

Lecture 4: Syntactic Analysis

COMP 524 Programming Language Concepts

Stephen Olivier

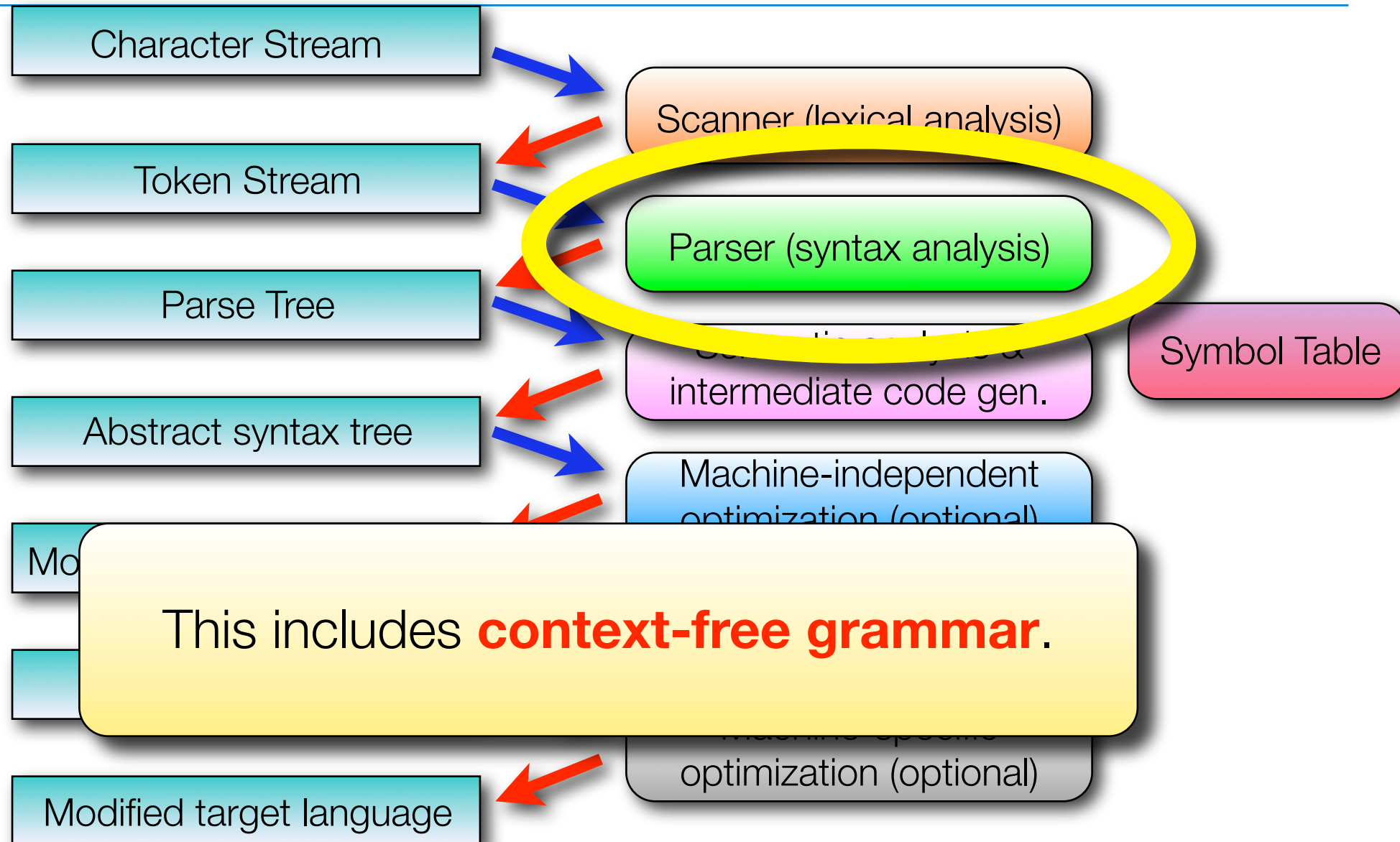
January 27, 2009

Based on notes by N. Fisher, F. Hernandez-Campos, and D. Stotts

The University of North Carolina at Chapel Hill



Goal of Lecture



Parsing

- The main task of parsing is to **identify the syntax**.



Review: Regular Expression Rules

- A RE consist of:
 - A character (e.g., “0”, “1”, ...)
 - The empty string (i.e., “ ϵ ”)
 - Two REs next to each other (e.g., “*non_negative_digit digit*”) to denote concatenation.
 - Two REs separated by “|” next to each other (e.g., “*non_negative_digit | digit*”) to denote one RE or the other.
 - An RE followed by “*” (called the **Kleene star**) to denote zero or more iterations of the RE.
 - Parentheses (in order to avoid ambiguity).



Review: Regular Expression Rules

- A RE consist of:
 - A character (e.g., “0”, “1”, ...)
 - The
 - Two
 - Two REs separated by “|” next to each other (e.g., “*non_negative_digit* | *digit*”) to denote one RE or the other.
 - An RE followed by “*” (called the **Kleene star**) to denote zero or more iterations of the RE.
 - Parentheses (in order to avoid ambiguity).

A RE is **NEVER** defined in terms of itself!
Thus, REs cannot **define recursive statements.**



Review: Regular Expression Rules

- A RE consist of:
 - A character (e.g., “0”, “1”, ...)
 - The
 - Two
 - Two REs separated by “|” next to each other (e.g., “*non_negative_digit* | *digit*”) to
 - An RE followed by “*” (called the **Kleene star**) to denote zero or more iterations of the RE.
 - Parentheses (in order to avoid ambiguity).

For example, REs cannot define arithmetic expressions with parentheses.



Context-Free Grammars

- **Context-Free Grammars (CFGs)** are similar to REs except that they can handle **recursion**.

Arithmetic expression with parentheses

$$expr \rightarrow id \mid number \mid - expr \mid (expr) \mid expr \ op \ expr$$
$$op \rightarrow + \mid - \mid * \mid /$$


Context-Free Grammars

- **Context-**

Each rule is called a **production**.

except that they can handle **recursion**.

Arithmetic expression with parentheses

$expr \rightarrow id \mid number \mid - expr \mid (expr) \mid expr op expr$

$op \rightarrow + \mid - \mid * \mid /$



Context

One of the nonterminals, usually the first one, is called the **start symbol**, and it defines the construct defined by the grammar.

- **Context** except that they can handle **recursion**.

Arithmetic expression with parentheses

$expr \rightarrow id \mid number \mid - expr \mid (expr) \mid expr \ op \ expr$

$op \rightarrow + \mid - \mid * \mid /$



Non-terminals are symbols that are defined by the CFG and can appear on both the **left** and **right** side of “ $\rightarrow$ ”.

- **Context-Free Grammars (CFGs)** are similar to REs except that they can handle **recursion**.

Arithmetic expression with parentheses

$expr \rightarrow id \mid number \mid - expr \mid (expr) \mid expr op expr$

$op \rightarrow + \mid - \mid * \mid /$



Co In the book, non-terminals are written in *italics*,
but in other literature they are written in
<brackets>.

- **Context-Free Grammars (CFGs)** are similar to REs except that they can handle **recursion**.

Arithmetic expression with parentheses

expr → id | number | - *expr* | (*expr*) | *expr* op *expr*

op → + | - | * | /



Terminals are the strings that define the grammar and can only appear on the right side of “ $\rightarrow$ ”.

- **Context-Free Grammars (CFGs)** are similar to REs except that they can handle **recursion**.

Arithmetic expression with parentheses

$expr \rightarrow id \mid number \mid - expr \mid (expr) \mid expr \ op \ expr$

$op \rightarrow + \mid - \mid * \mid /$



Co

In the book non-terminals are written in **typewriter** font, others write it in “normal” font.

- **Context-Free Grammars (CFGs)** are similar to REs except that they can handle **recursion**.

Arithmetic expression with parentheses

$expr \rightarrow id \mid number \mid - expr \mid (expr) \mid expr \ op \ expr$

$op \rightarrow + \mid - \mid * \mid /$



BNF

- Technically, the Kleene star (*) and parentheses are not allowed under the CFG rules, called **Backus-Naur Form (BNF)**.
- However, for convenience, we will use Extended BNF that includes the Kleene star, parentheses, and the **Kleene Plus** (+), which stands for “one or more iterations.”



EBNF Example.

- The Kleene star and parentheses can be written as follows

$id_list \rightarrow id (, id)^*$

$id_list \rightarrow id$

$id_list \rightarrow id_list, id$



Derivation

- A **derivation** is “a series of replacement operations that derive a string of terminals from the start symbol.”

slope * x + intercept

$\Rightarrow \text{expr op } \underline{\text{expr}} + \text{id}$

$\Rightarrow \text{expr } \underline{\text{op}} \text{ id} + \text{id}$

$\Rightarrow \underline{\text{expr}} * \text{id} + \text{id}$

$\Rightarrow \text{id} * \text{id} + \text{id}$

$\text{expr} \Rightarrow \text{expr op } \underline{\text{expr}}$

$\Rightarrow \text{expr } \underline{\text{op}}$ id

$\Rightarrow \underline{\text{expr}} + \text{id}$

(slope) * (x) + (intercept)

Derivation

- A **derivation** is “a series of replacement operations that derive a start symbol.”

$\Rightarrow$ denotes “derived from”

slope * x + intercept



$\Rightarrow \text{expr} \Rightarrow \text{expr op } \underline{\text{expr}}$

$\Rightarrow \text{expr op } \underline{\text{id}}$

$\Rightarrow \underline{\text{expr}} + \text{id}$

$\Rightarrow \text{expr op } \underline{\text{expr}} + \text{id}$

$\Rightarrow \text{expr op } \underline{\text{id}} + \text{id}$

$\Rightarrow \underline{\text{expr}} * \text{id} + \text{id}$

$\Rightarrow \text{id} * \text{id} + \text{id}$

(slope) * (x) + (intercept)

Deriv

This derivation replaces the **rightmost** non-terminal. Derivations with this behavior are called (surprisingly) **rightmost** or **canonical derivation**.

- A **d**erivation

erations that
mbol.”

slope * x + intercept

$\Rightarrow \text{expr op } \underline{\text{expr}} + \text{id}$

$\Rightarrow \text{expr } \underline{\text{op}} \text{ id} + \text{id}$

$\text{expr} \Rightarrow \text{expr op } \underline{\text{expr}}$

$\Rightarrow \underline{\text{expr}} * \text{id} + \text{id}$

$\Rightarrow \text{expr } \underline{\text{op}} \text{ id}$

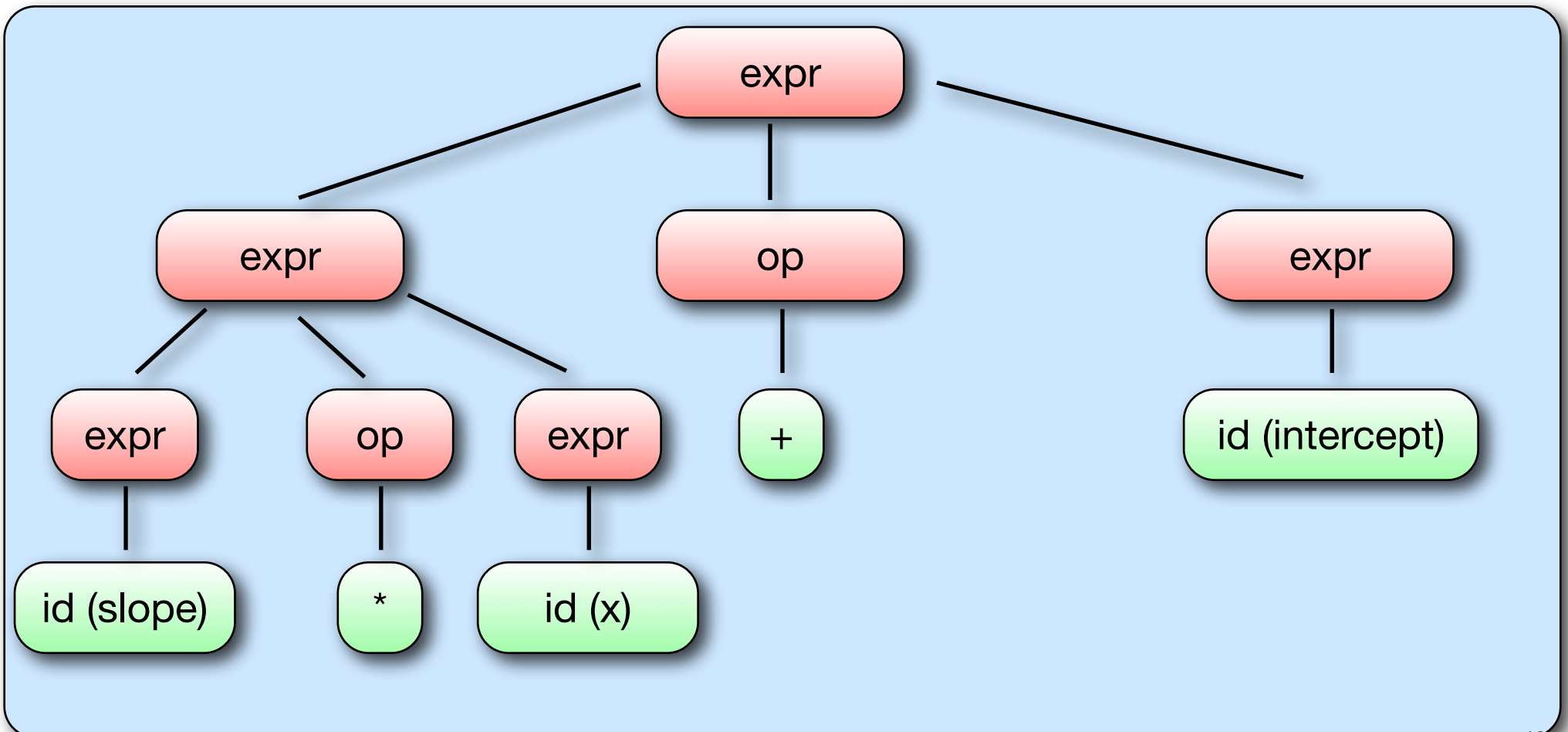
$\Rightarrow \text{id} * \text{id} + \text{id}$

$\Rightarrow \underline{\text{expr}} + \text{id}$

$(\text{slope}) * (\text{x}) + (\text{intercept})$

Parse Tree

- A parse tree is the graphical representation of the derivation.

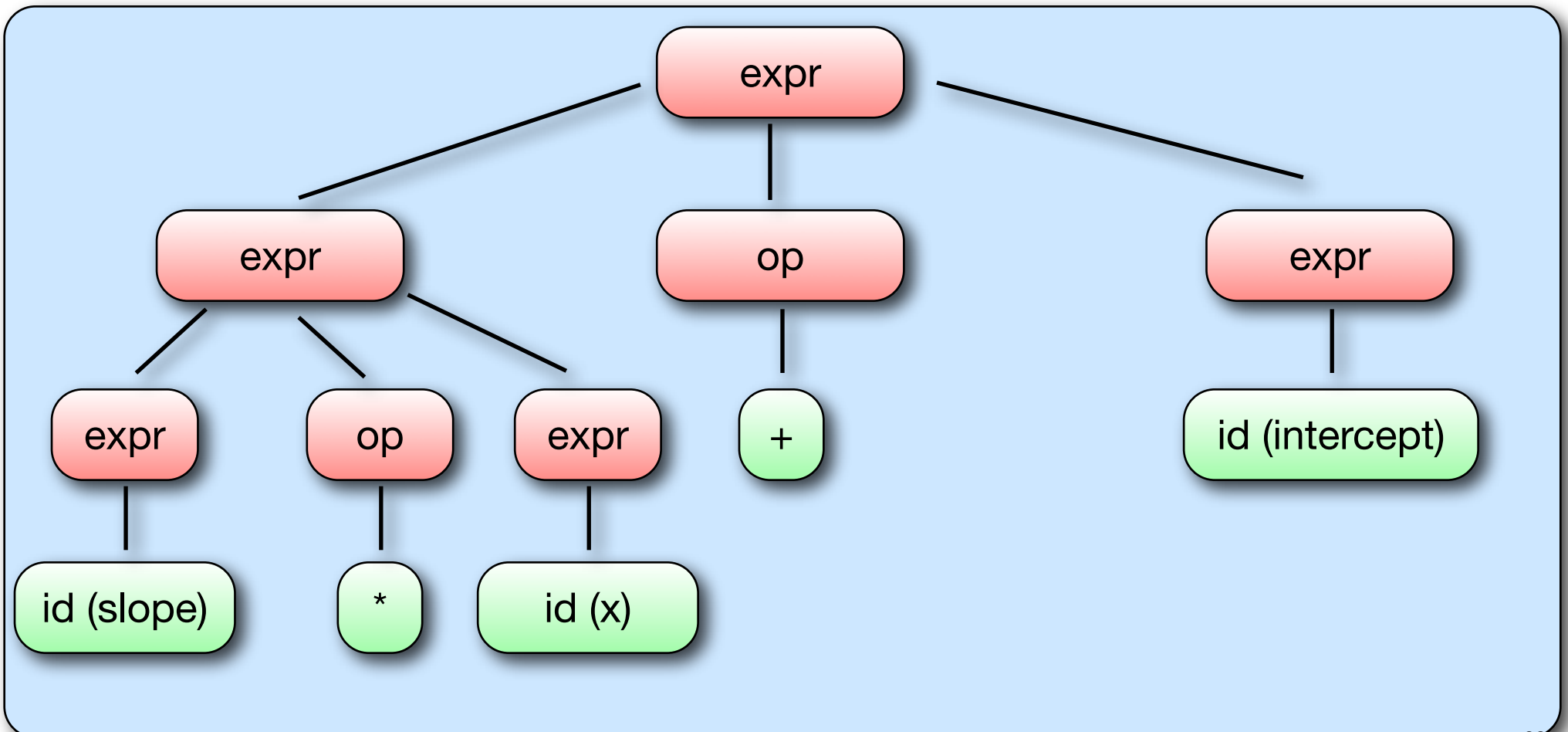


Pa

This parse tree constructs the formula

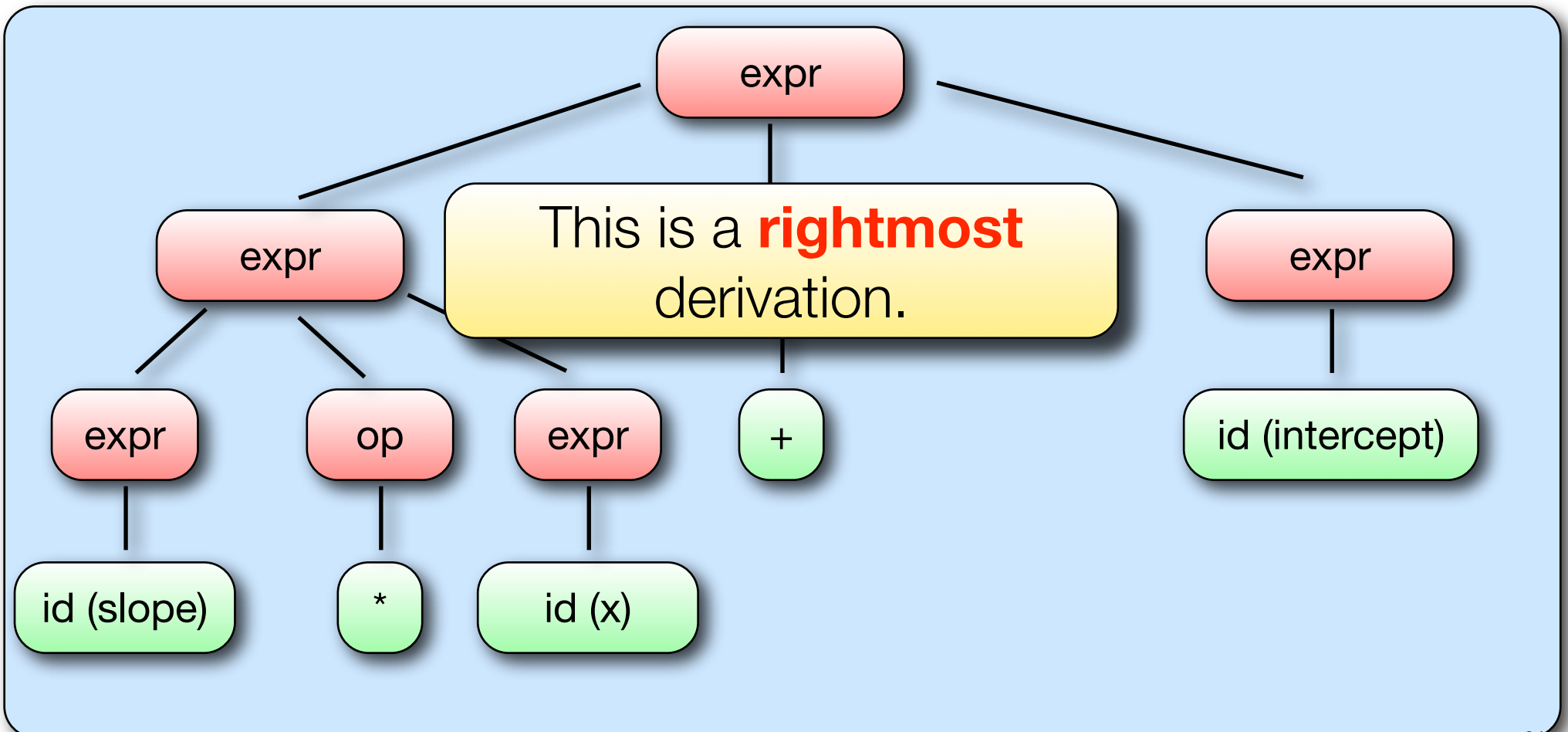
(slope*x) + intercept

- A parse tree is the graphical representation of the derivation.



Parse Tree

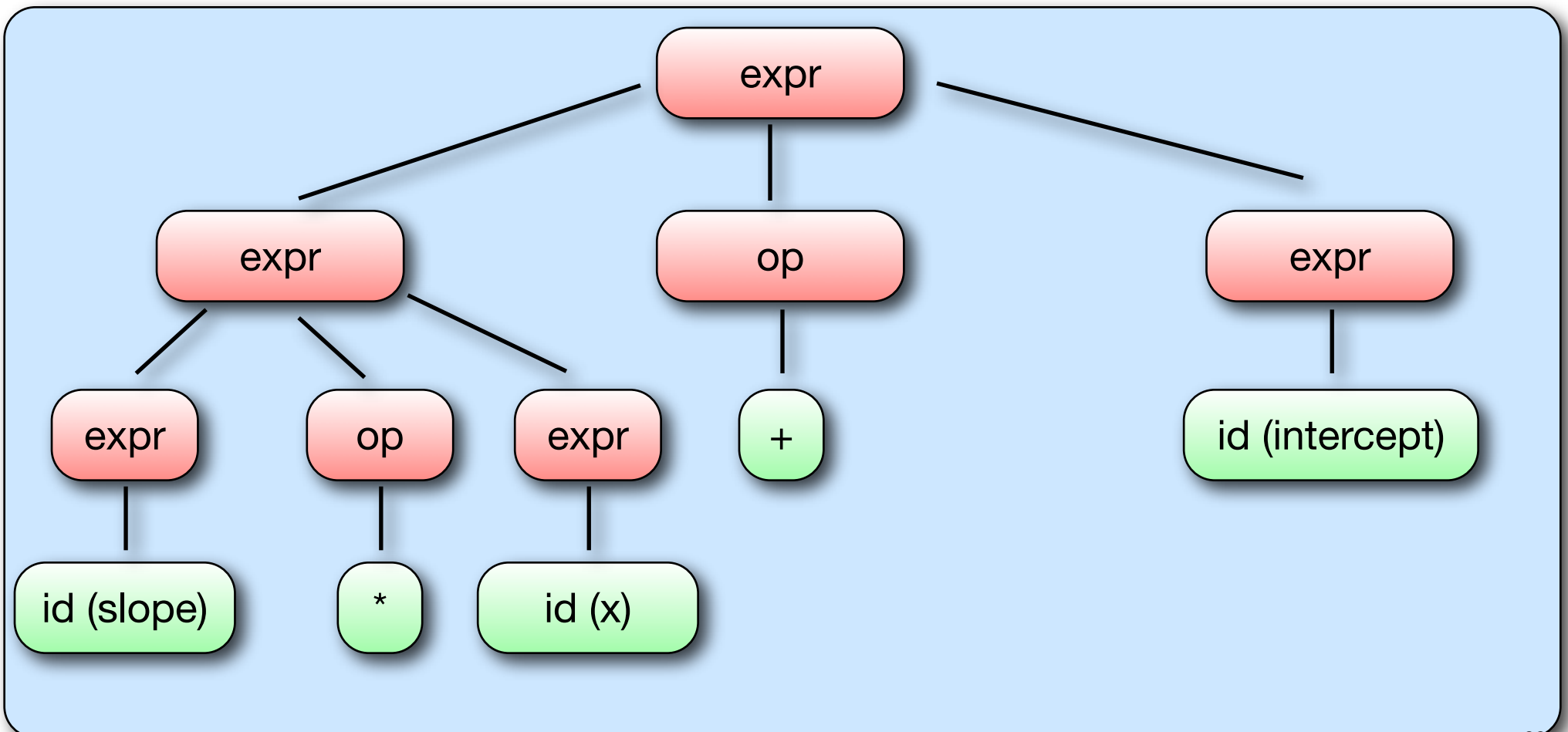
- A parse tree is the graphical representation of the derivation.



Lets try deriving “2*a*b+c”

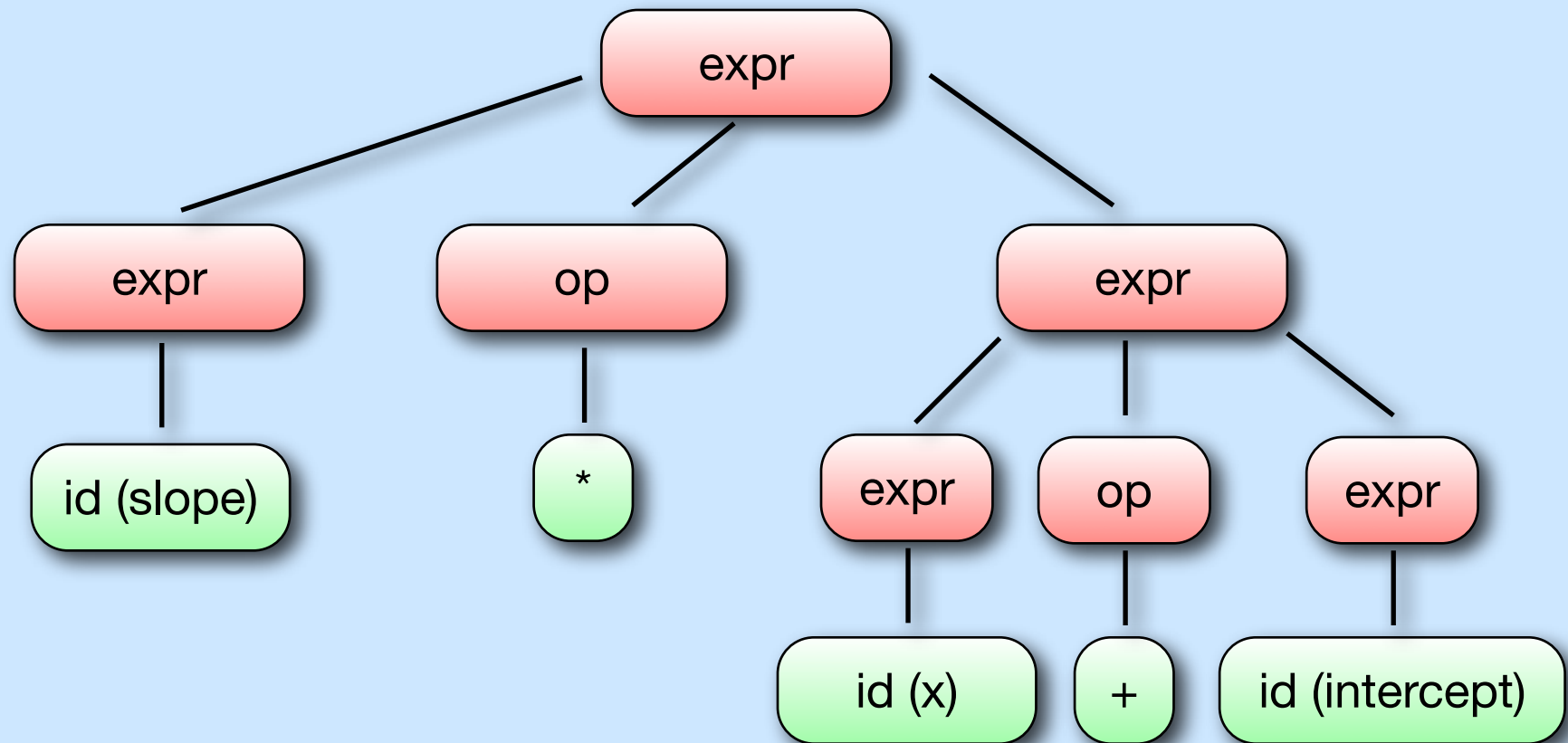
Parse Tree

- A parse tree is the graphical representation of the derivation.



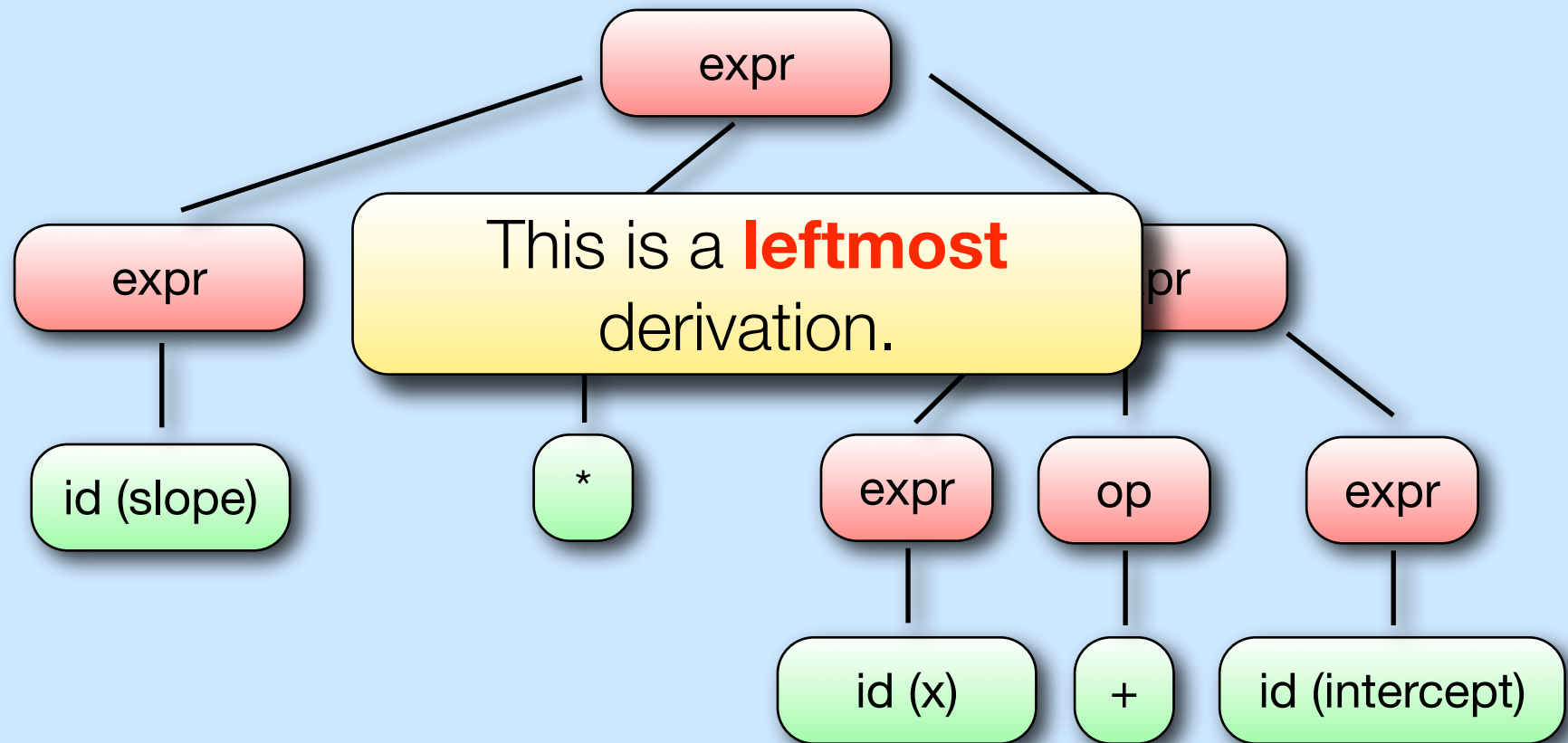
Parse Tree (Ambiguous)

- This grammar is ambiguous and can construct the following parse tree.



Parse Tree (Ambiguous)

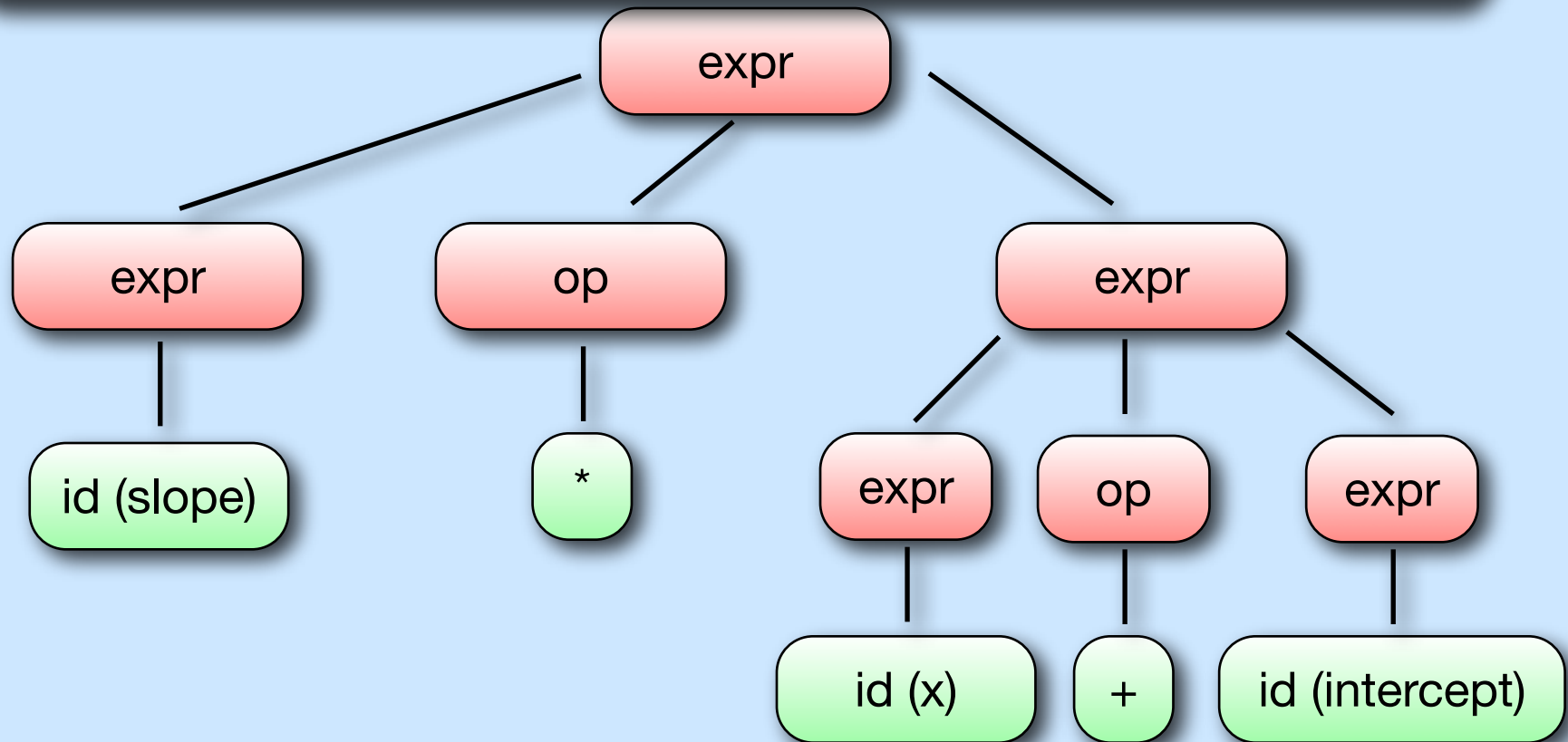
- This grammar is ambiguous and can construct the following parse tree.



Pa

• T
fo

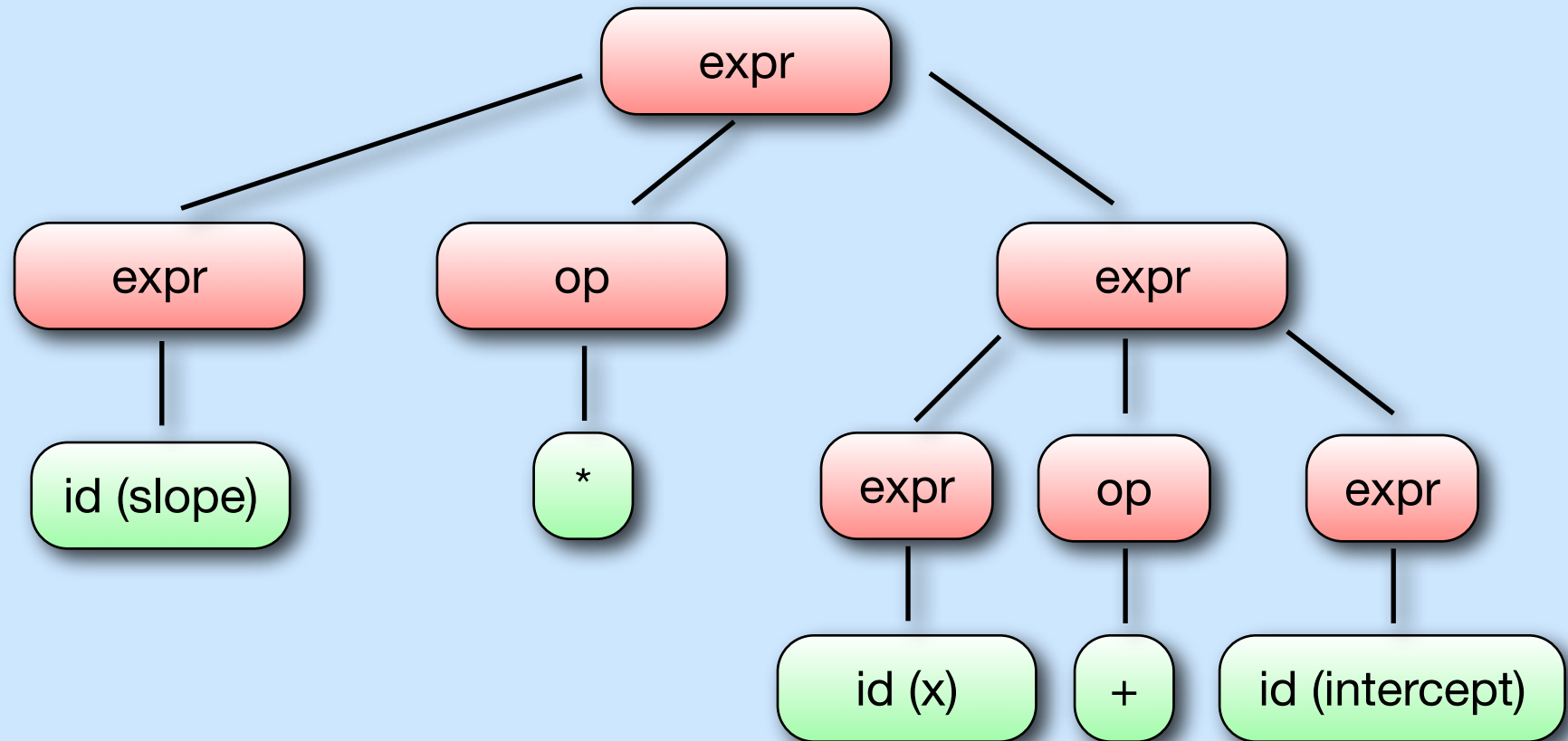
This parse tree constructs the formula
slope*(x+intercept) which is **not equal** to
slope*x + intercept



Lets try deriving “2*a*b+c”

Parse Tree (Ar

- This grammar, is ambiguous and can construct the following parse tree.



Disambiguating grammar

- The problem with our original grammar was that we did **not fully express the grammatical structure** (i.e., associativity and precedence).
- To create an unambiguous grammar, we need to fully specify the grammar.

$expr \rightarrow term \mid expr \text{ add_op } term$

$term \rightarrow factor \mid term \text{ mult_op } factor$

$factor \rightarrow id \mid number \mid - factor \mid (expr)$

$add_op \rightarrow + \mid -$

$mult_op \rightarrow * \mid /$



Disambiguating grammar

- The problem with our original grammar was that we did **not fully express the grammatical structure** (i.e. associativity and precedence)
- To create an unambiguous grammar, we need to fully specify the grammar.

Gives precedence to multiply

$expr \rightarrow term \mid expr \text{ add_op } term$

$term \rightarrow factor \mid term \text{ mult_op } factor$

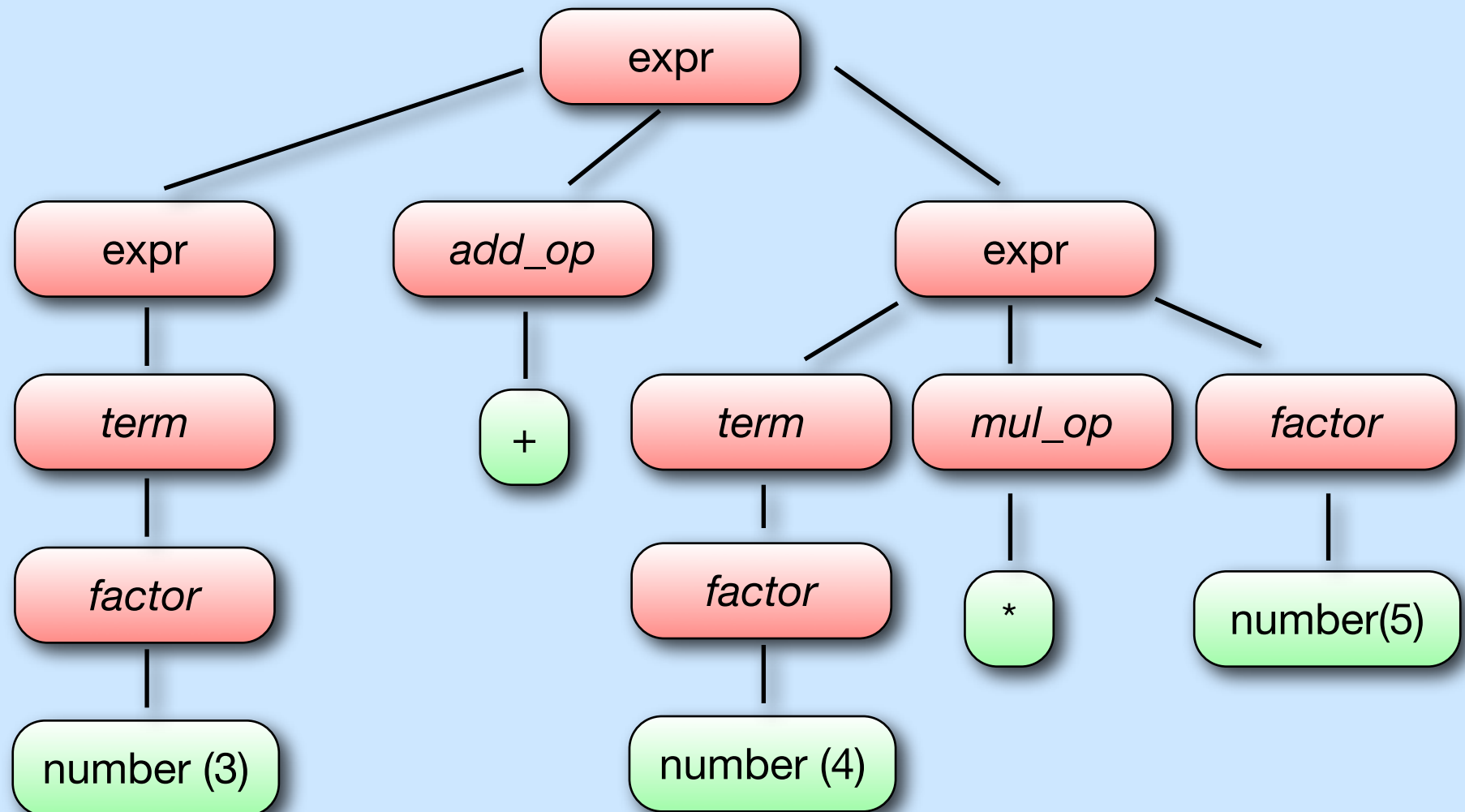
$factor \rightarrow id \mid number \mid - factor \mid (expr)$

$add_op \rightarrow + \mid -$

$mult_op \rightarrow * \mid /$



3+4*5



Disambiguating grammar

- The problem with the previous grammar was that it did **not fully** express the precedence and associativity and
- To create an unambiguous grammar, we need to fully specify the grammar.

Lets try deriving “3*4+5*6+7”

$expr \rightarrow term \mid expr \text{ add_op } term$

$term \rightarrow factor \mid term \text{ mult_op } factor$

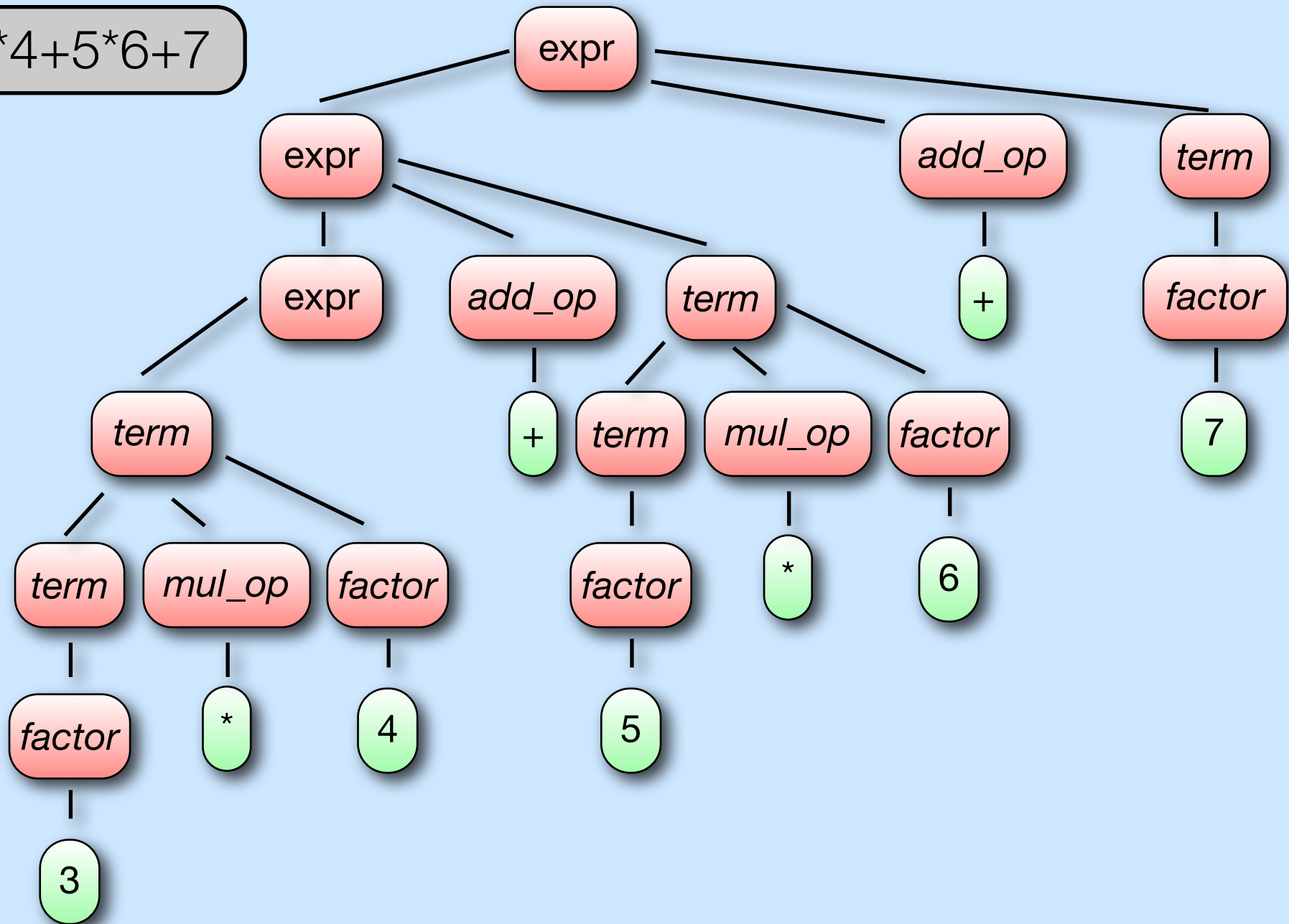
$factor \rightarrow id \mid number \mid - factor \mid (expr)$

$add_op \rightarrow + \mid -$

$mult_op \rightarrow * \mid /$



$3*4+5*6+7$



Disambiguating grammar

- The problem (the grammar did **not fully** express the precedence).
Lets try deriving “3*4” & “3+4”
- To create an unambiguous grammar, we need to fully specify the grammar.

$expr \rightarrow term \mid expr \text{ add_op } term$

$term \rightarrow factor \mid term \text{ mult_op } factor$

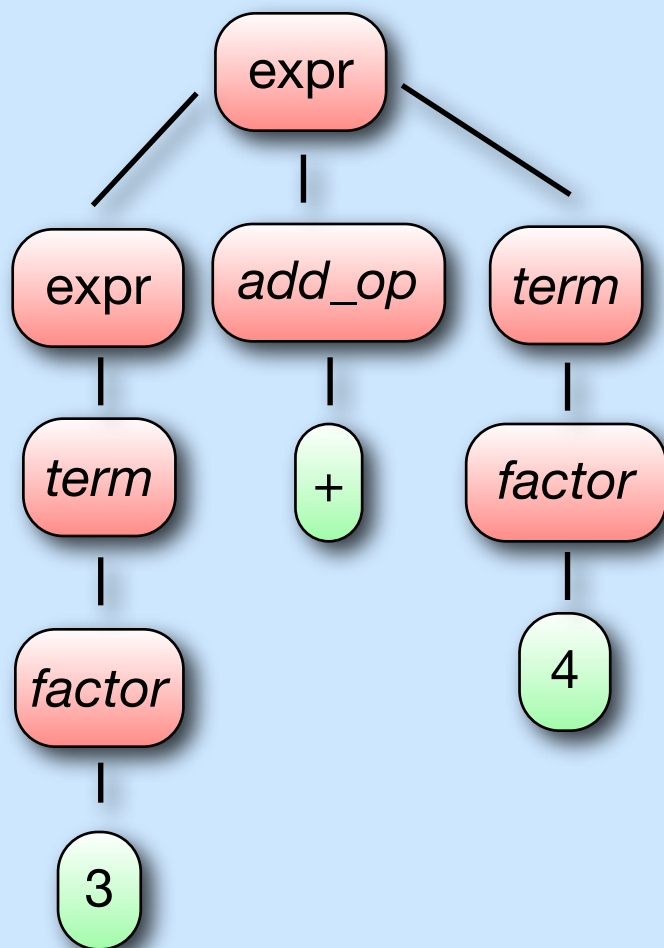
$factor \rightarrow id \mid number \mid - factor \mid (expr)$

$add_op \rightarrow + \mid -$

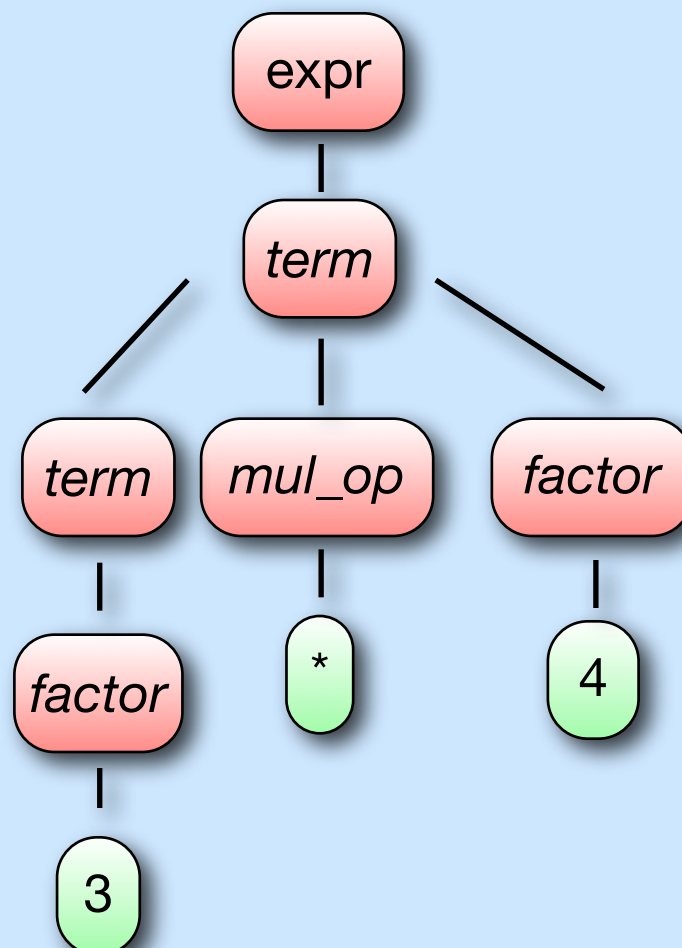
$mult_op \rightarrow * \mid /$



3+4

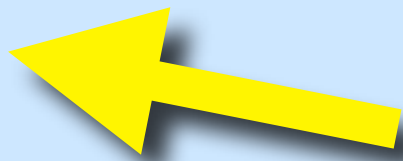


3*4

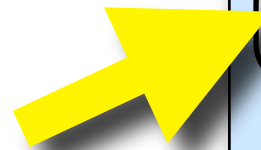


How can you derive these trees by examining one character at a time?

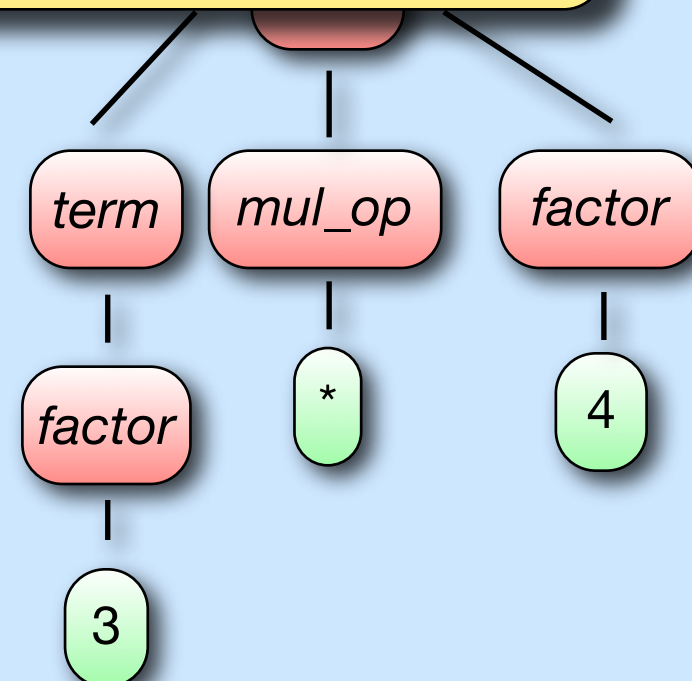
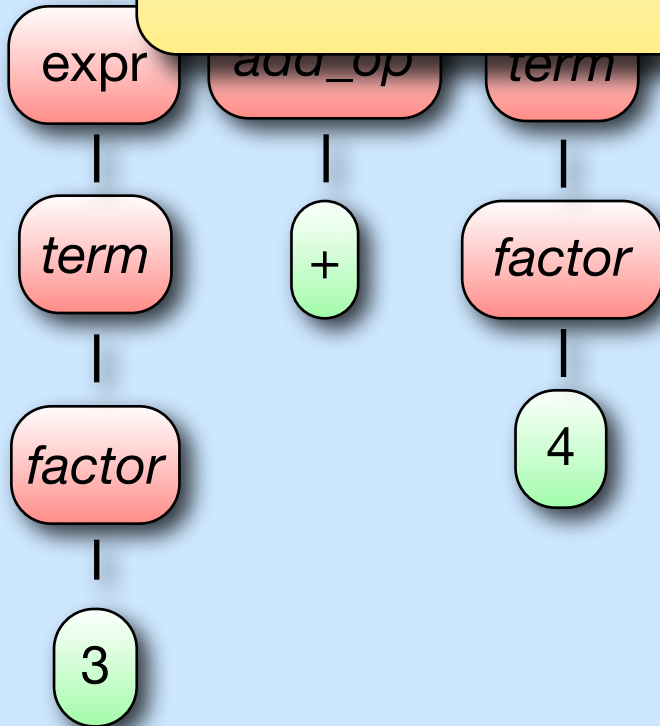
3+4



3*4



In order to derive these trees, the first character that we need to examine is the math symbol.



3+4

3*4

expr

expr

expr

term

factor

3

4

Thus, to parse these “sentences,” we first need to search through them to find the math symbols . . . then we need to sort out the multiplication from the addition. . . ugh. . .

factor

*

4

3

Java Spec

- Available on-line

- http://java.sun.com/docs/books/jls/second_edition/html/j.title.doc.html

- Examples

- Comments: http://java.sun.com/docs/books/jls/second_edition/html/lexical.doc.html#48125
 - Multiplicative Operators: http://java.sun.com/docs/books/jls/second_edition/html/expressions.doc.html#239829
 - Unary Operators: http://java.sun.com/docs/books/jls/second_edition/html/expressions.doc.html#4990



LL and LR Derivation

- CFGs can be parsed in **$O(n^3)$** time, where n is length of the program.
- This is too long for most code; however, there are two types of grammars that can be parsed in **linear time**, i.e., **$O(n)$** .
 - **LL**: “Left-to-right, Left-most derivation”
 - **LR**: “Left-to-right, Right-most derivation”



Top-down

- **LL parsers** are top-down, i.e., they identify the non-terminals first and terminals second.
- **LL grammars** are grammars that can be parsed by an LL parser.



Top-Down

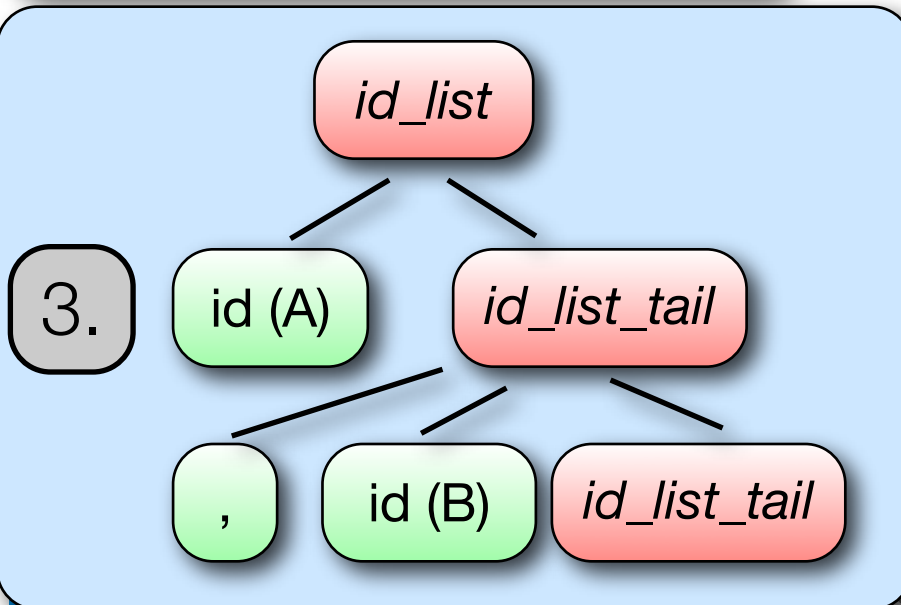
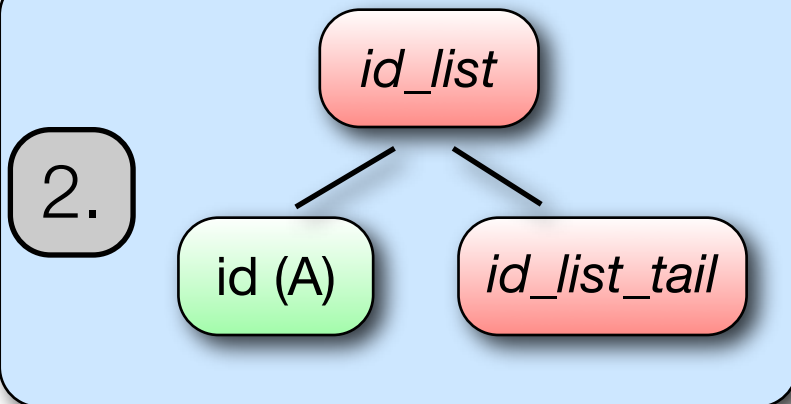
$id_list \rightarrow id\ id_list_tail$

$id_list_tail \rightarrow ;$

$id_list_tail \rightarrow ,\ id\ id_list_tail$

A, B, C;

1. id_list



Top-Down

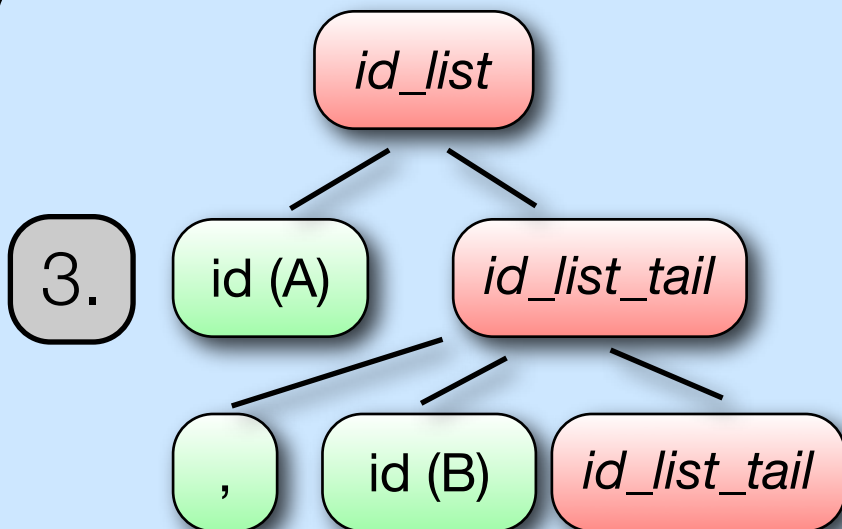
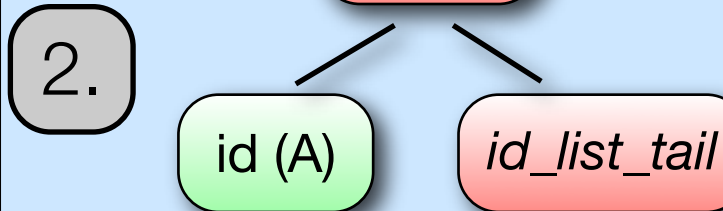
$id_list \rightarrow id\ id_list_tail$

$id_list_tail \rightarrow ;$

$id_list_tail \rightarrow ,\ id\ id_list_tail$

A, B, C;

1. id_list



LL parsers are some time called **predictive** because they “predict” the next state.



Top-Down

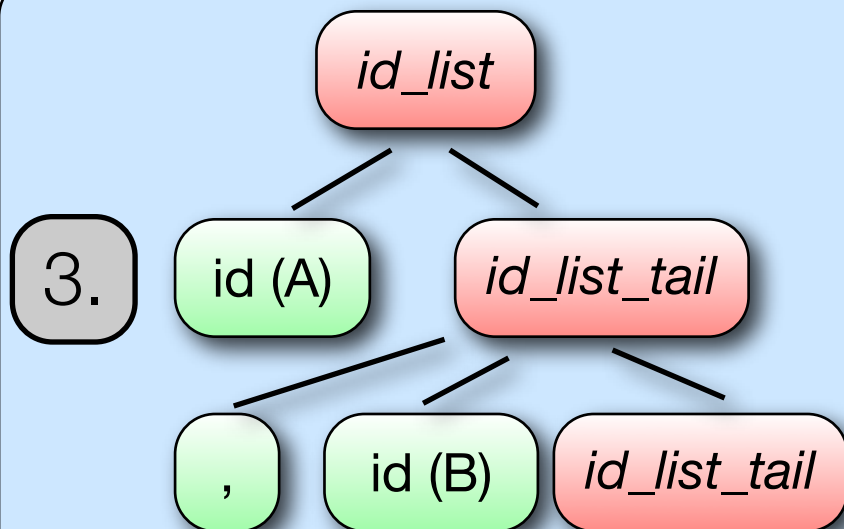
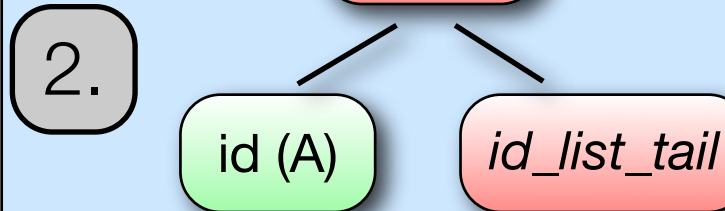
$id_list \rightarrow id\ id_list_tail$

$id_list_tail \rightarrow ;$

$id_list_tail \rightarrow ,\ id\ id_list_tail$

A, B, C;

1. id_list



For example, after $id(A)$ is “discovered”, the next state is “predicted as id_list_tail .”



Top-Down

$id_list \rightarrow id\ id_list_tail$

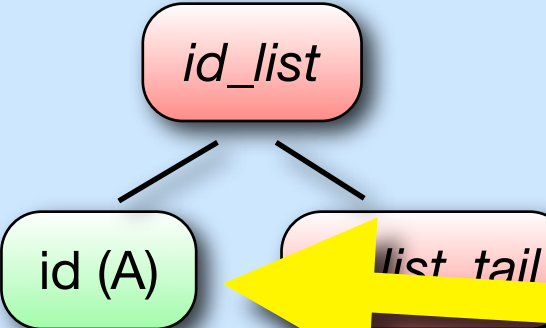
$id_list_tail \rightarrow ;$

$id_list_tail \rightarrow ,\ id\ id_list_tail$

A, B, C;

1. id_list

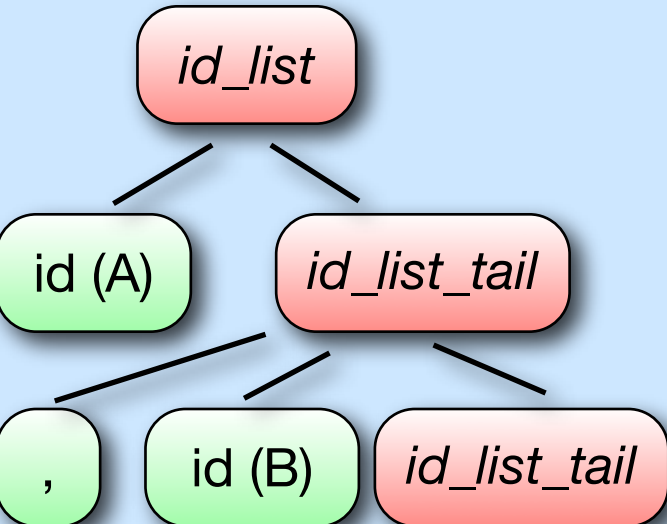
2.



```
graph TD; id_list[id_list] --- id_A[id (A)]; id_list --- list_tail[id_list tail];
```

A partial parse tree with root id_list (red) and two children: $id\ (A)$ (green) and $id_list\ tail$ (pink). A yellow arrow points from the text box to the $id_list\ tail$ node.

3.



```
graph TD; id_list[id_list] --- id_A[id (A)]; id_list --- id_list_tail[id_list_tail]; id_list_tail --- comma[,]; id_list_tail --- id_B[id (B)]; id_list_tail --- id_list_tail[id_list_tail];
```

A complete parse tree for the first two tokens. The root id_list (red) has children $id\ (A)$ (green) and id_list_tail (pink). The id_list_tail node has three children: a comma (green), $id\ (B)$ (green), and another id_list_tail (pink).

Notice that tokens are placed in the tree from the **left-most** to **right-most**.



Bottom-up

- **LR parsers** are bottom-up, i.e., they discover the terminals first and non-terminals second.
 - **LR grammars** are grammars that can be parsed by an LR parser.
 - All **LL grammars are LR grammars** but not vice versa.



Bottom-up

$id_list \rightarrow id\ id_list_tail$

$id_list_tail \rightarrow ;$

$id_list_tail \rightarrow ,\ id\ id_list_tail$

A, B, C;

1. id (A)

2. id (A) ,

6. id (A) , id (B) , id (C) ;

7. id (A) , id (B) , id (C) id_list_tail

;



Bottom-up

$id_list \rightarrow id\ id_list_tail$

$id_list_tail \rightarrow ;$

$id_list_tail \rightarrow ,\ id\ id_list_tail$

A, B, C;

8.

id (A)

,

id (B)

id_list_tail

id_list_tail

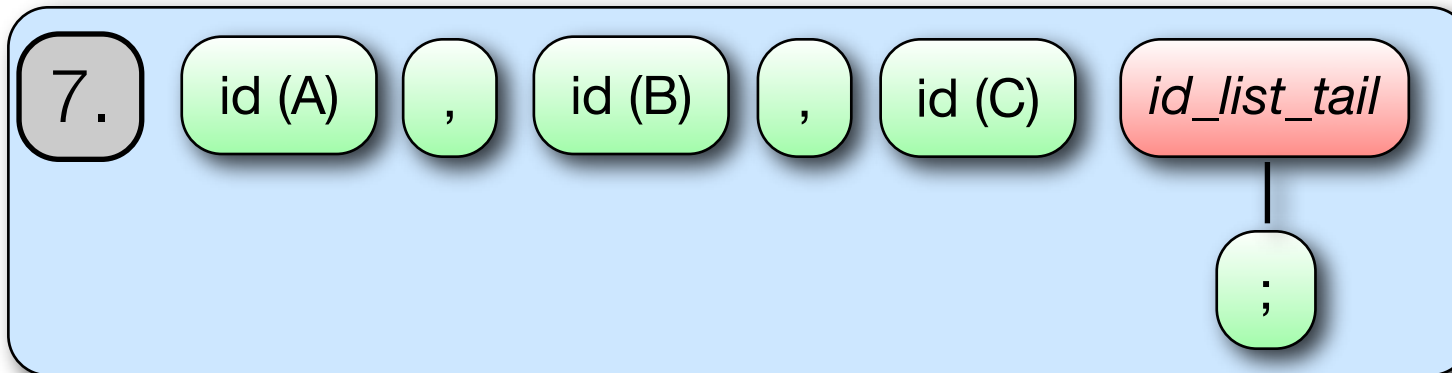
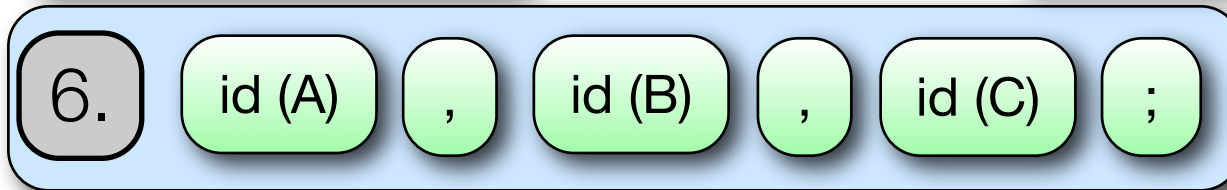
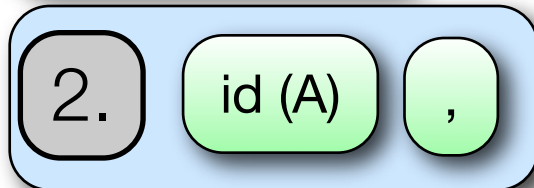
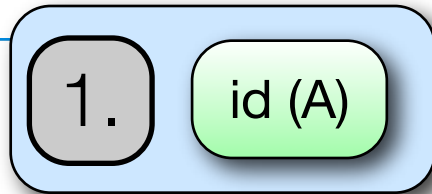
,

id (C)

;



Bottom-up



$id_list \rightarrow id\ id_list_tail$

$id_list_tail \rightarrow ;$

$id_list_tail \rightarrow ,\ id\ id_list_tail$

A B C:

LR parsers are sometimes called **shift** because they “shift” the states



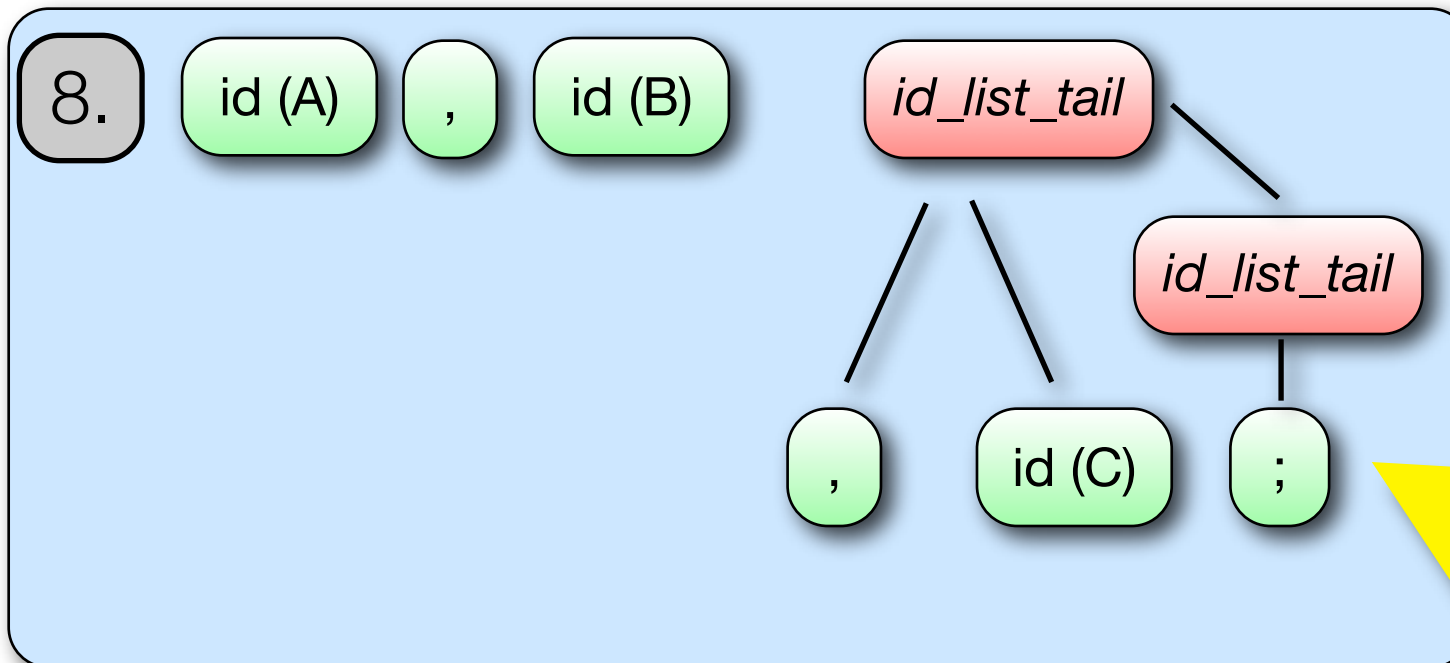
Bottom-up

$id_list \rightarrow id\ id_list_tail$

$id_list_tail \rightarrow ;$

$id_list_tail \rightarrow ,\ id\ id_list_tail$

A, B, C;



Notice that tokens are added to the tree from the **right-most** to the **left-most**.



Bottom-up

$id_list \rightarrow id\ id_list_tail$

$id_list_tail \rightarrow ;$

$id_list_tail \rightarrow ,\ id\ id_list_tail$

A, B, C;

8.

id (A)

,

id (B)

id_list_tail

id_list_tail

Sometimes you see LL and LR parsers written as **LL(n)** and **LR(n)** to denote the parser needs to look ahead **n** tokens .



Bottom-up

$id_list \rightarrow id\ id_list_tail$

$id_list_tail \rightarrow ;$

$id_list_tail \rightarrow ,\ id\ id_list_tail$

A, B, C;

8.

id (A)

,

id (B)

id_list_tail

id_list_tail

The problem with this grammar is that it can require an **arbitrarily large** number of terminals to be “shifted” before placing them into the tree



A better bottom-up grammar

This grammar limits the number of “suspended” non-terminals.

$id_list \rightarrow id_list_prefix ;$

$id_list_prefix \rightarrow id_list_prefix, id$

$id_list_prefix \rightarrow id$



A better bottom-up grammar

Lets try parsing “A, B, C;”

$id_list \rightarrow id_list_prefix ;$

$id_list_prefix \rightarrow id_list_prefix, id$

$id_list_prefix \rightarrow id$



A better bottom-up grammar

However, it **cannot** be parsed by LL (top-down) parser. Since when the parser discovers an “id” it does not “know” the number of “id_list_prefixs”

$id_list \rightarrow id_list_prefix ;$

$id_list_prefix \rightarrow id_list_prefix, id$

$id_list_prefix \rightarrow id$

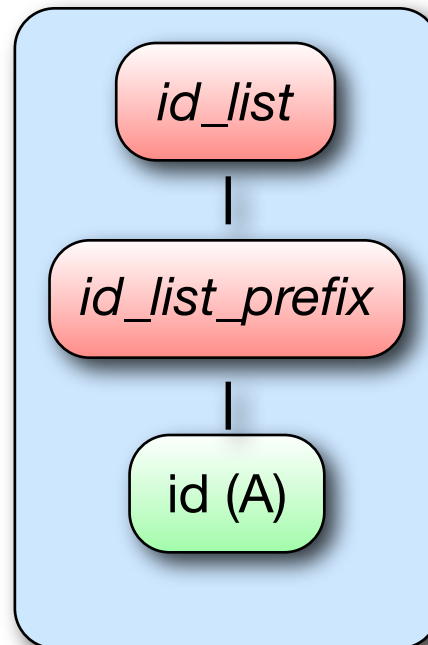
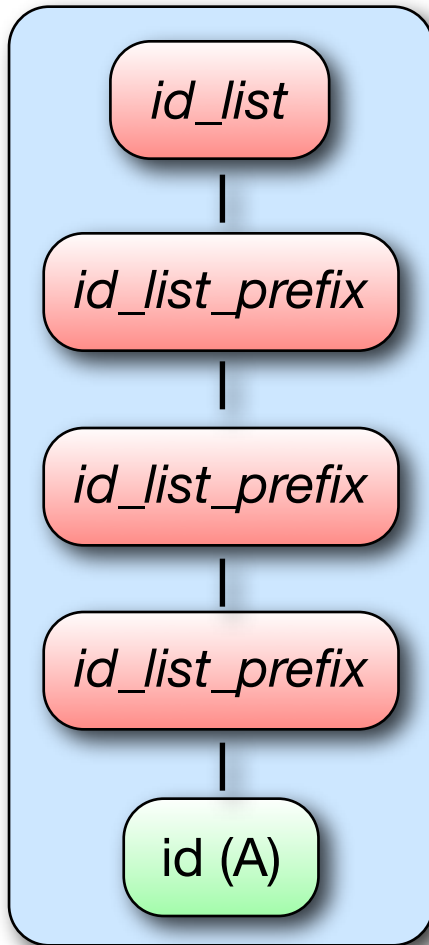


A better bottom-up grammar

$id_list \rightarrow id_list_prefix ;$

$id_list_prefix \rightarrow id_list_prefix, id$

$id_list_prefix \rightarrow id$



Both of these are valid break downs, but we don't know which one. Therefore, is **NOT** a valid **LL grammar**, but it is a valid **LR grammar**.



$program \rightarrow stmt_list \$\$$

$stmt_list \rightarrow stmt\ stmt_list \mid \epsilon$

$stmt \rightarrow id := expr \mid read\ id \mid write\ expr$

$expr \rightarrow term\ term_tail$

$term_tail \rightarrow add_op\ term\ term_tail \mid \epsilon$

$term \rightarrow factor\ factor_tail$

$factor_tail \rightarrow mult_op\ factor\ factor_tail \mid \epsilon$

$factor \rightarrow (expr) \mid id \mid literal$

$add_op \rightarrow + \mid -$

$mult_op \rightarrow * \mid /$

This grammar (for the calculator) unlike the previous calculator grammar is an LL grammar, because when an “id” is encountered **we know exactly where it belongs.**

$program \rightarrow stmt_list \$\$$

$stmt_list \rightarrow stmt\ stmt_list \mid \epsilon$

$stmt \rightarrow id := expr \mid read\ id \mid write\ expr$

$expr \rightarrow term\ term_tail$

$term_tail \rightarrow add_op\ term\ term_tail \mid \epsilon$

$term \rightarrow factor\ factor_tail$

$factor_tail \rightarrow mult_op\ factor\ factor_tail \mid \epsilon$

$factor \rightarrow (expr) \mid id \mid literal$

$add_op \rightarrow + \mid -$

$mult_op \rightarrow * \mid /$

Let's try "c := 2*A+B"

Let's try "2*A+B"

$expr \rightarrow term \mid expr \text{ add_op } term$

$term \rightarrow factor \mid term \text{ mult_op } factor$

$factor \rightarrow id \mid number \mid - factor \mid (expr)$

$add_op \rightarrow + \mid -$

$mult_op \rightarrow * \mid /$



Let's try "c := 2*A+B"

$expr \rightarrow id \mid number \mid - expr \mid (expr) \mid expr \ op \ expr$

$op \rightarrow + \mid - \mid * \mid /$



Recursive Descent & LL Parse Table

- There are two ways to code a parser for LL grammars:
 - **Recursive Descent**, which is a recursive program with case statements that correspond to each one-to-one to nonterminals.
 - **LL Parse Table**, which consist of an iterative driver program and a table that contains all of the nonterminals.



$program \rightarrow stmt_list \$\$$

$stmt_list \rightarrow stmt\ stmt_list \mid \epsilon$

$stmt \rightarrow id := expr \mid read\ id \mid write\ expr$

$expr \rightarrow term\ term_tail$

$term_tail \rightarrow add_op\ term\ term_tail \mid \epsilon$

$term \rightarrow factor\ factor_tail$

$factor_tail \rightarrow mult_op\ factor\ factor_tail \mid \epsilon$

$factor \rightarrow (expr) \mid id \mid literal$

$add_op \rightarrow + \mid -$

$mult_op \rightarrow * \mid /$

Recursive Descent

procedure program()

case in_tok **of**

id, read, write, \$\$:
 stmt_list()
 match(\$\$)

else

return error

procedure stmt()

case in_tok **of**

id: match(id); match(:=); expr()
read: match(read); match(id)
write: match(write); expr()

else

return error

procedure stmt_list()

case in_tok **of**

id, read, write:
 stmt(); stmt_list();

\$\$:

skip

else

return error

procedure match(expec)

if in_tok = expec

 consume in_tok

else

return error

Recursive Descent

```
procedure program
case in_tok of
  id, read, wr
    stmt_list()
    match($$)
else
  return error
```

- **in_tok** is a global variable that is the current token
- **consume** changes in_tok to the next token.

```
  n(:=); expr()
  match(id)
  e); expr()
```

return error

```
procedure stmt_list()
case in_tok of
  id, read, write:
    stmt(); stmt_list();
  $$:
    skip
else
  return error
```

```
procedure match(expec)
if in_tok = expec
  consume in_tok
else
  return error
```

Recursive Descent

procedure program()

case in_tok **of**

id, read, write, \$\$:
 stmt_list()
 match(\$\$)

else

return error

procedure stmt()

case in_tok **of**

id: match(id); match(:=); expr()
read: match(read); match(id)
write: match(write); expr()

else

return error

procedure stmt_list()

case in_tok **of**

id, read, write:
 stmt(); stmt_list();

\$\$:

skip

else

return error

procedure match(expec)

if in_tok = expec

 consume in_tok

else

return error

Recursive Descent

The question is how do we label the case statements?

```
procedure program()
```

```
  case in_tok of  
    id, read, write, $$:  
      stmt_list()  
      match($$)
```

```
  else  
    return error
```

```
procedure stmt()
```

```
  case in_tok of  
    id: match(id); match(:=); expr()  
    read: match(read); match(id)  
    write: match(write); expr()
```

```
  else  
    return error
```

```
procedure stmt_list()
```

```
  case in_tok of  
    id, read, write:  
      stmt_list()
```

```
    $$:  
      skip
```

```
  else  
    return error
```

```
procedure match(expec)
```

```
  if in_tok = expec  
    consume in_tok
```

```
  else  
    return error
```

First, Follow, and Predict

- Three functions allow us to label the branches
 - **FIRST(a)**: The terminals (and ϵ) that can be the first tokens of the non-terminal symbol **a**.
 - **FOLLOW(A)**: The terminals that can follow the terminal or non-terminal symbol **A**
 - **PREDICT(A $\rightarrow$ a)**: The terminals that can be the first tokens as a result of the production **A $\rightarrow$ a**



$program \rightarrow stmt_list \$\$$

$stmt_list \rightarrow stmt\ stmt_list \mid \epsilon$

$stmt \rightarrow id := expr \mid read\ id \mid write\ expr$

$expr \rightarrow term\ term_tail$

$term_tail \rightarrow add_op\ term\ term_tail \mid \epsilon$

$term \rightarrow factor\ factor_tail$

$factor_tail \rightarrow mult_op\ factor\ factor_tail \mid \epsilon$

$factor \rightarrow (expr) \mid id \mid literal$

$add_op \rightarrow + \mid -$

$mult_op \rightarrow * \mid /$

- $FIRST(program) = \{id, read, write, \$\}$
- $FOLLOW(program) = \{\epsilon\}$
- $PREDICT(program \rightarrow stmt_list \$\$) = \{id, read, write, \$\}$
- $FOLLOW(id) = \{+, -, *, /,), :=, id, read, write, \$\}$

$program \rightarrow stmt_list \$\$$

$stmt_list \rightarrow stmt\ stmt_list \mid \epsilon$

$stmt \rightarrow id := expr \mid read\ id \mid write\ expr$

$expr \rightarrow term\ term_tail$

$term_tail \rightarrow add_op\ term\ term_tail \mid \epsilon$

$term \rightarrow factor\ factor_tail$

$factor_tail \rightarrow mult_op\ factor\ factor_tail \mid \epsilon$

$factor \rightarrow (expr) \mid id \mid literal$

$add_op \rightarrow + \mid -$

$mult_op \rightarrow * \mid /$

- $FIRST(factor_tail) = \{ *, /, \epsilon \}$
- $FOLLOW(factor_tail) = \{ +, -,), id, read, write, \$\$ \}$
- $PREDICT(factor_tail \rightarrow m_op\ factor\ factor_tail) = \{ *, / \}$
- $PREDICT(factor_tail \rightarrow \epsilon) = \{ +, -,), id, read, write, \$\$ \}$

$program \rightarrow stmt_list \$\$$

$stmt_list \rightarrow stmt \mid stmt_list \mid \epsilon$

$stmt \rightarrow id := expr \mid \epsilon$

$expr \rightarrow term \mid term \mid \epsilon$

$term_tail \rightarrow add_op \ term \ term_tail \mid \epsilon$

$term \rightarrow factor \ factor_tail$

$factor_tail \rightarrow mult_op \ factor \ factor_tail \mid \epsilon$

$factor \rightarrow (expr) \mid id \mid literal$

$add_op \rightarrow + \mid -$

$mult_op \rightarrow * \mid /$

Since $factor_tail$ can be “transformed” into an empty statement **PREDICT($factor_tail \rightarrow \epsilon$)** equals **FOLLOW($factor_tail$)**

$PREDICT(factor_tail \rightarrow \epsilon) =$
 $\{+, -,), id, read,$
 $write, \$\$ \}$

- $PREDICT(factor_tail \rightarrow m_op \ factor \ factor_tail)$
 $= \{*, /\}$

- $PREDICT(factor_tail \rightarrow \epsilon)$
 $= \{+, -,), id, read,$
 $write, \$\$ \}$

Recursive Descent

These are all of the PREDICT() values from every production.

procedure program()

case in_tok of
id, read, write, \$\$:
stmt_list()
match(\$\$)
else
return error

procedure stmt()

case in_tok of
id: match(id); match(:=); expr()
read: match(read); match(id)
write: match(write); expr()
else
return error

procedure stmt_list()

case in_tok of
id, read, write:
stmt_list(),
\$\$:
skip
else
return error

procedure match(expec)

if in_tok = expec
consume in_tok
else
return error

Constructing FIRST, FOLLOW, and PREDICT

- To construct the FIRST, FOLLOW, and PREDICT tables we iterate through the grammar building on knowledge.
 - First, we define all of the “**obvious**” FIRST and FOLLOW values
 - For example, $\$ \$ \in \text{FOLLOW}(stmt_list)$ and $\{id, read, write\} \in \text{FIRST}(stmt)$
 - Next, **we build on this**,
 - For example, $\{id, read, write\} \in \text{FIRST}(stmt_list)$ since *stmt_list* can begin with *stmt* and $\{id, read, write\} \in \text{FIRST}(stmt)$
 - We then continue on until we get **no more knowledge**.



Let's try making tables
for this grammar

$expr \rightarrow term \mid expr \text{ add_op } term$

$term \rightarrow factor \mid term \text{ mult_op } factor$

$factor \rightarrow id \mid number \mid - factor \mid (expr)$

$add_op \rightarrow + \mid -$

$mult_op \rightarrow * \mid /$



Table-Driven

```
p_stack: stack of symbols;
p_stack.push(st_symbol);
loop
  exp_sym := p_stack.pop
  if exp_sym = terminal
    match(exp_sym);
    if exp_sym==$ return
  else
    if table[exp_sym,in_tok].action = error
      return error
    else
      prediction := table[exp_sym,in_tok].prod;
      foreach sym in prod_table[prediction]
        p_stack.push(sym)
```



Parse Stack	Input Stream
<i>program</i>	read A read B
<i>stmt_list</i> \$\$	read A read B
<i>stmt stmt_list</i> \$\$	read A read B
read id <i>stmt_list</i> \$\$	read A read B
id <i>stmt_list</i> \$\$	A read B

Writing an LL(1) Grammar--Left Recursion

- **Left recursion** is where the leftmost symbol is a recursive non-terminal symbol.

$id_list \rightarrow id_list_prefix;$

$id_list_prefix \rightarrow id_list_prefix, id$

$id_list_prefix \rightarrow id$

- This can cause a grammar **not to be LL(1)**.
- It is **desirable for LR grammars**.



Writing an LL(1) Grammar-- Eliminating Left Recursion

- To eliminate left recursion replace it with **right recursion**.

$id_list \rightarrow id_list_prefix;$

$id_list_prefix \rightarrow id_list_prefix, id$

$id_list_prefix \rightarrow id$

$id_list \rightarrow id\ id_list_tail$

$id_list_tail \rightarrow , id\ id_list_tail$

$id_list_tail \rightarrow ;$



Writing an LL(1) Grammar--Common Prefix

- **Common prefixes** occur when there is more than one prefix for a given nonterminal.

$stmt \rightarrow id := expr$

$stmt \rightarrow id (arguments)$

- Again, this causes a grammar **not to be LL(1)**.



Writing an LL(1) Grammar--Left factoring

- **Common prefixes** To get rid of common prefixes we use a technique called **left factoring**.

$stmt \rightarrow id := expr$

$stmt \rightarrow id (arguments)$

$stmt \rightarrow id stmt_list_tail$

$stmt_list_tail \rightarrow := expr \mid (arguments)$



Dangling else

- Even if left recursion and common prefixes don't exist a language may not be LL(1).
- In Pascal, there is the problem that an else statement in if-then-else statements is optional. Because **we don't know which if to match else to**.

```
if AAA then
    if BBB then
        CCC
else
    DDD
```



Dangling else

- In Pascal there is **NO LL(1)** parser that can handle this problem.
- Even though a proper LR(1) parser can handle this, it may not handle it in a **method the programmer desires**.

```
if AAA then
  if BBB then
    CCC
else
  DDD
```



Dangling else

- Thus, to write this code correctly (based on indentation) “begin” and “end” statements must be added.

```
if AAA then
  if BBB then
    CCC
else
  DDD
```

```
if AAA then
  begin
    if BBB then
      CCC
    end
  else
    DDD
```

