

For $A \rightarrow X_1 \dots X_n$, $\text{Predict}(A \rightarrow X_1 \dots X_n) = \text{Set of all initial (first) tokens derivable from } A \rightarrow X_1 \dots X_n$
 $= \{a \text{ in } V_t \mid A \rightarrow X_1 \dots X_n \Rightarrow^* a \dots\}$

For example, given

Stmt $\rightarrow$ **Label** **id** = **Expr** ;
 | **Label** **if** **Expr** **then** **Stmt** ;
 | **Label** **read** (**IdList**) ;
 | **Label** **id** (**Args**) ;
Label $\rightarrow$ **intlrit** :
 | λ

Production	Predict Set
Stmt $\rightarrow$ Label id = Expr ;	{ id , intlrit }
Stmt $\rightarrow$ Label if Expr then Stmt ;	{ if , intlrit }
Stmt $\rightarrow$ Label read (IdList) ;	{ read , intlrit }
Stmt $\rightarrow$ Label id (Args) ;	{ id , intlrit }

We now will match a production p only if the next unmatched token is in p 's predict set. We'll avoid trying productions that clearly won't work, so parsing will be faster.

But what is the predict set of a λ -production?

It can't be what's generated by λ (which is nothing!), so we'll define it as the tokens that can *follow* the use of a λ -production.

That is, $\text{Predict}(A \rightarrow \lambda) = \text{Follow}(A)$ where (by definition)

$\text{Follow}(A) = \{a \text{ in } V_t \mid S \Rightarrow^+ \dots Aa \dots\}$

In our example,
 $\text{Follow}(\text{Label} \rightarrow \lambda) = \{\text{id}, \text{if}, \text{read}\}$
 (since these terminals can immediately follow uses of **Label** in the given productions).

Now let's parse

id (**intlrit**) ;

Our start symbol is **Stmt** and the initial token is **id**.

id can predict

Stmt $\rightarrow$ **Label** **id** = **Expr** ;

id then predicts **Label** $\rightarrow \lambda$

The **id** is matched, but "(" doesn't match "=" so we backup and try a different production for **Stmt**.

id also predicts

Stmt $\rightarrow$ **Label** **id** (**Args**) ;

Again, **Label** $\rightarrow \lambda$ is predicted and used, and the input tokens can match the rest of the remaining production.

We had only one misprediction, which is better than before.

Now we'll rewrite the productions a bit to make predictions easier.

We remove the **Label** prefix from all the statement productions (now **intlrit** won't predict all four productions).

We now have

Stmt $\rightarrow$ **Label** **BasicStmt**

BasicStmt $\rightarrow$ **id** = **Expr** ;

 | **if** **Expr** **then** **Stmt** ;

 | **read** (**IdList**) ;

 | **id** (**Args**) ;

Label $\rightarrow$ **intlrit** :

 | λ

Now **id** predicts two different **BasicStmt** productions. If we rewrite these two productions into

BasicStmt $\rightarrow$ **id** **StmtSuffix**

StmtSuffix $\rightarrow$ = **Expr** ;

 | (**Args**) ;

we no longer have any doubt over which production id predicts.

We now have

Production	Predict Set
Stmt $\rightarrow$ Label BasicStmt	Not needed!
BasicStmt $\rightarrow$ id StmtSuffix	{id}
BasicStmt $\rightarrow$ if Expr then Stmt ;	{if}
BasicStmt $\rightarrow$ read (IdList) ;	{read}
StmtSuffix $\rightarrow$ (Args) ;	{ (}
StmtSuffix $\rightarrow$ = Expr ;	{ = }
Label $\rightarrow$ intlit :	{intlit}
Label $\rightarrow$ λ	{if, id, read}

This grammar generates the same statements as our original grammar did, but now prediction never fails!

Whenever we must decide what production to use, the predict sets for productions with the same lefthand side are always disjoint.

Any input token will predict a unique production or no production at all (indicating a syntax error).

If we never mispredict a production, we never backup, so parsing will be fast and absolutely accurate!

LL(1) GRAMMARS

A context-free grammar whose Predict sets are always disjoint (for the same non-terminal) is said to be *LL(1)*.

LL(1) grammars are ideally suited for top-down parsing because it is always possible to correctly predict the expansion of any non-terminal. No backup is ever needed.

Formally, let

$\text{First}(X_1 \dots X_n) =$

$\{a \text{ in } V_t \mid A \rightarrow X_1 \dots X_n \Rightarrow^* a \dots\}$

$\text{Follow}(A) = \{a \text{ in } V_t \mid S \Rightarrow^+ \dots A a \dots\}$

$\text{Predict}(A \rightarrow X_1 \dots X_n) =$

If $X_1 \dots X_n \Rightarrow^* \lambda$

Then $\text{First}(X_1 \dots X_n) \cup \text{Follow}(A)$

Else $\text{First}(X_1 \dots X_n)$

If some CFG, G, has the property that for all pairs of distinct productions with the same lefthand side,

$A \rightarrow X_1 \dots X_n$ and $A \rightarrow Y_1 \dots Y_m$

it is the case that

$\text{Predict}(A \rightarrow X_1 \dots X_n) \cap$

$\text{Predict}(A \rightarrow Y_1 \dots Y_m) = \phi$

then G is LL(1).

LL(1) grammars are easy to parse in a top-down manner since predictions are always correct.

Example

Production	Predict Set
$S \rightarrow A a$	$\{b, d, a\}$
$A \rightarrow B D$	$\{b, d, a\}$
$B \rightarrow b$	$\{b\}$
$B \rightarrow \lambda$	$\{d, a\}$
$D \rightarrow d$	$\{d\}$
$D \rightarrow \lambda$	$\{a\}$

Since the predict sets of both B productions and both D productions are disjoint, this grammar is LL(1).

RECURSIVE DESCENT PARSERS

An early implementation of top-down (LL(1)) parsing was recursive descent.

A parser was organized as a set of *parsing procedures*, one for each non-terminal. Each parsing procedure was responsible for parsing a sequence of tokens derivable from its non-terminal.

For example, a parsing procedure, A, when called, would call the scanner and match a token sequence derivable from A.

Starting with the start symbol's parsing procedure, we would then match the entire input, which must be derivable from the start symbol.

This approach is called recursive descent because the parsing procedures were typically *recursive*, and they *descended* down the input's parse tree (as top-down parsers always do).

Building A RECURSIVE DESCENT PARSER

We start with a procedure **Match**, that matches the current input token against a predicted token:

```
void Match(Terminal a) {  
    if (a == currentToken)  
        currentToken = Scanner();  
    else SyntaxError();  
}
```

To build a parsing procedure for a non-terminal A, we look at all productions with A on the lefthand side:

$$A \rightarrow X_1 \dots X_n \mid A \rightarrow Y_1 \dots Y_m \mid \dots$$

We use predict sets to decide which production to match (LL(1) grammars always have disjoint predict sets).

We match a production's righthand side by calling **Match** to

match terminals, and calling parsing procedures to match non-terminals.

The general form of a parsing procedure for

$A \rightarrow X_1 \dots X_n \mid A \rightarrow Y_1 \dots Y_m \mid \dots$ is

```
void A() {
  if (currentToken in Predict(A→X1...Xn))
    for(i=1;i<=n;i++)
      if (X[i] is a terminal)
        Match(X[i]);
      else X[i]();
  else
    if (currentToken in Predict(A→Y1...Ym))
      for(i=1;i<=m;i++)
        if (Y[i] is a terminal)
          Match(Y[i]);
        else Y[i]();
  else
    // Handle other A → ... productions
  else // No production predicted
    SyntaxError();
}
```

Usually this general form isn't used.

Instead, each production is "macro-expanded" into a sequence of **Match** and parsing procedure calls.

Example: CSX-Lite

Production	Predict Set
$\text{Prog} \rightarrow \{ \text{Stmts} \} \text{Eof}$	{
$\text{Stmts} \rightarrow \text{Stmt Stmts}$	id if
$\text{Stmts} \rightarrow \lambda$	}
$\text{Stmt} \rightarrow \text{id} = \text{Expr} ;$	id
$\text{Stmt} \rightarrow \text{if} (\text{Expr}) \text{Stmt}$	if
$\text{Expr} \rightarrow \text{id Etail}$	id
$\text{Etail} \rightarrow + \text{Expr}$	+
$\text{Etail} \rightarrow - \text{Expr}$	-
$\text{Etail} \rightarrow \lambda$	) ;

CSX-Lite Parsing Procedures

```
void Prog() {
  Match("{");
  Stmts();
  Match("}");
  Match(Eof);
}

void Stmts() {
  if (currentToken == id ||
      currentToken == if){
    Stmt();
    Stmts();
  } else {
    /* null */
  }
}

void Stmt() {
  if (currentToken == id){
    Match(id);
    Match("=");
    Expr();
    Match(";");
  } else {
    Match(if);
    Match("(");
    Expr();
    Match(")");
    Stmt();
  }
}
```

```

void Expr() {
    Match(id);
    Etail();
}

void Etail() {
    if (currentToken == "+") {
        Match("+");
        Expr();
    } else if (currentToken == "-") {
        Match("-");
        Expr();
    } else {
        /* null */
    }
}

```

Let's use recursive descent to parse
{ a = b + c; } Eof
 We start by calling **Prog()** since this
 represents the start symbol.

Calls Pending	Remaining Input
Prog()	{ a = b + c; } Eof
Match("{"); Stmts(); Match("}"); Match(Eof);	{ a = b + c; } Eof
Stmts(); Match("{"); Match(Eof);	a = b + c; } Eof
Stmt(); Stmts(); Match("}"); Match(Eof);	a = b + c; } Eof
Match(id); Match("="); Expr(); Match(";"); Stmts(); Match("}"); Match(Eof);	a = b + c; } Eof

Calls Pending	Remaining Input
Match("="); Expr(); Match(";"); Stmts(); Match("}"); Match(Eof);	= b + c; } Eof
Expr(); Match(";"); Stmts(); Match("}"); Match(Eof);	b + c; } Eof
Match(id); Etail(); Match(";"); Stmts(); Match("}"); Match(Eof);	b + c; } Eof
Etail(); Match(";"); Stmts(); Match("}"); Match(Eof);	+ c; } Eof

Calls Pending	Remaining Input
Match("+"); Expr(); Match(";"); Stmts(); Match("}"); Match(Eof);	+ c; } Eof
Expr(); Match(";"); Stmts(); Match("}"); Match(Eof);	c; } Eof
Match(id); Etail(); Match(";"); Stmts(); Match("}"); Match(Eof);	c; } Eof
Etail(); Match(";"); Stmts(); Match("}"); Match(Eof);	; } Eof
/* null */ Match(";"); Stmts(); Match("}"); Match(Eof);	; } Eof

Calls Pending	Remaining Input
<code>Match(";");</code> <code>Stmts();</code> <code>Match(")");</code> <code>Match(Eof);</code>	<code>; } Eof</code>
<code>Stmts();</code> <code>Match(")");</code> <code>Match(Eof);</code>	<code>} Eof</code>
<code>/* null */</code> <code>Match(")");</code> <code>Match(Eof);</code>	<code>} Eof</code>
<code>Match(")");</code> <code>Match(Eof);</code>	<code>} Eof</code>
<code>Match(Eof);</code>	<code>Eof</code>
Done!	All input matched

SYNTAX ERRORS IN RECURSIVE DESCENT PARSING

In recursive descent parsing, syntax errors are automatically detected. In fact, they are detected *as soon as possible* (as soon as the first illegal token is seen).

How? When an illegal token is seen by the parser, either it fails to predict any valid production or it fails to match an expected token in a call to **Match**.

Let's see how the following illegal CSX-lite program is parsed:

```
{ b + c = a; } Eof
```

(Where should the first syntax error be detected?)

Calls Pending	Remaining Input
<code>Prog()</code>	<code>{ b + c = a; } Eof</code>
<code>Match("{");</code> <code>Stmts();</code> <code>Match(")");</code> <code>Match(Eof);</code>	<code>{ b + c = a; } Eof</code>
<code>Stmts();</code> <code>Match(")");</code> <code>Match(Eof);</code>	<code>b + c = a; } Eof</code>
<code>Stmt();</code> <code>Stmts();</code> <code>Match(")");</code> <code>Match(Eof);</code>	<code>b + c = a; } Eof</code>
<code>Match(id);</code> <code>Match("=");</code> <code>Expr();</code> <code>Match(";");</code> <code>Stmts();</code> <code>Match(")");</code> <code>Match(Eof);</code>	<code>b + c = a; } Eof</code>

Calls Pending	Remaining Input
<code>Match("=");</code> <code>Expr();</code> <code>Match(";");</code> <code>Stmts();</code> <code>Match(")");</code> <code>Match(Eof);</code>	<code>+ c = a; } Eof</code>
Call to Match fails!	<code>+ c = a; } Eof</code>

Table-Driven Top-Down PARSERS

Recursive descent parsers have many attractive features. They are actual pieces of code that can be read by programmers and extended.

This makes it fairly easy to understand how parsing is done.

Parsing procedures are also convenient places to add code to build ASTs, or to do type-checking, or to generate code.

A major drawback of recursive descent is that it is quite inconvenient to change the grammar being parsed. Any change, even a minor one, may force parsing procedures to be

reprogrammed, as productions and predict sets are modified.

To a less extent, recursive descent parsing is less efficient than it might be, since subprograms are called just to match a single token or to recognize a righthand side.

An alternative to parsing procedures is to encode all prediction in a parsing table. A pre-programed driver program can use a parse table (and list of productions) to parse any LL(1) grammar.

If a grammar is changed, the parse table and list of productions will change, but the driver need not be changed.

LL(1) PARSE TABLES

An LL(1) parse table, T , is a two-dimensional array. Entries in T are production numbers or blank (error) entries.

T is indexed by:

- A , a non-terminal. A is the non-terminal we want to expand.
- CT , the current token that is to be matched.
- $T[A][CT] = A \rightarrow X_1 \dots X_n$
if CT is in $\text{Predict}(A \rightarrow X_1 \dots X_n)$
 $T[A][CT] = \text{error}$
if CT predicts no production with A as its lefthand side

CSX-lite Example

	Production	Predict Set
1	$\text{Prog} \rightarrow \{ \text{Stmts} \} \text{Eof}$	{
2	$\text{Stmts} \rightarrow \text{Stmt Stmts}$	id if
3	$\text{Stmts} \rightarrow \lambda$	}
4	$\text{Stmt} \rightarrow \text{id} = \text{Expr} ;$	id
5	$\text{Stmt} \rightarrow \text{if} (\text{Expr}) \text{Stmt}$	if
6	$\text{Expr} \rightarrow \text{id Etail}$	id
7	$\text{Etail} \rightarrow + \text{Expr}$	+
8	$\text{Etail} \rightarrow - \text{Expr}$	-
9	$\text{Etail} \rightarrow \lambda$	) ;

	{	}	if	(	)	id	=	+	-	;	eof
Prog	1										
Stmts		3	2			2					
Stmt			5			4					
Expr						6					
Etail					9			7	8	9	