means for specifying regular languages
- Example:
 - Letter_(letter_|digit)*
- Each regular expression is a pattern specifying the form of strings

Regular expressions

- ϵ is a regular expression, $L(\epsilon) = \{\epsilon\}$
- If a is a symbol in Σ the naive regular expression, $L(a) = \{a\}$
- $(r)|(s)$ is a regular expression denoting the language $L(r) \cup L(s)$
- $(r)(s)$ is a regular expression denoting the language $L(r)L(s)$
- $(r)^*$ is a regular expression denoting $(L(r))^*$
- (r) is a regular expression denoting $L(r)$

Regular definitions

$d_1 \rightarrow r_1$

$d_2 \rightarrow r_2$

...

$d_n \rightarrow r_n$

- Example:

$\text{letter_} \rightarrow A|B|\dots|Z|a|b|\dots|Z|_ \text{ digit} \rightarrow$

$0 | 1 | \dots | 9$

$\text{id} \rightarrow \text{letter_}(\text{letter_}|\text{digit})^*$

Extensions

- One or more instances: $(r)^+$
- Zero of one instances: $r^?$
- Character classes: $[abc]$
- Example:
 - $\text{letter_} \rightarrow [A-Za-z_]$
 - $\text{digit} \rightarrow [0-9]$
 - $\text{id} \rightarrow \text{letter_}(\text{letter}|\text{digit})^*$

Recognition of tokens

- Starting point is the language grammar to understand the tokens:

stmt-**>if** expr **then** stmt
 | **if** expr **then** stmt **else** stmt
 | ϵ

expr-**>term relop** term
 | term

term -**>id**
 | **number**

Recognition of tokens(cont.)

- The next step is to formalize the patterns:

digit ->[0-9]

Digits ->digit+

number->digit(.digits)?(E[+-]?Digit)?

letter->[A-Za-z_]

id ->letter(letter|digit)*

If ->if

Then ->then

Else ->else

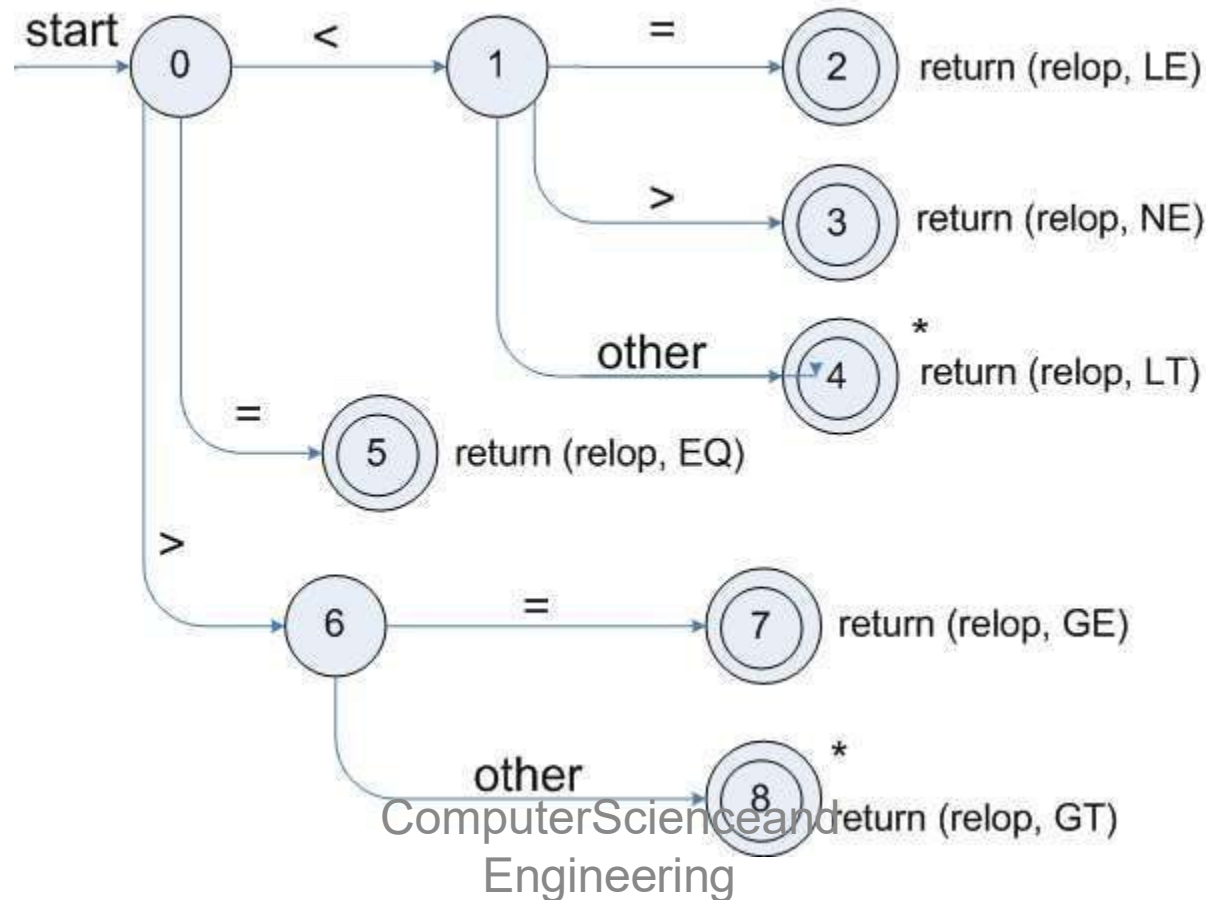
Relop ->< |>|<=|>=| =|<>

- We also need to handle whitespaces:

ws->(blank|tab|newline)+

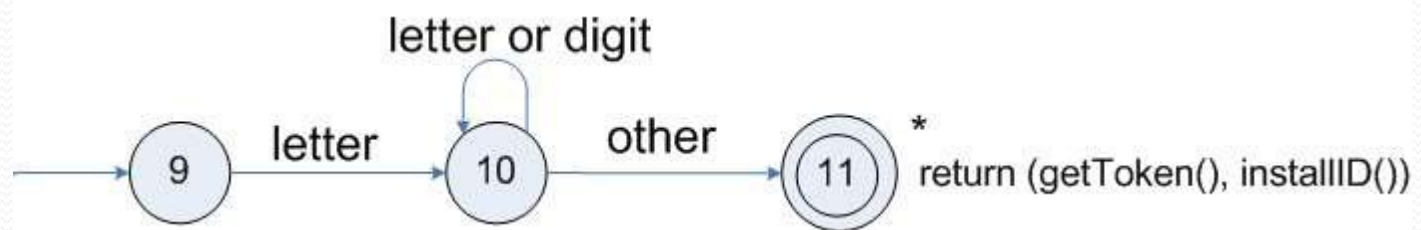
Transition diagrams

- Transition diagram for relop



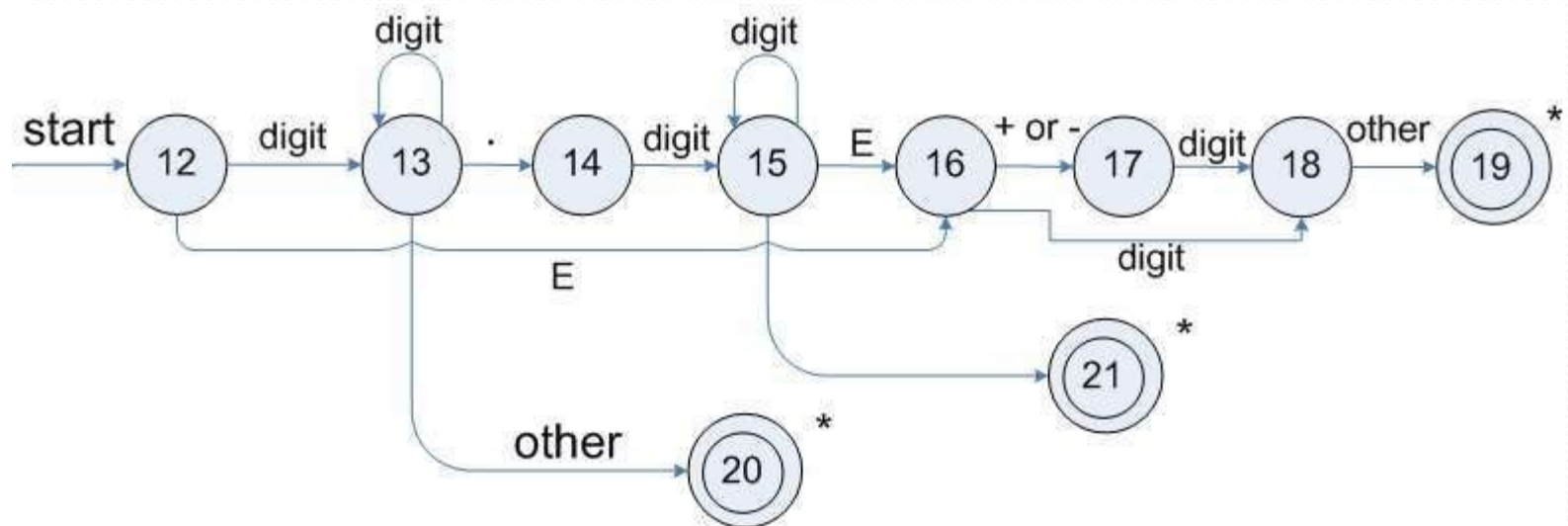
Transition diagrams(cont.)

- Transition diagram for reserved words and identifiers



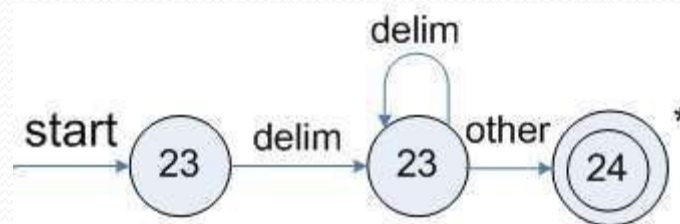
Transition diagrams(cont.)

- Transition diagram for unsigned numbers



Transition diagrams(cont.)

- Transition diagram for whitespace



Architecture of a transition-diagram-based lexical analyzer

```
TOKENgetRelop()
{
    TOKENretToken=new(RELOP)
    while(1){          /*repeatcharacterprocessinguntil a
                        returnorfailureoccurs*/
        switch(state){
            case0:c=nextchar();
                    if(c=='<')state=1;
                    else if (c=='=')state=5;
                    else if (c=='>')state=6;
                    else fail();/*lexemeisnotarelop*/ break;
            case1:...
            ...
            case8:retract();
                    retToken.attribute=GT;
                    return(retToken);
        }
    }
```

}



your roots to success...

Lexical Analyzer Generator-Lex

LexSourceprogram
lex.l



lex.yy.c



Input stream



Structure of Lex programs

declarations

%%

Translation rules

%%

Auxiliary
functions

Pattern {Action}

Example



your roots to success...

```
%{
    /*definitions of manifest constants
    LT,LE, EQ, NE, GT,GE,
    IF,THEN,ELSE,ID,NUMBER,RELOP*/
}%

/*regular definitions
delim      [\t\n]
ws         {delim}+
letter     [A-Za-z]
digit      [0-9]
id          {letter}({letter}|{digit})*
number     {digit}+(\.{digit}+)?(E[+-]?{digit}+)?

%%
{ws}       { /*no action and no return*/}
if         {return(IF);}
then       {return(THEN);}
else       {return(ELSE);}
{id}       {yylval=(int)install ID();return(ID);}
{number}   {yylval=(int)install Num();return(NUMBER);}
...

```

```
Int install ID() { /*function to
lexeme, whose first character is
pointed to by yy text, and whose
length is yy leng, into the symbol
table and return a pointer thereto
*/
}

```

```
Int install Num() { /* similar to install
ID, but puts numerical constants
into a separate table */
}

```

Finite Automata

- Regular expressions=specification
- Finite automata=implementation
- A finite automata consists of
 - An input alphabet Σ
 - A set of states S
 - A start state
 - A set of accepting states $F \subseteq S$
 - A set of transitions $\text{state} \rightarrow^{\text{input}} \text{state}$

Finite Automata

- Transition

$$s_1 \xrightarrow{a} s_2$$

- Is

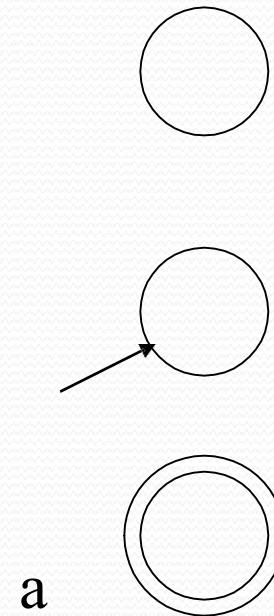
read In state s_1 on input “a” go to state s_2

- If end of input
 - If in accepting state=>accept, otherwise=>reject
- If no transition possible=>reject

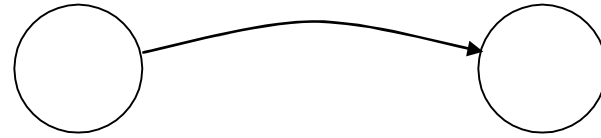


Finite Automata State Graphs

- A state
 - The start state
 - An accepting state

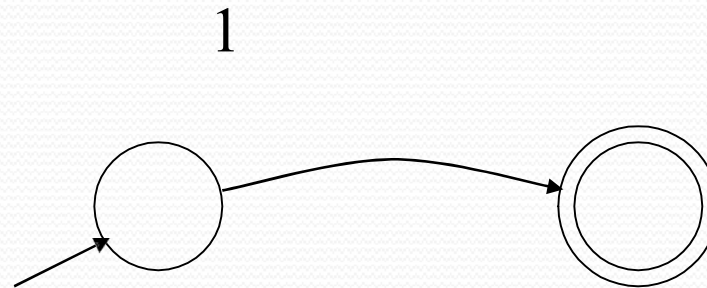


- A transition



A Simple Example

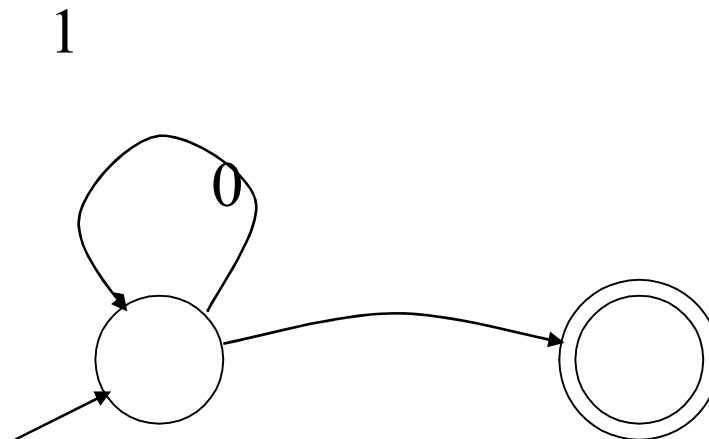
- A finite automaton that accepts only “1”



- A finite automaton accepts a string if we can follow transitions labeled with the characters in the string from the start to some accepting state

Another Simple Example

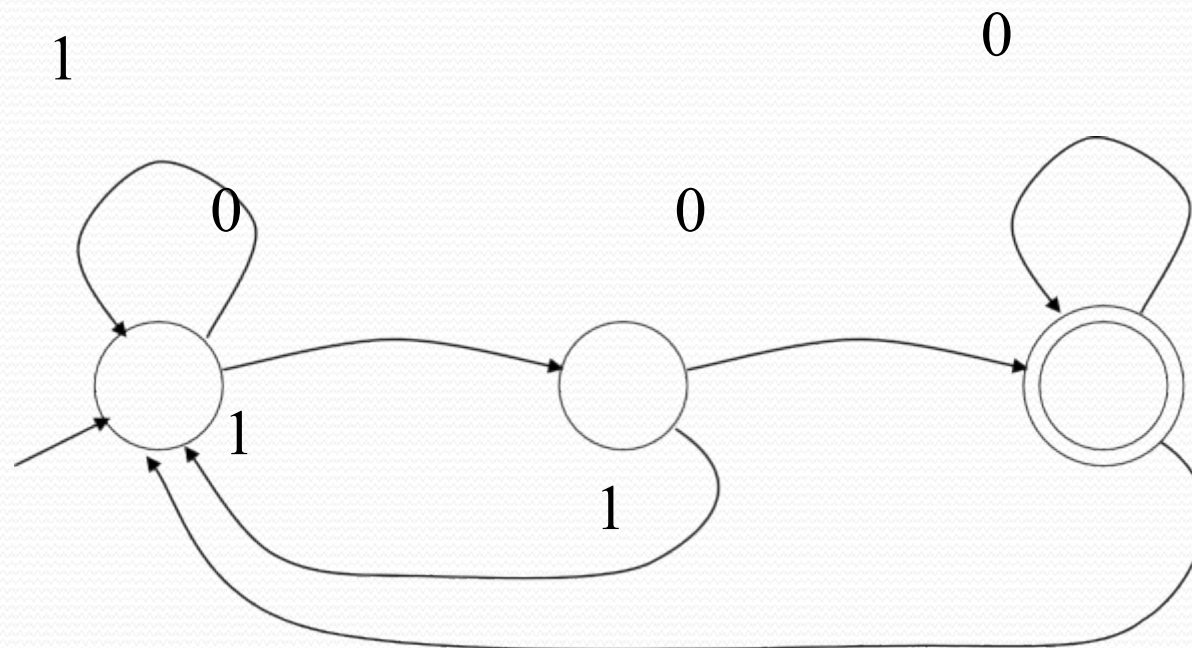
- A finite automaton accepting any number of 1's followed by a single 0
- Alphabet: $\{0,1\}$



- Check that “1110” is accepted but “110...” is not

And Another Example

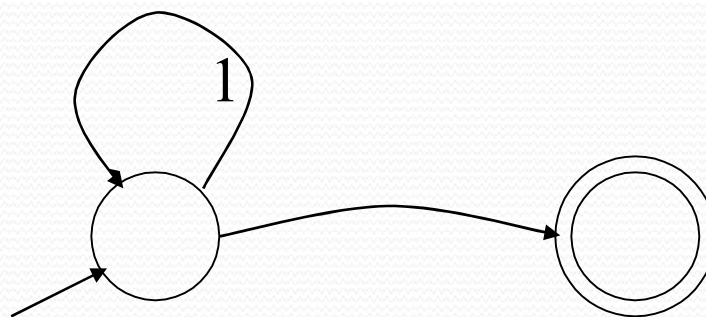
- Alphabet{0,1}
- What language does this recognize?



And Another Example

- Alphabet still $\{0,1\}$

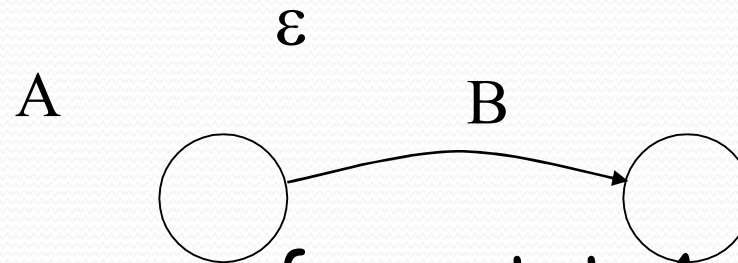
1



- The operation of the automaton is not completely defined by the input
 - On input “11” the automaton could be in either state

Epsilon Moves

- Another kind of transition: ϵ -moves



- Machine can move from state A to state B without reading input

Deterministic and Nondeterministic Automata



- Deterministic Finite Automata(DFA)
 - One transition per input per state
 - No ϵ -moves
- Non deterministic Finite Automata(NFA)
 - Can have multiple transitions for one input in a given state
 - Can have ϵ -moves
- *Finite* automata have *finite* memory
 - Need only to encode the current state



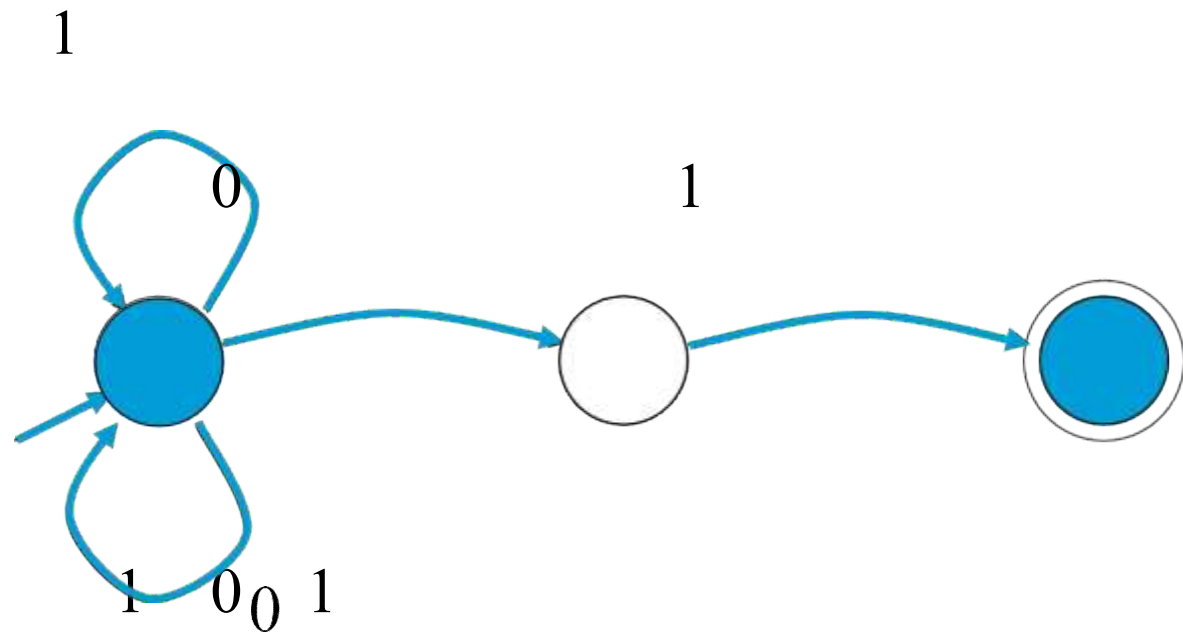
your roots to success...

Execution of Finite Automata

- A DFA can take only one path through the state graph
 - Completely determined by input
- NFAs can choose
 - Whether to make ϵ -moves
 - Which of multiple transitions for a single input to take

Acceptance of NFAs

- An NFA can get into multiple states



- Rule: NFA accepts if it can get in a final state



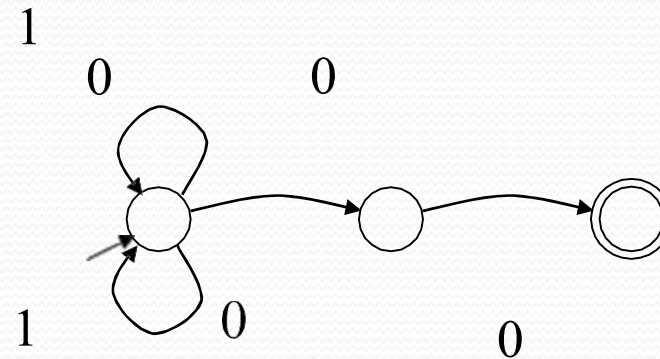
NFA vs. DFA (1)

- NFAs and DFAs recognize the same set of languages (regular languages)
- DFA are easier to implement
 - There are no choices to consider

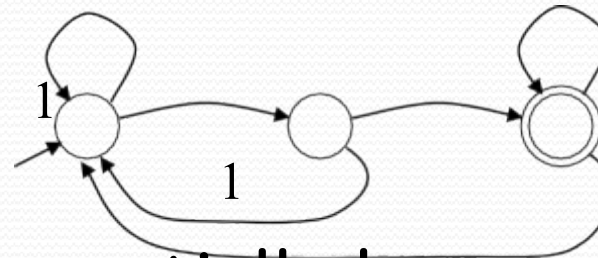
NFA vs. DFA(2)

- For a given language the NFA can be simpler than the DFA

NFA



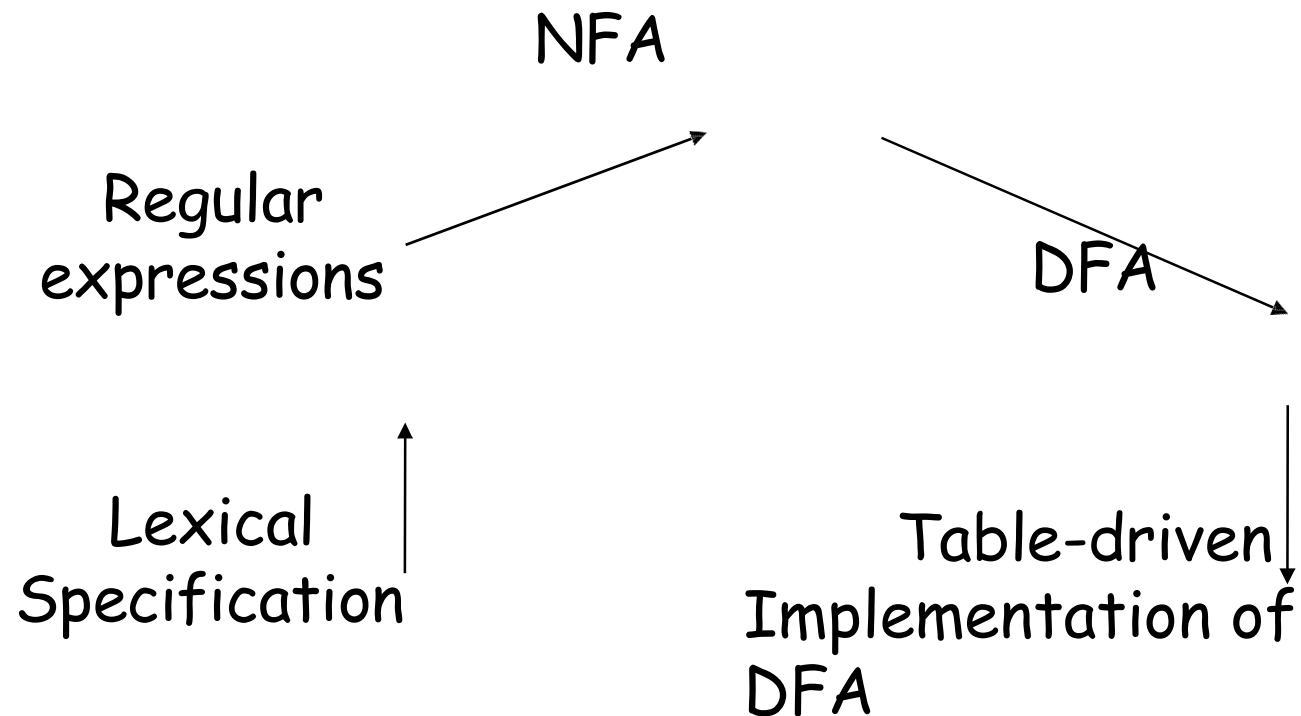
DFA



- DFA can be exponentially larger than NFA

Regular Expressions to Finite Automata

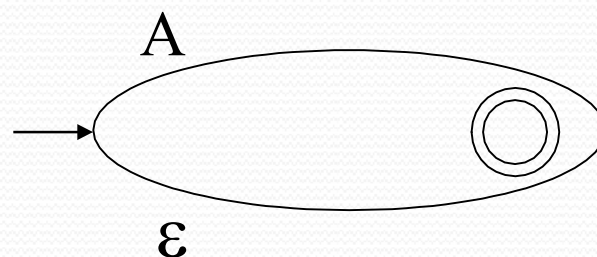
- High-level sketch



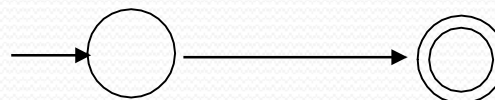
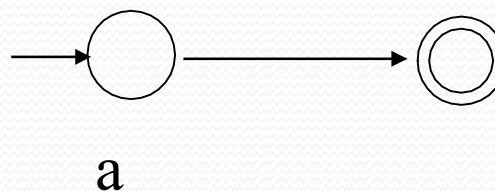
Regular Expressions to NFA(→,

- For each kind of rexp, define an NFA
 - Notation :NFA for rexpA

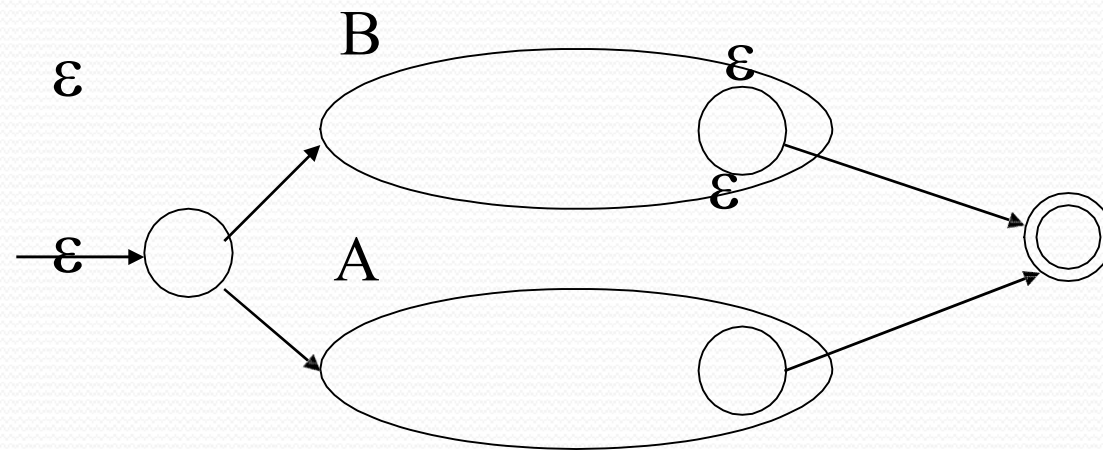
• For ϵ



• For input a

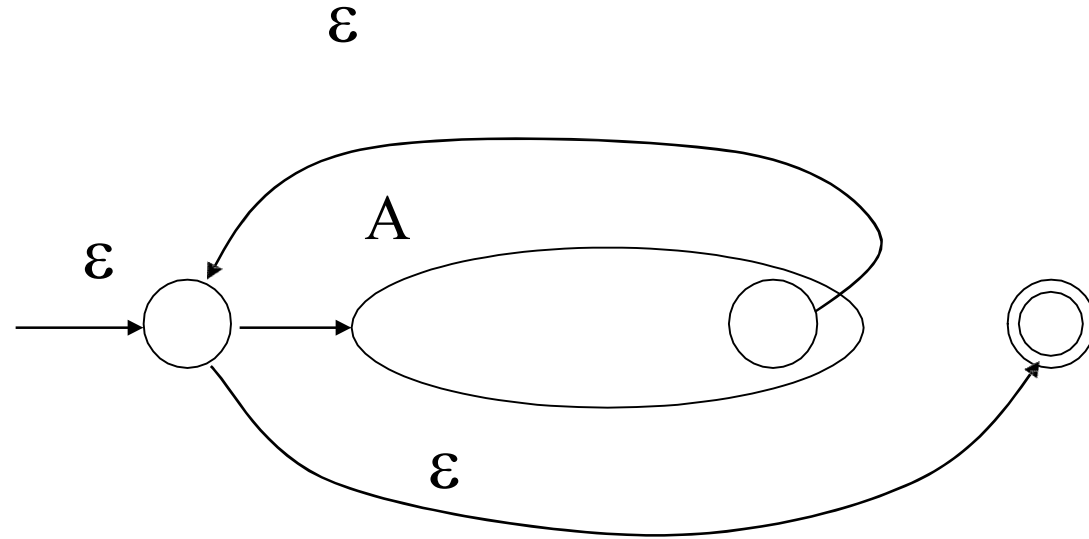


- For $A|B$



Regular Expressions to NFA(ϵ),

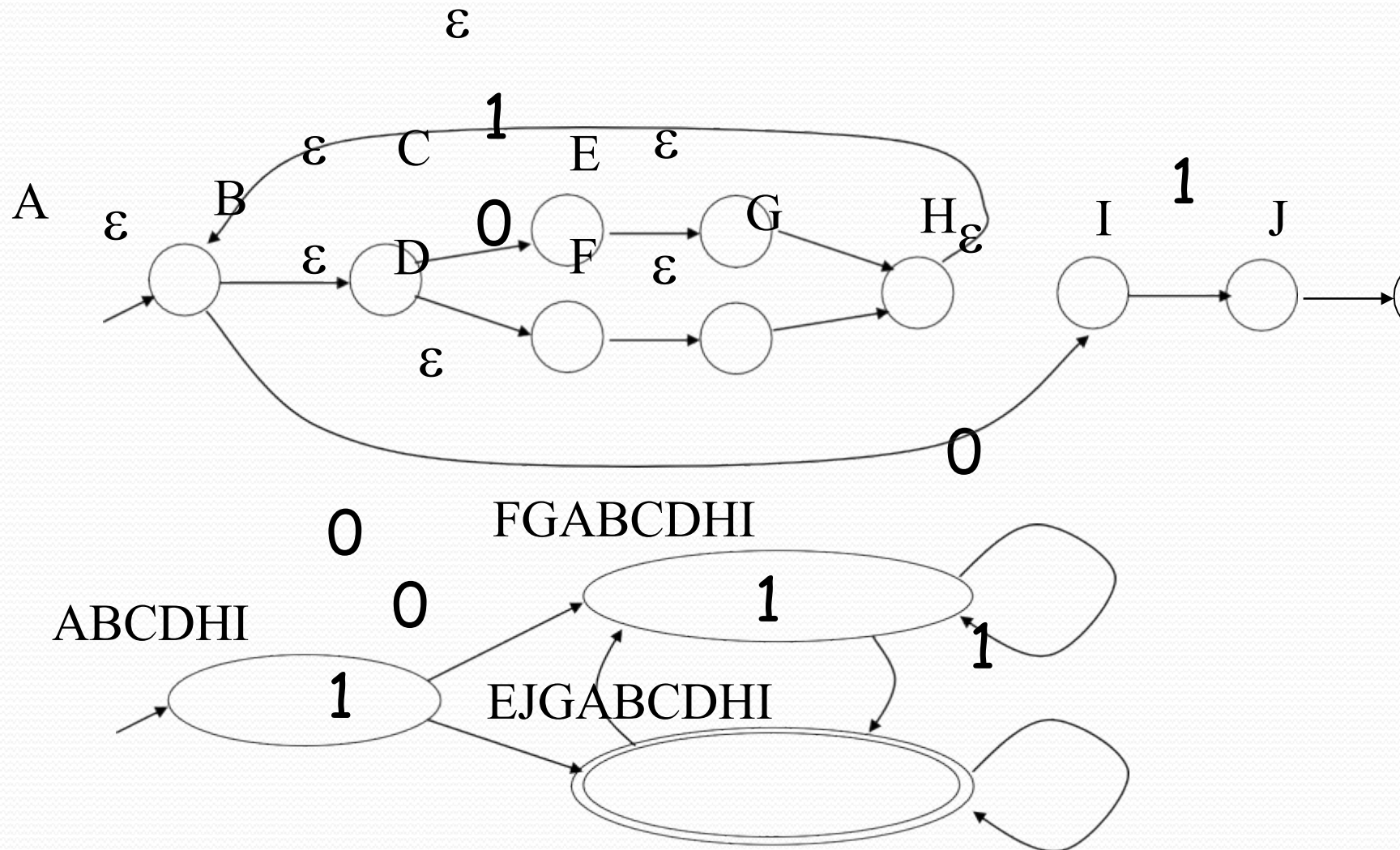
- For A^*



NFA to DFA. The Trick

- Simulate the NFA
- Each state of resulting DFA
= a non-empty subset of states of the NFA
- Start state
= the set of NFA states reachable through ϵ -moves from NFA start state
- Add a transition $S \xrightarrow{a} S'$ to DFA iff
 - S' is the set of NFA states reachable from the states in S After seeing the input a
 - considering ϵ -moves as well

NFA->DFA Example



$$x = *p$$

LDR_{1,p} //R_{1=p}

LDR_{2,0}(R₁) **//R₂=**

contents(o+contents(R1))

$$\mathbf{ST}_{\mathbf{x}, \mathbf{R}_2} \quad // \mathbf{x} = \mathbf{R}_2$$

conditional-jump three-address instruction

If $x < y$ goto L



LD $R_{1,x}$

LD $R_{2,y}$

SUBR $_{1,R_1,R_2}$ BLTZ R_1, M

//R1=x

//R2=y

//R1= R1-R2

//ifR1<0jumptoM



costs associated with the addressing modes

● LD R0,R1

● LD R0,M

● LD $R_1, *_{100}(R_2)$

cost=1

cost=2

cost=3

Addresses in the Target Code

- A statically determined area Code
- A statically determined data area Static
- A dynamically managed area Heap
- A dynamically managed area Stack

Procedure calls and returns

- Call callee
- Return
- Halt
- action

Target program for a sample call and return

	<i>// code for c</i>	100: ACTION ₁	<i>// code for c</i>
action ₁		120: ST 364, #140	<i>// code for action₁</i>
call p		132: BR 200	<i>// save return address 140 in location 364</i>
action ₂		140: ACTION ₂	<i>// call p</i>
halt		160: HALT	<i>// return to operating system</i>
		...	
	<i>// code for p</i>		<i>// code for p</i>
action ₃		200: ACTION ₃	
return		220: BR *364	<i>// return to address saved in location 364</i>
		...	
			<i>// 300-363 hold activation record for c</i>
		300:	<i>// return address</i>
		304:	<i>// local data for c</i>
		...	
			<i>// 364-451 hold activation record for p</i>
		364:	<i>// return address</i>
		368:	<i>// local data for p</i>

Stack Allocation

```
LD    SP, #stackStart           // initialize the stack
code for the first procedure

HALT                             // terminate execution
```

```
ADD    SP, SP, #caller.recordSize    // increment stack pointer
ST      *SP, #here + 16               // save return address
BR      callee.codeArea               // branch to callee procedure
// return to caller
```

Return to caller in caller:

In

Callee:

BR*0(SP)

SUB SP, Φ , #caller
r.recordsize

Target code for stack allocation

```

// code for m
action1
call q
action2
halt

// code for p
action3
return

// code for q
action4
call p
action5
call q
action6
call q
return

100: LD SP, #600           // initialize the stack
108: ACTION1             // code for action1
128: ADD SP, SP, #msize    // call sequence begins
136: ST *SP, #152          // push return address
144: BR 300                // call q
152: SUB SP, SP, #msize    // restore SP
160: ACTION12
180: HALT
...

200: ACTION3             // code for p
220: BR *0(SP)            // return
...

300: ACTION4             // code for q
                           // contains a conditional jump
336: BR 200                // call p
344: SUB SP, SP, #qsize    // push return address
352: ACTION5
372: ADD SP, SP, #qsize    // call q
380: BR *SP, #396          // push return address
388: BR 300                // call q
396: SUB SP, SP, #qsize
404: ACTION6
424: ADD SP, SP, #qsize
432: ST *SP, #440          // push return address
440: BR 300                // call q
448: SUB SP, SP, #qsize
456: BR *0(SP)            // return
600:                      // stack starts here

```

Basic blocks and flow graphs

- Partition the intermediate code into basic blocks
 - The flow of control can only enter the basic block through the first instruction in the block. That is, there are no jumps into the middle of the block.
 - Control will leave the block without halting or branching, except possibly at the last instruction in the block.
- The basic blocks become the nodes of a flowgraph

rules for finding leaders

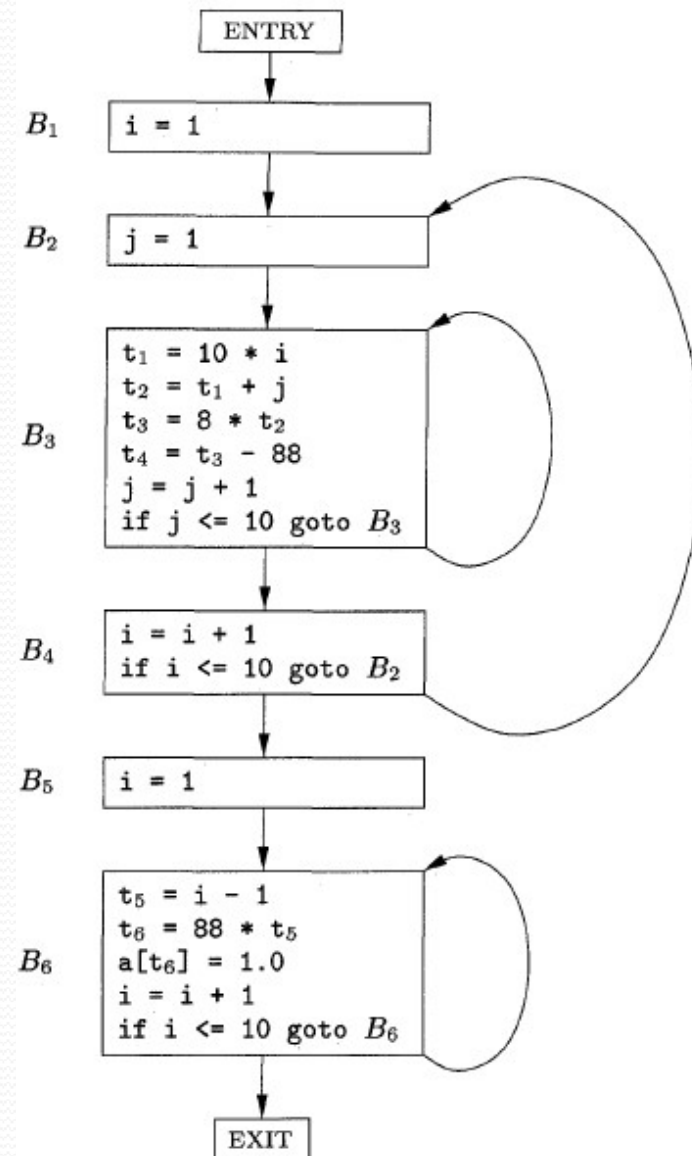
- The first three-address instruction in the intermediate code is a leader.
- Any instruction that is the target of a conditional or unconditional jump is a leader.
- Any instruction that immediately follows a conditional or unconditional jump is a leader.

Intermediate code to set a 10*10 matrix to an identity matrix

```
for i from 1 to 10 do
    for j from 1 to 10 do
         $a[i, j] = 0.0;$ 
for i from 1 to 10 do
     $a[i, i] = 1.0;$ 
```

```
1)  i = 1
2)  j = 1
3)  t1 = 10 * i
4)  t2 = t1 + j
5)  t3 = 8 * t2
6)  t4 = t3 - 88
7)  a[t4] = 0.0
8)  j = j + 1
9)  if j <= 10 goto (3)
10) i = i + 1
11) if i <= 10 goto (2)
12) i = 1
13) t5 = i - 1
14) t6 = 88 * t5
15) a[t6] = 1.0
16) i = i + 1
17) if i <= 10 goto (13)
```

Flowgraph based on BasicBlocks



liveness and next-use information

- We wish to determine for each three address statement $x=y+z$ what the next uses of x , y and z are.
- Algorithm:
 - Attach to statement the information currently found in the symbol table regarding then extuse and liveness of x , y , and z .
 - In the symbol table, set x to "not live" and "non extuse."
 - In the symbol table, set y and z to "live" and the next uses of y and z to i .

DAG representation of basic blocks

- There is a node in the DAG for each of the initial values of the variables appearing in the basic block.
- There is a node N associated with each statements within the block. The children of N are those nodes corresponding to statements that are the last definitions, prior to s , of the operands used by s .
- Node N is labeled by the operator applied at s , and also attached to N is the list of variables for which it is the last definition within the block.
- Certain nodes are designated *output nodes*. These are the nodes whose variables are *live on exit* from

the block.

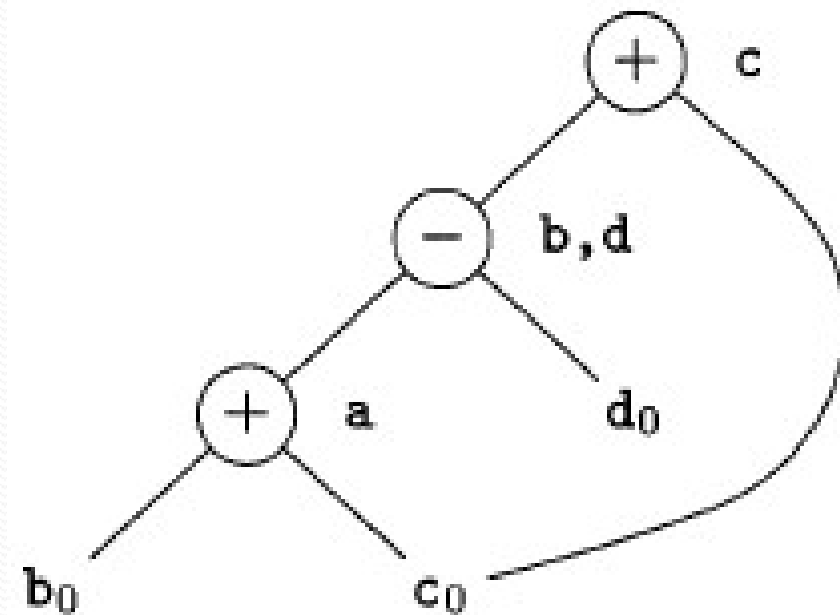
Code improving transformations

- We can eliminate *local common subexpressions*, that is, instructions that compute a value that has already been computed.
- We can eliminate *deadcode*, that is, instructions that compute a value that is never used.
- We can reorder statements that do not depend on one another; such reordering may reduce the time a temporary value needs to be preserved in a register.
- We can apply algebraic law store or deroper and s of three-address instructions, and sometimes thereby simplify the computation.

DAG for basicblock

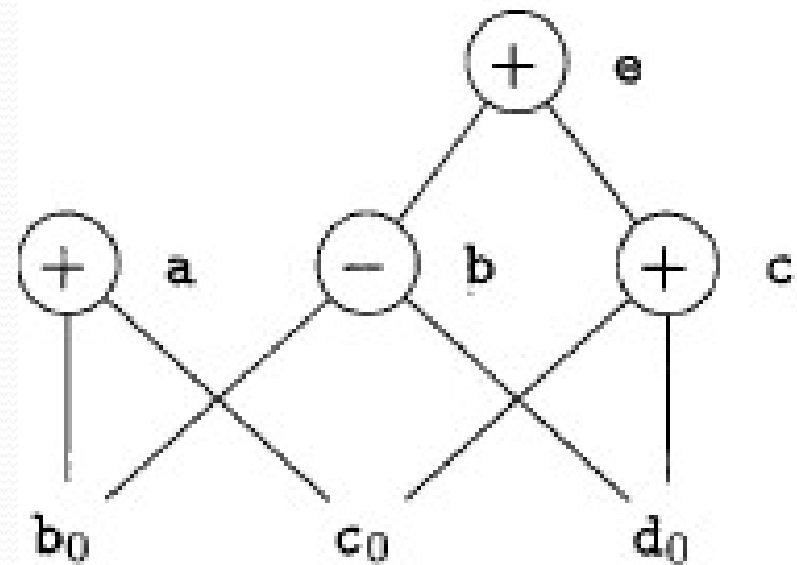
```

a = b + c
b = a - d
c = b + c
d = a - d
  
```



DAG for basicblock

```
a = b + c;
b = b - d
c = c + d
e = b + c
```

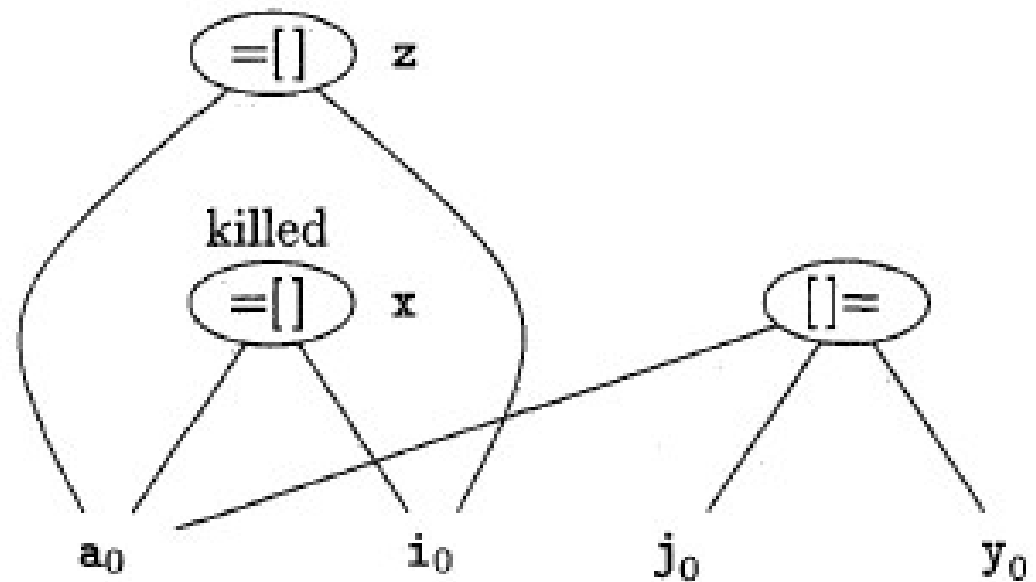


array accesses in a DAG

- An assignment from an array, like $x = a[i]$, is represented by creating a node with operator $=[]$ and two children representing the initial value of the array, a in this case, and the index i . Variable x becomes a label of this new node.
- An assignment to an array, like $a[j] = y$, is represented by a new node with operator $[] =$ and three children representing a , j , and y . There is no variable labeling this node. What is different is that the creation of this node *kills* all currently constructed nodes whose value depends on a . A node that has been killed cannot receive any more labels; that is, it cannot become a common subexpression.

DAG for a sequence of array assignments

```
x = a[i]
a[j] = y
z = a[i]
```



Rules for reconstructing the basic block from a DAG

- The order of instructions must respect the order of nodes in the DAG. That is, we cannot compute a node's value until we have computed a value for each of its children.
- Assignments to an array must follow all previous assignments to, or evaluations from, the same array, according to the order of these instructions in the original basic block.
- Evaluations of array elements must follow any previous (according to the original block) assignments to the same array. The only permutation allowed is that two evaluations from the same array may be done in either order, as long as neither crosses over an assignment to that array.
- Any use of a variable must follow all previous (according to the original block) procedure calls or indirect assignments through a pointer.
- Any procedure call or indirect assignment through a pointer must follow all previous (according to the original block) evaluations of any variable.

principal uses of registers

- In most machine architectures, some or all of the operands of an operation must be in registers in order to perform the operation.
- Registers make good temporaries—places to hold the result of a subexpression while a larger expression is being evaluated, or more generally, a place to hold a variable that is used only within a single basic block.
- Registers are often used to help with run-time storage management, for example, to manage the run-time stack, including the maintenance of stack pointers and possibly the top elements of the stack itself.

Descriptors for datastructure

- For each available register, a register descriptor keeps track of the variable names whose current value is in that register. Since we shall use only those registers that are available for local use within a basic block, we assume that initially, all register descriptors are empty. As the code generation progresses, each register will hold the value of zero or more names.
- For each program variable, an address descriptor keeps track of the location or locations where the current value of that variable can be found. The location might be a register, a memory address, a stack location, or some set of more than one of these. The information can be stored in the symbol-table entry for that variable name.

Machine Instructions for Operations

- Use `getReg(x=y+z)` to select registers for x , y , and z . Call these R_x , R_y and R_z .
- If y is not in R_y (according to the register descriptor for R_y), then issue an instruction `LD  $R_y, y'$` , where y' is one of the memory locations for y (according to the address descriptor for y).
- Similarly, if z is not in R_z , issue an instruction `LDR  $R_z, z'$` , where z' is a location for z .
- Issue the instruction `ADDR $_x, R_y, R_z$` .

Rules for updating the register and address descriptors

- For the instruction LDR, x
 - Change the register descriptor for register R_x so it holds only x .
 - Change the address descriptor for x by adding register R_x as an additional location.
- For the instruction STx, R , change the address descriptor for x to include its own memory location.
- For an operation such as $ADDR_x, R_y, R_z$ implementing a three-address instruction $x = y + x$
 - Change the register descriptor for R_x so that it holds only x .
 - Change the address descriptor for x so that its only location is R_x . Note that the memory location for x is *not* now in the address descriptor for x .
 - Remove R_x from the address descriptor of any variable other than x .
- When we process a copy statement $x = y$, after generating the load for y into register R_y , if needed, and after managing descriptors as for all load statements (per rule 1):
 - Add x to the register descriptor for R_y .
 - Change the address descriptor for x so that its only location is R_y .

Instructions generated and the changes in the register and address descriptors

```
a = d
LD R2, d
```

```
d = v + u
ADD R1, R3, R1
```

```
exit
ST a, R2
ST d, R1
```

```
ADD R3, R2, R1
```

R1	R2	R3	a	b	c	d	t	u	v
----	----	----	---	---	---	---	---	---	---

u	a, d	v	R2	b	c	d, R2			R1
---	------	---	----	---	---	-------	--	--	----

d	a	v	R2	b	c	R1			
---	---	---	----	---	---	----	--	--	--

d	a	v	a, R2	b	c	d, R1			
---	---	---	-------	---	---	-------	--	--	--

u	t	v	a	b	c	d	R2	R1	R3
---	---	---	---	---	---	---	----	----	----

Rules for picking register R_y for y

- If y is currently in a register, pick a register already containing y as R_y . Do not issue a machine instruction to load this register, as none is needed.
- If y is not in a register, but there is a register that is currently empty, pick one such register as R_y .
- The difficult case occurs when y is not in a register, and there is no register that is currently empty. We need to pick one of the allowable registers anyway, and we need to make it safe to reuse.

Possibilities for value of R

- If the address descriptor for v says that v is somewhere besides R , then we are OK.
- If v is x , the value being computed by instruction I , and x is not also one of the other operands of instruction I (in this example), then we are OK. The reason is that in this case, we know this value of x is never again going to be used, so we are free to ignore it.
- Otherwise, if v is not used later (that is, after the instruction I , there are no further uses of v , and if v is live on exit from the block, then v is recomputed within the block), then we are OK.
- If we are not OK by one of the first two cases, then we need to generate the store instruction ST_v, R to place a copy of v in its own memory location. This operation is called a spill.

Selection of the register Rx

1. Since a new value of x is being computed, a register that holds only x is always an acceptable choice for R_x .
2. If y is not used after instruction I , and R_y holds only y after being loaded, R_y can also be used as R_x . A similar option holds regarding z and R_x .

Possibilities for value of R

- If the address descriptor for v says that v is somewhere besides R , then we are OK.
- If v is x , the value being computed by instruction I , and x is not also one of the other operands of instruction I (in this example), then we are OK. The reason is that in this case, we know this value of x is never again going to be used, so we are free to ignore it.
- Otherwise, if v is not used later (that is, after the instruction I , there are no further uses of v , and if v is live on exit from the block, then v is recomputed within the block), then we are OK.
- If we are not OK by one of the first two cases, then we need to generate the store instruction ST_v, R to place a copy of v in its own memory location. This operation is called a spill.

Characteristic of peephole optimizations

- Redundant-instructionelimination
- Flow-of-controloptimizations
- Algebraicsimplifications
- Useofmachineidioms

Redundant-instruction elimination

- LDa, Ro
ST Ro, a
- ifdebug==1gotoL1
gotoL2
L1:printdebugginginformation
L2:

Flow-of-control optimizations

gotoL1

...

L1:gotoL2

Canbereplacedby:

goto L2

...

L1:gotoL2

ifa<bgotoL1

...

L1:gotoL2

Canbereplacedby: if

a<b goto L2

...

L1:gotoL2

Global register allocation

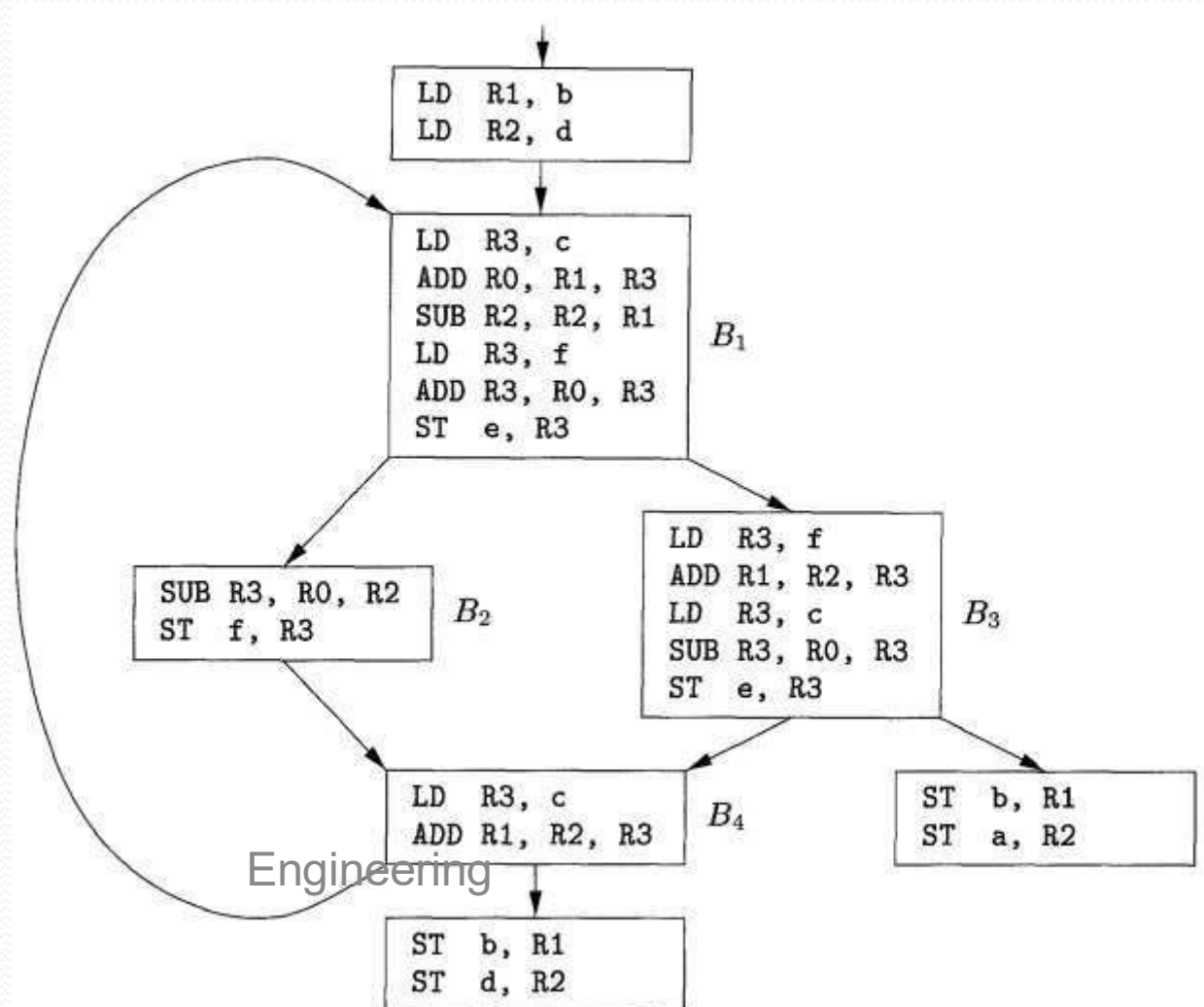
- Previously explained algorithm does local (block based) register allocation
- This resulted that all live variables be stored at the end of block
- To save some of these stores and their corresponding loads, we might arrange to assign registers to frequently used variables and keep these registers consistent across block boundaries (globally)
- Some options are:
 - Keep values of variables used in loops inside registers
 - Use graph coloring approach for more globally allocation

Usage counts

- For the loops we can approximate the saving by register allocation as:
 - $\text{Sum over all blocks } (B) \text{ in a loop } (L)$
 - For each use of x before any definition in the block we add one unit of saving
 - If x is live on exit from B and is assigned a value in B , then we add 2 units of saving

Flowgraph of an innerloop

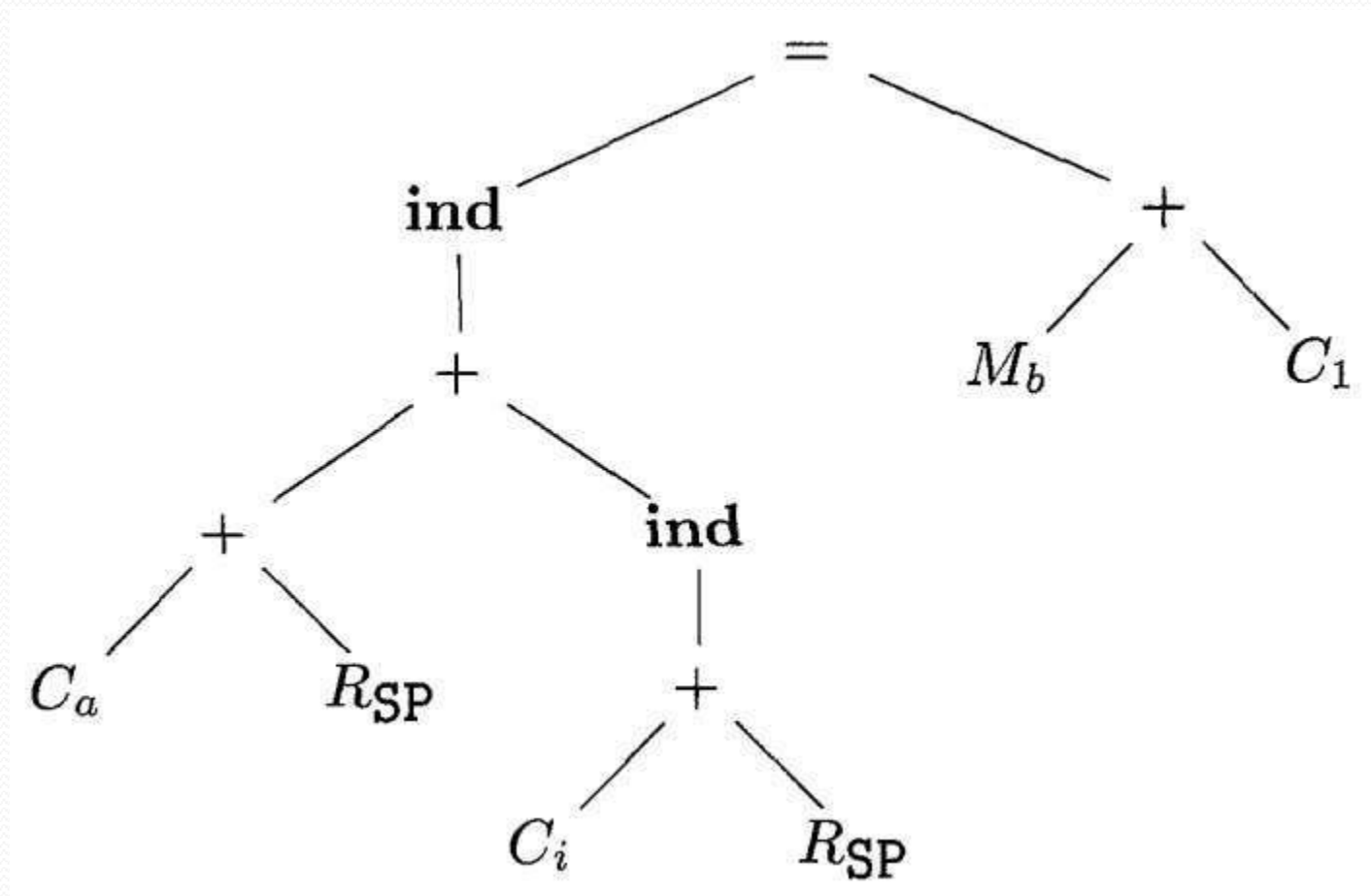
Code sequence using global register assignment



Register allocation by Graph coloring

- Two passes are used
 - Target-machine instructions are selected as though there are an infinite number of symbolic registers
 - Assign physical registers to symbolic ones
 - Create a register-interference graph
 - Nodes are symbolic registers and edges connect two nodes if one is live at a point where the other is defined.
 - For example in the previous example an edge connects a and d in the graph
 - Use a graph coloring algorithm to assign registers.

Intermediate-code tree for $a[i]=b+1$



Tree-rewriting rules

1)	$R_i \leftarrow C_a$	$\{ \text{LD } R_i, \#a \}$
2)	$R_i \leftarrow M_x$	$\{ \text{LD } R_i, x \}$
3)	$M \leftarrow \begin{array}{c} = \\ \swarrow \quad \searrow \\ M_x \quad R_i \end{array}$	$\{ \text{ST } x, R_i \}$
4)	$M \leftarrow \begin{array}{c} = \\ \swarrow \quad \searrow \\ \text{ind} \quad R_j \\ \\ R_i \end{array}$	$\{ \text{ST } *R_i, R_j \}$
5)	$R_i \leftarrow \begin{array}{c} \text{ind} \\ \\ + \\ \swarrow \quad \searrow \\ C_a \quad R_j \end{array}$	$\{ \text{LD } R_i, a(R_j) \}$
6)	$R_i \leftarrow \begin{array}{c} + \\ \swarrow \quad \searrow \\ R_i \quad \text{ind} \\ \quad \\ \quad + \\ \quad \swarrow \quad \searrow \\ \quad C_a \quad R_j \end{array}$	$\{ \text{ADD } R_i, R_i, a(R_j) \}$
7)	$R_i \leftarrow \begin{array}{c} + \\ \swarrow \quad \searrow \\ R_i \quad R_j \end{array}$	$\{ \text{ADD } R_i, R_i, R_j \}$
8)	$R_i \leftarrow \begin{array}{c} + \\ \swarrow \quad \searrow \\ R_i \quad C_1 \end{array}$	$\{ \text{INC } R_i \}$

Syntax-directed translation scheme

- | | | |
|-----|---|--------------------------------------|
| 1) | $R_i \rightarrow \mathbf{c}_a$ | $\{ \text{LD } R_i, \#a \}$ |
| 2) | $R_i \rightarrow M_x$ | $\{ \text{LD } R_i, x \}$ |
| 3) | $M \rightarrow = M_x R_i$ | $\{ \text{ST } x, R_i \}$ |
| 4) | $M \rightarrow = \mathbf{ind} R_i R_j$ | $\{ \text{ST } *R_i, R_j \}$ |
| 5) | $R_i \rightarrow \mathbf{ind} + \mathbf{c}_a R_j$ | $\{ \text{LD } R_i, a(R_j) \}$ |
| 6) | $R_i \rightarrow + R_i \mathbf{ind} + \mathbf{c}_a R_j$ | $\{ \text{ADD } R_i, R_i, a(R_j) \}$ |
| 7) | $R_i \rightarrow + R_i R_j$ | $\{ \text{ADD } R_i, R_i, R_j \}$ |
| 8) | $R_i \rightarrow + R_i \mathbf{c}_1$ | $\{ \text{INC } R_i \}$ |
| 9) | $R \rightarrow \mathbf{sp}$ | |
| 10) | $M \rightarrow \mathbf{m}$ | |

An instruction set for tree matching

$$R_i \leftarrow C_a$$

$$R_i \leftarrow \begin{array}{c} + \\ \swarrow \quad \searrow \\ R_i \quad \text{ind} \\ | \\ + \\ \swarrow \quad \searrow \\ C_a \quad R_j \end{array}$$

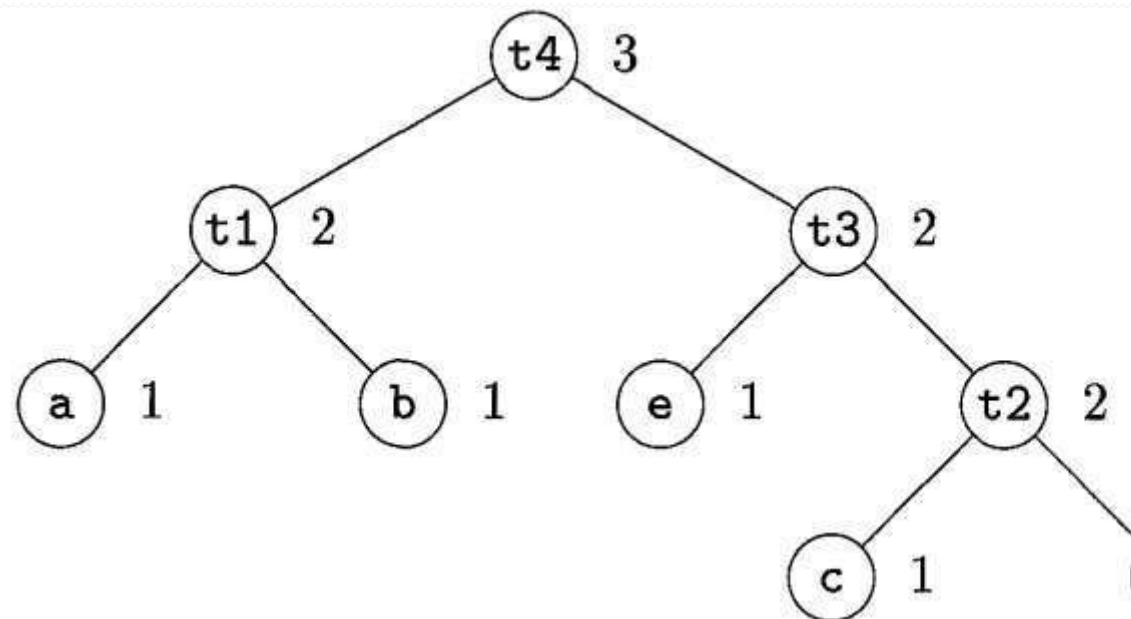
$$R_i \leftarrow \begin{array}{c} + \\ \swarrow \quad \searrow \\ R_i \quad R_j \end{array}$$

Ershov Numbers

- Label any leaf i .
- The label of an interior node with one child is the label of its child.
- The label of an interior node with two children is
 - The larger of the labels of its children, if those labels are different.
 - One plus the label of its children if the labels are the same.

A tree labeled with Ershov numbers

$t_1 = a - b$
 $t_2 = c + d$
 $t_3 = e * t_2$
 $t_4 = t_1 + t_3$



Generating code from a labeled expression tree

- To generate machine code for an interior node with label k and two children with equal labels (which must be $k-1$) do the following:
 - Recursively generate code for the right child, using base $b+1$. The result of the right child appears in register R_{b+k} .
 - Recursively generate code for the left child, using base b ; the result appears in R_{b+k-1} .
 - Generate the instruction $OP R_{b+k}, R_{b+k-1}, R_{b+k}$, where OP is the appropriate operation for the interior node in question.
- Suppose we have an interior node with label k and children with unequal labels. Then one of the children, which we'll call the "big" child, has label k , and the other child, the "little" child, has some label $m < k$. Do the following to generate code for this interior node, using base b :
 - Recursively generate code for the big child, using base b ; the result appears in register R_{b+k-1} .
 - Recursively generate code for the small child, using base b ; the result appears in register R_{b+m-1} . Note that since $m < k$, neither R_{b+k-1} nor any higher-numbered register is used.
 - Generate the instruction $OP R_{b+k-1}, R_{b+m-1}, R_{b+k-1}$ or the instruction $OP R_{b+k-1}, R_{b+k-1}, R_{b+m-1}$, depending on whether the big child is the right or left child, respectively.
- For a leaf representing operand x , if the base is b generate the instruction $LD R_b, x$.

Optimal three-register code

```
LD  R3, d
LD  R2, c
ADD R3, R2, R3
LD  R2, e
MUL R3, R2, R3
LD  R2, b
LD  R1, a
SUB R2, R1, R2
ADD R3, R2, R3
```

Evaluating Expressions with an Insufficient Supply of Registers

- Node N has at least one child with label r or greater. Pick the larger child (or either if their labels are the same) to be the "big" child and let the other child be the "little" child.
- Recursively generate code for the big child, using $base = 1$. The result of this evaluation will appear in register R_r .
- Generate the machine instruction ST_{t_k}, R_r , where t_k is a temporary variable used for temporary results used to help evaluate nodes with label k .
- Generate code for the little child as follows. If the little child has label r or greater, pick $base = 1$. If the label of the little child is $j < r$, then pick $b = r - j$. Then recursively apply this algorithm to the little child; the result appears in R_r .
- Generate the instruction LDR_{r-1}, t_k .
- If the big child is the right child of N , then generate the instruction $OP_{R_r, R_{r-1}}$. If the big child is the left child, generate $OP_{R_r, R_{r-1}}, R_r$.

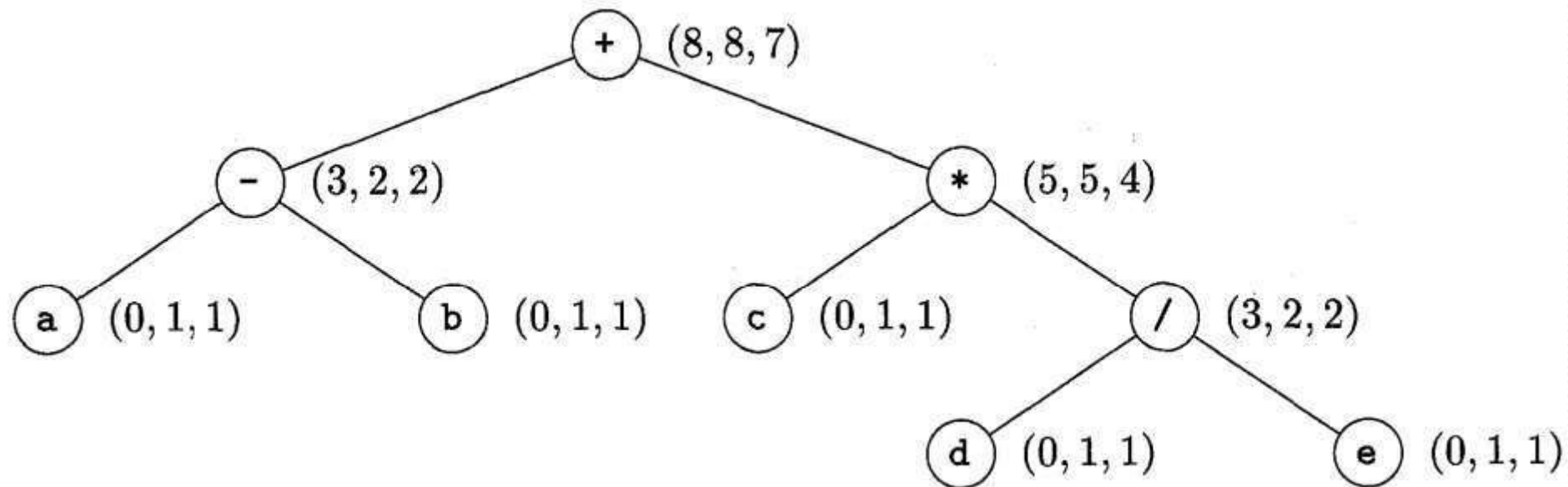
Optimal three-register code using only two registers

```
LD  R2, d
LD  R1, c
ADD R2, R1, R2
LD  R1, e
MUL R2, R1, R2
ST  t3, R2
LD  R2, b
LD  R1, a
SUB R2, R1, R2
LD  R1, t3
ADD R2, R2, R1
```

Dynamic Programming Algorithm

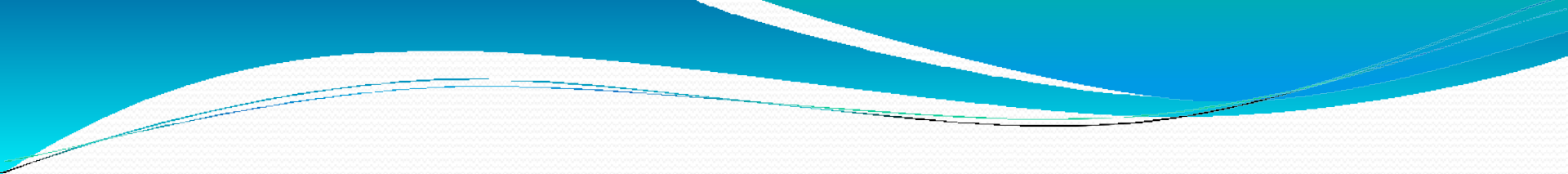
- Compute bottom-up for each node of the expression tree T an array C of costs, in which the i th component $C[i]$ is the optimal cost of computing the subtree S rooted at n into a register, assuming i registers are available for the computation, for
- Traverse T , using the cost vectors to determine which subtrees of T must be computed into memory.
- Traverse each tree using the cost vectors and associated instructions to generate the final target code. The code for the subtrees computed into memory locations is generated first.

Syntax tree for $(a-b)+c*(d/e)$ with cost vector at each node



minimum cost of evaluating the root with two registers available

- Computetheleftsubtreewithtwo registersavailableinto register R_0 , computetherightsubtreewithoneregister availableintoregister R_1 , andusetheinstruction $ADD R_0, R_0, R_1$ to computetheroot. Thissequencehascost $2+5+1=8$.
- Computetherightsubtreewithtwo registersavailableinto R_1 , computetheleftsubtreewithoneregisteravailable into R_0 , andusetheinstruction $ADD R_0, R_0, R_1$. This sequence hascost $4+2+1=7$.
- Computetherightsubtreeintomemorylocation M , computetheleftsubtreewithtwo registersavailableinto register R_0 , anduse $DD R_0, R_0, M$. This sequencehascost $5+2+1=8$.



UNIT-III

Syntax-Directed Translation

Outline

- Syntax Directed Definitions
- Evaluation Orders of SDD's
- Applications of Syntax Directed Translation
- Syntax Directed Translation Schemes

Introduction

- We can associate information with a language construct by attaching attributes to the grammar symbols.
- A syntax directed definition specifies the values of attributes by associating semantic rules with the grammar productions.

Production

$E \rightarrow E1 + T$

Semantic Rule

$E.code = E1.code || T.code || '+'$

- We may alternatively insert the semantic actions inside the grammar

$E\{$

Syntax Directed Definitions

- ASD is a context free grammar with attributes and rules
- Attributes are associated with grammar symbols and rules with productions
- Attributes may be of many kinds: numbers, types, table references, strings, etc.
- Synthesized attributes
 - A synthesized attribute at node N is defined only in terms of attribute values of children of N and N 's parent
- Inherited attributes
 - An inherited attribute at node N is defined only in terms of attribute values at N 's parent, N itself and N 's siblings

Example of S-attributed SDD

Production

- 1) $L \rightarrow E_n$
- 2) $E \rightarrow E_1 + T$
- 3) $E \rightarrow T$
- 4) $T \rightarrow T_1 * F$
- 5) $T \rightarrow F$
- 6) $F \rightarrow (E)$
- 7) $F \rightarrow \text{digit}$

Semantic Rules $L.val = E.val$

$E.val = E1.val + T.val$

$E.val = T.val$

$T.val = T1.val * F.val$ $T.val = F.val$

$F.val = E.val$

$F.val = \text{digit.lexval}$

Example of mixed attributes

Production

1) $T \rightarrow FT'$

2) $T' \rightarrow *FT'_1$

- 3) $T' \rightarrow \epsilon$
- 4) $F \rightarrow \text{digit}$

SemanticRules

$T'.inh = F.val$

$T.val = T'.syn$

$T'_1.inh = T'.inh * F.val$

$T'.syn = T'_1.syn$

$T'.syn = T'.inh$

$F.val = F.val = \text{digit.lexval}$

Evaluation orders for SDD's

- A dependency graph is used to determine the order of computation of attributes
- Dependency graph
 - For each parse tree node, the parse tree has a node for each attribute associated with that node
 - If a semantic rule defines the value of synthesized attribute $A.b$ in terms of the value of $X.c$ then the dependency graph has an edge from $X.c$ to $A.b$
 - If a semantic rule defines the value of inherited attribute $B.c$ in terms of the value of $X.a$ then the dependency graph has an edge from $X.c$ to $B.c$

Ordering the evaluation of attributes

- If dependency graph has an edge from M to N then M must be evaluated before the attribute of N
- Thus the only allowable orders of evaluation are those sequence of nodes $N_1, N_2, \dots, N_k$ such that if there is an edge from N_i to N_j then $i < j$
- Such an ordering is called a topological sort of a graph
- Example!

S-Attributed definitions

- An SDD is S-attributed if every attribute is synthesized
- We can have a post-order traversal of parse-tree to evaluate attributes in S-attributed definitions

```
postorder(N){  
    for(each child C of N, from the left) preorder(C);  
    evaluate the attributes associated with node N;  
}
```

- S-Attributed definitions can be implemented during bottom-up parsing without the need to explicitly create parse trees

L-Attributed definitions

- ASDDisL-Attributed if the edges in dependency graph goes from Left to Right but not from Right to Left.
- More precisely, each attribute must be either
 - Synthesized
 - Inherited, but if there is a production $A \rightarrow X_1 X_2 \dots X_n$ and there is an inherited attribute X_i computed by a rule associated with this production, then the rule may only use:
 - Inherited attributes associated with the head A
 - Either inherited or synthesized attributes associated with the occurrences of symbols $X_1, X_2, \dots, X_{i-1}$ located to the left of X_i
 - Inherited or synthesized attributes associated with this occurrence of X_i itself, but in such a way that there is no cycle in the graph

Application of Syntax Directed Translation

- Type checking and intermediate code generation (chapter 6)
- Construction of syntax trees
 - Leaf nodes: $\text{Leaf}(\text{op}, \text{val})$
 - Interior node: $\text{Node}(\text{op}, c_1, c_2, \dots, c_k)$
- Example:

Production

- 1) $E \rightarrow E_1 + T$
- 2) $E \rightarrow E_1 - T$
- 3) $E \rightarrow T$
- 4) $T \rightarrow (E)$
- 5) $T \rightarrow \text{id}$
- 6) $T \rightarrow \text{num}$

Semantic Rules

- $E.\text{node} = \text{new node}('+', E_1.\text{node}, T.\text{node})$
- $E.\text{node} = \text{new node}('-', E_1.\text{node}, T.\text{node})$
- $E.\text{node} = T.\text{node}$
- $T.\text{node} = E.\text{node}$
- $T.\text{node} = \text{new Leaf}(\text{id}, \text{id.entry})$
- $T.\text{node} = \text{new Leaf}(\text{num}, \text{num.val})$

Syntax tree for L-attributed definition

Production	SemanticRules ₊
1) $E \rightarrow TE'$	$E.\text{node} = E'.\text{syn}$ $E'.\text{inh} = T.\text{node}$
2) $E' \rightarrow +TE_1'$	$E_1'.\text{inh} = \text{newnode}('+', E'.\text{inh}, T.\text{node})$ $E'.\text{syn} = E_1'.\text{syn}$
3) $E' \rightarrow -TE_1'$	$E_1'.\text{inh} = \text{newnode}('-', E'.\text{inh}, T.\text{node})$ $E'.\text{syn} = E_1'.\text{syn}$
4) $E' \rightarrow \epsilon$	$E'.\text{syn} = E'.\text{inh}$
5) $T \rightarrow (E)$	$T.\text{node} = E.\text{node}$
6) $T \rightarrow \text{id}$	$T.\text{node} = \text{newLeaf}(\text{id}, \text{id}.\text{entry})$
7) $T \rightarrow \text{num}$	$T.\text{node} = \text{newLeaf}(\text{num}, \text{num}.\text{val})$

Syntax directed translation schemes

- An SDT is a Context Free grammar with program fragments embedded within production bodies
- Those program fragments are called semantic actions
- They can appear at any position within production body
- Any SDT can be implemented by first building a parse tree and then performing the actions in a left-to-right depth first order
- Typically SDT's are implemented during parsing without building a parse tree

Postfix translation schemes

- Simplest SDDs are those that we can parse the grammar bottom-up and the SDD is attributed
- For such cases we can construct SDT where each action is placed at the end of the production and is executed along with the reduction of the body to the head of that production
- SDT's with all actions at the right end of the production bodies are called postfix SDT's

Example of postfix SDT

- | | |
|---------------------------------|--|
| 1) $L \rightarrow En$ | $\{ \text{print}(E.\text{val}); \}$ |
| 2) $E \rightarrow E1 + T$ | $\{ E.\text{val} = E1.\text{val} + T.\text{val}; \}$ |
| 3) $E \rightarrow T$ | $\{ E.\text{val} = T.\text{val}; \}$ |
| 4) $T \rightarrow T1 * F$ | $\{ T.\text{val} = T1.\text{val} * F.\text{val}; \}$ |
| 5) $T \rightarrow F$ | $\{ T.\text{val} = F.\text{val}; \}$ |
| 6) $F \rightarrow (E)$ | $\{ F.\text{val} = E.\text{val}; \}$ |
| 7) $F \rightarrow \text{digit}$ | $\{ F.\text{val} = \text{digit}.\text{lexval}; \}$ |

Parse-Stack implementation of postfix SDT's

- In a shift-reduce parser we can easily implement semantic action using the parser stack
- For each non terminal(or state) on the stack we can associate a record holding its attributes
- Then in a reduction step we can execute the semantic action at the end of a production to evaluate the attribute(s) of then on-terminal at the left side of the production
- And put the value on the stack in replace of the right side of production

Example

L->En	{print(stack[top-1].val); top=top-1;}
E->E1+T E	{stack[top-2].val=stack[top-2].val+stack.val; top=top-2;}
-> T	
T->T1* F	{stack[top-2].val=stack[top-2].val+stack.val; top=top-2;}
T->F	
F -> (E)	{stack[top-2].val=stack[top-1].val top=top-2;}
F->digit	

SDT's with actions inside productions

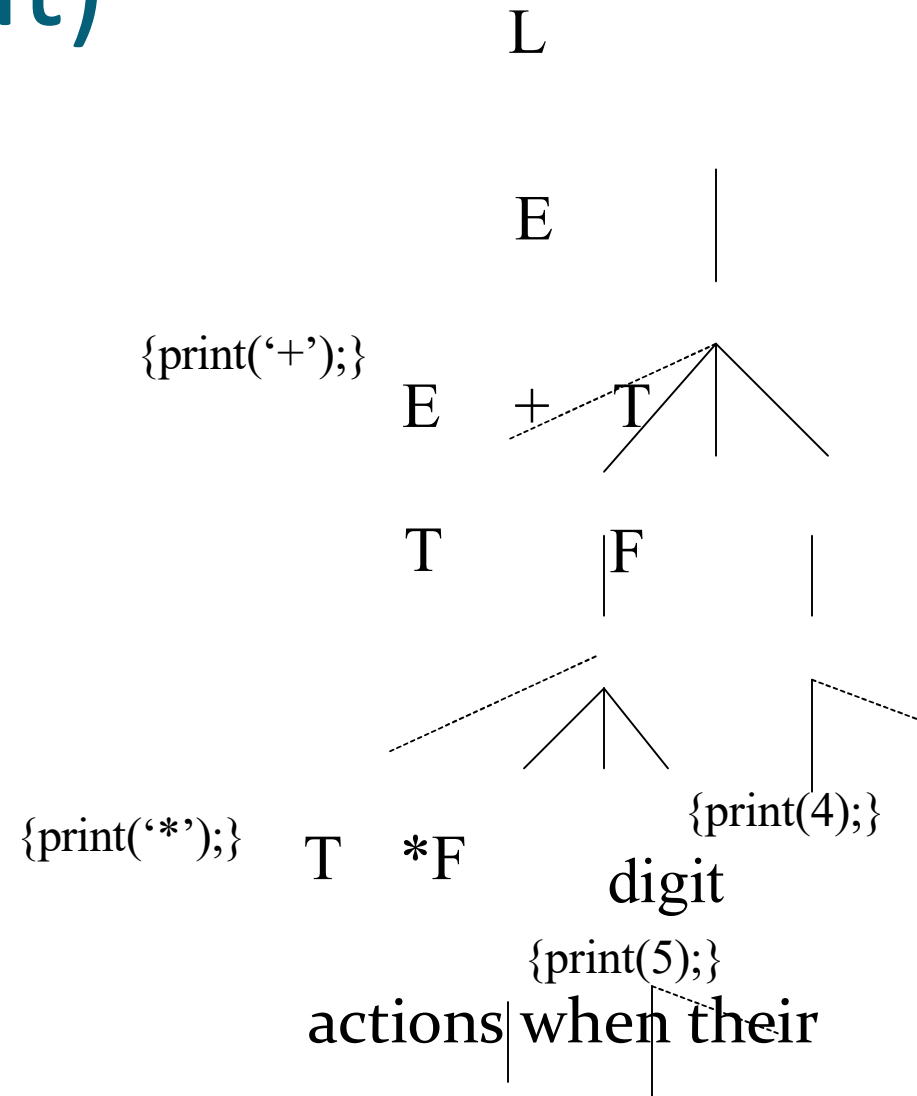
- For a production $B \rightarrow X\{a\}Y$
 - If the parse is bottom-up then we perform action “a” as soon as this occurrence of X appears on the top of the parser stack
 - If the parser is top down we perform “a” just before we expand Y
 - Sometimes we cant do things as easily as explained above
 - One example is when we are
- 1) $L \rightarrow En$
 - 2) $E \rightarrow \{\text{print}(' + '); \} E1 + T$
 - 3) $E \rightarrow T$
 - 4) $T \rightarrow \{\text{print}(' * '); \} T1 * F$
 - 5) $T \rightarrow F$
 - 6) $F \rightarrow (E)$

Computer Science and Engineering
parsing this SDT with a bottom-up 7) $F \rightarrow \text{digit} \{ \text{print}(\text{digit.lexval}); \}$

SDT's with actions inside productions (cont)

- Any SDT can be implemented as follows

- Ignore the actions and produce a parse tree
- Examine each interior Node N and add actions
As new children at the correct position
- Perform a postorder traversal and execute actions when their





nodes

are
visited

digit

F

digit

{print(3);}

SDT's for L-Attributed definitions

- We can convert an L-attributed SDD into an SDT using following two rules:
 - Embed the action that computes the inherited attributes for an on terminal A immediately before that occurrence of A. if several inherited attributes of A are dependent on one another in anacyclic fashion, order them so that those needed first are computed first
 - Place the action of a synthesized attribute for the head Of a production at the end of the body of the production

Example

```
S->while(C) S1      L1=new();  
                   L2=new();  
                   S1.next=L1;  
                   C.false=S.next;  
                   C.true=L2;  
                   S.code=label||L1||C.code||label||L2||S1.code
```

```
S->while( {L1=new();L2=new();C.false=S.next;C.true=L2;}  
C) {S1.next=L1;}  
S1 {S.code=label||L1||C.code||label||L2||S1.code;}
```

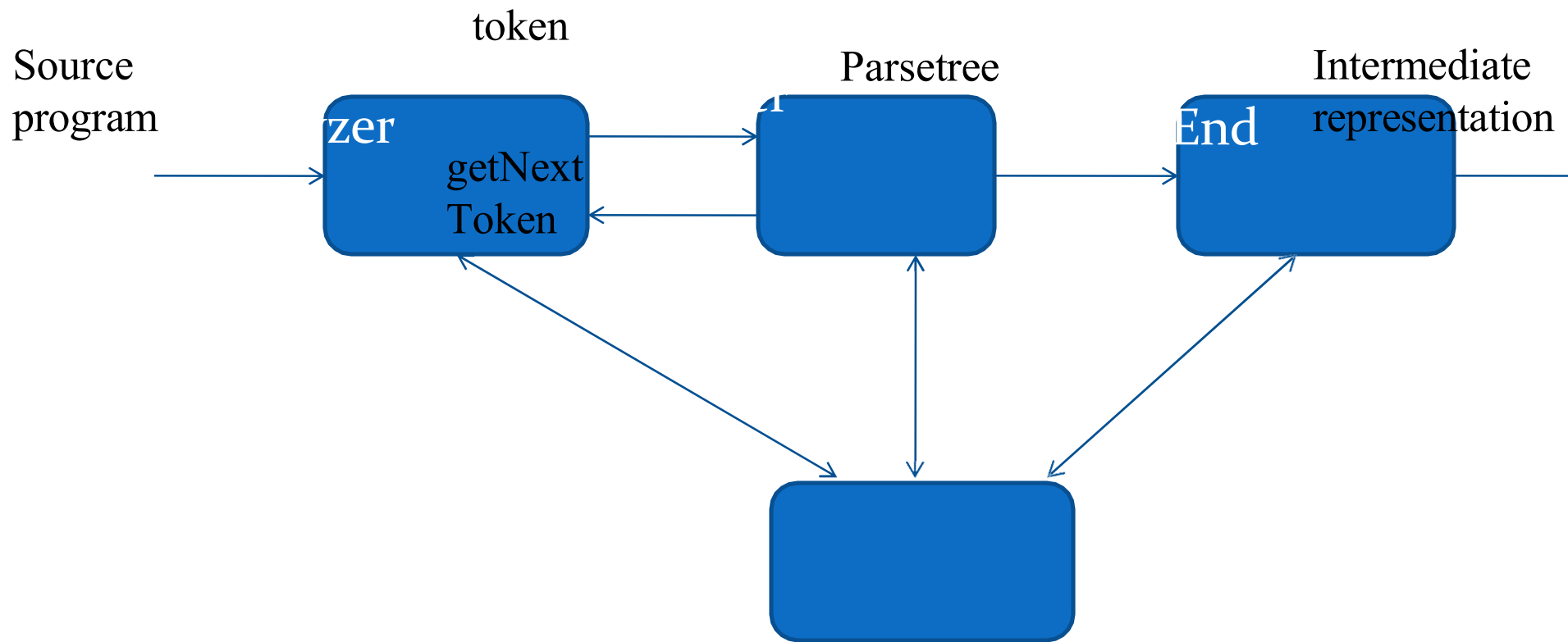
UNIT-II

Syntax Analyzer

Outline

- Role of parser
- Context free grammars
- Top down parsing
- Bottom up parsing
- Parser generators

The role of parser



Symbol
table



Uses of grammars

$E \rightarrow E + T \mid T$

$T \rightarrow T * F \mid F$

$F \rightarrow (E) \mid \text{id}$

$E \rightarrow TE'$

$E' \rightarrow +TE' \mid \varepsilon$

$T \rightarrow FT'$

$T' \rightarrow *FT' \mid \varepsilon$

$F \rightarrow (E) \mid \mathbf{id}$

Error handling

- Common programming errors
 - Lexical errors
 - Syntactic errors
 - Semantic errors
 - Lexical errors
- Error handler goals
 - Report the presence of errors clearly and accurately
 - Recover from each error quickly enough to detect subsequent errors
 - Add minimal over head to the processing of correct progrms

Error-recover strategies

- Panic mode recovery
 - Discard input symbol one at a time until one of designated set of synchronization tokens is found
- Phrase level recovery
 - Replacing a prefix of remaining input by some string That allows the parser to continue
- Error productions
 - Augment the grammar with productions that generate the erroneous constructs
- Global correction
 - Choosing minimal sequence of changes to obtain a globally least-cost correction

Context free grammars

- Terminals
- Non terminals
- Start symbol
- productions

$\text{expression} \rightarrow \text{expression} + \text{term}$
 $\text{expression} \rightarrow \text{expression} - \text{term}$
 $\text{expression} \rightarrow \text{term}$
 $\text{term} \rightarrow \text{term} * \text{factor}$
 $\text{term} \rightarrow \text{term} / \text{factor}$
 $\text{term} \rightarrow \text{factor}$
 $\text{factor} \rightarrow (\text{expression})$
 $\text{factor} \rightarrow \mathbf{id}$

Derivations

- Productions are treated as writing rules to generate a string
- Right most and left most derivations
 - $E \rightarrow E + E \mid E * E \mid -E \mid (E) \mid \text{id}$
 - Derivations for $-(\text{id} + \text{id})$
 - $E \Rightarrow -E \Rightarrow -(E) \Rightarrow -(E + E) \Rightarrow -(\text{id} + E) \Rightarrow -(\text{id} + \text{id})$

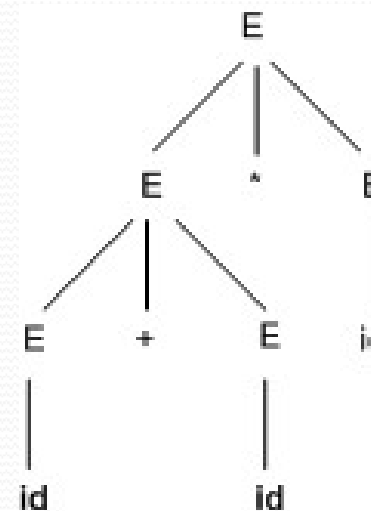
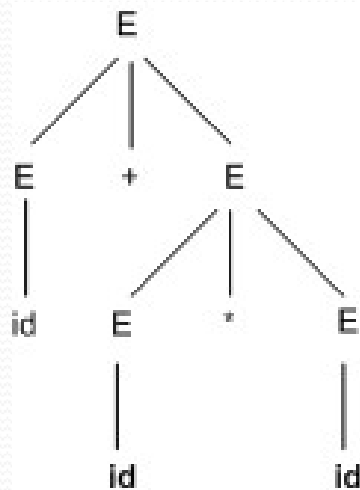
Parse trees

- **-(id+id)**

- $E \Rightarrow -E \Rightarrow -(E) \Rightarrow -(E+E) \Rightarrow -(\text{id}+E) \Rightarrow -(\text{id}+\text{id})$

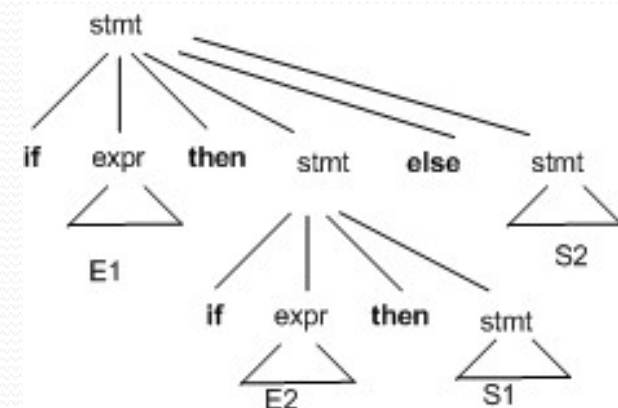
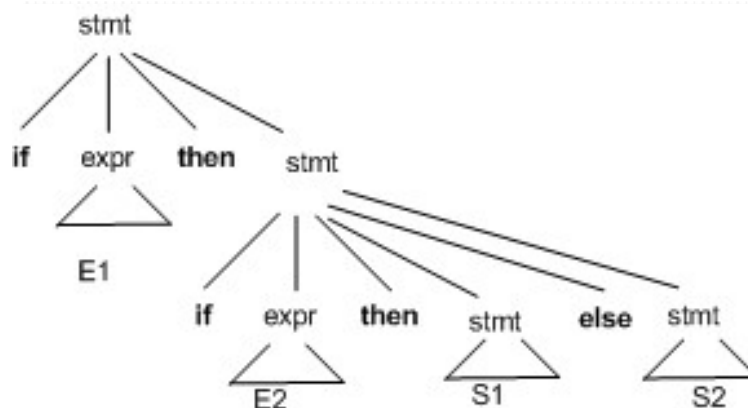
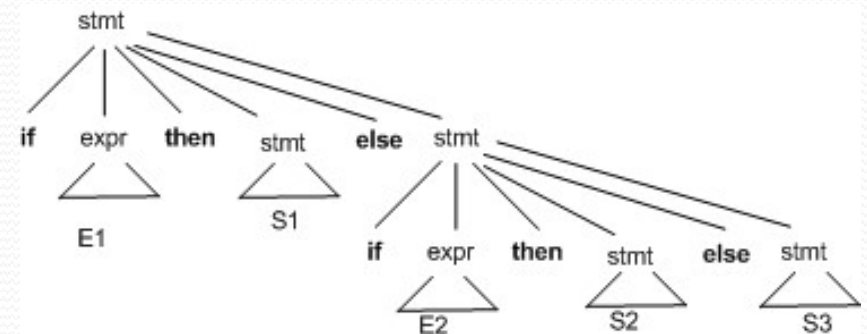
Ambiguity

- For some strings there exist more than one parse tree
- Or more than one left most derivation
- Or more than one right most derivation
- Example: $\text{id} + \text{id} * \text{id}$



Elimination of ambiguity

stmt $\rightarrow$ **if** expr **then** stmt
 | **if** expr **then** stmt **else** stmt
 | **other**



Elimination of ambiguity(cont.)

- Idea:

- A statement appearing between a **then** and an **else**
Must be matched

```
stmt  →  matched_stmt  
      |  open_stmt  
  
matched_stmt → If expr then matched_stmt else matched_stmt  
              | other  
  
open_stmt  →  If expr then stmt  
              |  If expr then matched_stmt else open_stmt
```

Elimination of left recursion

- A grammar is left recursive if it has a non-terminal A such that there is a derivation $A \xRightarrow{+} A \alpha$
- Top down parsing methods cant handle left-recursive grammars
- A simple rule for direct left recursion elimination:
 - For a rule like:
 - $A \rightarrow A \alpha \mid \beta$
 - We may replace it with
 - $A \rightarrow \beta A'$
 - $A' \rightarrow \alpha A' \mid \epsilon$

Left recursion elimination (cont.)

- There are cases like following
 - $S \rightarrow Aa|b$
 - $A \rightarrow Ac|Sd|\epsilon$
- Left recursion elimination algorithm:
 - Arrange the nonterminals in some order $A_1, A_2, \dots, A_n$.
 - For (each i from 1 to n) {
 - For (each j from 1 to $i-1$) {
 - Replace each production of the form $A_i \rightarrow A_j \gamma$ by the production $A_i \rightarrow \delta_1 \gamma | \delta_2 \gamma | \dots | \delta_k \gamma$ where $A_j \rightarrow \delta_1 | \delta_2 | \dots | \delta_k$ are all current A_j productions
 - }
 - Eliminate left recursion among the A_i -productions
 - }

Left factoring

- Left factoring is a grammar transformation that is useful for producing a grammar suitable for predictive or top-down parsing.
- Consider following grammar:
 - Stmt $\rightarrow$ **if** expr **then** stmt **else** stmt
 - | **if** expr **then** stmt
- On seeing input **if** it is not clear for the parser which production to use
- We can easily perform left factoring:
 - If we have $A \rightarrow \alpha \beta_1 \mid \alpha \beta_2$ Then we place it with
 - $A \rightarrow \alpha A'$
 - $A' \rightarrow \beta_1 \mid \beta_2$

Left factoring(cont.)

Algorithm

- For each non-terminal A , find the longest prefix α common to two or more of its alternatives. If $\alpha \neq \epsilon$, then replace all of A -productions $A \rightarrow \alpha \beta_1 \mid \alpha \beta_2 \mid \dots \mid \alpha \beta_n \mid \gamma$ by
 - $A \rightarrow \alpha A' \mid \gamma$
 - $A' \rightarrow \beta_1 \mid \beta_2 \mid \dots \mid \beta_n$

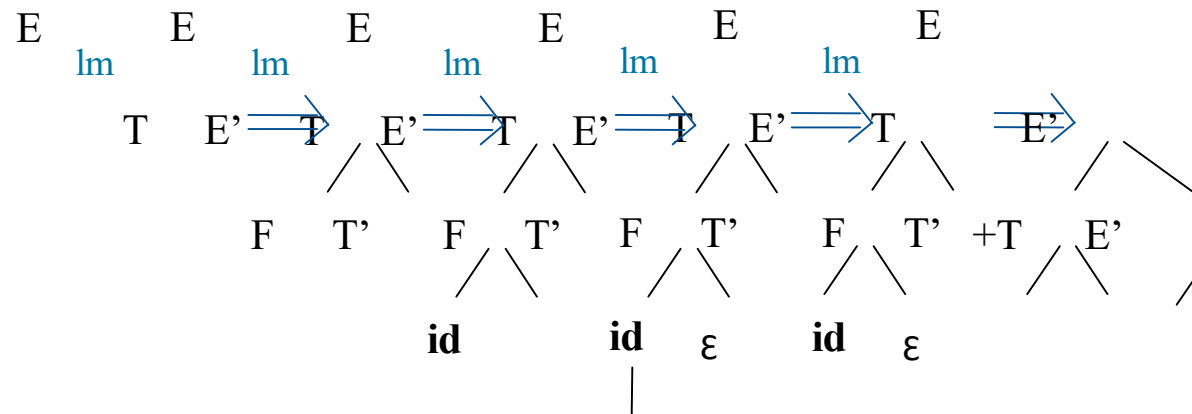
Example:

- $S \rightarrow IEtS \mid iEtSeS \mid a$
- $E \rightarrow b$

Top-down parser

- A Top-down parser tries to create a parse tree from the root towards the leaf canning input from
- Left to right
- It can be also viewed as finding a left most derivation for an input string
- Example: $\text{id}+\text{id}*\text{id}$

$E \rightarrow TE'$
 $E' \rightarrow +TE' | \epsilon$
 $T \rightarrow FT'$
 $T' \rightarrow *FT' | \epsilon$
 $F \rightarrow (E) | \text{id}$



Recursive descent parsing

- Consists of a set of procedures, one for each non terminal
- Execution begins with the procedure for start symbol
- A typical procedure for anon-terminal

```
voidA(){  
    chooseanA-production,A->X1X2..Xk for  
    (i=1 to k) {  
        if(Xiis anonterminal  
            callprocedureXi();  
        else if (Xi equals the current input symbol a)  
            advancetheinputtothenextsymbol;  
        else/*anerrorhasoccurred*/  
    }  
}
```

Recursive descent parsing(cont)

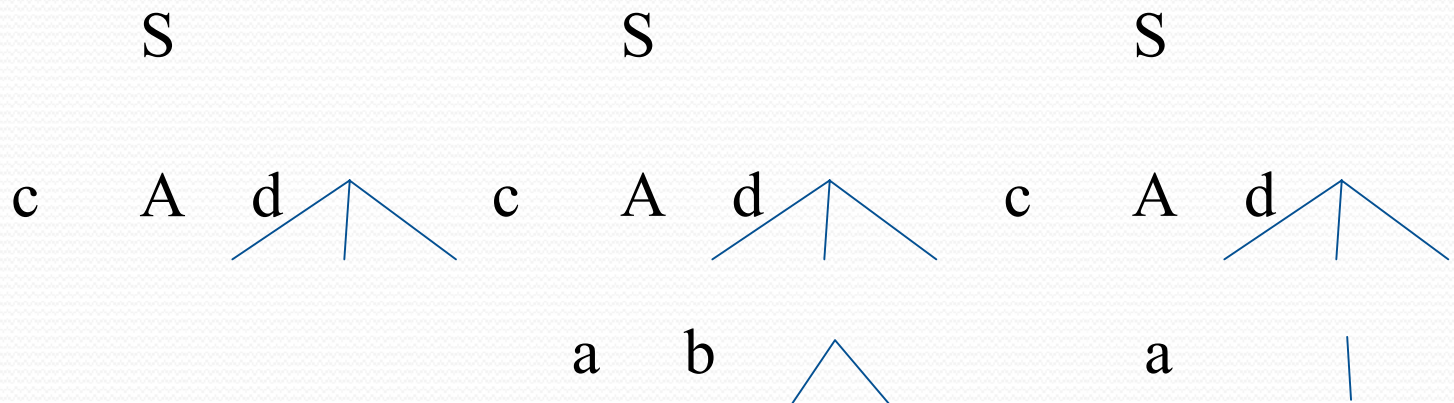
- General recursive descent may require back tracking
- The previous code needs to be modified to allow back tracking
- In general form it can't choose an A-production easily.
- So we need to try all alternatives
- If one failed the input pointer needs to be reset and another alternative should be tried
- Recursive descent parsers can't be used for left- recursive grammars

Example

$S \rightarrow cAd$

$A \rightarrow ab|a$

Input: cad



First and Follow

- First() is set of terminals that begins strings derived from
- If $\alpha \xRightarrow{*} \epsilon$ then ϵ is also in First(ϵ)
- In predictive parsing when we have $A \rightarrow \alpha \mid \beta$, if First(α) and First(β) are disjoint sets then we can select appropriate A-production by looking at the next input
- Follow(A), for any non terminal A, is set of terminals that can appear immediately after A in some sentential form
 - If we have $S \xRightarrow{*} \alpha A \beta$ for some α and β then α is in Follow(A)
- If A can be the right most symbol in some sentential form, then \$ is in Follow(A)

Computing First

- To compute $\text{First}(X)$ for all grammar symbols X , apply following* rules until no more terminals or ϵ can be added to any First set:
 1. If X is a terminal then $\text{First}(X) = \{X\}$.
 2. If X is a non terminal and $X \rightarrow Y_1 Y_2 \dots Y_k$ is a production for some $k \geq 1$, then place a in $\text{First}(X)$ if for some i a is in $\text{First}(Y_i)$ and ϵ is in all of $\text{First}(Y_1), \dots, \text{First}(Y_{i-1})$ that is $Y_1 \dots Y_{i-1}^* \Rightarrow \epsilon$. if ϵ is in $\text{First}(Y_j)$ for $j=1, \dots, k$ then add ϵ to $\text{First}(X)$.
 3. If $X \rightarrow \epsilon$ is a production then add ϵ to $\text{First}(X)$
- Example!

Computing follow

- To compute $\text{First}(A)$ for all non terminals A , apply following rules until nothing can be added to any follow set:
 1. Place $\$$ in $\text{Follow}(S)$ where S is the start symbol
 2. If there is a production $A \rightarrow \alpha B \beta$ then everything in $\text{First}(\beta)$ except ϵ is in $\text{Follow}(B)$.
 3. If there is a production $A \rightarrow B$ or a production $A \rightarrow \alpha B \beta$ where $\text{First}(\beta)$ contains ϵ , then everything In $\text{Follow}(A)$ is in $\text{Follow}(B)$
- Example!

LL(1) Grammars

- Predictive parsers are those recursive descent parsers needing no backtracking
- Grammars for which we can create predictive parsers are called LL(1)
 - The first L means scanning input from left to right
 - The second L means leftmost derivation
 - And i stands for using one input symbol for lookahead
- A grammar G is LL(1) if and only if whenever $A \rightarrow \alpha \mid \beta$ are two distinct productions of G , the following conditions hold:
 - For no terminal a do α and β both derive strings beginning with a
 - At most one of α or β can derive empty string
 - If $\alpha \Rightarrow^* \epsilon$ then β does not derive any string beginning with a terminal in $\text{Follow}(A)$.

Construction of predictive parsing table

- For each production $A \rightarrow \alpha$ in grammar do the following:
 1. For each terminal a in $\text{First}(\alpha)$ add $A \rightarrow$ in $M[A, a]$
 2. If ϵ is in $\text{First}(\alpha)$, then for each terminal b in $\text{Follow}(A)$ add $A \rightarrow \epsilon$ to $M[A, b]$. If ϵ is in $\text{First}(\alpha)$ and $\$$ is in $\text{Follow}(A)$, add $A \rightarrow \epsilon$ to $M[A, \$]$ as well
- If after performing the above, there is no production in $M[A, a]$ then set $M[A, a]$ to error

Example

	First	Follow
$E \rightarrow TE'$	{(,id}	{+,*,),}\$}
$E' \rightarrow +TE' \epsilon$	{(,id}	{+,),}\$}
$T \rightarrow FT'$	{(,id}	{),}\$}
$T' \rightarrow *FT' \epsilon$	{+, ϵ }	{),}\$}
$F \rightarrow (E) id$	{*, ϵ }	{+,),}\$}

Non-terminal		InputSymbol					
		id	+	*	(	)	\$
E	$E \rightarrow TE'$				$E \rightarrow TE'$		
E'			$E' \rightarrow +TE'$			$E' \rightarrow \epsilon$	$E' \rightarrow \epsilon$
T	$T \rightarrow FT'$				$T \rightarrow FT'$		
T'			$T' \rightarrow \epsilon$	$T' \rightarrow *FT'$		$T' \rightarrow \epsilon$	$T' \rightarrow \epsilon$
F	$F \rightarrow id$						

Bottom-up Parsing

- Construct a parse tree for an input string beginning at the leaves (the bottom) and working toward the root (the top)
- Example: $\text{id} * \text{id}$

$E \rightarrow E + T \mid T$

$T \rightarrow T * F \mid F$

$F \rightarrow (E) \mid \text{id}$

$\text{id} * \text{id}$

$F * \text{id}$

$T * \text{id}$

$T * F$

F

E

id

F

F

id

$T * F$

F

id

id

$T * F$

F

id

$F \text{ id id}$

id

Shift-reduce parser

- The general idea is to shift some symbol of input to the stack until a reduction can be applied
- At each reduction step, a specific substring matching the body of a production is replaced by the nonterminal at the head of the production
- The key decisions during bottom-up parsing are about when to reduce and about what production to apply
- A reduction is a reverse of a step in a derivation
- The goal of a bottom-up parser is to construct a derivation in reverse:
 - $E \Rightarrow T \Rightarrow T * F \Rightarrow T * id \Rightarrow F * id \Rightarrow id * id$

Shiftrreduceparsing

- A stack is used to hold grammar symbols
- Handle always appear on top of the stack
- Initial configuration:

Stack	Input
\$	w\$

- Acceptance configuration

Stack	Input
\$S	\$

Reduce/reduceconflict

stmt->id(parameter_list)

stmt -> expr:=expr

parameter_list->parameter_list, parameter

parameter_list->parameter

parameter->id

expr->id(expr_list)

expr->id

expr_list->expr_list,expr

expr_list->expr

Stack

... id(id

Input

,id) ...\$

LR Parsing

- The most prevalent type of bottom-up parsers
- LR(k), mostly interested on parsers with $k \leq 1$
- Why LR parsers?
 - Table driven
 - Can be constructed to recognize all programming language constructs
 - Most general non-backtracking shift-reduce parsing method
 - Can detect a syntactic error as soon as it is possible to do so
 - Class of grammars for which we can construct LR parsers are superset of those which we can construct LL parsers

States of an LR parser

- States represent set of items
- An LR(0) item of G is a production of G with the dot at some position of the body:
 - For $A \rightarrow XYZ$ we have following items
 - $A \rightarrow .XYZ$
 - $A \rightarrow X.YZ$
 - $A \rightarrow XY.Z$
 - $A \rightarrow XYZ.$
 - In a state having $A \rightarrow .XYZ$ we hope to see a string derivable from XYZ next on the input.
 - What about $A \rightarrow X.YZ$?

Constructing canonical LR(0) item sets

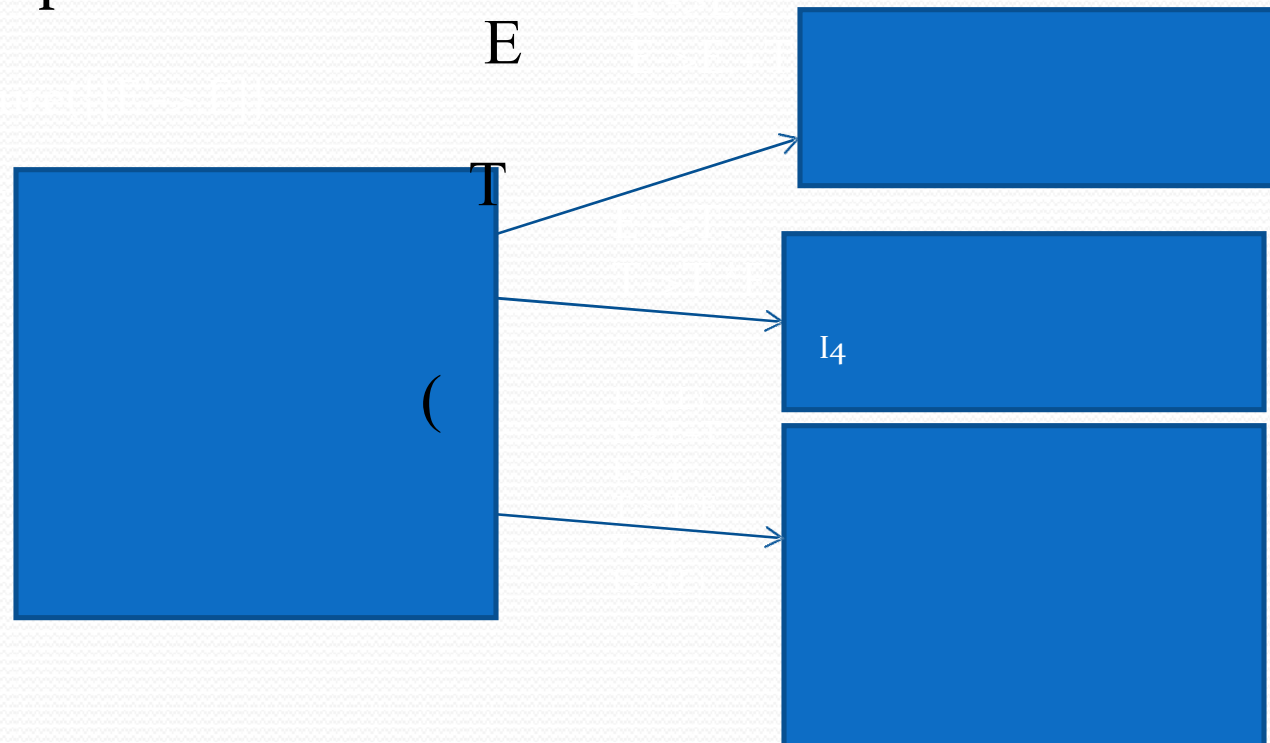
- Augmented grammar:
 - G with addition of a production: $S' \rightarrow S$
- Closure of item sets:
 - If I is a set of items, $\text{closure}(I)$ is a set of items constructed from I by the following rules:
 - Add every item in I to $\text{closure}(I)$
 - If $A \rightarrow \alpha.B\beta$ is in $\text{closure}(I)$ and $B \rightarrow \gamma$ is a production then add the item $B \rightarrow \gamma.$ to $\text{closure}(I)$.
- Example:

$$\begin{aligned} E' &\rightarrow E \\ E &\rightarrow E+T \mid T \\ T &\rightarrow T * F \mid F \\ F &\rightarrow (E) \mid \mathbf{id} \end{aligned}$$



Constructing canonical LR(0) item sets (cont.)

- $\text{Goto}(I, X)$ where I is an item set and X is a grammar symbol is closure of set of all items $[A \rightarrow \alpha X \beta]$ where $[A \rightarrow \alpha.X \beta]$ is in I
- Example

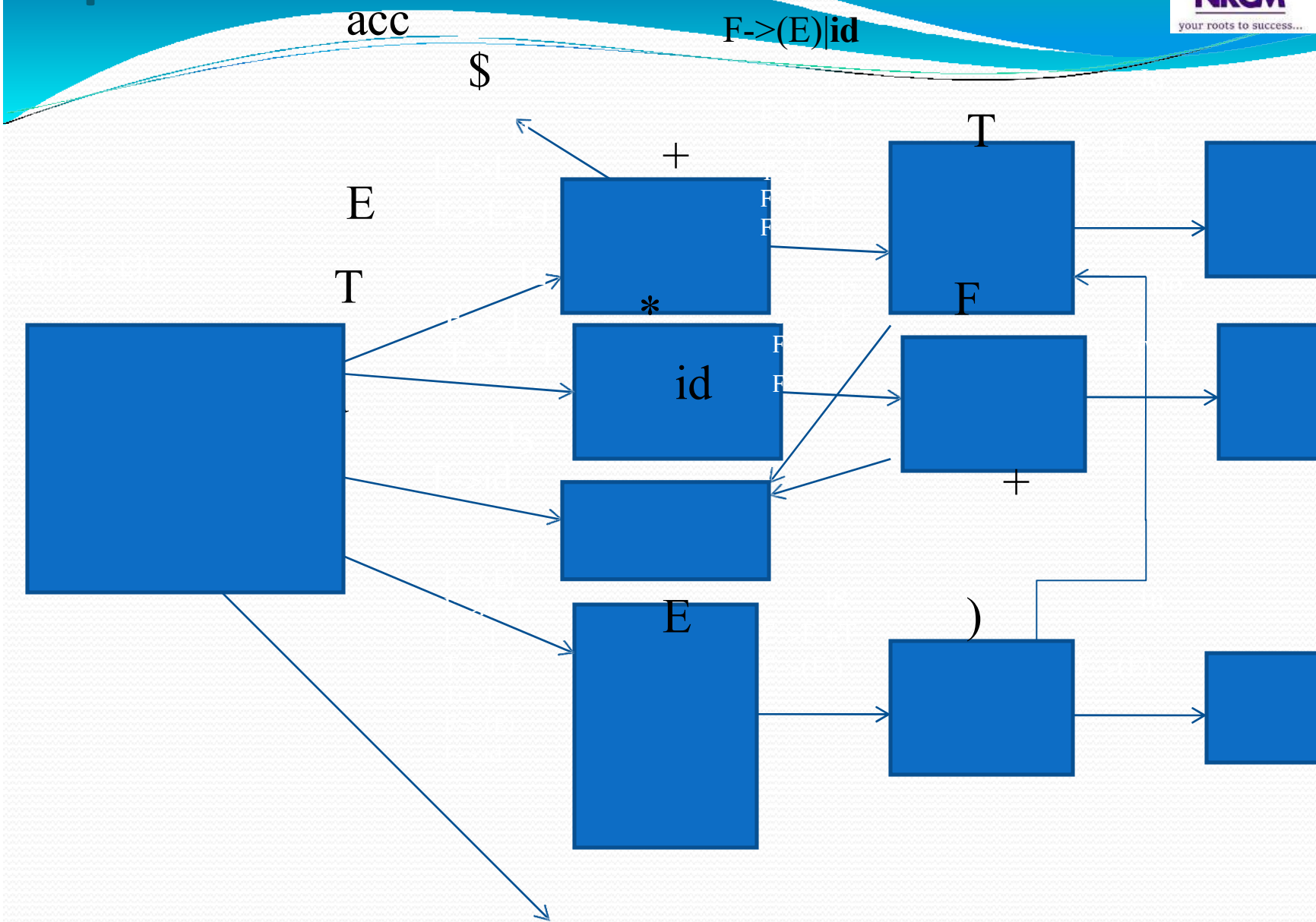


Canonical LR(0) items

```
Void items(G') {  
    C = CLOSURE({[S' -> .S]});  
    repeat  
        for (each set of items I in C)  
            for (each grammar symbol X)  
                if (GOTO(I, X) is not empty and not in C) add  
                    GOTO(I, X) to C;  
    until no new set of items are added to C on a round;  
}
```

Example

$E' \rightarrow E$
 $E \rightarrow E + T \mid T$
 $T \rightarrow T * F \mid F$
 $F \rightarrow (E) \mid id$

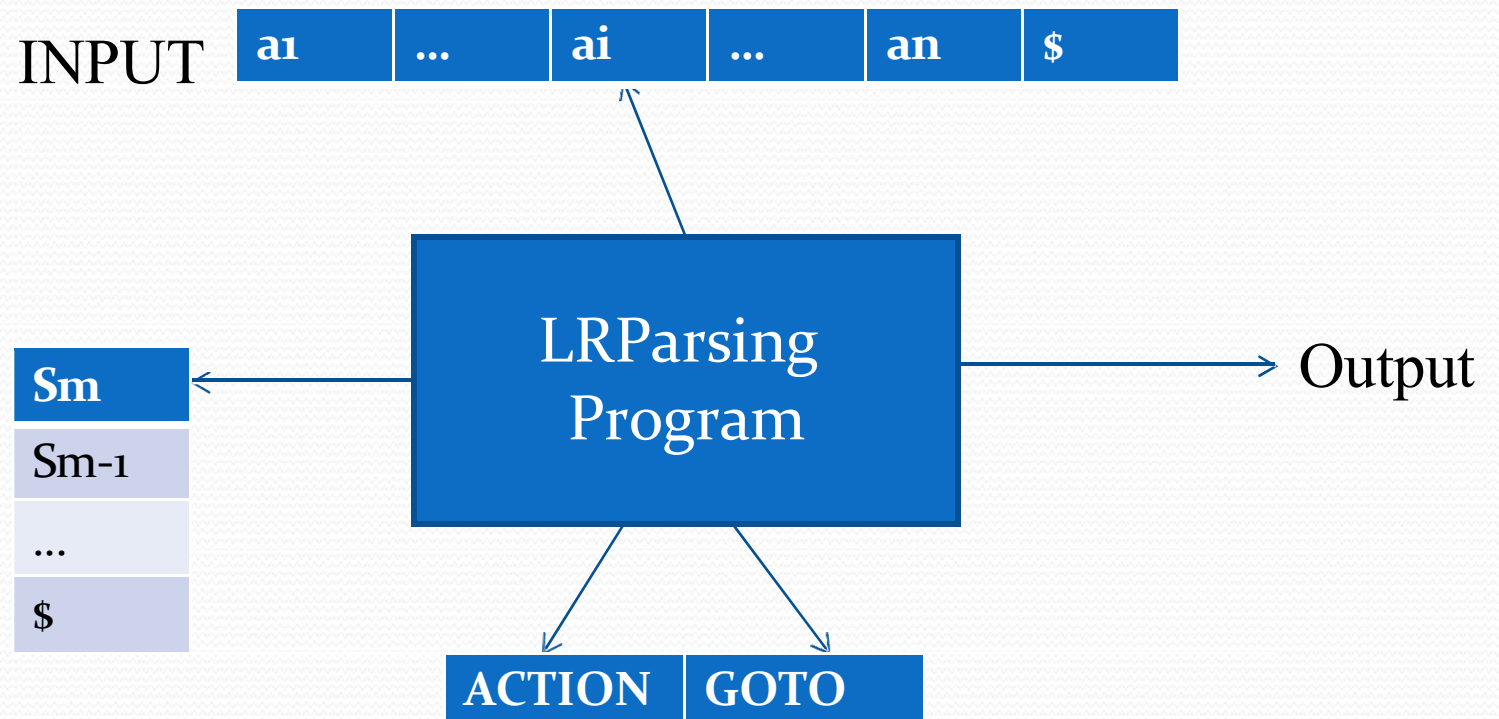


Use of LR(0) automaton

● Example: $\text{id} * \text{id}$

Line	Stack	Symbols	Input	Action
(1)	0	\$	id*id\$	Shift to 5
(2)	05	\$id	*id\$	Reduce by $F \rightarrow id$
(3)	03	\$F	*id\$	Reduce by $T \rightarrow F$
(4)	02	\$T	*id\$	Shift to 7
(5)	027	\$T*	id\$	Shift to 5
(6)	0275	\$T*id	\$	Reduce by $F \rightarrow id$
(7)	02710	\$T*F	\$	Reduce by $T \rightarrow T*F$
(8)	02	\$T	\$	Reduce by $E \rightarrow T$
(9)	01	\$E	\$	accept

LR-Parsingmodel



LR parsing algorithm

```
let a be the first symbol of w$;  
while(1){ /* repeat forever */  
    let s be the state on top of the stack; if  
    (ACTION[s,a] = shift t) {  
        push t onto the stack;  
        let a be the next input symbol;  
    } elseif (ACTION[s,a] = reduce A →  $\beta$ ) {  
        pop  $|\beta|$  symbols of the stack;  
        let state t now be on top of the stack;  
        push GOTO[t,A] onto the stack;  
        output the production A →  $\beta$ ;  
    } elseif (ACTION[s,a] = accept) break; /* parsing is done */ else  
        call error-recovery routine;  
}
```


Example

STATE	ACTION						GOTO		
	id	+	*	(	)	\$	E	T	F
0	S5			S4			1	2	3
1		S6				Acc			
2		R2	S7		R2	R2			
3		R4	R7		R4	R4			
4	S5			S4			8	2	3
5		R6	R6		R6	R6			
6	S5			S4				9	3
7	S5			S4					10
8		S6			S11				
9		R1	S7		R1	R1			
10		R3	R3		R3	R3			
11		R5	R5		R5	R5			

(0) $E' \rightarrow E$

(1) $E \rightarrow E+T$

(2) $E \rightarrow T$

(3) $T \rightarrow T * F$

(4) $T \rightarrow F$

(5) $F \rightarrow (E)$

(6) $F \rightarrow id$

$id * id + id?$

Constructing SLR parsing table

● Method

- Construct $C = \{I_0, I_1, \dots, I_n\}$, the collection of LR(0) items for G'
- State i is constructed from state l_i :
 - If $[A \rightarrow \alpha.a\beta]$ is in l_i and $\text{Goto}(l_i, a) = l_j$, then set $\text{ACTION}[i, a]$ to “shift j ”
 - If $[A \rightarrow \alpha.]$ is in l_i , then set $\text{ACTION}[i, a]$ to “reduce $A \rightarrow \alpha$ ” for all a in $\text{follow}(A)$
 - If $[S' \rightarrow .S]$ is in l_i , then set $\text{ACTION}[i, \$]$ to “Accept”
- If any conflicts appear then we say that the grammar is not SLR(1).
- If $\text{GOTO}(l_i, A) = l_j$ then $\text{GOTO}[i, A] = j$
- All entries not defined by above rules are made “error”
- The initial state of the parser is the one constructed from the set of items containing $[S' \rightarrow .S]$

Example grammar which is not SLR(1)

$S \rightarrow L=R \mid R$

$L \rightarrow *R \mid id$

$R \rightarrow L$

I0
 $S' \rightarrow .S$
 $S \rightarrow .L=R$
 $S \rightarrow .R$
 $L \rightarrow .*R \mid$
 $L \rightarrow .id$
 $R \rightarrow .L$

I1
 $S' \rightarrow S.$

I3
 $S \rightarrow R.$

I5
 $L \rightarrow id.$

I7
 $L \rightarrow *R.$

I2
 $S \rightarrow L.=R$
 $R \rightarrow L.$

I4
 $L \rightarrow *.R$
 $R \rightarrow .L$
 $L \rightarrow .*R$
 $L \rightarrow .id$

I6
 $S \rightarrow L=.R$
 $R \rightarrow .L$
 $L \rightarrow .*R$
 $L \rightarrow .id$

I8
 $R \rightarrow L.$

I9
 $S \rightarrow L=R.$

Action

=

2

Shift6

ReduceR->

Computer Engineering

More powerful LR parsers

- Canonical-LR or just LR method
 - Use lookahead symbols for items: LR(1) items
 - Results in a large collection of items
- LALR: lookaheads are introduced in LR(0) items

Canonical LR(1) items

- In LR(1) item set each item is in the form: $[A \rightarrow \alpha.\beta, a]$
- An LR(1) item $[A \rightarrow \alpha.\beta, a]$ is valid for a viable prefix y if there is a derivation $S \xRightarrow{*} \delta A w \xRightarrow{rm} \delta \alpha \beta w$, where

- $\Gamma = \delta \alpha$

- Either a is the first symbol of w , or $w = \epsilon$ and a is $\$$

- Example:

- $S \rightarrow BB$

- $B \rightarrow aB | b$

$$S \xRightarrow{*} aaBab \xRightarrow{rm} aaaBab$$

Item $[B \rightarrow a.B, a]$ is valid for $\gamma = aaa$ and $w = ab$

Constructing LR(1) set of item

```
SetOfItemsClosure(I){  
    repeat  
        for(each item  $[A \rightarrow \alpha.B\beta, a]$  in I)  
            for(each production  $B \rightarrow \gamma$  in  $G'$ )  
                for(each terminal  $b$  in  $\text{First}(\beta a)$ )  
                    add  $[B \rightarrow \gamma.b]$  to set I;  
  
    until no more items are added to I;  
    return I;  
}
```

```
SetOfItemsGoto(I, X){  
    initialize J to be the empty set;  
    for(each item  $[A \rightarrow \alpha.X\beta, a]$  in I)  
        add item  $[A \rightarrow \alpha.X.\beta, a]$  to set J;  
    return closure(J);  
}
```

```
void items( $G'$ ){  
    initialize C to Closure( $\{[S' \rightarrow .S, \$]\}$ );  
    repeat  
        for(each set of items  $I$  in C)  
            for(each grammar symbol  $X$ )  
                if ( $\text{Goto}(I, X)$  is not empty and not in C)  
                    add  $\text{Goto}(I, X)$  to C;  
  
    until no new set of items are added to C;  
}
```


Example

$S' \rightarrow S$

$S \rightarrow CC$

$C \rightarrow cC$

$C \rightarrow d$

Canonical LR(1) parsing table

● Method

- Construct $C = \{I_0, I_1, \dots, I_n\}$, the collection of LR(1) items for G'
- State i is constructed from state I_i :
 - If $[A \rightarrow \alpha.a\beta, b]$ is in I_i and $\text{Goto}(I_i, a) = I_j$, then set $\text{ACTION}[i, a]$ to "shift j "
 - If $[A \rightarrow \alpha., a]$ is in I_i , then set $\text{ACTION}[i, a]$ to "reduce $A \rightarrow \alpha$ "
 - If $[S' \rightarrow .S, \$]$ is in I_i , then set $\text{ACTION}[i, \$]$ to "Accept"
- If any conflicts appear then we say that the grammar is not LR(1).
- If $\text{GOTO}(I_i, A) = I_j$ then $\text{GOTO}[i, A] = j$
- All entries not defined by above rules are made "error"
- The initial state of the parser is the one constructed from the set of items containing $[S' \rightarrow .S, \$]$

Example

$S' \rightarrow S$

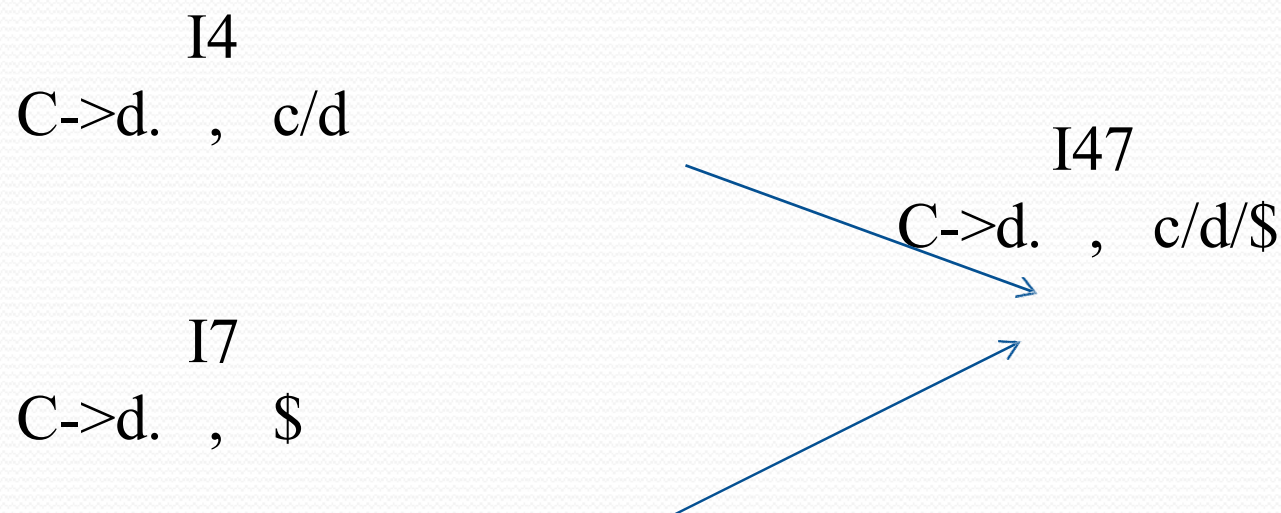
$S \rightarrow CC$

$C \rightarrow cC$

$C \rightarrow d$

LALR Parsing Table

- For the previous example we had:



- State merge can't produce Shift-Reduce conflicts. Why?
- But it may produce reduce-reduce conflict

Using ambiguous grammars

$E \rightarrow E + E$

$E \rightarrow E * E$

$E \rightarrow (E)$

$E \rightarrow id$

I0: $E' \rightarrow .E$

$E \rightarrow .E + E$

$E \rightarrow .E * E$

$E \rightarrow .(E)$

$E \rightarrow .id$

I1: $E' \rightarrow E.$

$E \rightarrow E. + E$

$E \rightarrow E. * E$

I4: $E \rightarrow E + .E$

$E \rightarrow E. + E$

$E \rightarrow E. * E$

$E \rightarrow .(E)$

$E \rightarrow .id$

I2: $E \rightarrow (.E)$

$E \rightarrow .E + E$

$E \rightarrow .E * E$

$E \rightarrow .(E)$

$E \rightarrow .id$

I5: $E \rightarrow E * .E$

$E \rightarrow (.E)$

$E \rightarrow .E + E$

$E \rightarrow .E * E$

$E \rightarrow .(E)$

$E \rightarrow .id$

STATE	ACTION						
	id	+	*	(	)	\$	E
0	S ₃			S ₂			1
1		S ₄	S ₅			Acc	
2	S ₃		S ₂				6
3		R ₄	R ₄		R ₄	R ₄	
4	S ₃			S ₂			7
5	S ₃			S ₂			8
6		S ₄	S ₅				
7		R ₁	S ₅		R ₁	R ₁	
8		R ₂	R ₂		R ₂	R ₂	
9		R ₃	R ₃		R ₃	R ₃	

I6: $E \rightarrow (E.)$

$E \rightarrow E. + E$

$E \rightarrow E. * E$

I8: $E \rightarrow E * E.$

$E \rightarrow E. + E$
 $E \rightarrow E. * E$

I7: $E \rightarrow E + E.$

$E \rightarrow E. + E$

$E \rightarrow E. * E$

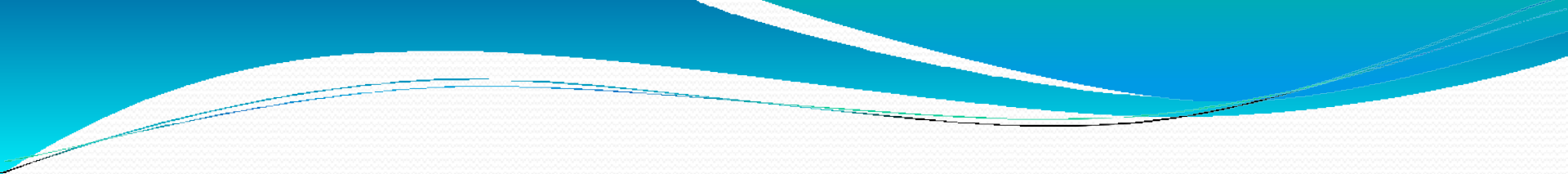
I9: $E \rightarrow (E).$

Computer Science and
Engineering

UNIT-IV



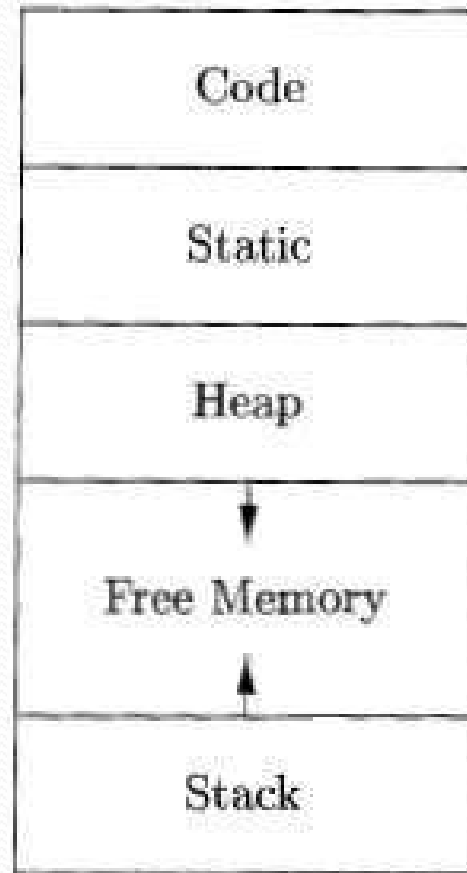
Run-Time Environments



Outline

- Compiler must do the storage allocation and provide access to variables and data
- Memory management
 - Stack allocation
 - Heap management
 - Garbage collection

Storage Organization



Static vs. Dynamic Allocation

- Static: Compile time, Dynamic: Runtime allocation
- Many compilers use some combination of following
 - Stack storage: for local variables, parameters and so on
 - Heap storage: Data that may outlive the call to the procedure that created it
- Stack allocation is a valid allocation for procedures since procedure calls are nested

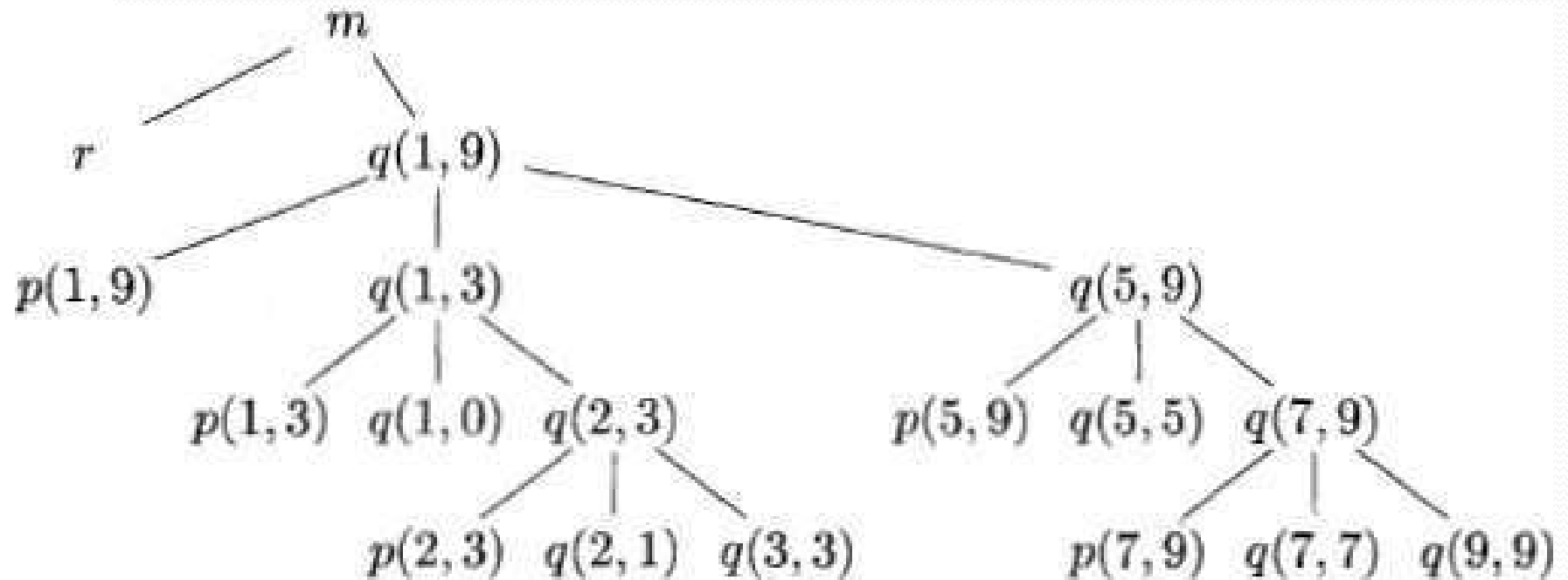
Sketch of a quicksort program

```
int a[11];
void readArray() { /* Reads 9 integers into a[1], ..., a[9]. */
    int i;
    ...
}
int partition(int m, int n) {
    /* Picks a separator value  $v$ , and partitions  $a[m..n]$  so that
        $a[m..p-1]$  are less than  $v$ ,  $a[p] = v$ , and  $a[p+1..n]$  are
       equal to or greater than  $v$ . Returns  $p$ . */
    ...
}
void quicksort(int m, int n) {
    int i;
    if (n > m) {
        i = partition(m, n);
        quicksort(m, i-1);
        quicksort(i+1, n);
    }
}
main() {
    readArray();
    a[0] = -9999;
    a[10] = 9999;
    quicksort(1,9);
}
```

Activation for Quicksort

```
enter main()
  enter readArray()
  leave readArray()
  enter quicksort(1,9)
    enter partition(1,9)
    leave partition(1,9)
    enter quicksort(1,3)
    ...
    leave quicksort(1,3)
    enter quicksort(5,9)
    ...
    leave quicksort(5,9)
  leave quicksort(1,9)
leave main()
```

Activation tree representing calls during execution of quicksort



Activation records

- Procedure calls and returns are usually managed by a run-time stack called the control stack.
- Each live activation has an activation record (sometimes called a frame)
- The root of activation tree is at the bottom of the stack
- The current execution path specifies the content of the stack with the last activation has record in the top of the stack.

AGeneralActivationRecord

Actual parameters
Returned values
Control link
Access link
Saved machine status
Local data
Temporaries

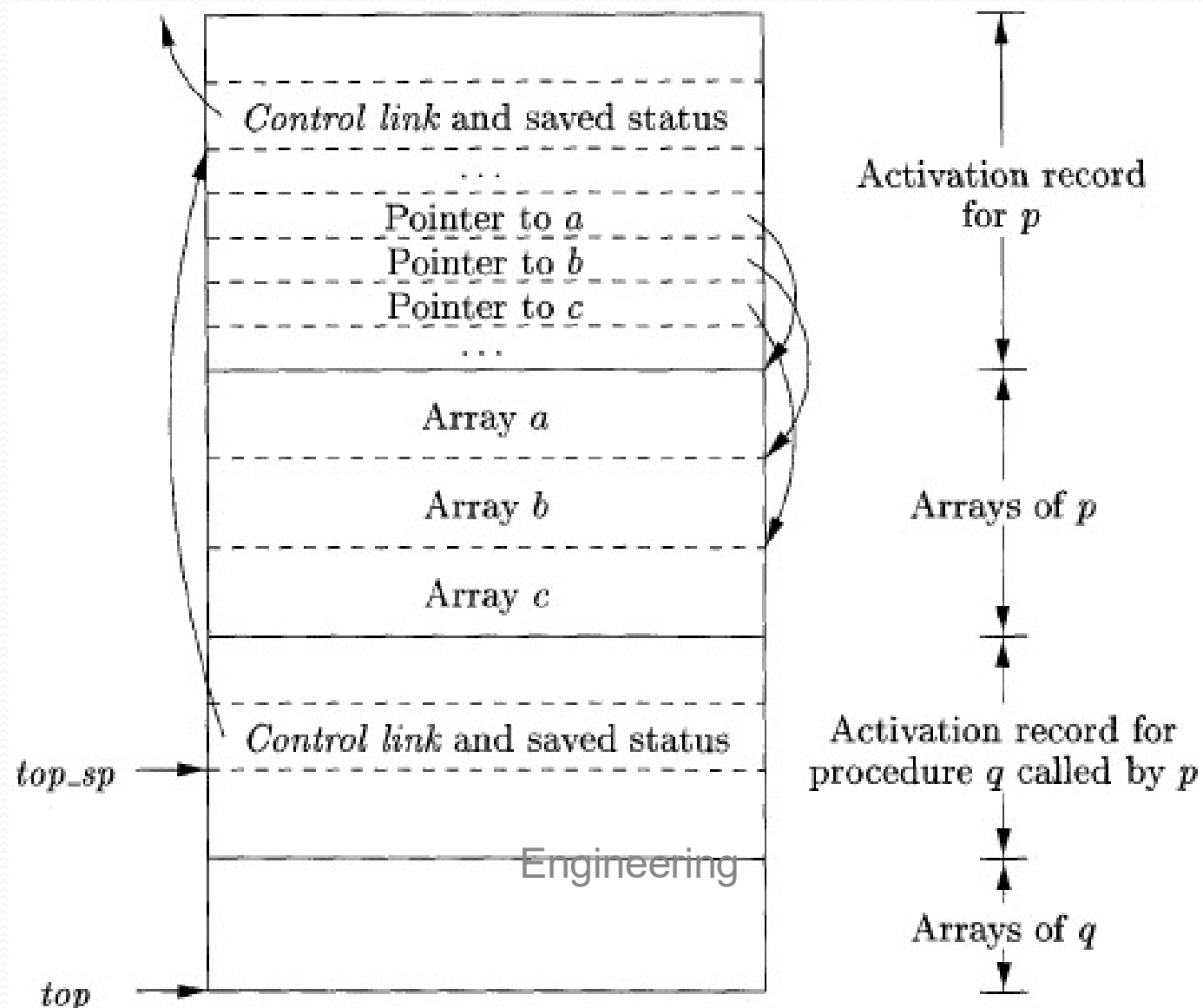
ActivationRecord

- Temporary values
- Local data
- A saved machine status
- An “access link”
- A control link
- Space for the return value of the called function
- The actual parameters used by the calling procedure

Access to dynamically allocated arrays



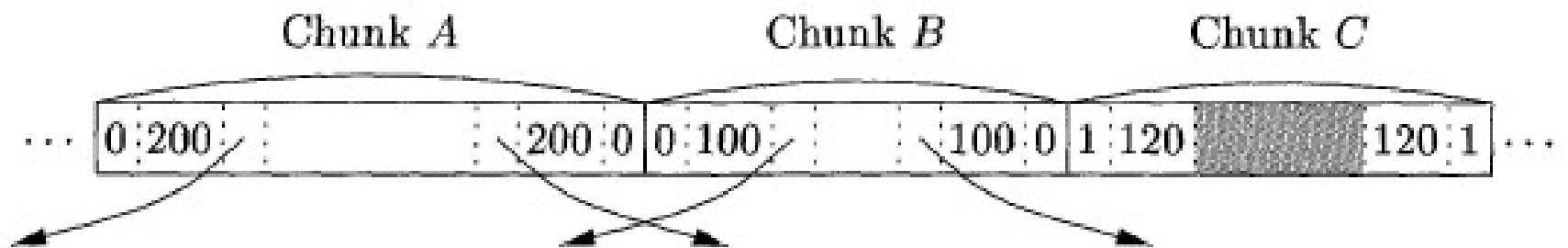
your roots to success...



MemoryManager

- Two basic functions:
 - Allocation
 - Deallocation
- Properties of memory managers:
 - Space efficiency
 - Program efficiency
 - Low overhead

Part of a Heap



Introduction

- The final phase of a compiler is code generator
- It receives an intermediate representation (IR) with supplementary information in symbol table
- Produces a semantically equivalent target program
- Code generator main tasks:
 - Instruction selection
 - Register allocation and assignment
 - Instruction ordering



Issues in the Design of Code Generator

- The most important criterion is that it produces correct code
- Input to the code generator
 - IR + Symbol table
 - We assume frontend produces low-level IR, i.e. values of names in it can be directly manipulated by the machine instructions.
 - Syntactic and semantic errors have been already detected
- The target program
 - Common target architectures are: RISC, CISC and Stack based machines
 - In this chapter we use a very simple RISC-like computer with addition of some CISC-like addressing modes

complexity of mapping

- the level of the IR
- the nature of the instruction-set architecture
- the desired quality of the generated code.

$x = y + z$

```
LD    R0,y
ADD   R0,R0,z
ST    x,R0
```

$a = b + c$

$d = a + e$

```
LD    R0,b
ADD   R0,R0,c
ST    a, R0
LD    R0,a
ADD   R0,R0,e
      d, R0
```

Register allocation

- Two sub problems
 - Register allocation: selecting the setoff variables that will reside in register sat each point in the program
 - Resister assignment: selecting specific register that a variable reside in
- Complications imposed by the hardware architecture
 - Example: register pairs for multiplication and division

t=a+b
t=t*c
T=t/d

L	R1, a
A	R1,b
M	R0, c
D	R0,d
ST	R1, t

t=a+b
t=t+c
T=t/d

L	R0,a
A	R0,b
M	R0,c
SRDA	R0,32
ST	R1,t

A simple target machine model

- Load operations: $LD_{r,x}$ and $LD_{r1,r2}$
- Store operations: $ST_{x,r}$
- Computation operations: $OP_{dst,src1,src2}$
- Unconditional jumps: BRL
- Conditional jumps: $Bcond_{r,L}$ like $BLTZ_{r,L}$

Addressing Modes

- Variable name: x
- Indexed address: $a(r)$ like $LDR\ R_1, a(R_2)$ means
 $R_1 = \text{contents}(a + \text{contents}(R_2))$
- Integer indexed by a register: like $LDR\ R_1, 100(R_2)$
- Indirect addressing mode: $*r$ and $*100(r)$
- Immediate constant addressing mode: like $LDR\ R_1, \#100$



```
//R1=i
```

```
//R1=R1*8
```

```
//R2=contents(a+contents(R1))
```

//b=R2



```
//R1=c
```

//R2=j

```
//R2=R2*8
```

```
//contents(a+contents(R2))=R1
```

UNIT-V

Machine-Independent Optimization

Machine independent optimization attempts to improve the intermediate code to get a better target code. The part of the code which is transformed here does not involve any absolute memory location or any CPU registers.

Code Optimization can perform in the following different ways:

(1) Compile Time Evaluation:

(a) $z = 5 * (45.0 / 5.0) * r$

Perform $5 * (45.0 / 5.0) * r$ at compile time.

(b) $x = 5.7$

$y = x / 3.6$

Evaluate $x / 3.6$ as $5.7 / 3.6$ at compile time.

(2) Variable Propagation:

Before Optimization the code is:

```
c = a * b
```

```
x=a
```

```
till
```

```
d=x*b+4
```

After Optimization the code is:

```
c = a * b
```

```
x=a
```

```
till
```

```
    d=a*b+4
```

(3) Dead code elimination:

Before elimination the code is:

```
c = a * b
```

```
x=b
```

```
till
```

```
d=a*b+4
```

After elimination the code is:

```
c = a * b
```

```
till
```

```
d=a*b+4
```

(4) Code Motion:

It reduces the evaluation frequency of expression.

It brings loop in variant statements out of the loop. do

```
{  
    Item =10;  
    value value=value+ item;  
} while (value<100);
```


//This code can be further optimized as

```
item = 10;  
do  
{  
    Value value=value+ item;  
} while (value<100);
```

(5) Induction Variable and Strength Reduction: Strength reduction is used to replace the high strength operator by the low strength.

An induction variable is used in loop for the following kind of assignment like $i = i + \text{constant}$.

Before reduction the code is:

After Reduction the code is:

```
i = 1
```

```
t=4
```

```
{
```

```
    while(t<40)
```

```
        y = t;
```

```
        t = t + 4;
```

```
}
```

Data Flow Analysis



To efficiently optimize the code compiler collects all the information about the program and distribute this information to each block of the flow graph. This process is known as data-flow graph analysis.

Certain optimization can only be achieved by examining the entire program. It can't be achieved by examining just a portion of the program.

consider the following code:

```
x = a + b;  
x = 6 * 3
```

Some optimization needs more global information. For example, consider the following code:

```
a=1;  
b=2;  
c=3;  
if (...)x=a+ 5;  
else x=b+4; c  
= x + 1;
```