

4th Unit

Intermediate Code Generation:

Variants of Syntax tree

✓ Three address code

Types & Declarations

Translation of Expressions

✓ type checking

✓ Control Flow

Storage Organization

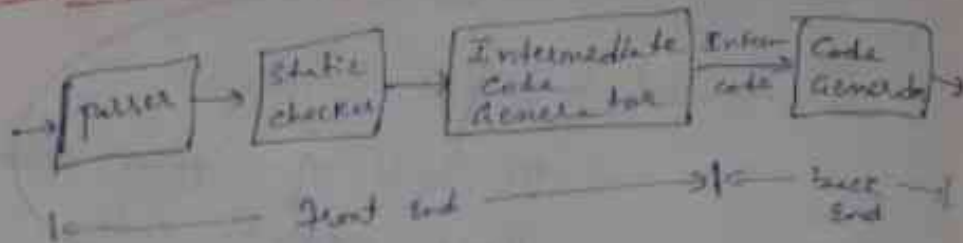
✓ Stack Allocation of space

Access to Non-local data on
the Stack

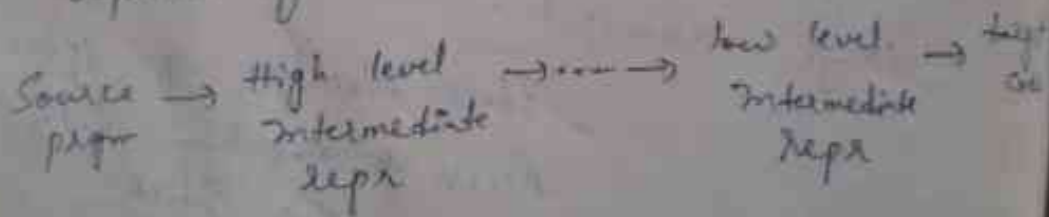
Heap Management

Garbage Collection

Intermediate Code Generation



→ In the process of translating a program in a given source lang into code for a given target machine, a compiler may construct a sequence of intermediate repr.



→ High-level repr are close to the source lang & low-level repr are close to the target machine.

Variants of Syntax trees

{ DAG for Expr
value - no method for Construct
DAGs

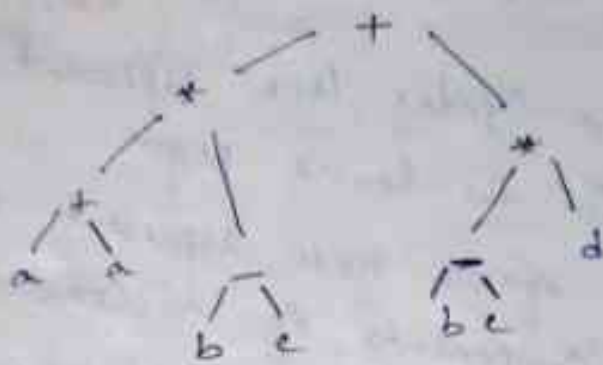
- Nodes in a Syntax tree represent constructs in the source program
- the children of a node represents the meaningful components of a construct.
- The DAG (Directed Acyclic Graph) for an Expr identifies the Common Sub-Expr of the Expr.

DAG for Expression

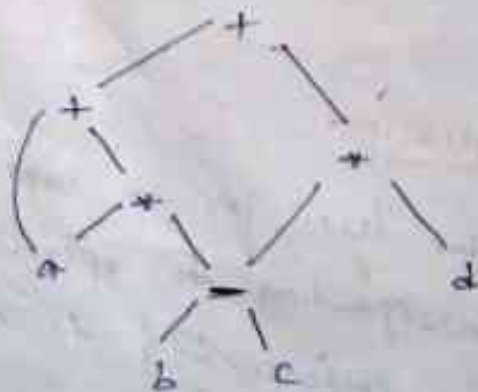
- Like the Syntax Tree for an Expr, A DAG has leaves corresponding to operands & interior nodes corresponding to operators.
- the difference is that a node N in a DAG has more than one parent if N represents a Common SubExpression.
- in a Syntax tree, the tree for the Common SubExpression would be Replicated as many times as the sub-Expr appears in the original Expression.

eg- $a + a * (b - c) + (b - c) * d$

Draw DAG & Syntax tree



SYNTAX tree



DAG

→ SDD to produce Syntax trees or DAG's
Syntax-Directed Definition

SDD

Production	Semantic Rules
① $E \rightarrow E_1 + T$	$E.\text{node} = \text{new Node}('+', E_1.\text{node}, T.\text{node})$
② $E \rightarrow E_1 - T$	$E.\text{node} = \text{new Node}('-', E_1.\text{node}, T.\text{node})$
③ $E \rightarrow T$	$E.\text{node} = T.\text{node}$
④ $T \rightarrow (E)$	$T.\text{node} = E.\text{node}$
⑤ $T \rightarrow \text{id}$	$T.\text{node} = \text{new leaf}(\text{id}, \text{id}.\text{entry})$
⑥ $T \rightarrow \text{num}$	$T.\text{node} = \text{new leaf}(\text{num}, \text{num.val})$

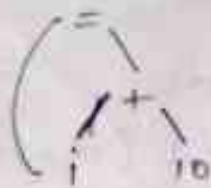
→ steps for constructing ~~the~~ syntax tree

1. $P_1 = \text{leaf}(\text{id}, \text{entry}-a) \rightarrow P_1$
2. $P_2 = \text{leaf}(\text{id}, \text{entry}-a) \rightarrow P_1$
3. $P_3 = \text{leaf}(\text{id}, \text{entry}-b)$
4. $P_4 = \text{leaf}(\text{id}, \text{entry}-c)$
5. $P_5 = \text{Node}('-', P_3, P_4) \rightarrow \{b - c\}$
6. $P_6 = \text{Node}('*', P_1, P_5) \rightarrow \{a * (b - c)\}$
7. $P_7 = \text{Node}('+', P_2, P_6) \rightarrow \{a + a * (b - c)\}$
8. $P_8 = \text{leaf}(\text{id}, \text{entry}-b) \rightarrow P_3$
9. $P_9 = \text{leaf}(\text{id}, \text{entry}-c) \rightarrow P_4$
10. $P_{10} = \text{Node}('-', P_8, P_9) \rightarrow P_5$
11. $P_{11} = \text{leaf}(\text{id}, \text{entry}-d)$
12. $P_{12} = \text{Node}('*', P_{10}, P_{11}) \rightarrow \{(b - c) * d\}$
13. $P_{13} = \text{Node}('+', P_7, P_{12}) \rightarrow \{a + a * (b - c) + (b - c) * d\}$

The Value-no. Method for Constructing DAG

→ the nodes of a Syntax tree or DAG are stored in an array of records.

eg:- $i = i + 10$



DAG

1	id			→ Entry for i
2	nom	10		
3	+	1	2	
4	=	1	3	

Array

→ Each row of the array represents one record, & therefore one node.

→ In each record, the first field is an operation code, indicating the label of the node.

→ In array, leaves have one additional field, which holds the lexical value (either a symbol-table pointer or a constant) and

interior nodes have two additional fields indicating the left & right children.

Three - Address Code

Addresses & Instructions
quadruples ✓
triples ✓
indirect triples ✓
SSA

→ In three-address code, there is at most one operator on the right side of an instruction.

$$\rightarrow x + y * z$$

$$t_1 = y * z$$

$$t_2 = x + t_1$$

Ex:- $a + a * (b - c) + (b - c) * d$

draw DAG of three-address code



DAG

$$t_1 = b - c$$

$$t_2 = a + t_1$$

$$t_3 = a + t_2$$

$$t_4 = t_1 * d$$

$$t_5 = t_3 + t_4$$

Three-address code

Addresses & Instructions

→ Three-address code is built from two concepts:

Addresses

Instructions

→ 3-address code can be implemented using records with fields for the addresses.

Records called

Quadruples

triples &

Indirect - triples

→ An address can be one of the following

A Name

A Constant

A Compiler-generated temporary

→ A Name :-

For convenience, we allow source-
program names to appear as address in
three-address code.

→ Constant :- Compiler must deal with many
diff types of constants & variables.

→ Compiler-generated temporary :-

It is useful, especially in
optimizing compiler, to create a distinct
name each time a temporary is needed.

COMMON THREE-ADDRESS INSTRUCTIONS FORMS:

1. Assignment Instruction

$x = y \text{ op } z$

2. Copy instructions of the form

$$x = y$$

3. An Unconditional jump goto L

4. Conditional jumps of the form if x goto L

5. Indexed copy inst of the form

$$x = y[i] \quad ;$$

$$x[i] = y$$

6. Address & pointer assignments of the form

$$x = \&y,$$

$$x = *y$$

$$*x = y.$$

eg:-

do $i = i + 1$;

while ($a[i] < v$);

Show types of labels for 3-adde struct.

L : $t_1 = i + 1$

$i = t_1$

$t_2 = i + 8$

$t_3 = a[t_2]$

while $t_3 < v$ goto L

100: $t_1 = i + 1$

101: $i = t_1$

102: $t_2 = i + 8$

103: $t_3 = a[t_2]$

104: while $t_3 < v$
goto 100

SYMBOLIC LABEL

POSITION NUMBERS
OR
NUMERIC LABEL

QUADRUPLES

→ In a Compiler, Instructions can be implemented as objects or as records with fields for the operator and the operands.

→ A Quadruple (Quad) has four fields, which are op, arg1, arg2, result

Eg:- $a = b * -c + b * -c$
Write 3-address code & Quadruples

$t_1 = \text{minus } c$

$t_2 = b * t_1$

$t_3 = \text{minus } c$

$t_4 = b * t_3$

$t_5 = t_2 + t_4$

$a = t_5$

3-address code

op	arg1	arg2	result
minus	c		t_1
*	b	t_1	t_2
minus	c		t_3
*	b	t_3	t_4
+	t_2	t_4	t_5
=	t_5		a

Quadruples

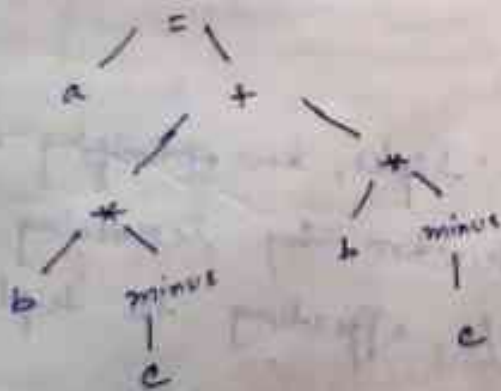
Triples

→ A triple has only three fields
op
arg1
arg2

→ using triples, we refer to the result of an operation x op y by its position, rather than by an explicit temporary name.

→ A triple repr would refer to position(o).

Ex - $a = b * -c + b * -c$
Draw syntax tree & Triples



SYNTAX TREE

	op	arg1	arg2
0	minus	c	
1	*	b	(0)
2	minus	c	
3	*	b	(2)
4	+	(1)	(3)
5	=	a	(4)

Triples

→ Disadvantage

With triples, the result of an operation is referred to by its position. So moving an instruction may require us to change all references to that result.

INDIRECT TRIPLES

→ It consist of a listing of pointers to triples, rather than a listing of triples themselves.

$$a = b * - c + b * - c$$

Instruction list	
55	(0)
56	(1)
57	(2)
58	(3)
59	(4)
60	(5)

	op	arg1	arg2
0	minus	c	
1	*	b	(0)
2	minus	c	
3	*	b	(2)
4	+	(1)	(3)
5	=	a	(4)

→ Advantage

With indirect triples, an optimizing compiler can move an instruction by reordering the instruction list, w/o affecting the triples themselves.

⇒ STATIC SINGLE-ASSIGNMENT FORM

→ Static Single-Assignment form (SSA) is an intermediate repⁿ that facilitates certain code optimizations.

→ two distinctive aspects distinguish SSA from 3-address code.

→ the first is that all assignments in SSA are to variables with distinct names.

→ subscripts distinguish each definition of variables.

$((a+b)-c) * d$

$p = a + b$

$q = p - c$

$r = q * d$

3-address code

$p_1 = a + b$

$q_1 = p_1 - c$

$r = q_1 * d$

SSA form

⇒ Types & Declarations { type expr
type equivalence
declarations
storage layout for local names

The applications of types can be grouped under checking & translation.

Type checking uses logical rules to reason about the behavior of a program at run time. Especially, it ensures that the types of the operands match the type expected by an operator.

Translation Applications :-

From the type of a name, a compiler can determine the storage that will be needed for that name at run-time.

type info is also needed to calculate the address denoted by an array reference, to insert explicit type conversion

⇒ Type Expressions

→ types have structure, which we shall represent using Type Expressions:

a type expr is either a basic type or is formed by applying an operator called a Type Constructor to a type expr.

The sets of basic types and constructors depend on the lang to be checked.

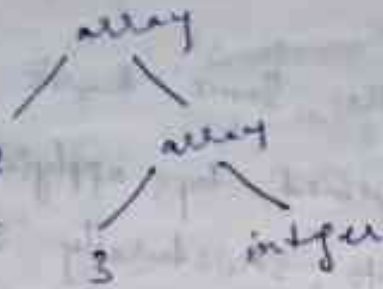
eg:-

`int [2][3]`

and as "array of 2 arrays of 3 integers each"

written as a type expr

array(2, array(3, integer)).



→ A basic type is a type expr.
boolean, char, integer, float & void

→ A type name is a type expr.

→ A type expr can be formed by using the

type constructor → for for types.

we write $s \rightarrow t$ for
for from type s to type t .

⇒ Type Equivalence

→ when are two type expr equivalent?

→ the key issue is whether a name in
a type expr stands for itself or whether
it is an abbreviation for another type
expr.

→ When type Expr are represented by graphs, two types are structurally equivalent if and only if one of the following condition is true:

- ① they are the same basic type.
- ② they are formed by applying the same constructor to structurally equivalent types
- ③ one is a type name that denotes the other

⇒ Declarations

→ types & Declarations using a simplified grammar that declares just one name at a time; declarations with list of names can be handled as.

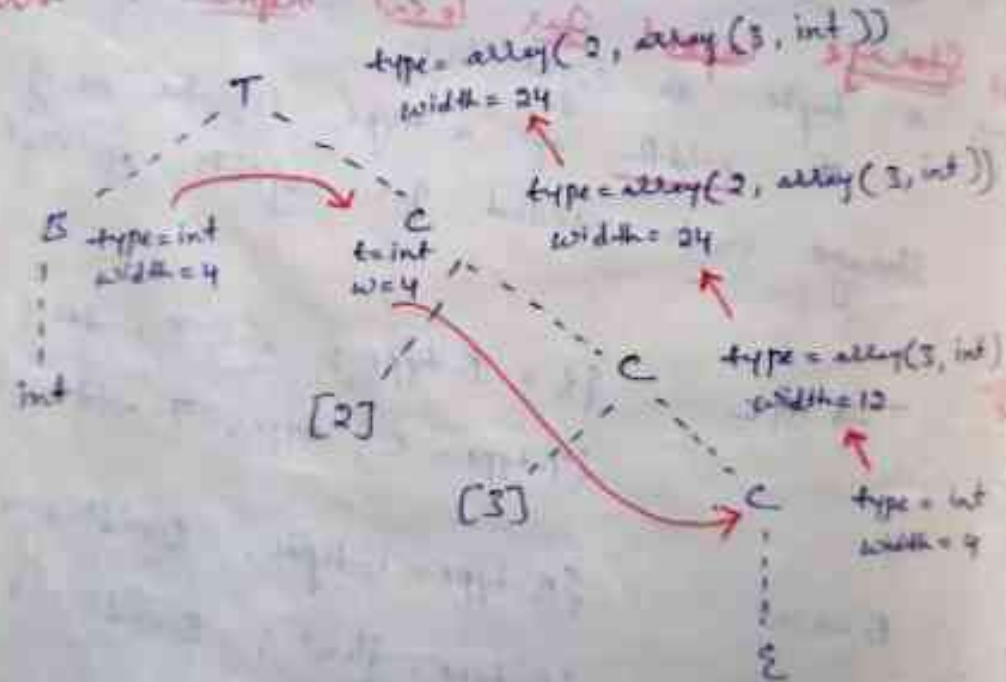
the grammar is

$D \rightarrow T \text{ id}; \text{ D/2}$ Existential
 $T \rightarrow B \text{ c} \mid \text{record } \{ D \}$ Type
 $B \rightarrow \text{int} \mid \text{float}$
 $C \rightarrow \text{z} \mid [\text{num}] \text{ c}$

Non-terminal D generates a sequence of declarations.

- The translation Scheme (SDT) computes types & their widths for basic & Array types;
- SDT uses synthesized attributes TYPE & WIDTH for each nonterminal & two var t & w to pass type & width info.

parse tree for type int[2][3]



SDT of array types

attributes passed & computed postorder

Translation of Expressions

operations within Expr

Incremental translation

Addressing array elements

translation of array references

⇒ Operations within Expressions

→ There - add code for an Alignment start

$\{$ S using attribute code for S
attribute addr, code for an

$\{$ Expr E.

S-add code for Expr

Semantic Rules

Production

$S \rightarrow id = E;$

$S.code = E.code \parallel$

$gen(top.get(id.lexeme) '=' E.addr)$

$E \rightarrow E_1 + E_2$

$E.addr = new Temp()$

$E.code = E_1.code \parallel E_2.code \parallel$

$gen(E.addr '=' E_1.addr '+' E_2.addr)$

$E \rightarrow E_1 - E_2$

$E.addr = new Temp()$

$E.code = E_1.code \parallel$

$gen(E.addr '=' E_1.addr '-' E_2.addr)$

$| (E_1)$

$E.addr = E_1.addr$

$E.code = E_1.code$

$| id$

$E.addr = top.get(id.lexeme)$

$E.code =$

⇒ Incremental Translation

→ the translation scheme generates the same code as the syntax directed def

→ with the incremental approach, the **Code** attribute is not used, since there is a single sequence of inst that is created by successive calls to **gen**

$S \rightarrow id = E;$ $\{ \text{gen}(\text{top.get}(id.\text{lexeme})) = E.\text{addr} \}$

$E \rightarrow E_1 + E_2$ $\{ E.\text{addr} = \text{new Temp}();$
 $\text{gen}(E.\text{addr} = E_1.\text{addr} + E_2.\text{addr}) \}$

$| - E_1$ $\{ E.\text{addr} = \text{new Temp}();$
 $\text{gen}(E.\text{addr} = \text{'minus'} E_1.\text{addr}) \}$

$| (E_1)$ $\{ E.\text{addr} = E_1.\text{addr}; \}$

$| id$ $\{ E.\text{addr} = \text{top.get}(id.\text{lexeme}); \}$

E-addr code for expr incrementally

→ Addressing Array Elements

→ Array elements can be accessed quickly if they are stored in a block of consecutive locations.

→ If the width of each array element is w , then the i th element of array A begins in location $3D$

$$\text{base} + i \times w$$

where base = relative addr of the storage allocated for the array.

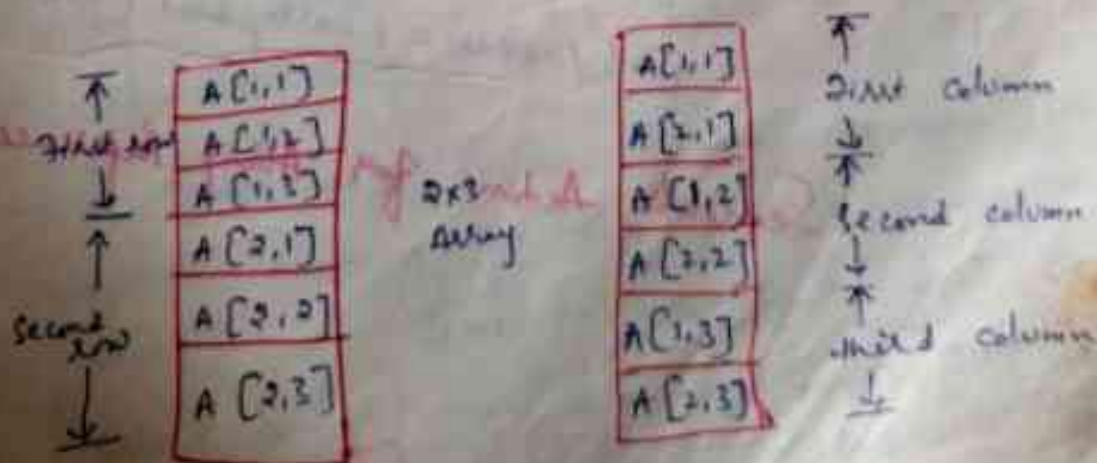
$2D$

$$\text{base} + i_1 \times w_1 + i_2 \times w_2$$

K Dimensions

$$\text{base} + i_1 \times w_1 + i_2 \times w_2 + \dots + i_K \times w_K$$

→ A $2D$ array is normally stored in one of the two forms, either **ROW-MAJOR** (row-by-row) or **COLUMN-MAJOR** (column-by-column)



ROW MAJOR

COLUMN MAJOR

⇒ Translation of Array References

→ the chief problem in generating code for array references is to relate the address calculation formulas to a grammar for array references.

→ L generate an array name followed by a sequence of index expressions.

$L \rightarrow L[E] \mid \text{ID} \mid [E]$

L = array

→ 3-addr code for Expr with array reference

$S \rightarrow \text{id} = E;$ $\text{gen}(\text{top.get}(\text{id.lexeme}) = E.\text{addr});$

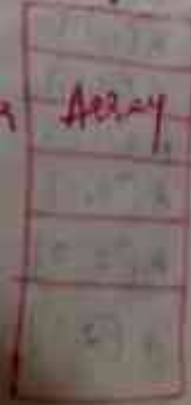
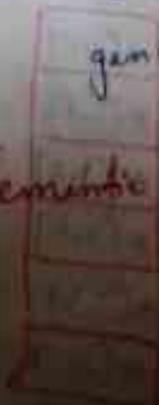
$L = E;$ $\text{gen}(L.\text{array.base}[L.\text{addr}] = E.\text{addr});$

$E \rightarrow E_1 + E_2$ $E.\text{addr} = \text{new Temp}();$
 $\text{gen}(E.\text{addr} = E_1.\text{addr} + E_2.\text{addr});$

id $E.\text{addr} = \text{top.get}(\text{id.lexeme});$

ID $E.\text{addr} = \text{new Temp}();$
 $\text{gen}(E.\text{addr} = L.\text{array.base}[L.\text{addr}]);$

Semantic Actions for Array References



TYPE CHECKING

Rules for Type checking

Type conversions

overloading of fns & operators

type inference & polymorphic fns

Type checking

- to do type checking a compiler needs to assign a type expr to each component of the source program.
- the compiler must then determine that these type expr conform to a collection of logical rules that is called TYPE SYSTEM for the source language.
- Type checking has the potential for catching errors in programs.
- In principle, any check can be done dynamically if the target code carries the type of an element along with the value of the element.
- A SOUND type system eliminates the need for dynamic checking for type errors, because it allows us to determine statistically that these errors cannot occur when the target program runs.

→ An implementation of a language is STRONGLY TYPED if a compiler guarantees that the program it accepts will run w/o type errors.

→ Type checking have been used to improve the Security of systems that allow software modules to be imported & executed.

⇒ Rules for Type checking

→ Type checking can take on two forms:
SYNTHESIS &
INFERENCE

→ TYPE SYNTHESIS builds up the type of an expr from the types of its subexpressions. It requires names to be declared before they are used.

Ex - the type of $E_1 + E_2$
before using this expr,
first declare exprs E_1, E_2

→ TYPE INFERENCE determines the type of a language construct from the way it is used. type inference is needed for languages

like ML, which check types, but do not require names to be declared.

⇒ Type Conversions

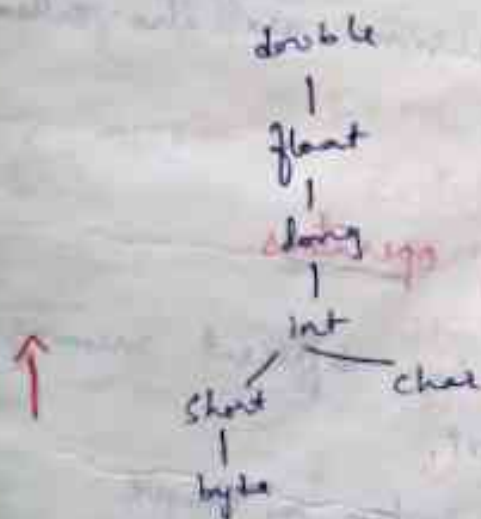
→ Type Conversion rules vary from lang to lang.

→ In Java, there are two types of conversions

WIDENING Conversions

NARROWING Conversions

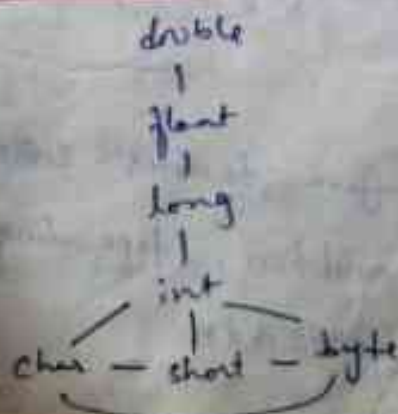
→ WIDENING Conversions



- Intended to preserve information

- Conversion by Compiler

→ NARROWING Conversions



- which can lose information

- Conversion by programmer

→ Conversion from one type to another is said to be IMPLICIT if it is done automatically by the compiler.

IMPLICIT type conversions, also called as COERCIONS

coercions, are limited in many languages to widening conversions.

→ Conversion is said to be EXPLICIT if the programmer must write something to cause the conversion.

Explicit conversions are also called CASTS / TYPE CASTING.

⇒ overloading of fn & operators

→ An overloaded symbol has different meanings depending on its context.

→ overloading is resolved when a unique meaning is determined for each occurrence of a name.

Eg:- + operator in Java denotes either string concatenation or addition, depending on the types of its operands.

⇒ Type Inference & Polymorphic fns

→ Type Inference is useful for a language like ML, which is strongly typed, but doesn't require names to be declared before they are used.

type Inference ensures that names are used consistently.

→ The term polymorphic refers to any code fragment that can be executed with arguments of different types.

CONTROL

FLOW ✓

Boolean Expressions

Short-circuit code

Flow of control statements

Avoiding Redundant Codes

Boolean values & jumping code

⇒ Control Flow tools

The translation of statements such as

if-else statement

while statement

is tied to the translation of

boolean expressions.

In programming lang, boolean expressions are often used to

1. Alter the flow of control:

Boolean Expr are used as Conditional Expr in stmts that alter the flow of control.

The value of such boolean Expr is implicit in a position reached in a program.

2. Compute logical values:

A boolean Expr can represent T or F as values.

⇒ Boolean Expressions

→ Boolean Expr are composed of the boolean operators ($\&\&$, $\|\$, $!$) (AND, OR, NOT)

applied to elements that are boolean variable or relational Expr.

→ using the attribute rel.op to indicate which of the six comparison operators $<$, $<=$, $=$, $!=$, $>$, $>=$ is represented by rel.

⇒ Short-Circuit Code

⇒ In SHORT-CIRCUIT (or JUMPING) code, the boolean operators `&&`, `||`, `!` translate into jumps.

→ The operators themselves do not appear in the code; instead the value of a boolean expr is represented by a position in the code sequence.

Eg:- The stmt

`if (x < 100 || x > 200 && x != y)`

might be translated into the code

`if x < 100 goto L2`
`if false x > 200 goto L1`
`if false x != y goto L1`

`L2: x = 0`

`L1:`

JUMPING CODE

⇒ FLOW-OF-CONTROL Statements

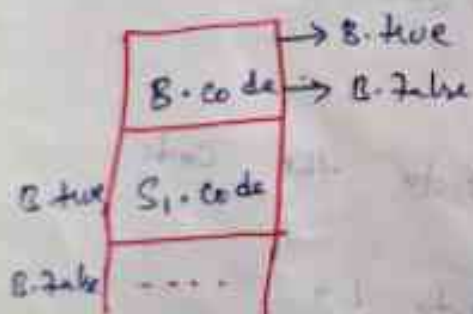
→ The translation of boolean expr into 3-addr code in the context of stmts such as those generated by the following grammar;

$S \rightarrow \text{if } (B) S_1$

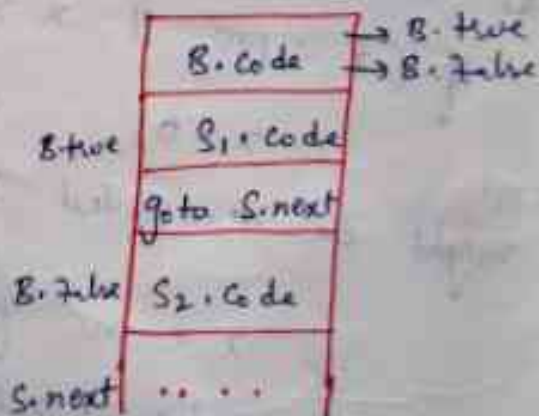
$S \rightarrow \text{if } (B) S_1 \text{ else } S_2$

$S \rightarrow \text{while } (B) S_1$

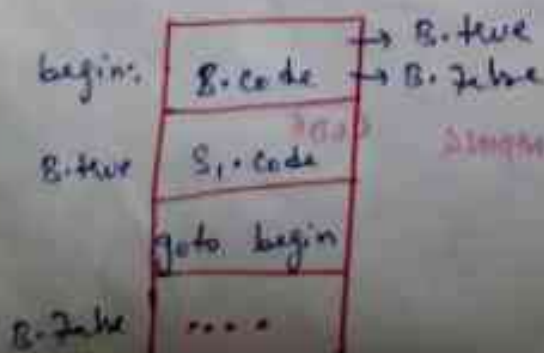
The translation



if



if-else



while

⇒ Avoiding Redundant Gotos

if $x > 200$ goto L_1

L_1 : goto L_4

Instead, consider this:

if $x > 200$ goto L_4

⇒ Boolean values, Jumping Code

1. Use two passes :-

Construct a complete Syntax tree for the g/p, and then walk the tree in depth-first order, computing the translation specified by the semantic rules.

2. Use one pass for stmt,
but two passes for expr :-

RUN-TIME ENVIRONMENTS

Storage Organization / memory organization

Stack Allocation of Space

Access to Non-local Data on the Stack

Heap Management

Garbage Collection

Parameter Passing
1. Memory organization

1. STORAGE ORGANIZATION

Static Vs Dynamic Storage Allocation

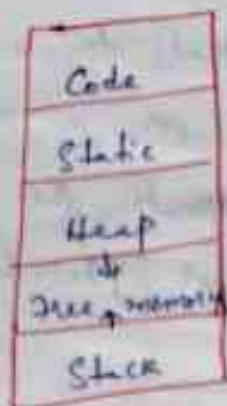
→ the Executing target program runs in its own logical address space in which each program value has a location.

→ The management & organization of this logical address space is shared b/w the Compiler, OS & target machine.

→ The OS maps the logical addresses into physical addresses, which are usually spread throughout memory.

→ The Run-time representation of an object program in the logical address space consists of data

and program area as



⇒ Static vs Dynamic Storage Allocation

→ Storage-allocation decision is STATIC, if it can be made by the compiler looking only at the text of the program, not at what the program does when it executes.

→ A decision is DYNAMIC if it can be decided only while the program is running.

→ Many compilers use some combination of the following two strategies for Dynamic Storage Allocation:

✓ Allocation:

Stack Storage

Heap Storage

1. STACK Storage

Names local to a procedure are allocated space on a stack. Stack supports the normal call/return policy for procedures.

2. HEAP Storage

Data that may outlive the call to the procedure that created it is usually allocated on a Heap of reusable storage.

Heap is an area of virtual memory that allows objects or other data elements to obtain storage when they are created and to return that storage when they are invalidated.

3. STACK ALLOCATION OF SPACE

Activation Records

Calling Sequences

Variable-length data on stack

→ Each time a procedure is called, space for its local variables is pushed onto a stack, and when the procedure terminates, that space is popped off the stack.

→ Activation Records

→ procedure calls & return are usually managed by a run-time stack called the CONTROL STACK.

→ Each live Activation has an ACTIVATION RECORD (frames) on the control stack, with the root of the activation tree at the bottom, and the entire sequence of activation records on the stack corresponding to the paths in the activation tree to the activation where control currently resides.

→ the latter activation has its record at the top of the stack.

→ The contents of Activation records vary with the lang being implemented.

Actual Parameters
Returned Values
Control Link
Access Link
Saved machine Status
Local data
Temporaries

Activation Record

1. Temporary values, such as those arising from the evaluation of expr.
2. Local data belonging to the procedure whose activation record this is.
3. Saved Machine Status, with inf about the state of the machine just before the call to the procedure.
inf like return addr, contents of registers.

4. Access link may be needed to locate data needed by the called procedure but found elsewhere, in another activation record.

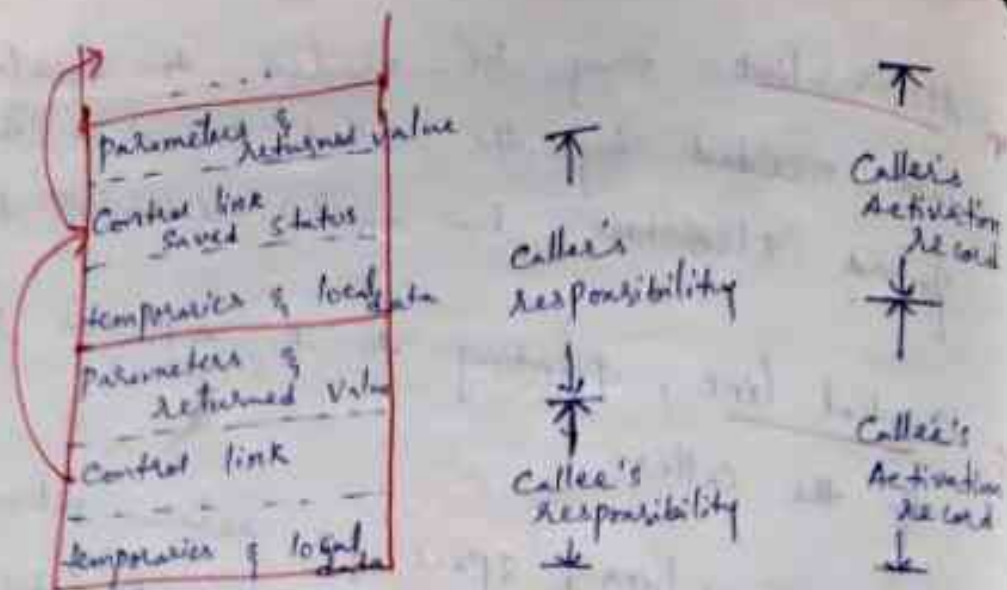
5. Control link, pointing to the activation record of the caller.

6. Returned values, space for return values if any. Again, not all called procedures return a value, and if one does, we may prefer to place that value in a register for efficiency.

7. Actual parameters used by the calling procedure. Commonly, these values are not placed in the activation record but rather in registers.

⇒ Calling Sequences are implemented by what are known as CALLING SEQUENCES, which consists of code that allocates an activation record on the stack and enters info into its fields.

⇒ Return Sequence is similar code to restore the state of the machine so the calling procedure can continue its execution after the call.



Division of tasks b/w caller & callee

⇒ Variable-length Data on the Stack

→ A common strategy for allocating variable-length arrays (i.e., arrays whose size depends on the value of one or more parameters of the called procedure)

Access To Non-Local DATA ON THE STACK

topic covered in 3rd Unit.

4. HEAP MANAGEMENT / Dynamic Memory

The Memory Manager

Memory Hierarchy of a Computer

Locality in programs

Reducing Fragmentation

Manual Reallocation Requests

HEAP Management

→ The heap is the portion of the store that is used for data that lives indefinitely, or until the program explicitly deletes it.

→ While local variables typically become inaccessible when their procedures end, many languages enable us to create objects or other data whose existence is not tied to the procedure activation that creates them.

→ For eg, both C++ & Java give the programmer new to create objects that may be passed, or pointers to them may be passed, from procedure to procedure, so they continue to exist long after the procedure that created them is gone.

→ Such objects are stored on a heap.

→ MEMORY MANAGER

→ The memory manager keeps track of all the free space in heap storage at all times.

→ It performs two basic fns:
Allocation
Deallocation

- Allocation:- When a program requests memory for a variable or object, the memory manager produces a chunk of continuous heap memory of the requested size.
- Deallocation:- The Memory manager returns deallocated space to the pool of free space. So it can reuse the space to satisfy other allocation requests.

PROPERTIES of Memory Managers:-

- Space Efficiency:- A Memory manager should minimize the total heap space needed by a program.

- Program Efficiency:- A Memory manager should make good use of the memory subsystem to allow programs to run faster.

⇒ Memory Hierarchy of a Computer

> 2GB	<u>Virtual Memory</u>	3-15ms
256MB-2GB	<u>Physical Memory</u>	100-150ns
128KB-4MB	<u>2nd level cache</u>	40-60ns
16-64KB	<u>1st level cache</u>	5-10ns
32 words	<u>Registers (processors)</u>	1ns

⇒ Locality In Programs

→ Most programs exhibit a high degree of LOCALITY; that ~~initials, input~~ they spend most of ~~their time~~ executing a relatively small fraction of the code & touching only a small fraction of the data.

→ Program has two types of locality
TEMPORAL locality
SPATIAL locality

→ TEMPORAL locality:-

we say that a program has temporal locality if the memory location it accesses are likely to be accessed again within a short period of time.

→ SPATIAL locality:-

we say that a program has spatial locality if memory locations close to the location accessed are likely also to be accessed within a short period of time.

→ prgrms spend 90% of their time executing 10% of the code.

→ Reducing Fragmentation

→ At the beginning of prgm execution, the heap is one contiguous unit of free space.

→ As the prgm allocates & deallocates memory, this space is broken up into free & used

chunks of memory, and the free chunks need not reside in a contiguous area of the heap.

→ we refer to the free chunks of memory as HOLES.

→ with each allocation request, the memory manager must PLACE the requested chunk of memory into a large-enough hole.

→ unless a hole of exactly the right size is found, we need to split some hole, creating a yet smaller hole.

→ with each deallocation request, the freed chunks of memory are added back to the pool of free space.

→ we COALESCE contiguous holes into larger holes, as the holes can only get smaller otherwise.

→ if we are not careful, the free memory may end up getting fragmented, consisting of large no. of small, noncontiguous holes.

→ It is then possible that no hole is large enough to satisfy a future request, even though there may be sufficient aggregate free space.

Best-fit & Next-fit object placement

Best-fit algorithm tends to spare the large holes to satisfy subsequent, larger requests.

Managing & coalescing free space

There are two data structures that are useful to support coalescing of adjacent free blocks:

- Boundary Tags
- Doubly linked, Embedded Free list

• Boundary Tags:

At both ends, we keep a free/used bit that tells whether or not the block is currently allocated (used) or available (free).

Adjacent to each free/used bit is a count of the total no. of bytes in the chunk.

• Doubly linked, Embedded Free list:-

The free chunks are also linked in a doubly linked list.

The pointers for this list are within the blocks themselves, say adjacent to the boundary tags at either end.

Thus, no additional space is needed for the free list, although its existence does place a lower bound on how small chunks can get; they must accommodate two boundary tags & two pointers, even if the object is a single byte.

⇒ Manual Deallocation Requests

Problems with Manual Deallocation:-
manual memory mgt is error-prone.

The common mistake take two forms:

Failing ever to delete data that cannot be referenced is called a MEMORY LEAK error.

~~Referencing~~ deleted data is a
referencing

DANGLING- POINTER - REFERENCE error.

Programming Conventions & Tools

Tools are

Object Ownership

Reference Counting

Region-based Allocation

OBJECT OWNERSHIP

it is useful when an object's lifetime can be statically reasoned about. The idea is to associate an OWNER with each object at all times.

The owner is a pointer to that object, presumably belonging to some function invocation.

The owner is responsible for either deleting the object or for passing the object to another owner.

However, it doesn't help solve the dangling-pointer-reference problem, because it is possible to follow a non-owning pointer to an object that has been deleted.

Reference Counting

it is useful when an object's lifetime needs to be determined dynamically.

The idea is to associate a count with each dynamically allocated object.

Whenever a reference to the object is created, we increment the reference count.

whenever a reference to the object is created, we increment the reference count;

whenever a reference is removed, we decrement the reference count.

When the count goes to zero, the object can no longer be referenced and can therefore be deleted.

Region-Based Allocation

it is useful for collections of objects whose lifetimes are tied to specific phases in a computation.

When objects are created to be used only within some steps of a computation, we can allocate all such objects in the same region.

Region-based allocation technique has limited applicability.

GARBAGE COLLECTION

Design Goals for Garbage Collectors

Reachability

Reference Counting Garbage Collectors

→ Data that cannot be referenced is generally known as GARBAGE.

→ Many high-level programming languages remove the burden of manual memory mgt from the programmer by offering automatic garbage collection, which deallocates unreachable data.

⇒ Design Goals

→ Garbage Collection is the reclamation of chunks of storage holding objects that can no longer be accessed by a program.

→ we need to assume that objects have a type that can be determined by the garbage collector at run time.

→ From the type information, we can tell how large the object is and which components of the object contain references to other objects.

→ A user program, which we shall refer to as the MUTATOR, modifies the collection of objects in the heap.

→ The mutator creates objects by acquiring space from the memory manager, and the mutator may introduce and drop references to existing objects.

→ objects become garbage when the mutator program cannot reach them.

⇒ Type Safety

type safety

performance metrics -
reachability

→ A language in which the type of any data component can be determined is said to be TYPE SAFE.

→ There are type-safe languages like ML, for which we can determine types at compile time.

→ Java, type will be determined at run-time, called DYNAMICALLY typed language.

→ lang which is neither statically nor dynamically type safe, then it is said to be UNSAFE.

Performance Metrics

performance metrics that must be considered when designing a garbage collector are:

overall Execution time

Space usage

pause time

program locality

Reachability

→ we refer to all the data that can be accessed directly by a program, w/o having to dereference any pointer, as the ROOT set.

→ Reachability becomes a bit more complex when the program has been optimized by the compiler.

There are four basic operations that a mutator performs to change the set of reachable objects

Object Allocations

parameter passing & Return values

Reference Assignments

Procedure Returns

→ Reference Counting Garbage Collectors

→ Garbage Collector. Identifies garbage as an object changes from being reachable to unreachable;

→ the object can be deleted when its count drop to zero.

→ With a reference-counting garbage collector, every object must have a field for the reference count.

→ Reference counts can be maintained as follows:

Object Allocation:- the reference count of the new object is set to 1.

Parameter passing:- the reference count of each object passed into a procedure is incremented.

Reference Assignments:- for stmt $U = V$, where U & V are references, the reference count of the object referred to by V goes up by one, and the count for the old object referred to by U goes down by one.

Procedure Returns :- As a procedure exist, objects referred to by the local variables in its activation record have their counts decremented.

Transitive loss of Reachability :- whenever the reference count of an object becomes zero, we must also decrement the count of each object pointed to by a reference within the object.

PDF Created Using



Camera Scanner

Easily Scan documents & Generate PDF



<https://play.google.com/store/apps/details?id=photo.pdf.maker>