



மனோன்மணியம் சுந்தரனார் பல்கலைக்கழகம்

MANONMANIAM SUNDARANAR UNIVERSITY

TIRUNELVELI – 627 012

தொலைநிலை தொடர் கல்வி இயக்ககம்

**DIRECTORATE OF DISTANCE AND
CONTINUING EDUCATION**



B.Sc., MATHEMATICS

III YEAR

DISCRETE MATHEMATICS

Sub. Code: JEMA51

Prepared by

Dr. R. THAYALARAJAN

Assistant Professor

Department of Mathematics

M.S. University College, Govindaperi – 627 414.



II YEAR - B.SC., MATHEMATICS

DISCRETE MATHEMATICS

SYLLABUS

Unit I:

Mathematical logic: Statements and Notations – Connectives – Negation – Conjunction - Statement formulas and truth table – Conditional and Bi-conditional - Well-formed formulas – Tautologies.

Chapter 1: Sec 1.1, 1.2.1 to 1.2.4, 1.2.6 to 1.2.8

Unit II:

Normal forms – Disjunctive Normal forms – Conjunctive Normal forms – Principal Disjunctive Normal forms – Principal conjunctive Normal forms – Ordering and Uniqueness of Normal forms – Validity using truth tables – Rules of inference.

Chapter 1: Sec 1.3.1 to 1.3.5, 1.4.1, 1.4.2

Unit III:

The Predicate calculus – Predicates – The statement function, Variables and quantifiers – Predicate formulas – Free and bound variables – The Universe of discourse – Inference theory of the predicate calculus – Valid formulas and Equivalence – Some valid formulas involving quantifiers – Theory of inference for the Predicate calculus.

Chapter 1: Sec 1.5.1 to 1.5.5 and 1.6.1 to 1.6.4

Unit IV:

Relations and Ordering – Relations – Properties of Binary relations in a set – Partial ordering – Partially ordered set: Representation and Associated terminology – Functions: Definition and Introduction – Composition of functions – Inverse functions.

Chapter 2: Sec 2.3.1, 2.3.2, 2.3.8, 2.3.9, 2.4.1 to 2.4.3

Unit V:

Lattices as partially ordered sets: Definition and examples – Some properties of Lattices – Sub lattices , Direct product and Homomorphism – Boolean algebra: Definition and examples – Sub Algebra, Direct product and Homomorphism.

Chapter 4: Sec 4.1.1, 4.1.2, 4.1.4, 4.2.1, 4.2.2

Text Book:

1.J.P.Tremblay, R. Manohar, Discrete Mathematics Structures with Applications to Computer Science, Tata Mc Grawhill, 2001.

UNIT – I

MATHEMATICAL LOGIC

Introduction:

Mathematical logic is the foundation of discrete mathematics, providing the rules and language for reasoning about discrete objects like integers and propositions. It uses symbols to represent statements and connectives to build compound statements, which are then evaluated for their truth values to prove theorems and determine the validity of arguments. Key areas include propositional logic (dealing with simple statements) and predicate logic, and the study of logic underpins many computer science concepts such as algorithm design and verification.

- ❖ It is the study of formal reasoning and how to prove mathematical statements.
- ❖ It provides the rules and techniques for determining if a statement is true or false, and if an argument is valid.
- ❖ It establishes a theoretical basis for many areas of mathematics and computer science.

Statements and Notation:

In discrete mathematics, a statement (or proposition) is a declarative sentence that is definitively either true or false. Notation is used to represent these statements and logical relationships between them, using symbols for logical connectives ($\wedge, \vee, \rightarrow, \neg, \leftrightarrow$), quantifiers ($\forall, \exists$) and single letters to represent entire statements (P, Q, R). These are the primary statements. Only those declarative sentences will be admitted in the object language which have one and only one of two possible values called “truth values”. The two truth values are *true* and *false* and are denoted by the symbols T and F respectively. Occasionally they are also denoted by the symbols 1 and 0. The truth values have nothing to do with our feelings of the truth or falsity of these admissible sentences because these feelings are

subjective and depend upon context. For our purpose, it is enough to assume that it is possible to assign one and only one of the two possible values to a declarative sentence. We are concerned in our study with the effect of assigning any particular truth value to declarative sentences rather than with the actual truth value of these sentences. Since we admit only two possible truth values, our logic is sometimes called a two-valued logic.

We develop a mechanism by which we can construct in our object language other declarative sentences having one of the two possible truth values. Note that we do not admit any other types of sentence, such as exclamatory, interrogative, etc., in the object language.

Declarative sentences in the object language are of two types. The first type includes those sentences which are considered to be primitive in the object language. These will be denoted by distinct symbols selected from the capital letters $A, B, C, \dots, P, Q, \dots$, while declarative sentences of the second type are obtained from the primitive ones by using certain symbols, called connectives, and certain punctuation marks, such as parentheses, to join primitive sentences. In any case, all the declarative sentences to which it is possible to assign one and only one of the two possible truth values are called *statements*. These statements which do not contain any of the connectives are called *atomic (primary, primitive) statements*.

We shall now give examples of sentences and show why some of them are not admissible in the object language and, hence, will not be symbolized.

- 1 Canada is a country.
- 2 Moscow is the capital of Spain.
- 3 This statement is false.
- 4 $1 + 101 = 110$.
- 5 Close the door.
- 6 Toronto is an old city.
- 7 Man will reach Mars by 1980.

Obviously Statements (1) and (2) have truth values *true* and *false* respectively. Statement (3) is not a statement according to our definition, because we cannot properly assign to it a definite truth value. If we assign the value *true*, then Sentence (3) says that Statement (3) is false. On the other hand, if we assign it the value *false*, then Sentence (3) implies that Statement (3) is true. This example illustrates a semantic paradox. In (4) we have a statement whose truth value depends upon the context; viz., if we are talking about numbers in the decimal system, then it is a false statement. On the other hand, for numbers in binary, it is a true statement. The truth value of a statement often depends

upon its context, which is generally unstated but nonetheless understood. We shall soon see that we are not going to be preoccupied with the actual truth value of a statement. We shall be interested only in the fact that it *has* a truth value. In this sense (4), (6), and (7) are all statements. Note that Statement (6) is considered true in some parts of the world and false in certain other parts. The truth value of (7) could be determined only in the year 1980, or earlier if a man reaches Mars before that date. But this aspect is not of interest to us. Note that (5) is not a statement; it is a command.

Once we know those atomic statements which are admissible in the object language, we can use symbols to denote them. Methods of constructing and analyzing statements constructed from one or more atomic statements are discussed in Sec. 1-2, while the method of symbolizing atomic statements will be described here after we discuss some conventions regarding the use and mention of names in statements.

It is customary to use the *name* of an object, not the object itself, when making a statement about the object. As an example, consider the statement

8 This table is big.

The expression "this table" is used as a name of the object. The actual object, namely a particular table, is not used in the statement. It would be inconvenient to put the actual table in place of the expression "this table." Even for the case of small objects, where it may be possible to insert the actual object in place of its name, this practice would not permit us to make two simultaneous statements about the same object without using its name at one place or the other. For this reason it may be agreed that a statement about an object would contain never the object itself but only its name. In fact, we are so familiar with this convention that we take it for granted.

Consider, now, a situation in which we wish to discuss something about a *name*, so that the name is the object about which a statement is to be made. According to the rule just stated, we should use not the name itself in the statement but some name of that name. How does one give a name to a name? A usual method is to enclose the name in quotation marks and to treat it as a name for the name. For example, let us look at the following two statements.

9 Clara is smart.

10 "Clara" contains five letters.

In (9) something is said about a person whose name is Clara. But Statement (10) is not about a person but about a name. Thus "Clara" is used as a name of this name. By enclosing the name of a person in quotation marks it is made clear that the statement made in (10) is about a name and not about a person.

This convention can be explained alternatively by saying that we *use* a certain word in a sentence when that word serves as the name of an object under consideration. On the other hand, we *mention* a word in a sentence when that

word is acting not as the name of an object but as the name of the word itself. To “mention” a word means that the word itself has been converted into an object of our consideration.

Throughout the text we shall be making statements not only about what we normally consider objects but also about other statements. Thus it would be necessary to name the statements under consideration. The same device used for naming names could also be used for naming statements. A statement enclosed in quotation marks will be used as the *name* of the statement. More generally, any expression enclosed in quotation marks will be used as the name of that expression. In other words, any expression that is mentioned is placed in quotation marks. The following statement illustrates the above discussion.

11 “Clara is smart” contains “Clara.”

Statement (11) is a statement about Statement (9) and the word “Clara.” Here Statement (9) was named first by enclosing it in quotation marks and then by using this name in (11) along with the name “Clara”!

In this discussion we have used certain other devices to name statements. One such device is to display a statement on a line separated from the main text. This method of display is assumed to have the same effect as that obtained

by using quotation marks to delimit a statement within the text. Further, we have sometimes numbered these statements by inserting a number to the left of the statement. In a later reference this number is used as a name of the statement. This number is written within the text without quotation marks. Such a display and the numbering of statements permit some reduction in the number of quotation marks. Combinations of these different devices will be used throughout the text in naming statements. Thus the statement

12 “Clara is smart” is true.

could be written as “(9) is true,” or equivalently,

12a (9) is true.

A particular person or an object may have more than one name. It is an accepted principle that one may substitute for the name of an object in a given statement any other name of the same object without altering the meaning of the statement. This principle was used in Statements (12) and (12a).

We shall be using the name-forming devices just discussed to form the names of statements. Very often such distinctions are not made in mathematical writings, and generally the difference between the name and the object is assumed to be clear from the context. However, this practice sometimes leads to confusion.

A situation analogous to the name-object concept just discussed exists in many programming languages. In particular, the distinction between the name of a variable and its value is frequently required when a procedure (function or

subroutine) is invoked (called). The arguments (also called actual parameters) in the statement which invokes the procedure are associated with the (formal) parameters of the procedure either by name or by value. If the association is made by value, then only the value of an argument is passed to its corresponding parameter. This procedure implies that we cannot change the value of the argument from within the function since it is not known where this argument is stored in the computer memory. On the other hand, a call-by-name association makes the name or address of the argument available to the procedure. Such an association allows the value of an argument to be changed by instructions in the procedure. We shall now discuss how call-by-name and call-by-value associations are made in a number of programming languages.

In certain versions of FORTRAN compilers (such as IBM's FORTRAN H and G) the name of an argument, not its value, is passed to a function or subroutine. This convention also applies to the case of an argument's being a constant. The address of a constant (stored in some symbol table of the compiler) is passed to the corresponding parameter of the function. This process could lead to catastrophic results. For example, consider the simple function FUN described by the following sequence of statements:

```

INTEGER FUNCTION FUN(I)
  I = 5
  FUN = I
  RETURN
END

```

Suppose that the main program, which invokes FUN, consists of the trivial statements

```

      K = 3
      J = FUN(K) * 3
      L = FUN(3) * 3
      PRINT 10, J, L
10 FORMAT(1H , I3, I3)
      STOP
      END

```

This program yields values of 15 and 20 for variables J and L respectively. In the evaluation of J, the address of K is known within the function. K is changed in the function to a value of 5 by the statement $I = 5$. The functional value returned by the function is 5, and a value of 15 for J results. The computation of L, however, is quite different. The address of 3 is passed to the function. Since the corresponding parameter I is changed to 5, the value of 3 in the symbol table in the main program will also be changed to 5. Note that since all references to the symbol table entry for constant 3 were made at compile time, all such future references in the remainder of the main program still refer to that entry or location, but the value will now be 5, not 3. More specifically, the name 3 in the right operand of the multiplication of L has a value of 5.

In other versions of FORTRAN compilers, such results are prevented by creating a dummy variable for each argument that is a constant. These internal (dummy) variables are not accessible to the programmer. A change in parameter corresponding to a dummy variable changes the value of that variable, but it does not change the value of the original argument from which it was constructed.

WATFIV permits the passing of arguments by value by merely enclosing such arguments in slashes. For example, in the function call

TEST(I,/K/,5)

the value of K is passed to the function TEST.

In PL/I arguments can be passed by value or by name. An argument is passed by value if it is enclosed within parentheses; otherwise it is passed by name. In the function call

TEST(I,(K),5)

the arguments I and K are passed by name and by value respectively.

As mentioned earlier, we shall use the capital letters $A, B, \dots, P, Q, \dots$ (with the exception of T and F) as well as subscripted capital letters to represent statements in symbolic logic. As an illustration, we write

13 P : It is raining today.

In Statement (13) we are including the information that " P " is a statement in symbolic logic which corresponds to the statement in English, "It is raining today." This situation is similar to the translation of the same statement into French as "Aujourd'hui il pleut." Thus " P " in (13)—"It is raining today"—and "Aujourd'hui il pleut" are the names of the same statement. Note that " P " and not P is used as the name of a statement.

CONNECTIVES

The notions of a statement and of its truth value have already been introduced. In the case of simple statements, their truth values are fairly obvious. However, it is possible to construct rather complicated statements from simpler statements by using certain connecting words or expressions known as "sentential connectives." Several such connectives are used in the English language. Because they are used with a variety of meanings, it is necessary to define a set of connectives with definite meanings. It is convenient to denote these new connectives by means of symbols. We define these connectives in this section and then develop methods to determine the truth values of statements that are formed by using them. Various properties of these statements and some relationships between them are also discussed. In addition, we show that the statements along with the connectives define an algebra that satisfies a set of properties. These properties enable us to do some calculations by using statements as objects. The algebra developed here has interesting and important applications in the field of switching theory and logical design of computers, as is shown in Sec. 1-2.15. Some of these results are also used in the theory of inference discussed in Sec. 1-4.

The statements that we consider initially are simple statements, called *atomic* or *primary statements*. As already indicated, new statements can be formed from atomic statements through the use of sentential connectives. The resulting statements are called *molecular* or *compound statements*. Thus the atomic statements are those which do not have any connectives.

In our everyday language we use connectives such as “and,” “but,” “or,” etc., to combine two or more statements to form other statements. However, their use is not always precise and unambiguous. Therefore, we will not symbolize these connectives in our object language; however, we will define connectives which have some resemblance to the connectives in the English language.

The idea of using the capital letters $P, Q, \dots, P_1, P_2, \dots$ to denote statements was already introduced in Sec. 1-1. Now the same symbols, namely, the capital letters with or without subscripts, will also be used to denote arbitrary statements. In this sense, a statement “ P ” either denotes a particular statement or serves as a placeholder for any statement whatsoever. This dual use of the same symbol to denote either a definite statement, called a *constant*, or an arbitrary statement, called a *variable*, does not cause any confusion as its use will be clear from the context. The truth value of “ P ” is the truth value of the actual statement which it represents. It should be emphasized that when “ P ” is used as a statement variable, it has no truth value and as such does not represent a statement in symbolic logic. We understand that if it is to be replaced, then its replacement must be a statement. Then the truth value of P could be determined. It is convenient to call “ P ” in this case a “statement formula.” We discuss the notion of “statement formula” in Sec. 1-2.4. However, in the sections that follow, we often abbreviate the term “statement formula” simply by “statement.” This abbreviation keeps our discussion simple and emphasizes the meaning of the connectives introduced.

As an illustration, let

P : It is raining today.

Q : It is snowing.

and let R be a statement variable whose possible replacements are P and Q . If no replacement for R is specified, it remains a statement variable and has no truth value. On the other hand, the truth values of P and Q can be determined because they are statements.

Negation

The *negation* of a statement is generally formed by introducing the word “not” at a proper place in the statement or by prefixing the statement with the phrase “It is not the case that.” If “ P ” denotes a statement, then the negation of “ P ” is written as “ $\neg P$ ” and read as “not P .” If the truth value of “ P ” is T , then the truth value of “ $\neg P$ ” is F . Also if the truth value of “ P ” is F , then the truth value of “ $\neg P$ ” is T . This definition of the negation is summarized by Table 1-2.1.

Notice that we have not used the quotation marks to denote the names of the statements in the table. This practice is in keeping with the one adopted

earlier, when a statement was separated from the main text and written on a separate line. From now on we shall drop the quotation marks even within the text when we use symbolic names for the statements, except in the case where this practice may lead to confusion. We now illustrate the formation of the negation of a statement.

Consider the statement

P : London is a city.

Then $\neg P$ is the statement

$\neg P$: It is not the case that London is a city.

Normally $\neg P$ can be written as

$\neg P$: London is not a city.

Although the two statements "It is not the case that London is a city" and "London is not a city" are not identical, we have translated both of them by $\neg P$. The reason is that both these statements have the same meaning in English. A given statement in the object language is denoted by a symbol, and it may correspond to several statements in English. This multiplicity happens because in a natural language one can express oneself in a variety of ways.

As an illustration, if a statement is

P : I went to my class yesterday.

then $\neg P$ is any one of the following

1 I did not go to my class yesterday.

**Table 1-2.1 TRUTH TABLE FOR
NEGATION**

P	$\neg P$
T	F
F	T

2 I was absent from my class yesterday.

3 It is not the case that I went to my class yesterday.

The symbol " $\neg$ " has been used here to denote the negation. Alternate symbols used in the literature are " $\sim$," a bar, or " NOT ," so that $\neg P$ is written as $\sim P$, $\bar{P}$, or $NOT P$. Note that a negation is called a connective although it only modifies a statement. In this sense, negation is a unary operation which operates on a single statement or a variable. The word "operation" will be explained in Chap. 2. For the present it is sufficient to note that an operation on statements generates other statements. We have chosen " $\neg$ " to denote negation because this symbol is commonly used in the textbooks on logic and also in several programming languages, one of which will be used here.

Conjunction

The *conjunction* of two statements P and Q is the statement $P \wedge Q$ which is read as " P and Q ." The statement $P \wedge Q$ has the truth value T whenever both P and Q have the truth value T ; otherwise it has the truth value F . The conjunction is defined by Table 1-2.2.

EXAMPLE 1 Form the conjunction of

P : It is raining today.

Q : There are 20 tables in this room.

SOLUTION It is raining today and there are 20 tables in this room. ////

Normally, in our everyday language the conjunction "and" is used between two statements which have some kind of relation. Thus a statement "It is raining today and $2 + 2 = 4$ " sounds odd, but in logic it is a perfectly acceptable statement formed from the statements "It is raining today" and " $2 + 2 = 4$."

EXAMPLE 2 Translate into symbolic form the statement

Jack and Jill went up the hill.

SOLUTION In order to write it as a conjunction of two statements, it is necessary first to paraphrase the statement as

Jack went up the hill and Jill went up the hill.

Table 1-2.2 TRUTH TABLE FOR CONJUNCTION

P	Q	$P \wedge Q$
T	T	T
T	F	F
F	T	F
F	F	F

If we now write

P : Jack went up the hill.

Q : Jill went up the hill.

then the given statement can be written in symbolic form as $P \wedge Q$.

So far we have seen that the symbol $\wedge$ is used as a translation of the connective "and" appearing in English. However, the connective "and" is sometimes used in a different sense, and in such cases it cannot be translated by the symbol $\wedge$ defined above. In order to see this difference, consider the statements:

- 1 Roses are red and violets are blue.
- 2 He opened the book and started to read.
- 3 Jack and Jill are cousins.

In Statement (1) the conjunction “and” is used in the same sense as the symbol $\wedge$. In (2) the word “and” is used in the sense of “and then,” because the action described in “he started to read” occurs after the action described in “he opened the book.” In (3) the word “and” is not a conjunction. Note that our definition of conjunction is symmetric as far as P and Q are concerned; that is to say, the truth values of $P \wedge Q$ and of $Q \wedge P$ are the same for specific values of P and Q . Obviously the truth value of (1) will not change if we write it as

Violets are blue and roses are red.

On the other hand, we cannot write (2) as

He started to read and opened the book.

These examples show that the symbol $\wedge$ has a specific meaning which corresponds to the connective “and” in general, although “and” may also be used with some other meanings. Some authors use the symbol $\&$, or a dot, or “AND” to denote the conjunction. Note that the conjunction is a binary operation in the sense that it connects two statements to form a new statement.

Disjunction

The *disjunction* of two statements P and Q is the statement $P \vee Q$ which is read as “ P or Q .” The statement $P \vee Q$ has the truth value F only when both P and Q have the truth value F ; otherwise it is *true*. The disjunction is defined by Table 1-2.3.

Table 1-2.3 TRUTH TABLE FOR DISJUNCTION

P	Q	$P \vee Q$
T	T	T
T	F	T
F	T	T
F	F	F

The connectives $\neg$ and $\wedge$ defined earlier have the same meaning as the words “not” and “and” in general. However, the connective $\vee$ is not always the same as the word “or” because of the fact that the word “or” in English is commonly used both as an “exclusive OR” and as an “inclusive OR.” For example, consider the following statements:

- 1 I shall watch the game on television or go to the game.
- 2 There is something wrong with the bulb or with the wiring.
- 3 Twenty or thirty animals were killed in the fire today.

In Statement (1), the connective “or” is used in the exclusive sense; that is to say, one or the other possibility exists but not both. In (2) the intended meaning is clearly one or the other or both. The connective “or” used in (2) is the “inclusive OR.” In (3) the “or” is used for indicating an approximate number of animals, and it is not used as a connective.

From the definition of disjunction it is clear that $\vee$ is “inclusive OR.” The symbol $\vee$ comes from the Latin word “vel” which is the “inclusive OR.” It is not necessary to introduce a new symbol for “exclusive OR,” since there are other ways to express it in terms of the symbols already defined. We demonstrate this point in Sec. 1-2.14.

Normally in our everyday language, the disjunction “or” is used between two statements which have some kind of relationship between them. It is not necessary in logic that there be any relationship between them according to the definition of disjunction. The truth value of $P \vee Q$ depends only upon the truth values of P and Q . As before, it may be necessary to paraphrase given statements in English before they can be translated into symbolic form. Similarly, translations of statements from symbolic logic into statements in English may require paraphrasing in order to make them grammatically acceptable.

Statement Formulas and Truth Tables

We have defined the connectives $\neg$, $\wedge$, and $\vee$ so far. Other connectives will be defined subsequently. We shall occasionally distinguish between two types of statements in our symbolic language. Those statements which do not contain any connectives are called *atomic* or *primary* or *simple statements*. On the other hand, those statements which contain one or more primary statements and some connectives are called *molecular* or *composite* or *compound statements*. As an example, let P and Q be any two statements. Some of the compound statements formed by using P and Q are

$$\neg P \quad P \vee Q \quad (P \wedge Q) \vee (\neg P) \quad P \wedge (\neg Q) \quad (1)$$

The compound statements given above are statement formulas derived from the statement variables P and Q . Therefore, P and Q may be called the components of the statement formulas. Observe that in addition to the connectives we have also used parentheses in some cases in order to make the formula unambiguous. We discuss the rules of constructing statement formulas in Sec. 1-2.7.

Recall that a statement formula has no truth value. It is only when the statement variables in a formula are replaced by definite statements that we get a statement. This statement has a truth value which depends upon the truth values of the statements used in replacing the variables.

In the construction of formulas, the parentheses will be used in the same sense in which they are used in elementary arithmetic or algebra or sometimes in a computer programming language. This usage means that the expressions in the innermost parentheses are simplified first. With this convention in mind, $\neg(P \wedge Q)$ means the negation of $P \wedge Q$. Similarly $(P \wedge Q) \vee (Q \wedge R)$ means

the disjunction of $P \wedge Q$ and $Q \wedge R$. $((P \wedge Q) \vee R) \wedge (\neg P)$ means the conjunction of $\neg P$ and $(P \wedge Q) \vee R$, while $(P \wedge Q) \vee R$ means the disjunction of $P \wedge Q$ and R .

In order to reduce the number of parentheses, we will assume that the negation affects as little as possible of what follows. Thus $\neg P \vee Q$ is written for $(\neg P) \vee Q$, and the negation means the negation of the statement immediately following the symbol $\neg$. On the other hand, according to our convention, $\neg(P \wedge Q) \vee R$ stands for the disjunction of $\neg(P \wedge Q)$ and R . The negation affects $P \wedge Q$ but not R .

Truth tables have already been introduced in the definitions of the connectives. Our basic concern is to determine the truth value of a statement formula for each possible combination of the truth values of the component statements. A table showing all such truth values is called the *truth table* of the formula. In Table 1-2.1 we constructed the truth table for $\neg P$. There is only one component or atomic statement, namely P , and so there are only two possible truth values to be considered. Thus Table 1-2.1 has only two rows. In Tables 1-2.2 and 1-2.3 we constructed truth tables for $P \wedge Q$ and $P \vee Q$ respectively. These statement formulas have two component statements, namely P and Q , and there are 2^2 possible combinations of truth values that must be considered. Thus each of the two tables has 2^2 rows. In general, if there are n distinct components in a statement formula, we need to consider 2^n possible combinations of truth values in order to obtain the truth table.

Two methods of constructing truth tables are shown in the following examples.

EXAMPLE 1 Construct the truth table for the statement formula $P \vee \neg Q$.

SOLUTION It is necessary to consider all possible truth values of P and Q . These values are entered in the first two columns of Table 1-2.4 for both methods. In the table which is arrived at by method 1, the truth values of $\neg Q$ are entered

Table 1-2.4a

P	Q	$\neg Q$	$P \vee \neg Q$
T	T	F	T
T	F	T	T
F	T	F	F
F	F	T	T

Method 1

Table 1-2.4b

P	Q	P	$\vee$	$\neg$	Q
T	T	T	T	F	T
T	F	T	T	T	F
F	T	F	F	F	T
F	F	F	T	T	F
Step					
Number	1	3	2	1	

Method 2

in the third column, and the truth values of $P \vee \neg Q$ are entered in the fourth column. In method 2, as given in Table 1-2.4b, a column is drawn for each statement as well as for the connectives that appear. The truth values are entered step by step. The step numbers at the bottom of the table show the sequence followed in arriving at the final step. ////

EXAMPLE 2 Construct the truth table for $P \wedge \neg P$.

SOLUTION See Table 1-2.5. Note that the truth value is F for every possible truth value of P . In this special case, the truth value of $P \wedge \neg P$ is independent of the truth value of P . ////

EXAMPLE 3 Construct the truth table for $(P \vee Q) \vee \neg P$.

SOLUTION See Table 1-2.6. In this case the truth value of the formula $(P \vee Q) \vee \neg P$ is independent of the truth values of P and Q . This independence is due to the special construction of the formula, as we shall see in Sec. 1-2.8.

Table 1-2.5

P	$\neg P$	$P \wedge \neg P$
T	F	F
F	T	F

Method 1

P	P	$\wedge$	$\neg$	P
T	T	F	F	T
F	F	F	T	F

Step Num- ber	1	3	2	1
---------------------	---	---	---	---

Method 2

Table 1-2.6

P	Q	$P \vee Q$	$\neg P$	$(P \vee Q) \vee \neg P$
T	T	T	F	T
T	F	T	F	T
F	T	T	T	T
F	F	F	T	T

Method 1

P	Q	$(P$	$\vee$	$Q)$	$\vee$	$\neg$	P
T	T	T	T	T	T	F	T
T	F	T	T	F	T	F	T
F	T	F	T	T	T	T	F
F	F	F	F	F	T	T	F

Step Number	1	2	1	3	2	1
----------------	---	---	---	---	---	---

Method 2

Observe that if the truth values of the component statements are known, then the truth value of the resulting statement can be readily determined from the truth table by reading along the row which corresponds to the correct truth values of the component statements.

EXERCISES 1-2.4

1 Using the statements

R : Mark is rich.

H : Mark is happy.

write the following statements in symbolic form:

- Mark is poor but happy.
- Mark is rich or unhappy.
- Mark is neither rich nor happy.
- Mark is poor or he is both rich and unhappy.

- 2 Construct the truth tables for the following formulas.
- $\neg(\neg P \vee \neg Q)$
 - $\neg(\neg P \wedge \neg Q)$
 - $P \wedge (P \vee Q)$
 - $P \wedge (Q \wedge P)$
 - $(\neg P \wedge (\neg Q \wedge R)) \vee (Q \wedge R) \vee (P \wedge R)$
 - $(P \wedge Q) \vee (\neg P \wedge Q) \vee (P \wedge \neg Q) \vee (\neg P \wedge \neg Q)$
- 3 For what truth values will the following statement be true? "It is not the case that houses are cold or haunted and it is false that cottages are warm or houses ugly." (Hint: There are four atomic statements.)
- 4 Given the truth values of P and Q as T and those of R and S as F , find the truth values of the following:
- $P \vee (Q \wedge R)$
 - $(P \wedge (Q \wedge R)) \vee \neg((P \vee Q) \wedge (R \vee S))$
 - $(\neg(P \wedge Q) \vee \neg R) \vee (((\neg P \wedge Q) \vee \neg R) \wedge S)$

Conditional and Biconditional

If P and Q are any two statements, then the statement $P \rightarrow Q$ which is read as "If P , then Q " is called a *conditional* statement. The statement $P \rightarrow Q$ has a truth value F when Q has the truth value F and P the truth value T ; otherwise it has the truth value T . The conditional is defined by Table 1-2.8.

The statement P is called the *antecedent* and Q the *consequent* in $P \rightarrow Q$. Again, according to the definition, it is not necessary that there be any kind of relation between P and Q in order to form $P \rightarrow Q$.

Table 1-2.8 TRUTH TABLE FOR CONDITIONAL

P	Q	$P \rightarrow Q$
T	T	T
T	F	F
F	T	T
F	F	T

EXAMPLE 1 Express in English the statement $P \rightarrow Q$ where

P : The sun is shining today.

Q : $2 + 7 > 4$.

SOLUTION If the sun is shining today, then $2 + 7 > 4$.

The conditional often appears very confusing to a beginner, particularly when one tries to translate a conditional in English into symbolic form. A variety of expressions are used in English which can be appropriately translated by the symbol $\rightarrow$. It is customary to represent any one of the following expressions by $P \rightarrow Q$:

- 1 Q is necessary for P .
- 2 P is sufficient for Q .
- 3 Q if P .
- 4 P only if Q .
- 5 P implies Q .

We shall avoid the translation “implies.” Although, in mathematics, the statements “If P , then Q ” and “ P implies Q ” are used interchangeably, we want to use the word “implies” in a different way.

In our everyday language, we use the conditional statements in a more restricted sense. It is customary to assume some kind of relationship or implication or feeling of cause and effect between the antecedent and the consequent in using the conditional. For example, the statement “If I get the book, then I shall read it tonight” sounds reasonable because the second statement “I shall read it (the book) tonight” refers to the book mentioned in the first part of the statement. On the other hand, a statement such as “If I get the book, then this room is red” does not make sense to us in our conventional language. However, according to our definition of the conditional, the last statement is perfectly acceptable and has a truth value which depends on the truth values of the two statements being connected.

The first two entries in Table 1-2.8 are similar to what we would expect in our everyday language. Thus, if P is true and Q is true, then $P \rightarrow Q$ is true. Similarly, if P is true and Q is false, then “If P , then Q ” appears to be false. Consider, for example, the statement “If I get the money, then I shall buy the car.” If I actually get the money and buy the car, then the statement appears to be correct or true. On the other hand, if I do not buy the car even though I get the money, then the statement is false. Normally, when a conditional statement is made, we assume that the antecedent is true. Because of this convention in English, the first two entries in the truth table do not appear strange. Referring to the above statement again, if I do not get the money and I still buy the car, it is not so clear whether the statement made earlier is true or false. Also, if I do not buy the car and I do not get the money, then it is not intuitively clear whether the statement made is true or false. It may be possible to justify entries in the last two rows of the truth table by considering special examples or even by emphasizing certain aspects of the statements given in the above examples. However, it is best to consider Table 1-2.8 as the definition of the conditional in which the entries in the last two rows are arbitrarily assigned in order to avoid any ambiguity. Any other choice for the last two entries would correspond to some other connective which has either been defined or will be defined. In general, the use of “If ..., then ...” in English has only partial resemblance to the use of the conditional $\rightarrow$ as defined here.

EXAMPLE 2 Write the following statement in symbolic form.

If either Jerry takes Calculus or Ken takes Sociology,
then Larry will take English.

SOLUTION Denoting the statements as

J : Jerry takes Calculus.

K : Ken takes Sociology.

L : Larry takes English.

the above statement can be symbolized as

$$(J \vee K) \rightarrow L$$

EXAMPLE 3 Write in symbolic form the statement

The crop will be destroyed if there is a flood.

SOLUTION Let the statements be denoted as

C : The crop will be destroyed.

F : There is a flood.

Note that the given statement uses “if” in the sense of “If ..., then ...” It is better to rewrite the given statement as “If there is a flood, then the crop will be destroyed.” Now it is easy to symbolize it as

$$F \rightarrow C \quad \text{////}$$

EXAMPLE 4 Construct the truth table for $(P \rightarrow Q) \wedge (Q \rightarrow P)$.

SOLUTION See Table 1-2.9. Note that the given formula has the truth value T whenever both P and Q have identical truth values. ////

If P and Q are any two statements, then the statement $P \Leftrightarrow Q$, which is read as “ P if and only if Q ” and abbreviated as “ P iff Q ,” is called a *biconditional statement*. The statement $P \Leftrightarrow Q$ has the truth value T whenever both P and

Table 1-2.9

P	Q	$P \rightarrow Q$	$Q \rightarrow P$	$(P \rightarrow Q) \wedge (Q \rightarrow P)$
T	T	T	T	T
T	F	F	T	F
F	T	T	F	F
F	F	T	T	T

Table 1-2.10 TRUTH TABLE FOR BICONDITIONAL

P	Q	$P \rightleftharpoons Q$
T	T	T
T	F	F
F	T	F
F	F	T

Table 1-2.11

P	Q	$P \wedge Q$	$\neg(P \wedge Q)$	$\neg P$	$\neg Q$	$\neg P \vee \neg Q$	$\neg(P \wedge Q) \rightleftharpoons (\neg P \vee \neg Q)$
T	T	T	F	F	F	F	T
T	F	F	T	F	T	T	T
F	T	F	T	T	F	T	T
F	F	F	T	T	T	T	T

Q have identical truth values. Table 1-2.10 defines the biconditional. The statement $P \rightleftharpoons Q$ is also translated as “ P is necessary and sufficient for Q .” Note that the truth values of $(P \rightarrow Q) \wedge (Q \rightarrow P)$ given in Table 1-2.9 are identical to the truth values of $P \rightleftharpoons Q$ defined here.

EXAMPLE 5 Construct the truth table for the formula

$$\neg(P \wedge Q) \rightleftharpoons (\neg P \vee \neg Q)$$

SOLUTION See Table 1-2.11. Note that the truth values of the given formula are T for all possible truth values of P and Q . ////

EXERCISES 1-2.6

- Show that the truth values of the following formulas are independent of their components.
 - $(P \wedge (P \rightarrow Q)) \rightarrow Q$
 - $(P \rightarrow Q) \rightleftharpoons (\neg P \vee Q)$
 - $((P \rightarrow Q) \wedge (Q \rightarrow R)) \rightarrow (P \rightarrow R)$
 - $(P \rightleftharpoons Q) \rightleftharpoons ((P \wedge Q) \vee (\neg P \wedge \neg Q))$
- Construct the truth tables of the following formulas.
 - $(Q \wedge (P \rightarrow Q)) \rightarrow P$
 - $\neg(P \vee (Q \wedge R)) \rightleftharpoons ((P \vee Q) \wedge (P \vee R))$
- A connective denoted by ∇ is defined by Table 1-2.12. Find a formula using P , Q , and the connectives $\wedge$, $\vee$, and $\neg$ whose truth values are identical to the truth values of $P \nabla Q$.
- Given the truth values of P and Q as T and those of R and S as F , find the truth values of the following:
 - $(\neg(P \wedge Q) \vee \neg R) \vee ((Q \rightleftharpoons \neg P) \rightarrow (R \vee \neg S))$
 - $(P \rightleftharpoons R) \wedge (\neg Q \rightarrow S)$
 - $(P \vee (Q \rightarrow (R \wedge \neg P))) \rightleftharpoons (Q \vee \neg S)$

Table 1-2.12

P	Q	$P \nabla Q$
T	T	F
T	F	T
F	T	T
F	F	F

Well-formed Formulas

The notion of a statement formula has already been introduced. A statement formula is not a statement (although, for the sake of brevity, we have often called it a statement); however, a statement can be obtained from it by replacing the variables by statements. A *statement formula* is an expression which is a string consisting of variables (capital letters with or without subscripts), parentheses, and connective symbols. Not every string of these symbols is a formula. We shall now give a recursive definition of a statement formula, often called a well-formed formula (wff). A *well-formed formula* can be generated by the following rules:

- 1 A statement variable standing alone is a well-formed formula.
- 2 If A is a well-formed formula, then $\neg A$ is a well-formed formula.
- 3 If A and B are well-formed formulas, then $(A \wedge B)$, $(A \vee B)$, $(A \rightarrow B)$, and $(A \rightleftharpoons B)$ are well-formed formulas.
- 4 A string of symbols containing the statement variables, connectives, and parentheses is a well-formed formula, iff it can be obtained by finitely many applications of the rules 1, 2, and 3.

According to this definition, the following are well-formed formulas:

$$\neg(P \wedge Q) \quad \neg(P \vee Q) \quad (P \rightarrow (P \vee Q)) \quad (P \rightarrow (Q \rightarrow R)) \\ ((P \rightarrow Q) \wedge (Q \rightarrow R)) \rightleftharpoons (P \rightarrow R))$$

The following are not well-formed formulas.

- 1 $\neg P \wedge Q$. Obviously P and Q are well-formed formulas. A wff would be either $(\neg P \wedge Q)$ or $\neg(P \wedge Q)$.
- 2 $(P \rightarrow Q) \rightarrow (\wedge Q)$. This is not a wff because $\wedge Q$ is not.
- 3 $P \rightarrow Q$. Note that $(P \rightarrow Q)$ is a wff.
- 4 $(P \wedge Q) \rightarrow Q$. The reason for this not being a wff is that one of the parentheses in the beginning is missing. $((P \wedge Q) \rightarrow Q)$ is a wff, while $(P \wedge Q) \rightarrow Q$ is still not a wff.

It is possible to introduce some conventions so that the number of parentheses used can be reduced. In fact, there are conventions which, when followed, allow one to dispense with all the parentheses. We shall not discuss these conven-

tions here. For the sake of convenience we shall omit the outer parentheses. Thus we write $P \wedge Q$ in place of $(P \wedge Q)$, $(P \wedge Q) \rightarrow Q$ in place of $((P \wedge Q) \rightarrow Q)$, and $((P \rightarrow Q) \wedge (Q \rightarrow R)) \rightleftharpoons (P \rightarrow R)$ instead of $((P \rightarrow Q) \wedge (Q \rightarrow R)) \rightleftharpoons (P \rightarrow R)$. Since the only formulas we will encounter are well-formed formulas, we will refer to well-formed formulas as formulas.

Tautologies

Well-formed formulas have been defined. We also know how to construct the truth table of a given formula. Let us consider what a truth table represents. If definite statements are substituted for the variables in a formula, there results a statement. The truth value of this resulting statement depends upon the truth values of the statements substituted for the variables. Such a truth value appears as one of the entries in the final column of the truth table. Observe that this entry will not change even if any of the definite statements that replace particular variables are themselves replaced by other statements, as long as the truth values associated with all variables are unchanged. In other words, an entry in the final column depends only on the truth values of the statements assigned to the variables rather than on the statements themselves. Different rows correspond to different sets of truth value assignments. A truth table is therefore a summary of the truth values of the resulting statements for all possible assignments of values to the variables appearing in a formula. It must be emphasized that a statement⁴ formula does not have a truth value. In our discussion which follows we shall, for the sake of simplicity, use the expression "the truth value of a statement formula" to mean the entries in the final column of the truth table of the formula.

In general, the final column of a truth table of a given formula contains both T and F . There are some formulas whose truth values are always T or always F regardless of the truth value assignments to the variables. This situation occurs because of the special construction of these formulas. We have already seen some examples of such formulas.

Consider, for example, the statement formulas $P \vee \neg P$ and $P \wedge \neg P$ in Table 1-2.13. The truth values of $P \vee \neg P$ and $P \wedge \neg P$, which are T and F respectively, are independent of the statement by which the variable P may be replaced.

A statement formula which is true regardless of the truth values of the statements which replace the variables in it is called a *universally valid formula* or a *tautology* or a *logical truth*. A statement formula which is false regardless of the truth values of the statements which replace the variables in it is called a *contradiction*. Obviously, the negation of a contradiction is a tautology. We may say that a statement formula which is a tautology is *identically true* and a formula which is a contradiction is *identically false*.

A straightforward method to determine whether a given formula is a tautology is to construct its truth table. This process can always be used but often becomes tedious, particularly when the number of distinct variables is

large or when the formula is complicated. Recall that the numbers of rows in a truth table is 2^n , where n is the number of distinct variables in the formula. Later,

Table 1-2.13

P	$\neg P$	$P \vee \neg P$	$P \wedge \neg P$
T	F	T	F
F	T	T	F

alternative methods will be developed that will be able to determine whether a statement formula is a tautology without having to construct its truth table.

A simple fact about tautologies is that the conjunction of two tautologies is also a tautology. Let us denote by A and B two statement formulas which are tautologies. If we assign any truth values to the variables of A and B , then the truth values of both A and B will be T . Thus the truth value of $A \wedge B$ will be T , so that $A \wedge B$ will be a tautology.

A formula A is called a *substitution instance* of another formula B if A can be obtained from B by substituting formulas for some variables of B , with the condition that the same formula is substituted for the same variable each time it occurs. We now illustrate this concept. Let

$$B: P \rightarrow (J \wedge P)$$

Substitute $R \rightleftharpoons S$ for P in B , and we get

$$A: (R \rightleftharpoons S) \rightarrow (J \wedge (R \rightleftharpoons S))$$

Then A is a substitution instance of B . Note that

$$(R \rightleftharpoons S) \rightarrow (J \wedge P)$$

is not a substitution instance of B because the variable P in $J \wedge P$ was not replaced by $R \rightleftharpoons S$. It is possible to substitute more than one variable by other formulas, provided that all substitutions are considered to occur simultaneously. For example, substitution instances of $P \rightarrow \neg Q$ are

- 1 $(R \wedge \neg S) \rightarrow \neg(J \vee M)$
- 2 $(R \wedge \neg S) \rightarrow \neg(R \wedge \neg S)$
- 3 $(R \wedge \neg S) \rightarrow \neg P$
- 4 $Q \rightarrow \neg(P \wedge \neg Q)$

In (2) both P and Q have been replaced by $R \wedge \neg S$. In (4), P is replaced by Q and Q by $P \wedge \neg Q$.

Next, consider the following formulas which result from $P \rightarrow \neg Q$.

- 1 Substitute $P \vee Q$ for P and R for Q to get the substitution instance $(P \vee Q) \rightarrow \neg R$.

2 First substitute $P \vee Q$ for P to obtain the substitution instance $(P \vee Q) \rightarrow \neg Q$. Next, substitute R for Q in $(P \vee Q) \rightarrow \neg Q$, and we get $(P \vee R) \rightarrow \neg R$. This formula is a substitution instance of $(P \vee Q) \rightarrow \neg Q$, but it is not a substitution instance of $P \rightarrow \neg Q$ under the substitution $(P \vee Q)$ for P and R for Q . This statement is true because we did not substitute simultaneously as we did in (1).

It may be noted that in constructing substitution instances of a formula, substitutions are made for the atomic formula and never for the molecular formula. Thus $P \rightarrow Q$ is not a substitution instance of $P \rightarrow \neg R$, because it is R which must be replaced and not $\neg R$.

The importance of the above concept lies in the fact that any substitution instance of a tautology is a tautology. Consider the tautology $P \vee \neg P$. Regardless of what is substituted for P , the truth value of $P \vee \neg P$ is always T . Therefore, if we substitute any statement formula for P , the resulting formula will be a tautology. Hence the following substitution instances of $P \vee \neg P$ are tautologies.

$$\begin{aligned}(R \rightarrow S) \vee \neg(R \rightarrow S) \\ ((P \vee S) \wedge R) \vee \neg((P \vee S) \wedge R) \\ (((P \vee \neg Q) \rightarrow R) \leftrightarrow S) \vee \neg(((P \vee \neg Q) \rightarrow R) \leftrightarrow S)\end{aligned}$$

Thus, if it is possible to detect whether a given formula is a substitution instance of a tautology, then it is immediately known that the given formula is also a tautology. Similarly, one can start with a tautology and write a large number of formulas which are substitution instances of this tautology and hence are themselves tautologies.

EXERCISES 1-2.8

- 1 From the formulas given below select those which are well-formed according to the definition in Sec. 1-2.7, and indicate which ones are tautologies or contradictions.
 - (a) $(P \rightarrow (P \vee Q))$
 - (b) $((P \rightarrow (\neg P)) \rightarrow \neg P)$
 - (c) $((\neg Q \wedge P) \wedge Q)$
 - (d) $((P \rightarrow (Q \rightarrow R)) \rightarrow ((P \rightarrow Q) \rightarrow (P \rightarrow R)))$
 - (e) $((\neg P \rightarrow Q) \rightarrow (Q \rightarrow P))$
 - (f) $((P \wedge Q) \leftrightarrow P)$
- 2 Produce the substitution instances of the following formulas for the given substitutions.
 - (a) $((P \rightarrow Q) \rightarrow P) \rightarrow P$; substitute $(P \rightarrow Q)$ for P and $((P \wedge Q) \rightarrow R)$ for Q .
 - (b) $((P \rightarrow Q) \rightarrow (Q \rightarrow P))$; substitute Q for P and $(P \wedge \neg P)$ for Q .
- 3 Determine the formulas which are substitution instances of other formulas in the list and give the substitutions.
 - (a) $(P \rightarrow (Q \rightarrow P))$
 - (b) $((((P \rightarrow Q) \wedge (R \rightarrow S)) \wedge (P \vee R)) \rightarrow (Q \vee S))$
 - (c) $(Q \rightarrow ((P \rightarrow P) \rightarrow Q))$
 - (d) $(P \rightarrow ((P \rightarrow (Q \rightarrow P)) \rightarrow P))$
 - (e) $((((R \rightarrow S) \wedge (Q \rightarrow P)) \wedge (R \vee Q)) \rightarrow (S \vee P))$

NORMAL FORMS

Let $A(P_1, P_2, \dots, P_n)$ be a statement formula where $P_1, P_2, \dots, P_n$ are the atomic variables. If we consider all possible assignments of the truth values to $P_1, P_2, \dots, P_n$ and obtain the resulting truth values of the formula A , then we get the truth table for A . Such a truth table contains 2^n rows. The formula may have the truth value T for all possible assignments of the truth values to the variables $P_1, P_2, \dots, P_n$. In this case, A is said to be identically true, or a tautology. If A has the truth value F for all possible assignments of the truth values to $P_1, P_2, \dots, P_n$, then A is said to be identically false, or a contradiction. Finally, if A has the truth value T for at least one combination of truth values assigned to $P_1, P_2, \dots, P_n$, then A is said to be *satisfiable*.

The problem of determining, in a finite number of steps, whether a given statement formula is a tautology or a contradiction or at least satisfiable is known as a *decision problem*. Obviously, the construction of truth tables involves a finite number of steps, and, as such, a decision problem in the statement calculus always has a solution. Similarly, decision problems can be posed for other logical systems, particularly for the predicate calculus. However, in the latter case, the solution of the decision problem may not be simple.

As was mentioned earlier, the construction of truth tables may not be practical, even with the aid of a computer. We therefore consider other procedures known as reduction to normal forms.

Disjunctive Normal Forms

It will be convenient to use the word “product” in place of “conjunction” and “sum” in place of “disjunction” in our current discussion.

A product of the variables and their negations in a formula is called an *elementary product*. Similarly, a sum of the variables and their negations is called an *elementary sum*.

Let P and Q be any two atomic variables. Then $P, \neg P \wedge Q, \neg Q \wedge P \wedge \neg P, P \wedge \neg P$, and $Q \wedge \neg P$ are some examples of elementary products. On the other hand, $P, \neg P \vee Q, \neg Q \vee P \vee \neg P, P \vee \neg P$, and $Q \vee \neg P$ are examples of elementary sums of the two variables. Any part of an elementary sum or product which is itself an elementary sum or product is called a *factor* of the original elementary sum or product. Thus $\neg Q, P \wedge \neg P$, and $\neg Q \wedge P$ are some of the factors of $\neg Q \wedge P \wedge \neg P$. The following statements hold for elementary sums and products.

A necessary and sufficient condition for an elementary product to be identically false is that it contain at least one pair of factors in which one is the negation of the other.

A necessary and sufficient condition for an elementary sum to be identically true is that it contain at least one pair of factors in which one is the negation of the other.

We shall now prove the first of these two statements. The proof of the second will follow along the same lines.

We know that for any variable P , $P \wedge \neg P$ is identically false. Hence if $P \wedge \neg P$ appears in the elementary product, then the product is identically false. On the other hand, if an elementary product is identically false and does not contain at least one factor of this type, then we can assign truth values T and F to variables and negated variables, respectively, that appear in the product. This assignment would mean that the elementary product has the truth value T . But that statement is contrary to our assumption. Hence the statement follows.

A formula which is equivalent to a given formula and which consists of a sum of elementary products is called a *disjunctive normal form* of the given formula.

We shall first discuss a procedure by which one can obtain a disjunctive normal form of a given formula. It has already been shown that if the connectives $\rightarrow$ and $\Leftrightarrow$ appear in the given formula, then an equivalent formula can be obtained in which " $\rightarrow$ " and " $\Leftrightarrow$ " are replaced by $\wedge$, $\vee$, and $\neg$. This statement would be true of any other connective yet undefined. The truth of this statement will become apparent after our discussion of principal disjunctive normal forms. Therefore, there is no loss of generality in assuming that the given formula contains the connectives $\wedge$, $\vee$, and $\neg$ only.

If the negation is applied to the formula or to a part of the formula and not to the variables appearing in it, then by using De Morgan's laws an equivalent formula can be obtained in which the negation is applied to the variables only. After this step, the formula obtained may still fail to be in disjunctive normal form because there may be some parts of it which are products of sums. By repeated application of the distributive laws we obtain the required normal form.

EXAMPLE 1 Obtain disjunctive normal forms of (a) $P \wedge (P \rightarrow Q)$; (b) $\neg(P \vee Q) \Leftrightarrow (P \wedge Q)$.

SOLUTION

$$(a) P \wedge (P \rightarrow Q) \Leftrightarrow P \wedge (\neg P \vee Q) \Leftrightarrow (P \wedge \neg P) \vee (P \wedge Q)$$

$$(b) \neg(P \vee Q) \Leftrightarrow (P \wedge Q)$$

$$\Leftrightarrow (\neg(P \vee Q) \wedge (P \wedge Q)) \vee ((P \vee Q) \wedge \neg(P \wedge Q))$$

$$[\text{using } R \Leftrightarrow S \Leftrightarrow (R \wedge S) \vee (\neg R \wedge \neg S)]$$

$$\Leftrightarrow (\neg P \wedge \neg Q \wedge P \wedge Q) \vee ((P \vee Q) \wedge (\neg P \vee \neg Q))$$

$$\Leftrightarrow (\neg P \wedge \neg Q \wedge P \wedge Q) \vee ((P \vee Q) \wedge \neg P)$$

$$\vee ((P \vee Q) \wedge \neg Q)$$

$$\begin{aligned} \Leftrightarrow (\neg P \wedge \neg Q \wedge P \wedge Q) \vee (P \wedge \neg P) \vee (Q \wedge \neg P) \\ \vee (P \wedge \neg Q) \vee (Q \wedge \neg Q) \end{aligned}$$

which is the required disjunctive normal form.

The disjunctive normal form of a given formula is not unique. In fact, different disjunctive normal forms can be obtained for a given formula if the distributive laws are applied in different ways. Apart from this fact, the factors in each elementary product, as well as the factors in the sum, can be commuted. However, we shall not consider as distinct the various disjunctive normal forms obtained by reordering the factors either in the elementary products or in the sums.

Consider the formula $F \vee (Q \wedge R)$. Here the formula is already in the disjunctive normal form. However, we may write

$$\begin{aligned} P \vee (Q \wedge R) &\Leftrightarrow (P \vee Q) \wedge (P \vee R) \Leftrightarrow (P \wedge P) \\ &\vee (P \wedge Q) \vee (P \wedge R) \vee (Q \wedge R) \end{aligned}$$

the last equivalent formula being another equivalent disjunctive normal form. Of course, different disjunctive normal forms of the same formula are equivalent. In order to arrive at a unique normal form of a given formula, we introduce the principal disjunctive normal form in Sec. 1-3.3.

Finally, we remark that a given formula is identically false if every elementary product appearing in its disjunctive normal form is identically false. For the assumption to be true, every elementary product would have to have at least two factors, of which one is the negation of the other.

Conjunctive Normal Forms

A formula which is equivalent to a given formula and which consists of a product of elementary sums is called a *conjunctive normal form* of the given formula.

The method for obtaining a conjunctive normal form of a given formula is similar to the one given for disjunctive normal forms and will be demonstrated by examples. Again, the conjunctive normal form is not unique. Furthermore, a given formula is identically true if every elementary sum in its conjunctive normal form is identically true. This would be the case if every elementary sum appearing in the formula had at least two factors, of which one is the negation of the other.

EXAMPLE 1 Obtain a conjunctive normal form of each of the formulas given in Example 1 of Sec. 1-3.1.

SOLUTION

- (a) $P \wedge (P \rightarrow Q) \Leftrightarrow P \wedge (\neg P \vee Q)$. Hence $P \wedge (\neg P \vee Q)$ is a required form.
 (b) $\neg(P \vee Q) \Leftrightarrow (P \wedge Q) \Leftrightarrow (\neg(P \vee Q) \rightarrow (P \wedge Q)) \wedge ((P \wedge Q)$

$$\begin{aligned}
& \rightarrow \neg(P \vee Q)) \\
& [\text{using } R \Leftrightarrow S \Leftrightarrow (R \rightarrow S) \wedge (S \rightarrow R)] \\
& \Leftrightarrow ((P \vee Q) \vee (P \wedge Q)) \wedge (\neg(P \wedge Q) \\
& \quad \vee (\neg P \wedge \neg Q)) \\
& \Leftrightarrow ((P \vee Q \vee P) \wedge (P \vee Q \vee Q)) \\
& \quad \wedge ((\neg P \vee \neg Q) \vee (\neg P \wedge \neg Q)) \\
& \Leftrightarrow (P \vee Q \vee P) \wedge (P \vee Q \vee Q) \\
& \quad \wedge (\neg P \vee \neg Q \vee \neg P) \wedge (\neg P \vee \neg Q \vee \neg Q) \\
& \hspace{15em} ////
\end{aligned}$$

EXAMPLE 2 Show that the formula $Q \vee (P \wedge \neg Q) \vee (\neg P \wedge \neg Q)$ is a tautology.

SOLUTION First we obtain a conjunctive normal form of the given formula.

$$\begin{aligned}
Q \vee (P \wedge \neg Q) \vee (\neg P \wedge \neg Q) & \Leftrightarrow Q \vee ((P \vee \neg P) \wedge \neg Q) \\
& \Leftrightarrow (Q \vee (P \vee \neg P)) \wedge (Q \vee \neg Q) \\
& \Leftrightarrow (Q \vee P \vee \neg P) \wedge (Q \vee \neg Q)
\end{aligned}$$

Since each of the elementary sums is a tautology, the given formula is a tautology. ////

1-3.3 Principal Disjunctive Normal Forms

Let P and Q be two statement variables. Let us construct all possible formulas which consist of conjunctions of P or its negation and conjunctions of Q or its negation. None of the formulas should contain both a variable and its negation. Furthermore, any formula which is obtained by commuting the formulas in the conjunction is not included in the list because such a formula will be equivalent to one included in the list. For example, either $P \wedge Q$ or $Q \wedge P$ is included, but not both. For two variables P and Q , there are 2^2 such formulas given by

$$P \wedge Q \quad P \wedge \neg Q \quad \neg P \wedge Q \quad \text{and} \quad \neg P \wedge \neg Q$$

These formulas are called *minterms* or Boolean conjunctions of P and Q . From the truth tables of these minterms, it is clear that no two minterms are equivalent. Each minterm has the truth value T for exactly one combination of the truth values of the variables P and Q . This fact is shown in Table 1-3.1.

We assert that if the truth table of any formula containing only the variables P and Q is known, then one can easily obtain an equivalent formula which consists of a disjunction of some of the minterms. This statement is demonstrated as follows.

For every truth value T in the truth table of the given formula, select the

Table 1-3.1

P	Q	$P \wedge Q$	$P \wedge \neg Q$	$\neg P \wedge Q$	$\neg P \wedge \neg Q$
T	T	T	F	F	F
T	F	F	T	F	F
F	T	F	F	T	F
F	F	F	F	F	T

minterm which also has the value T for the same combination of the truth values of P and Q . The disjunction of these minterms will then be equivalent to the given formula.

This discussion provides the basis for a proof that a formula containing any connective is equivalent to a formula containing $\wedge$, $\vee$, and $\neg$.

For a given formula, an equivalent formula consisting of disjunctions of minterms only is known as its *principal disjunctive normal form*. Such a normal form is also called the *sum-of-products canonical form*.

EXAMPLE 1 The truth tables for $P \rightarrow Q$, $P \vee Q$, and $\neg(P \wedge Q)$ are given in Table 1-3.2. Obtain the principal disjunctive normal forms of these formulas.

SOLUTION

$$P \rightarrow Q \Leftrightarrow (P \wedge Q) \vee (\neg P \wedge Q) \vee (\neg P \wedge \neg Q)$$

$$P \vee Q \Leftrightarrow (P \wedge Q) \vee (P \wedge \neg Q) \vee (\neg P \wedge Q)$$

$$\neg(P \wedge Q) \Leftrightarrow (P \wedge \neg Q) \vee (\neg P \wedge Q) \vee (\neg P \wedge \neg Q) \quad ////$$

Note that the number of minterms appearing in the normal form is the same as the number of entries with the truth value T in the truth table of the given formula. Thus every formula which is not a contradiction has an equivalent principal disjunctive normal form. Further, such a normal form is unique, except for the rearrangements of the factors in the disjunctions as well as in each of the minterms. One can get a unique normal form by imposing a certain order in which the variables appear in the minterms as well as a definite order in which the minterms appear in the disjunction. In that case, if two given formulas are equivalent, then both of them must have identical principal disjunctive normal forms. Therefore, if we can devise a method other than the construction of truth tables to obtain the principal disjunctive normal form of a given formula, then

Table 1-3.2

P	Q	$P \rightarrow Q$	$P \vee Q$	$\neg(P \wedge Q)$
T	T	T	T	F
T	F	F	T	T
F	T	T	T	T
F	F	T	F	T

the same method can be used to determine whether two given formulas are equivalent.

Although our discussion of the principal disjunctive normal form was restricted to formulas containing only two variables, it is possible to define the minterms for three or more variables. Minterms for the three variables P , Q , and R are

$$\begin{array}{cccc}
 P \wedge Q \wedge R & P \wedge Q \wedge \neg R & P \wedge \neg Q \wedge R & P \wedge \neg Q \wedge \neg R \\
 \neg P \wedge Q \wedge R & \neg P \wedge Q \wedge \neg R & \neg P \wedge \neg Q \wedge R & \neg P \wedge \neg Q \wedge \neg R
 \end{array}$$

These minterms satisfy properties similar to those given for two variables. An equivalent principal disjunctive normal form of any formula which depends upon the variables P , Q , and R can be obtained. Note that there are 2^3 minterms for three variables or, more generally, 2^n minterms for n variables. For any formula containing n variables which are denoted by $P_1, P_2, \dots, P_n$, an equivalent disjunctive normal form can be obtained by selecting appropriate minterms out of the set of 2^n possible minterms.

If a formula is a tautology, then obviously all the minterms appear in its principal disjunctive normal form; it is also possible to determine whether a given formula is a tautology by obtaining its principal disjunctive normal form.

In order to obtain the principal disjunctive normal form of a given formula without constructing its truth table, one may first replace the conditionals and biconditionals by their equivalent formulas containing only $\wedge$, $\vee$, and $\neg$. Next, the negations are applied to the variables by using De Morgan's laws followed by the application of distributive laws, as was done earlier in obtaining the disjunctive or conjunctive normal forms. Any elementary product which is a contradiction is dropped. Minterms are obtained in the disjunctions by introducing the missing factors. Identical minterms appearing in the disjunctions are deleted. This procedure is demonstrated by means of examples.

EXAMPLE 2 Obtain the principal disjunctive normal forms of (a) $\neg P \vee Q$; (b) $(P \wedge Q) \vee (\neg P \wedge R) \vee (Q \wedge R)$.

SOLUTION

$$\begin{aligned}
 (a) \quad \neg P \vee Q &\Leftrightarrow (\neg P \wedge (Q \vee \neg Q)) \vee (Q \wedge (P \vee \neg P)) \\
 &\quad (A \wedge T \Leftrightarrow A) \\
 &\Leftrightarrow (\neg P \wedge Q) \vee (\neg P \wedge \neg Q) \vee (Q \wedge P) \vee (Q \wedge \neg P)
 \end{aligned}$$

(distributive laws)

$$\Leftrightarrow (\neg P \wedge Q) \vee (\neg P \wedge \neg Q) \vee (P \wedge Q)$$

(commutative law and $P \vee P \Leftrightarrow P$)

(See Example 1.)

$$(b) (P \wedge Q) \vee (\neg P \wedge R) \vee (Q \wedge R)$$

$$\begin{aligned} &\Leftrightarrow (P \wedge Q \wedge (R \vee \neg R)) \vee (\neg P \wedge R \wedge (Q \vee \neg Q)) \\ &\quad \vee (Q \wedge R \wedge (P \vee \neg P)) \end{aligned}$$

$$\begin{aligned} &\Leftrightarrow (P \wedge Q \wedge R) \vee (P \wedge Q \wedge \neg R) \vee (\neg P \wedge Q \wedge R) \\ &\quad \vee (\neg P \wedge \neg Q \wedge R) \end{aligned}$$

////

EXAMPLE 3 Show that the following are equivalent formulas.

$$(a) P \vee (P \wedge Q) \Leftrightarrow P$$

$$(b) P \vee (\neg P \wedge Q) \Leftrightarrow P \vee Q$$

SOLUTION We write the principal disjunctive normal form of each formula and compare these normal forms.

$$(a) P \vee (P \wedge Q) \Leftrightarrow (P \wedge (Q \vee \neg Q)) \vee (P \wedge Q) \Leftrightarrow (P \wedge Q) \vee (P \wedge \neg Q)$$

$$P \Leftrightarrow P \wedge (Q \vee \neg Q) \Leftrightarrow (P \wedge Q) \vee (P \wedge \neg Q)$$

$$(b) P \vee (\neg P \wedge Q) \Leftrightarrow (P \wedge (Q \vee \neg Q)) \vee (\neg P \wedge Q)$$

$$\Leftrightarrow (P \wedge Q) \vee (P \wedge \neg Q) \vee (\neg P \wedge Q)$$

$$P \vee Q \Leftrightarrow (P \wedge (Q \vee \neg Q)) \vee (Q \wedge (P \vee \neg P))$$

$$\Leftrightarrow (P \wedge Q) \vee (P \wedge \neg Q) \vee (\neg P \wedge Q)$$

////

EXAMPLE 4 Obtain the principal disjunctive normal form of

$$P \rightarrow ((P \rightarrow Q) \wedge \neg(\neg Q \vee \neg P))$$

SOLUTION Using $P \rightarrow Q \Leftrightarrow \neg P \vee Q$ and De Morgan's law, we obtain

$$P \rightarrow ((P \rightarrow Q) \wedge \neg(\neg Q \vee \neg P))$$

$$\Leftrightarrow \neg P \vee ((\neg P \vee Q) \wedge (Q \wedge P))$$

$$\Leftrightarrow \neg P \vee (\neg P \wedge (Q \wedge P)) \vee (Q \wedge (Q \wedge P))$$

$$\Leftrightarrow \neg P \vee (Q \wedge P)$$

$$\Leftrightarrow (\neg P \wedge (Q \vee \neg Q)) \vee (Q \wedge P)$$

$$\Leftrightarrow (\neg P \wedge Q) \vee (\neg P \wedge \neg Q) \vee (P \wedge Q)$$

////

The procedure described above becomes tedious if the given formula is complicated and contains more than two or three variables. When the number of variables is large, even a comparison of two principal disjunctive normal forms becomes cumbersome. In Sec. 1-3.5, we describe an ordering procedure for the

variables and a notation which make such a comparison easy. We also discuss in Chap. 2 a computer program to obtain the sum-of-products canonical form for a given formula.

Principal Conjunctive Normal Forms

In order to define the principal conjunctive normal form, we first define formulas which are called maxterms. For a given number of variables, the *maxterm* consists of disjunctions in which each variable or its negation, but not both, appears only once. Thus the maxterms are the duals of minterms. Either from the duality principle or directly from the truth tables, it can be ascertained that each of the maxterms has the truth value F for exactly one combination of the truth values of the variables. Also different maxterms have the truth value F for different combinations of the truth values of the variables.

For a given formula, an equivalent formula consisting of conjunctions of the maxterms only is known as its *principal conjunctive normal form*. This normal form is also called the *product-of-sums canonical form*. Every formula which is not a tautology has an equivalent principal conjunctive normal form which is unique except for the rearrangement of the factors in the maxterms as well as in their conjunctions. The method for obtaining the principal conjunctive normal form for a given formula is similar to the one described previously for the principal disjunctive normal form. In fact, all the assertions made for the principal disjunctive normal forms can also be made for the principal conjunctive normal forms by the duality principle.

If the principal disjunctive (conjunctive) normal form of a given formula A containing n variables is known, then the principal disjunctive (conjunctive) normal form of $\neg A$ will consist of the disjunction (conjunction) of the remaining minterms (maxterms) which do not appear in the principal disjunctive (conjunctive) normal form of A . From $A \Leftrightarrow \neg \neg A$ one can obtain the principal conjunctive (disjunctive) normal form of A by repeated applications of De Morgan's laws to the principal disjunctive (conjunctive) normal form of $\neg A$. This procedure will be illustrated by an example.

In order to determine whether two given formulas A and B are equivalent, one can obtain any of the principal normal forms of the two formulas and compare them. It is not necessary to assume that both formulas have the same variables. In fact, each formula can be assumed to depend upon all the variables that appear in both formulas, by introducing the missing variables and then reducing them to their principal normal forms.

EXAMPLE 1 Obtain the principal conjunctive normal form of the formula S given by $(\neg P \rightarrow R) \wedge (Q \Leftrightarrow P)$.

SOLUTION

$$(\neg P \rightarrow R) \wedge (Q \Leftrightarrow P)$$

$$\begin{aligned}
&\Leftrightarrow (P \vee R) \wedge ((Q \rightarrow P) \wedge (P \rightarrow Q)) \\
&\Leftrightarrow (P \vee R) \wedge (\neg Q \vee P) \wedge (\neg P \vee Q) \\
&\Leftrightarrow (P \vee R \vee (Q \wedge \neg Q)) \wedge (\neg Q \vee P \vee (R \wedge \neg R)) \\
&\quad \wedge (\neg P \vee Q \vee (R \wedge \neg R)) \\
&\Leftrightarrow (P \vee Q \vee R) \wedge (P \vee \neg Q \vee R) \wedge (P \vee \neg Q \vee \neg R) \\
&\quad \wedge (\neg P \vee Q \vee R) \wedge (\neg P \vee Q \vee \neg R)
\end{aligned}$$

Now the conjunctive normal form of $\neg S$ can easily be obtained by writing the conjunction of the remaining maxterms; thus, $\neg S$ has the principal conjunctive normal form

$$(P \vee Q \vee \neg R) \wedge (\neg P \vee \neg Q \vee R) \wedge (\neg P \vee \neg Q \vee \neg R)$$

By considering $\neg \neg S$, we obtain

$$\begin{aligned}
&\neg(P \vee Q \vee \neg R) \vee \neg(\neg P \vee \neg Q \vee R) \vee \neg(\neg P \vee \neg Q \vee \neg R) \\
&\Leftrightarrow (\neg P \wedge \neg Q \wedge R) \vee (P \wedge Q \wedge \neg R) \vee (P \wedge Q \wedge R)
\end{aligned}$$

which is the principal disjunctive normal form of S . ////

EXAMPLE 2 The truth table for a formula A is given in Table 1-3.3. Determine its disjunctive and conjunctive normal forms.

SOLUTION By choosing the minterms corresponding to each T value of A , we obtain

$$\begin{aligned}
A \Leftrightarrow (P \wedge \neg Q \wedge R) \vee (\neg P \wedge Q \wedge R) \vee (\neg P \wedge Q \wedge \neg R) \\
\vee (\neg P \wedge \neg Q \wedge \neg R)
\end{aligned}$$

Similarly

$$\begin{aligned}
A \Leftrightarrow (\neg P \vee \neg Q \vee \neg R) \wedge (\neg P \vee \neg Q \vee R) \wedge (\neg P \vee Q \vee R) \\
\wedge (P \vee Q \vee \neg R)
\end{aligned}$$

Here the maxterms appearing in the normal form correspond to the F values of A . The maxterms are written down by including the variable if its truth value is F and its negation if the value is T . ////

Ordering and Uniqueness of Normal Forms

Given any n statement variables, let us first arrange them in some fixed order. If capital letters are used to denote the variables, then they may be arranged in alphabetical order. If subscripted letters are also used, then the following is an illustration of the order that may be used:

$$A, B, \dots, Z, A_1, B_1, \dots, Z_1, A_2, B_2, \dots$$

As an example, if the variables are P_1, Q, R_3, S_1, T_2 , and Q_3 , then they may be arranged as

$$Q, P_1, S_1, T_2, Q_3, R_3$$

Once the variables have been arranged in a particular order, it is possible to designate them as the first variable, second variable, and so on.

Let us assume that n variables are given and are arranged in a particular order. The 2^n minterms corresponding to the n variables can be designated by $m_0, m_1, \dots, m_{2^n-1}$. If we write the subscript of any particular minterm in binary and add an appropriate number of zeros on the left (if necessary) so that the number of digits in the subscript is exactly n , then we can obtain the corresponding minterm in the following manner. If in the i th location from the left there

Table 1-3.3

P	Q	R	A
T	T	T	F
T	T	F	F
T	F	T	T
T	F	F	F
F	T	T	T
F	T	F	T
F	F	T	F
F	F	F	T

appears 1, then the i th variable appears in the conjunction. If 0 appears in the i th location from the left, then the negation of the i th variable appears in the conjunction forming the minterm. Thus each of $m_0, m_1, \dots, m_{2^n-1}$ corresponds to a unique minterm, which can be determined from the binary representation of the subscript. Conversely, given any minterm, one can find which of $m_0, m_1, \dots, m_{2^n-1}$ designates it.

Let P, Q , and R be three variables arranged in that order. The corresponding minterms are denoted by $m_0, m_1, \dots, m_7$. We can write the subscript 5 in binary as 101, and the minterm m_5 is $P \wedge \neg Q \wedge R$. Similarly m_0 corresponds to $\neg P \wedge \neg Q \wedge \neg R$. To obtain the minterm m_3 , we write 3 as 11 and append a zero on the left to get 011, and m_3 is $\neg P \wedge Q \wedge R$.

If we have six variables $P_1, P_2, \dots, P_6$, then there are $2^6 = 64$ minterms denoted by $m_0, m_1, \dots, m_{63}$. To get a minterm, m_{38} say, we write 38 in binary as 100110; then the minterm m_{38} is $P_1 \wedge \neg P_2 \wedge \neg P_3 \wedge P_4 \wedge P_5 \wedge \neg P_6$.

Having developed a notation for the representation of the minterms, which can be further simplified by writing only the subscripts of $m_0, m_1, \dots, m_{2^n-1}$, we designate the disjunction (sum) of minterms by the compact notation $\sum$. Using such a notation, the sum-of-products canonical form representing the disjunction of m_i, m_j , and m_k can be written down as $\sum i, j, k$. As an example, it is known that

$$(P \wedge Q) \vee (\neg P \wedge R) \vee (Q \wedge R) \Leftrightarrow (\neg P \wedge \neg Q \wedge R) \\ \vee (\neg P \wedge Q \wedge R) \vee (P \wedge Q \wedge \neg R) \vee (P \wedge Q \wedge R)$$

Thus we denote the principal disjunctive normal form of

$$(P \wedge Q) \vee (\neg P \wedge R) \vee (Q \wedge R)$$

as $\sum 1,3,6,7$. With this type of representation and simplification of notation, it is easy to compare two principal disjunctive normal forms.

We now proceed to obtain a similar representation for the product-of-sums (principal conjunctive normal) forms. We denote the maxterms of n variables by $M_0, M_1, \dots, M_{2^n-1}$. Again, the maxterm corresponding to M_j , say, is obtained by writing j in binary and appending the required number of zeros to the left in order to get n digits. If 0 appears in the i th location from the left of this binary number, then the i th variable appears in the disjunction, while if 1 appears in the i th location, then the negation of the i th variable appears. Thus the binary representation of the subscript uniquely determines the maxterm, and, conversely, every binary representation of numbers between 0 and $2^n - 1$ determines a maxterm. Note that the convention regarding 1 and 0 here is the opposite of what was used for minterms. This convention is adopted with a view to connect the two principal normal forms of any given formula.

The maxterms, $M_0, M_1, \dots, M_7$, corresponding to three variables P, Q , and R , are

$$\begin{array}{cccc} P \vee Q \vee R & P \vee Q \vee \neg R & P \vee \neg Q \vee R & P \vee \neg Q \vee \neg R \\ \neg P \vee Q \vee R & \neg P \vee Q \vee \neg R & \neg P \vee \neg Q \vee R & \neg P \vee \neg Q \vee \neg R \end{array}$$

As before, further simplification is introduced by using $\prod$ to denote the conjunction (product) of maxterms. Thus $\prod i,j,k$ represents the conjunction of maxterms M_i, M_j, M_k .

To illustrate this discussion, we consider $(P \wedge Q) \vee (\neg P \wedge R)$. We obtain its principal conjunctive normal form as follows.

$$\begin{aligned} (P \wedge Q) \vee (\neg P \wedge R) \\ &\Leftrightarrow ((P \wedge Q) \vee \neg P) \wedge ((P \wedge Q) \vee R) \\ &\Leftrightarrow (P \vee \neg P) \wedge (Q \vee \neg P) \wedge (P \vee R) \wedge (Q \vee R) \\ &\Leftrightarrow (Q \vee \neg P \vee (R \wedge \neg R)) \wedge (P \vee R \vee (Q \wedge \neg Q)) \\ &\quad \wedge (Q \vee R \vee (P \wedge \neg P)) \\ &\Leftrightarrow (Q \vee \neg P \vee R) \wedge (Q \vee \neg P \vee \neg R) \wedge (P \vee R \vee Q) \\ &\quad \wedge (P \vee R \vee \neg Q) \wedge (Q \vee R \vee P) \wedge (Q \vee R \vee \neg P) \\ &\Leftrightarrow (\neg P \vee Q \vee R) \wedge (\neg P \vee Q \vee \neg R) \wedge (P \vee Q \vee R) \\ &\quad \wedge (P \vee \neg Q \vee R) \end{aligned}$$

Thus the product-of-sums canonical form of $(P \wedge Q) \vee (\neg P \wedge R)$ can be represented as $\prod 0,2,4,5$. Note that its disjunctive normal form is

$$(P \wedge Q \wedge R) \vee (P \wedge Q \wedge \neg R) \vee (\neg P \wedge Q \wedge R) \\ \vee (\neg P \wedge \neg Q \wedge R) \Leftrightarrow \sum 1,3,6,7$$

More generally, given any formula containing n variables and using the above notations to represent the equivalent principal disjunctive and conjunctive normal forms, we see clearly that all numbers lying between 0 and $2^n - 1$ which do not appear in one normal form will appear in the other. This statement follows from the principle of duality and the discussion given earlier regarding the relation between these two principal normal forms.

EXERCISES 1-3.5

1 Write equivalent forms for the following formulas in which negations are applied to the variables only.

- (a) $\neg(P \vee Q)$ (d) $\neg(P \nleftrightarrow Q)$
 (b) $\neg(P \wedge Q)$ (e) $\neg(P \uparrow Q)$
 (c) $\neg(P \rightarrow Q)$ (f) $\neg(P \downarrow Q)$

Obtain the principal conjunctive normal forms of (a), (c), and (d).

2 Obtain the product-of-sums canonical forms of the following formulas.

- (a) $(P \wedge Q \wedge R) \vee (\neg P \wedge R \wedge Q) \vee (\neg P \wedge \neg Q \wedge \neg R)$
 (b) $(\neg S \wedge \neg P \wedge R \wedge Q) \vee (S \wedge P \wedge \neg R \wedge \neg Q) \vee (\neg S \wedge P \wedge R \wedge \neg Q) \vee$
 $(Q \wedge \neg P \wedge \neg R \wedge S) \vee (P \wedge \neg S \wedge \neg R \wedge Q)$
 (c) $(P \wedge Q) \vee (\neg P \wedge Q) \vee (P \wedge \neg Q)$
 (d) $(P \wedge Q) \vee (\neg P \wedge Q \wedge R)$

3 Obtain the principal disjunctive and conjunctive normal forms of the following formulas.

- (a) $(\neg P \vee \neg Q) \rightarrow (P \nleftrightarrow \neg Q)$ (d) $(P \rightarrow (Q \wedge R)) \wedge (\neg P \rightarrow (\neg Q \wedge \neg R))$
 (b) $Q \wedge (P \vee \neg Q)$ (e) $P \rightarrow (P \wedge (Q \rightarrow P))$
 (c) $P \vee (\neg P \rightarrow (Q \vee (\neg Q \rightarrow R)))$ (f) $(Q \rightarrow P) \wedge (\neg P \wedge Q)$

Which of the above formulas are tautologies?

Validity Using Truth Tables

Let A and B be two statement formulas. We say that " B logically follows from A " or " B is a valid conclusion (consequence) of the premise A " iff $A \rightarrow B$ is a tautology, that is, $A \Rightarrow B$.

Just as the definition of implication was extended to include a set of formulas rather than a single formula, we say that from a set of premises $\{H_1, H_2, \dots, H_m\}$ a conclusion C follows logically iff

$$H_1 \wedge H_2 \wedge \dots \wedge H_m \Rightarrow C \quad (1)$$

Given a set of premises and a conclusion, it is possible to determine whether the conclusion logically follows (we shall simply say "follows") from the given premises by constructing truth tables as follows.

Let $P_1, P_2, \dots, P_n$ be all the atomic variables appearing in the premises $H_1, H_2, \dots, H_m$ and the conclusion C . If all possible combinations of truth values are assigned to $P_1, P_2, \dots, P_n$ and if the truth values of $H_1, H_2, \dots, H_m$ and C are entered in a table, then it is easy to see from such a table whether (1) is true. We look for the rows in which all $H_1, H_2, \dots, H_m$ have the value T . If, for every such row, C also has the value T , then (1) holds. Alternatively, we may look for the rows in which C has the value F . If, in every such row, at least one of the values of $H_1, H_2, \dots, H_m$ is F , then (1) also holds. We call such a method a "truth table technique" for the determination of the validity of a conclusion and demonstrate this technique by examples.

EXAMPLE 1 Determine whether the conclusion C follows logically from the premises H_1 and H_2 .

- (a) $H_1: P \rightarrow Q \quad H_2: P \quad C: Q$
- (b) $H_1: P \rightarrow Q \quad H_2: \neg P \quad C: Q$
- (c) $H_1: P \rightarrow Q \quad H_2: \neg(P \wedge Q) \quad C: \neg P$
- (d) $H_1: \neg P \quad H_2: P \rightleftharpoons Q \quad C: \neg(P \wedge Q)$
- (e) $H_1: P \rightarrow Q \quad H_2: Q \quad C: P$

SOLUTION We first construct the appropriate truth table, as shown in Table 1.4.1. For (a) we observe that the first row is the only row in which both

Table 1-4.1

P	Q	$P \rightarrow Q$	$\neg P$	$\neg Q$	$\neg(P \wedge Q)$	$P \rightleftharpoons Q$
T	T	T	F	F	F	T
T	F	F	F	T	T	F
F	T	T	T	F	T	F
F	F	T	T	T	T	T

the premises have the value T . The conclusion also has the value T in that row. Hence it is valid. In (b) observe the third and fourth rows. The conclusion Q is true only in the third row, but not in the fourth, and hence the conclusion is not valid. Similarly, we can show that the conclusions are valid in (c) and (d) but not in (e).

The conclusion P in (e) does not follow logically from the premises $P \rightarrow Q$ and Q , no matter which statements in English are translated as P and Q or what the truth value of the conclusion P may be. As a particular case, consider the argument

H_1 : If Canada is a country, then New York is a city.: ($P \rightarrow Q$)
 H_2 : New York is a city.: (Q)
 Conclusion C : Canada is a country.: (P)

Note that both the premises and the conclusion have the truth value T . However, the conclusion does not follow logically from the premises. This example is chosen to emphasize the fact that we are not so much concerned with the conclusion's being true or false as we are with determining whether the conclusion follows from the premises, i.e., whether the argument is valid or invalid. ////

Theoretically, it is possible to determine in a finite number of steps whether a conclusion follows from a given set of premises by constructing the appropriate truth table. However, this method becomes tedious when the number of atomic variables present in all the formulas representing the premises and conclusion is large. This disadvantage, coupled with the fact that the inference theory is applicable in more general situations where the truth table technique is no longer available, suggests that we should investigate other possible methods. This investigating will be done in the following sections.

EXERCISES 1-4.1

1 Show that the conclusion C follows from the premises $H_1, H_2, \dots$ in the following cases.

$$(a) H_1: P \rightarrow Q \quad C: P \rightarrow (P \wedge Q)$$

$$(b) H_1: \neg P \vee Q \quad H_2: \neg(Q \wedge \neg R) \quad H_3: \neg R \quad C: \neg P$$

$$(c) H_1: \neg P \quad H_2: P \vee Q \quad C: Q$$

$$(d) H_1: \neg Q \quad H_2: P \rightarrow Q \quad C: \neg P$$

$$(e) H_1: P \rightarrow Q \quad H_2: Q \rightarrow R \quad C: P \rightarrow R$$

$$(f) H_1: R \quad H_2: P \vee \neg P \quad C: R$$

2 Determine whether the conclusion C is valid in the following, when $H_1, H_2, \dots$ are the premises.

$$(a) H_1: P \rightarrow Q \quad H_2: \neg Q \quad C: P$$

$$(b) H_1: P \vee Q \quad H_2: P \rightarrow R \quad H_3: Q \rightarrow R \quad C: R$$

$$(c) H_1: P \rightarrow (Q \rightarrow R) \quad H_2: P \wedge Q \quad C: R$$

$$(d) H_1: P \rightarrow (Q \rightarrow R) \quad H_2: R \quad C: P$$

$$(e) H_1: \neg P \quad H_2: P \vee Q \quad C: P \wedge Q$$

3 Without constructing a truth table, show that $A \wedge E$ is not a valid consequence of

$$A \rightleftharpoons B \quad B \rightleftharpoons (C \wedge D) \quad C \rightleftharpoons (A \vee E) \quad A \vee E$$

Also show that $A \vee C$ is not a valid consequence of

$$A \rightleftharpoons (B \rightarrow C) \quad B \rightleftharpoons (\neg A \vee \neg C) \quad C \rightleftharpoons (A \vee \neg B) \quad B$$

4 Show that $L \vee M$ follows from

$$P \wedge Q \wedge R \quad (Q \rightleftharpoons R) \rightarrow (L \vee M)$$

5 Show without constructing truth tables that the following statements cannot all be true simultaneously.

$$(a) P \rightleftharpoons Q \quad Q \rightarrow R \quad \neg R \vee S \quad \neg P \rightarrow S \quad \neg S$$

$$(b) R \vee M \quad \neg R \vee S \quad \neg M \quad \neg S$$

Rules of Inference

We now describe the process of derivation by which one demonstrates that a particular formula is a valid consequence of a given set of premises. Before we do this, we give two rules of inference which are called rules **P** and **T**.

Rule P: A premise may be introduced at any point in the derivation.

Rule T: A formula S may be introduced in a derivation if S is tautologically implied by any one or more of the preceding formulas in the derivation.

Before we proceed with the actual process of derivation, we list some important implications and equivalences that will be referred to frequently.

Not all the implications and equivalences listed in Tables 1-4.2 and 1-4.3 respectively are independent of one another. One could start with only a minimum number of them and derive the others by using the above rules of inference. Such an axiomatic approach will not be followed. We list here most of the important implications and equivalences and show how some of them are used in a derivation. Those which are used more often than the others are given special names because of their importance.

EXAMPLE 1 Demonstrate that R is a valid inference from the premises $P \rightarrow Q$, $Q \rightarrow R$, and P .

SOLUTION

{1}	(1)	$P \rightarrow Q$	Rule P
{2}	(2)	P	Rule P
{1, 2}	(3)	Q	Rule T, (1), (2), and I_{11} (modus ponens)
{4}	(4)	$Q \rightarrow R$	Rule P
{1, 2, 4}	(5)	R	Rule T, (3), (4), and I_{11}

The second column of numbers designates the formula as well as the line of derivation in which it occurs. The set of numbers in braces (the first column) for each line shows the premises on which the formula in the line depends. On the right, **P** or **T** represents the rule of inference, followed by a comment showing from which formulas and tautology that particular formula has been obtained. For example, if we follow this notation, the third line shows that the formula in this line is numbered (3) and has been obtained from premises in (1) and (2). The comment on the right says that the formula Q has been introduced using rule **T** and also indicates the details of the application of rule **T**. ////

Table 1-4.2 IMPLICATIONS

I_1	$P \wedge Q \Rightarrow P$	(simplification)
I_2	$P \wedge Q \Rightarrow Q$	
I_3	$P \Rightarrow P \vee Q$	(addition)
I_4	$Q \Rightarrow P \vee Q$	
I_5	$\neg P \Rightarrow P \rightarrow Q$	
I_6	$Q \Rightarrow P \rightarrow Q$	
I_7	$\neg(P \rightarrow Q) \Rightarrow P$	
I_8	$\neg(P \rightarrow Q) \Rightarrow \neg Q$	
I_9	$P, Q \Rightarrow P \wedge Q$	
I_{10}	$\neg P, P \vee Q \Rightarrow Q$	(disjunctive syllogism)
I_{11}	$P, P \rightarrow Q \Rightarrow Q$	(modus ponens)
I_{12}	$\neg Q, P \rightarrow Q \Rightarrow \neg P$	(modus tollens)
I_{13}	$P \rightarrow Q, Q \rightarrow R \Rightarrow P \rightarrow R$	(hypothetical syllogism)
I_{14}	$P \vee Q, P \rightarrow R, Q \rightarrow R \Rightarrow R$	(dilemma)

Table 1-4.3 EQUIVALENCES

E_1	$\neg \neg P \Leftrightarrow P$	(double negation)
E_2	$P \wedge Q \Leftrightarrow Q \wedge P$	(commutative laws)
E_3	$P \vee Q \Leftrightarrow Q \vee P$	
E_4	$(P \wedge Q) \wedge R \Leftrightarrow P \wedge (Q \wedge R)$	(associative laws)
E_5	$(P \vee Q) \vee R \Leftrightarrow P \vee (Q \vee R)$	
E_6	$P \wedge (Q \vee R) \Leftrightarrow (P \wedge Q) \vee (P \wedge R)$	(distributive laws)
E_7	$P \vee (Q \wedge R) \Leftrightarrow (P \vee Q) \wedge (P \vee R)$	
E_8	$\neg(P \wedge Q) \Leftrightarrow \neg P \vee \neg Q$	(De Morgan's laws)
E_9	$\neg(P \vee Q) \Leftrightarrow \neg P \wedge \neg Q$	
E_{10}	$P \vee P \Leftrightarrow P$	
E_{11}	$P \wedge P \Leftrightarrow P$	
E_{12}	$R \vee (P \wedge \neg P) \Leftrightarrow R$	
E_{13}	$R \wedge (P \vee \neg P) \Leftrightarrow R$	
E_{14}	$R \vee (P \vee \neg P) \Leftrightarrow T$	
E_{15}	$R \wedge (P \wedge \neg P) \Leftrightarrow F$	
E_{16}	$P \rightarrow Q \Leftrightarrow \neg P \vee Q$	
E_{17}	$\neg(P \rightarrow Q) \Leftrightarrow P \wedge \neg Q$	
E_{18}	$P \rightarrow Q \Leftrightarrow \neg Q \rightarrow \neg P$	
E_{19}	$P \rightarrow (Q \rightarrow R) \Leftrightarrow (P \wedge Q) \rightarrow R$	
E_{20}	$\neg(P \Leftrightarrow Q) \Leftrightarrow P \Leftrightarrow \neg Q$	
E_{21}	$P \Leftrightarrow Q \Leftrightarrow (P \rightarrow Q) \wedge (Q \rightarrow P)$	
E_{22}	$(P \Leftrightarrow Q) \Leftrightarrow (P \wedge Q) \vee (\neg P \wedge \neg Q)$	

EXAMPLE 2 Show that $R \vee S$ follows logically from the premises $C \vee D$, $(C \vee D) \rightarrow \neg H$, $\neg H \rightarrow (A \wedge \neg B)$, and $(A \wedge \neg B) \rightarrow (R \vee S)$.

SOLUTION

{1}	(1)	$(C \vee D) \rightarrow \neg H$	P
{2}	(2)	$\neg H \rightarrow (A \wedge \neg B)$	P
{1, 2}	(3)	$(C \vee D) \rightarrow (A \wedge \neg B)$	T, (1), (2), and I_{12}
{4}	(4)	$(A \wedge \neg B) \rightarrow (R \vee S)$	P
{1, 2, 4}	(5)	$(C \vee D) \rightarrow (R \vee S)$	T, (3), (4), and I_{13}
{6}	(6)	$C \vee D$	P
{1, 2, 4, 6}	(7)	$R \vee S$	T, (5), (6), and I_{11}

The two tautologies frequently used in the above derivations are I_{13} , known as hypothetical syllogism, and I_{11} , known as modus ponens. ////

EXAMPLE 3 Show that $S \vee R$ is tautologically implied by $(P \vee Q) \wedge (P \rightarrow R) \wedge (Q \rightarrow S)$.

SOLUTION

{1}	(1)	$P \vee Q$	P
{1}	(2)	$\neg P \rightarrow Q$	T, (1), E_1 , and E_{18}
{3}	(3)	$Q \rightarrow S$	P
{1, 3}	(4)	$\neg P \rightarrow S$	T, (2), (3), and I_{13}
{1, 3}	(5)	$\neg S \rightarrow P$	T, (4), E_{18} , and E_1
{6}	(6)	$P \rightarrow R$	P
{1, 3, 6}	(7)	$\neg S \rightarrow R$	T, (5), (6), and I_{13}
{1, 3, 6}	(8)	$S \vee R$	T, (7), E_{16} , and E_1 ////

EXAMPLE 4 Show that $R \wedge (P \vee Q)$ is a valid conclusion from the premises $P \vee Q$, $Q \rightarrow R$, $P \rightarrow M$, and $\neg M$.

SOLUTION

{1}	(1)	$P \rightarrow M$	P
{2}	(2)	$\neg M$	P
{1, 2}	(3)	$\neg P$	T, (1), (2), and I_{12}
{4}	(4)	$P \vee Q$	P
{1, 2, 4}	(5)	Q	T, (3), (4), and I_{10}
{6}	(6)	$Q \rightarrow R$	P
{1, 2, 4, 6}	(7)	R	T, (5), (6), and I_{11}
{1, 2, 4, 6}	(8)	$R \wedge (P \vee Q)$	T, (4), (7), and I_9 ////

EXAMPLE 5 Show I_{12} : $\neg Q, P \rightarrow Q \Rightarrow \neg P$.

SOLUTION

{1}	(1)	$P \rightarrow Q$	P	
{1}	(2)	$\neg Q \rightarrow \neg P$	T , (1), and E_{18}	
{3}	(3)	$\neg Q$	P	
{1, 3}	(4)	$\neg P$	T , (2), (3), and I_{11}	////

We shall now introduce a third inference rule, known as rule **CP** or rule of conditional proof.

Rule CP If we can derive S from R and a set of premises, then we can derive $R \rightarrow S$ from the set of premises alone.

Rule **CP** is not new for our purpose here because it follows from the equivalence E_{19} which states that

$$(P \wedge R) \rightarrow S \Leftrightarrow P \rightarrow (R \rightarrow S)$$

Let P denote the conjunction of the set of premises and let R be any formula. The above equivalence states that if R is included as an additional premise and S is derived from $P \wedge R$, then $R \rightarrow S$ can be derived from the premises P alone.

Rule **CP** is also called the *deduction theorem* and is generally used if the conclusion is of the form $R \rightarrow S$. In such cases, R is taken as an additional premise and S is derived from the given premises and R .

EXAMPLE 6 Show that $R \rightarrow S$ can be derived from the premises $P \rightarrow (Q \rightarrow S)$, $\neg R \vee P$, and Q .

SOLUTION Instead of deriving $R \rightarrow S$, we shall include R as an additional premise and show S first.

{1}	(1)	$\neg R \vee P$	P	
{2}	(2)	R	P (assumed premise)	
{1, 2}	(3)	P	T , (1), (2), and I_{10}	
{4}	(4)	$P \rightarrow (Q \rightarrow S)$	P	
{1, 2, 4}	(5)	$Q \rightarrow S$	T , (3), (4), and I_{11}	
{6}	(6)	Q	P	
{1, 2, 4, 6}	(7)	S	T , (5), (6), and I_{11}	
{1, 4, 6}	(8)	$R \rightarrow S$	CP	////

These examples show that a derivation consists of a sequence of formulas, each formula in the sequence being either a premise or tautologically implied by formulas appearing before.

In Sec. 1-3.1 we discussed the decision problem in terms of determining whether a given formula is a tautology. We can extend this notion to the determination of validity of arguments. Accordingly, if one can determine in a finite number of steps whether an argument is valid, then the decision problem for validity is solvable.

One solution to the decision problem for validity is provided by the truth table method. Use of this method is often not practical. The method of derivation just discussed provides only a partial solution to the decision problem, because if an argument is valid, then it is possible to show by this method that the argument is valid. On the other hand, if an argument is not valid, then it is very difficult to decide, after a finite number of steps, that this is the case. There are other methods of derivation which do allow one to determine, after a finite number of steps, whether an argument is or is not valid. One such method is described in Sec. 1-4.4, and its computer implementation is given later in Sec. 2-7.

We shall now give some examples of derivation involving statements in English. We first symbolize the given statements and then use the method of derivation just discussed.

EXAMPLE 7 “If there was a ball game, then traveling was difficult. If they arrived on time, then traveling was not difficult. They arrived on time. Therefore, there was no ball game.” Show that these statements constitute a valid argument.

SOLUTION Let

P : There was a ball game.

Q : Traveling was difficult.

R : They arrived on time.

We are required to show that from the premises $P \rightarrow Q$, $R \rightarrow \neg Q$, and R the conclusion $\neg P$ follows. (Complete the rest of the derivation.) /////

EXAMPLE 8 If A works hard, then either B or C will enjoy themselves. If B enjoys himself, then A will not work hard. If D enjoys himself, then C will not. Therefore, if A works hard, D will not enjoy himself.

SOLUTION Let A : A works hard; B : B will enjoy himself; C : C will enjoy himself; D : D will enjoy himself. Show that $A \rightarrow \neg D$ follows from $A \rightarrow B \vee C$, $B \rightarrow \neg A$, and $D \rightarrow \neg C$. Since the conclusion is given in the form of a condition $A \rightarrow \neg D$, include A as an additional premise and show that $\neg D$ follows logically from all the premises including A . Finally, use rule **CP** to obtain the result. /////

UNIT – III

THE PREDICATE CALCULUS

So far our discussion of symbolic logic has been limited to the consideration of statements and statement formulas. The inference theory was also restricted in the sense that the premises and conclusions were all statements. The symbols $P, Q, R, \dots, P_1, Q_1, \dots$ were used for statements or statement variables. The statements were taken as basic units of statement calculus, and no analysis of any atomic statement was admitted. Only compound formulas were analyzed, and this analysis was done by studying the forms of the compound formulas, i.e., the connections between the constituent atomic statements. It was not possible to express the fact that any two atomic statements have some features in common. In order to investigate questions of this nature, we introduce the concept of a predicate in an atomic statement. The logic based upon the analysis of predicates in any statement is called *predicate logic*.

Predicates

Let us first consider the two statements

John is a bachelor.

Smith is a bachelor.

Obviously, if we express these statements by symbols, we require two different symbols to denote them. Such symbols do not reveal the common features of these two statements; viz., both are statements about two different individuals who are bachelors. If we introduce some symbol to denote “is a bachelor” and a method to join it with symbols denoting the names of individuals, then we will have a symbolism to denote statements about any individual’s being a bachelor. The part “is a bachelor” is called a *predicate*. Another consideration which leads to some similar device for the representation of statements is suggested by the following argument.

All human beings are mortal.

John is a human being.

Therefore, John is a mortal.

Such a conclusion seems intuitively true. However, it does not follow from the inference theory of the statement calculus developed earlier. The reason for this deficiency is the fact that the statement “All human beings are mortal” cannot be analyzed to say anything about an individual. If we could separate the part “are mortal” from the part “All human beings,” then it might be possible to consider any particular human being.

We shall symbolize a predicate by a capital letter and the names of individuals or objects in general by small letters. We shall soon see that using capital letters to symbolize statements as well as predicates will not lead to any confusion. Every predicate describes something about one or more objects (the word "object" is being used in a very general sense to include individuals as well). Therefore, a statement could be written symbolically in terms of the predicate letter followed by the name or names of the objects to which the predicate is applied.

We again consider the statements

- 1 John is a bachelor.
- 2 Smith is a bachelor.

Denote the predicate "is a bachelor" symbolically by the predicate letter B , "John" by j , and "Smith" by s . Then Statements (1) and (2) can be written as $B(j)$ and $B(s)$ respectively. In general, any statement of the type " p is Q " where Q is a predicate and p is the subject can be denoted by $Q(p)$.

A statement which is expressed by using a predicate letter must have at least one name of an object associated with the predicate. When an appropriate number of names are associated with a predicate, then we get a statement. Using a capital letter to denote a predicate may not indicate the appropriate number of names associated with it. Normally, this number is clear from the context or from the notation being used. This numbering can also be accomplished by at-

taching a superscript to a predicate letter, indicating the number of names that are to be appended to the letter. A predicate requiring m ($m > 0$) names is called an m -place predicate. For example, B in (1) and (2) is a 1-place predicate. Another example is that " L : is less than" is a 2-place predicate. In order to extend our definition to $m = 0$, we shall call a statement a 0-place predicate because no names are associated with a statement.

Let R denote the predicate "is red" and let p denote "This painting." Then the statement

- 3 This painting is red.

can be symbolized by $R(p)$. Further, the connectives described earlier can now be used to form compound statements such as "John is a bachelor, and this painting is red," which can be written as $B(j) \wedge R(p)$. Other connectives can also be used to form statements such as

$$B(j) \rightarrow R(p) \quad \neg R(p) \quad B(j) \vee R(p) \quad \text{etc.}$$

Consider, now, statements involving the names of two objects, such as

- 4 Jack is taller than Jill.
- 5 Canada is to the north of the United States.

The predicates "is taller than" and "is to the north of" are 2-place predicates because names of two objects are needed to complete a statement involving these predicates. If the letter G symbolizes "is taller than," j_1 denotes "Jack," and j_2 denotes "Jill," then Statement (4) can be translated as $G(j_1, j_2)$. Note that the order in which the names appear in the statement as well as in the predicate is important. Similarly, if N denotes the predicate "is to the north of," c : Canada, and s : United States, then (5) is symbolized as $N(c, s)$. Obviously, $N(s, c)$ is the statement "The United States is to the north of Canada."

Examples of 3-place predicates and 4-place predicates are:

6 Susan sits between Ralph and Bill.

7 Green and Miller played bridge against Johnston and Smith.

In general, an n -place predicate requires n names of objects to be inserted in fixed positions in order to obtain a statement. The position of these names is important. If S is an n -place predicate letter and $a_1, a_2, \dots, a_n$ are the names of objects, then $S(a_1, a_2, \dots, a_n)$ is a statement. If we use this convention, every predicate symbol is followed by an appropriate number of letters, which are the names of objects, enclosed in parentheses and separated by commas. Occasionally, the parentheses and the commas are dropped. The definition does not require that the names be chosen from any fixed set. For example, if B denotes the predicate "is a bachelor" and t denotes "This table," then $B(t)$ symbolizes "This table is a bachelor." In our everyday language, the only admissible name in this context would be that of an individual. However, such restrictions are not necessary according to the rules given above. We show a method of imposing such restrictions in Sec. 1-5.5.

The Statement Function, Variables, and Quantifiers

Let H be the predicate "is a mortal," b the name "Mr. Brown," c "Canada," and s "A shirt." Then $H(b)$, $H(c)$, and $H(s)$ all denote statements. In fact, these statements have a common form. If we write $H(x)$ for " x is mortal," then $H(b)$, $H(c)$, $H(s)$, and others having the same form can be obtained from $H(x)$ by replacing x by an appropriate name. Note that $H(x)$ is not a statement, but it results in a statement when x is replaced by the name of an object. The letter x used here is a placeholder. From now on we shall use small letters as individual or object variables as well as names of objects.

A *simple statement function* of one variable is defined to be an expression consisting of a predicate symbol and an individual variable. Such a statement function becomes a statement when the variable is replaced by the name of any object. The statement resulting from a replacement is called a *substitution instance* of the statement function and is a formula of statement calculus.

The word "simple" in the above definition distinguishes the simple statement function from those statement functions which are obtained from combining one or more simple statement functions and the logical connectives. For example, if we let $M(x)$ be " x is a man" and $H(x)$ be " x is a mortal," then we can form *compound statement functions* such as

$$M(x) \wedge H(x) \quad M(x) \rightarrow H(x) \quad \neg H(x) \quad M(x) \vee \neg H(x) \quad \text{etc.}$$

An extension of this idea to the statement functions of two or more variables is straightforward. Consider, for example, the statement function of two variables:

$$1 \quad G(x, y): x \text{ is taller than } y.$$

If both x and y are replaced by the names of objects, we get a statement. If m represents Mr. Miller and f Mr. Fox, then we have

$$G(m, f): \text{Mr. Miller is taller than Mr. Fox.}$$

and

$$G(f, m): \text{Mr. Fox is taller than Mr. Miller.}$$

It is possible to form statement functions of two variables by using statement functions of one variable. For example, given

$$M(x): x \text{ is a man.}$$

$$H(y): y \text{ is a mortal.}$$

then we may write

$$M(x) \wedge H(y): x \text{ is a man and } y \text{ is a mortal.}$$

It is not possible, however, to write every statement function of two variables using statement functions of one variable.

One way of obtaining statements from any statement function is to replace the variables by the names of objects. There is another way in which statements can be obtained. In order to understand this alternative method, we first consider some familiar equations in elementary algebra.

$$2 \quad x + 2 = 5$$

$$3 \quad x^2 + 1 = 0$$

$$4 \quad (x - 1) * (x - \frac{1}{2}) = 0$$

$$5 \quad x^2 - 1 = (x - 1) * (x + 1)$$

In algebra, it is conventional to assume that the variable x is to be replaced by numbers (real, complex, rational, integer, etc.). In the above equations, we would not normally consider substituting for x the name of a person or object instead of numbers. We may state this idea by saying that the *universe* of the variable x is the set of real numbers or complex numbers or integers, etc. The restriction depends upon the problem under consideration. For example, we may be interested in only the real solution or the positive solution in a particular case. In Statement (2), if x is replaced by a real number, we get a statement. The resulting statement is true when 3 is substituted for x , while, for every other substitution, the resulting statement is false. In (3) there is no real number which, when substituted for x , gives a true statement. If, however, the universe of x includes complex numbers as well, then we find that there are two substitution

instances which give true statements. In (4), if the universe of x is assumed to be integers, then there is only one number which produces a true statement when substituted. The situation is slightly different in (5) in the sense that if any number is substituted for x , then the resulting statement is true. Therefore, we may say that

$$6 \quad \text{For any number } x, x^2 - 1 = (x - 1) * (x + 1).$$

Note that (6) is a statement and not a statement function even though a variable x appears in it. In fact, the addition of the phrase "For any number x ," has changed the situation. The letter x , as used in (6), is different from the variable x used in Statements (2) to (5). In (6) the variable x need not be replaced by any name to obtain a statement. In mathematics this distinction is often not made. Occasionally when a statement involves an equality, a distinction is made by using the symbol $\equiv$ instead of the equality sign to show that it is a statement. In this case, (6) would be written as $x^2 - 1 \equiv (x - 1) * (x + 1)$. A similar situation occurs when a statement function does not involve an equality, and a distinction is necessary in logic between these two different uses of the variables.

Let us first consider the following statements. Each one is a statement about all individuals or objects belonging to a certain set.

- 7 All men are mortal.
- 8 Every apple is red.
- 9 Any integer is either positive or negative.

Let us paraphrase these in the following manner.

- 7a For all x , if x is a man, then x is a mortal.
- 8a For all x , if x is an apple, then x is red.
- 9a For all x , if x is a integer, then x is either positive or negative.

We have already shown how statement functions such as " x is a man," " x is an apple," or " x is red" can be written by using predicate symbols. If we introduce a symbol to denote the phrase "For all x ," then it would be possible to symbolize Statements (7a), (8a), and (9a).

We symbolize "For all x " by the symbol " $(\forall x)$ " or by " (x) " with an understanding that this symbol be placed before the statement function to which this phrase is applied. Using

$$M(x) : x \text{ is man.} \quad H(x) : x \text{ is a mortal.}$$

$$A(x) : x \text{ is an apple.} \quad R(x) : x \text{ is red.}$$

$$N(x) : x \text{ is an integer.} \quad P(x) : x \text{ is either positive or negative.}$$

we write (7a), (8a), and (9a) as

$$7b \quad (x)(M(x) \rightarrow H(x))$$

$$8b \quad (x)(A(x) \rightarrow R(x))$$

$$9b \quad (x)(N(x) \rightarrow P(x))$$

Sometimes $(x)(M(x) \rightarrow H(x))$ is also written as $(\forall x)(M(x) \rightarrow H(x))$. The symbols (x) or $(\forall x)$ are called *universal quantifiers*. Strictly speaking, the quantification symbol is “ $()$ ” or “ $(\forall)$,” and it contains the variable which is to be quantified. It is now possible for us to quantify any statement function of one variable to obtain a statement. Thus $(x)M(x)$ is a statement which can be translated as

- 10 For all x , x is a man.
- 10a For every x , x is a man.
- 10b Everything is a man.

In order to determine the truth values of any one of these statements involving a universal quantifier, one may be tempted to consider the truth values of the statement function which is quantified. This method is not possible for two reasons. First, statement functions do not have truth values. When the variables are replaced by the names of objects, we get statements which have a truth value. Second, in most cases there is an infinite number of statements that can be produced by such substitutions.

Note that the particular variable appearing in the statements involving a quantifier is not important because the statements remain unchanged if x is replaced by y throughout. Thus the statements $(x)(M(x) \rightarrow H(x))$ and $(y)(M(y) \rightarrow H(y))$ are equivalent.

Sometimes it is necessary to use more than one universal quantifier in a statement. For example consider

$$G(x, y): x \text{ is taller than } y.$$

We can state that “For any x and any y , if x is taller than y , then y is not taller than x ” or “For any x and y , if x is taller than y , then it is not true that y is taller than x .” This statement can now be symbolized as

$$(x)(y)(G(x, y) \rightarrow \neg G(y, x))$$

The universal quantifier was used to translate expressions such as “for all,” “every,” and “for any.” Another quantifier will now be introduced to symbolize expressions such as “for some,” “there is at least one,” or “there exists some” (note that “some” is used in the sense of “at least one”).

Consider the following statements:

- 11 There exists a man.
- 12 Some men are clever.
- 13 Some real numbers are rational.

The first statement can be expressed in various ways, two such ways being

- 11a There exists an x such that x is a man.
- 11b There is at least one x such that x is a man.

Similarly, (12) can be written as

12a There exists an x such that x is a man and x is clever.

12b There exists at least one x such that x is a man and x is clever.

Such a rephrasing allows us to introduce the symbol " $(\exists x)$," called the *existential quantifier*, which symbolizes expressions such as "there is at least one x such that" or "there exists an x such that" or "for some x ." Writing

$M(x)$: x is a man.

$C(x)$: x is clever.

$R_1(x)$: x is a real number.

$R_2(x)$: x is rational.

and using the existential quantifier, we can write the Statements (11) to (13) as

11c $(\exists x)(M(x))$

12c $(\exists x)(M(x) \wedge C(x))$

13c $(\exists x)(R_1(x) \wedge R_2(x))$

It may be noted that a conjunction has been used in writing the statements of the form "Some A are B ," while a conditional was used in writing statements of the form "All A are B ." To a beginner this usage may appear confusing. We show in Sec. 1-5.5 why these connectives are the right ones to be used in these cases.

Predicate Formulas

Recall that capital letters were first used to denote some definite statements. Subsequently they were used as placeholders for the statements, and, in this sense, they were called statement variables. These statement variables were also considered as special cases of statement formulas.

In Secs. 1-5.1 and 1-5.2 the capital letters were introduced as definite predicates. It was suggested that a superscript n be used along with the capital letters in order to indicate that the capital letter is used as an n -place predicate. However, this notation was not necessary because an n -place predicate symbol must be followed by n object variables. Such variables are called *object* or *individual variables* and are denoted by lowercase letters. When used as an n -place predicate, the capital letter is followed by n individual variables which are enclosed in parentheses and separated by commas. For example, $P(x_1, x_2, \dots, x_n)$ denotes an n -place predicate formula in which the letter P is an n -place predicate and

$x_1, x_2, \dots, x_n$ are individual variables. In general, $P(x_1, x_2, \dots, x_n)$ will be called an *atomic formula* of predicate calculus. It may be noted that our symbolism includes the atomic formulas of the statement calculus as special cases ($n = 0$). The following are some examples of atomic formulas.

R $Q(x)$ $P(x, y)$ $A(x, y, z)$ $P(a, y)$ and $A(x, a, z)$

A well-formed formula of predicate calculus is obtained by using the following rules.

- 1 An atomic formula is a well-formed formula.
- 2 If A is a well-formed formula, then $\neg A$ is a well-formed formula.
- 3 If A and B are well-formed formulas, then $(A \wedge B)$, $(A \vee B)$, $(A \rightarrow B)$, and $(A \rightleftharpoons B)$ are also well-formed formulas.
- 4 If A is a well-formed formula and x is any variable, then $(x)A$ and $(\exists x)A$ are well-formed formulas.
- 5 Only those formulas obtained by using rules (1) to (4) are well-formed formulas.

Since we will be concerned with only well-formed formulas, we shall use the term “formula” for “well-formed formula.” We shall follow the same conventions regarding the use of parentheses as was done in the case of statement formulas.

Free and Bound Variables

Given a formula containing a part of the form $(x)P(x)$ or $(\exists x)P(x)$, such a part is called an x -bound part of the formula. Any occurrence of x in an x -bound part of a formula is called a *bound occurrence* of x , while any occurrence of x or of any variable that is not a bound occurrence is called a *free occurrence*. Further, the formula $P(x)$ either in $(x)P(x)$ or in $(\exists x)P(x)$ is described as the *scope* of the quantifier. In other words, the scope of a quantifier is the formula immediately following the quantifier. If the scope is an atomic formula, then no parentheses are used to enclose the formula; otherwise parentheses are needed. As illustrations, consider the following formulas:

$$(x)P(x, y) \quad (1)$$

$$(x)(P(x) \rightarrow Q(x)) \quad (2)$$

$$(x)(P(x) \rightarrow (\exists y)R(x, y)) \quad (3)$$

$$(x)(P(x) \rightarrow R(x)) \vee (x)(P(x) \rightarrow Q(x)) \quad (4)$$

$$(\exists x)(P(x) \wedge Q(x)) \quad (5)$$

$$(\exists x)P(x) \wedge Q(x) \quad (6)$$

In (1), $P(x, y)$ is the scope of the quantifier, and both occurrences of x are bound occurrences, while the occurrence of y is a free occurrence. In (2), the scope of the universal quantifier is $P(x) \rightarrow Q(x)$, and all occurrences of x are bound. In (3), the scope of (x) is $P(x) \rightarrow (\exists y)R(x, y)$, while the scope of $(\exists y)$ is $R(x, y)$. All occurrences of both x and y are bound occurrences. In (4), the scope of the first quantifier is $P(x) \rightarrow R(x)$, and the scope of the second is

$P(x) \rightarrow Q(x)$. All occurrences of x are bound occurrences. In (5), the scope of $(\exists x)$ is $P(x) \wedge Q(x)$. However, in (6) the scope of $(\exists x)$ is $P(x)$, and the last occurrence of x in $Q(x)$ is free.

It is useful to note that in the bound occurrence of a variable, the letter which is used to represent the variable is not important. In fact, any other letter can be used as a variable without affecting the formula, provided that the new letter is not used elsewhere in the formula. Thus the formulas

$$(x)P(x, y) \quad \text{and} \quad (z)P(z, y)$$

are the same. Further, the bound occurrence of a variable cannot be substituted by a constant; only a free occurrence of a variable can be. For example, $(x)P(x) \wedge Q(a)$ is a substitution instance of $(x)P(x) \wedge Q(y)$. In fact, $(x)P(x) \wedge Q(a)$ can be expressed in English as "Every x has the property P , and a has the property Q ." A change of variables in the bound occurrence is not a substitution instance. Sometimes it is useful to change the variables in order to avoid confusion. In (6), it is better to write $(y)P(y) \wedge Q(x)$ instead of $(x)P(x) \wedge Q(x)$, so as to separate the free and bound occurrences of variables. Occasionally, one may come across a formula of the type $(x)P(y)$ in which the occurrence of y is free and the scope of (x) does not contain an x ; in such a case, we have a vacuous use of (x) . Finally, it may be mentioned that in a statement every occurrence of a variable must be bound, and no variable should have a free occurrence. In the case where a free variable occurs in a formula, then we have a statement function.

EXAMPLE 1 Let

$P(x) : x$ is a person.

$F(x, y) : x$ is the father of y .

$M(x, y) : x$ is the mother of y .

Write the predicate " x is the father of the mother of y ."

SOLUTION In order to symbolize the predicate, we name a person called z as the mother of y . Obviously we want to say that x is the father of z and z the mother of y . It is assumed that such a person z exists. We symbolize the predicate as

$$(\exists z)(P(z) \wedge F(x, z) \wedge M(z, y)) \quad \text{////}$$

EXAMPLE 2 Symbolize the expression "All the world loves a lover."

SOLUTION First note that the quotation really means that everybody loves a lover. Now let

$P(x) : x$ is a person.

$L(x) : x$ is a lover.

$R(x, y) : x$ loves y .

The required expression is

$$(x)(P(x) \rightarrow (y)(P(y) \wedge L(y) \rightarrow R(x, y))) \quad ////$$

The Universe of Discourse

Example 2 in Sec. 1-5.4 shows that the process of symbolizing a statement in predicate calculus can be quite complicated. However, some simplification can be introduced by limiting the class of individuals or objects under consideration. This limitation means that the variables which are quantified stand for only those objects which are members of a particular set or class. Such a restricted class is called the *universe of discourse* or the *domain* of individuals or simply the *universe*. If the discussion refers to human beings only, then the universe of discourse is the class of human beings. In elementary algebra or number theory, the universe of discourse could be numbers (real, complex, rational, etc.).

EXAMPLE 1 Symbolize the statement "All men are giants."

SOLUTION Using

$G(x)$: x is a giant.

$M(x)$: x is a man.

the given statement can be symbolized as $(x)(M(x) \rightarrow G(x))$. However, if we restrict the variable x to the universe which is the class of men, then the statement is

$$(x)G(x) \quad ////$$

EXAMPLE 2 Consider the statement "Given any positive integer, there is a greater positive integer." Symbolize this statement with and without using the set of positive integers as the universe of discourse.

SOLUTION Let the variables x and y be restricted to the set of positive integers. Then the above statement can be paraphrased as follows: For all x , there exists a y such that y is greater than x . If $G(x, y)$ is " x is greater than y ," then the given statement is $(x)(\exists y)G(y, x)$. If we do not impose the restriction on the universe of discourse and if we write $P(x)$ for " x is a positive integer," then we can symbolize the given statement as $(x)(P(x) \rightarrow (\exists y)(P(y) \wedge G(y, x)))$.

////

The universe of discourse, if any, must be explicitly stated, because the truth value of a statement depends upon it. For instance, consider the predicate

$Q(x)$: x is less than 5.

and the statements $(x)Q(x)$ and $(\exists x)Q(x)$. If the universe of discourse is given by the sets

- 1 $\{-1, 0, 1, 2, 4\}$
- 2 $\{3, -2, 7, 8, -2\}$
- 3 $\{15, 20, 24\}$

then $(x)Q(x)$ is true for the universe of discourse (1) and false for (2) and (3). The statement $(\exists x)Q(x)$ is true for both (1) and (2), but false for (3).

It may be noted that there are two ways of obtaining a 0-place predicate from an n -place predicate. The first way is to substitute names of objects from of the predicate calculus with the understanding that the atomic variables there stand for prime predicate formulas.

Some Valid Formulas over Finite Universes

In this and in the following section we denote predicate formulas by capital letters such as $A, B, C, \dots$ followed by object variables $x, y, \dots$. Thus $A(x), A(x, y), B(y)$, and $C(x, y, z)$ are examples of predicate formulas. Some clarification is necessary at this stage. In the formula $A(x)$, we wish to say that A is a predicate formula in which x is one of the free variables. This variable x is of interest to us, and we want to emphasize the dependence of A on it. For example, we may write $B(x)$ for $(y)P(y) \vee Q(x)$.

If in a formula $A(x)$ we replace each free occurrence of the variable x by another variable y , then we say that y is *substituted* for x in the formula, and the resulting formula is denoted by $A(y)$. For such a substitution, the formula $A(x)$ must be free for y . A formula $A(x)$ is said to be *free for y* if no free occurrence of x is in the scope of the quantifiers (y) or $(\exists y)$. If $A(x)$ is not free for y , then it is necessary to change the variable y , appearing as a bound variable, to another variable before substituting y for x . If y is to be substituted, then it is usually a good idea to make all the bound variables different from y . The following examples show what $A(y)$ is for a given $A(x)$.

$A(x)$	$A(y)$
$P(x, y) \wedge (\exists y)Q(y)$	$P(y, y) \wedge (\exists y)Q(y)$ or $P(y, y) \wedge (\exists z)Q(z)$
$(S(x) \wedge S(y)) \vee (x)R(x)$	$(S(y) \wedge S(y)) \vee (x)R(x)$ or $(S(y) \wedge S(y)) \vee (z)R(z)$

The following formulas are not free for y .

$$P(x, y) \wedge (y)Q(x, y) \quad (y)(S(y) \rightarrow S(x))$$

In order to substitute y in place of the variable x in these formulas, it is necessary to first make them free for y as follows:

$A(x)$	$A(y)$
$P(x, y) \wedge (z)Q(x, z)$	$P(y, y) \wedge (z)Q(y, z)$
$(z)(S(x) \rightarrow S(x))$	$(z)(S(z) \rightarrow S(y))$

If the universe of discourse is a finite set, then all possible substitutions of the object variables can be enumerated. However, it is not possible to enumerate all possible substitutions if the universe of discourse is infinite. We shall now give some equivalences which hold for a finite universe. Later we show that these equivalences also hold for an arbitrary universe.

Theory of Inference for The Predicate Calculus

The method of derivation involving predicate formulas uses the rules of inference given for the statement calculus and also certain additional rules which are required to deal with the formulas involving quantifiers. The rules **P** and **T**, regarding the introduction of a premise at any stage of derivation and the introduction of any formula which follows logically from the formulas already introduced, remain the same. If the conclusion is given in the form of a conditional, we shall also use the rule of conditional proof called **CP**. Occasionally, we may use the indirect method of proof in introducing the negation of the conclusion as an additional premise in order to arrive at a contradiction.

The equivalences and implications of the statement calculus can be used in the process of derivation as before, except that the formulas involved are generalized to predicates. But these formulas do not have any quantifiers in them, while some of the premises or the conclusion may be quantified. In order to use the equivalences and implications, we need some rules on how to eliminate quantifiers during the course of derivation. This elimination is done by *rules of specification* called rules **US** and **ES**. Once the quantifiers are eliminated, the derivation proceeds as in the case of the statement calculus, and the conclusion is reached. It may happen that the desired conclusion is quantified. In this case, we need *rules of generalization* called rules **UG** and **EG**, which can be used to attach a quantifier.

The rules of generalization and specification follow. Here $A(x)$ is used to denote a formula with a free occurrence of x . $A(y)$ denotes a formula obtained by the substitution of y for x in $A(x)$. Recall that for such a substitution $A(x)$ must be free for y .

Rule US (Universal Specification) From $(x)A(x)$ one can conclude $A(y)$.

Rule ES (Existential Specification) From $(\exists x)A(x)$ one can conclude $A(y)$ provided that y is not free in any given premise and also not free in any prior step of the derivation. These requirements can easily be met by choosing a new variable each time **ES** is used. (The conditions of **ES** are more restrictive than ordinarily required, but they do not affect the possibility of deriving any conclusion.)

Rule EG (Existential Generalization) From $A(x)$ one can conclude $(\exists y)A(y)$.

Rule $\mathbf{UG}$ (Universal Generalization) From $A(x)$ one can conclude $(y)A(y)$ provided that x is not free in any of the given premises and provided that if x is free in a prior step which resulted from use of **ES**, then no variables introduced by that use of **ES** appear free in $A(x)$.

We shall now show, by means of an example, how an invalid conclusion could be arrived at if the second restriction on rule **UG** were not imposed. The other restrictions on **ES** and **UG** are easy to understand.

Let $D(u, v)$: u is divisible by v . Assume that the universe of discourse is $\{5, 7, 10, 11\}$, so that the statement $(\exists u)D(u, 5)$ is true because both $D(5, 5)$

RELATIONS AND ORDERING

The concept of a relation is a basic concept in everyday life as well as in mathematics. We have already used various relations. Associated with a relation is the act of comparing objects which are related to one another. The ability of a computer to perform different tasks based upon the result of a comparison is one of its most important attributes which is used several times during the execution of a typical program. In this section we first formalize the concept of a relation and then discuss methods of representing a relation by using a matrix or its graph. The relation matrix is useful in determining the properties of a relation and also in representing a relation on a computer. Various basic properties of relations are given, and certain important classes of relations are introduced. Among these, the compatibility relation and the equivalence relation have useful applications in the design of digital computers and other sequential machines. Partial ordering and its associated terminology are introduced next. The material in Chap. 4 is based upon these notions. Several relations given as examples in this section are used throughout the book. Algorithms to determine certain properties of relations are also given.

member of each of the following sets:

$$R_4 = R_1 \cap R_2 \cap R_3$$

$$= \{ \langle x, y \rangle \mid \langle x, y \rangle \in \mathbf{R} \times \mathbf{R} \wedge x * y \geq 1 \wedge x^2 + y^2 \leq 9 \wedge y^2 < x \}$$

$$R_5 = R_2 \cap (R_1 \cup R_3) \cap \sim(R_1 \cap R_3)$$

$$= \{ \langle x, y \rangle \mid \langle x, y \rangle \in \mathbf{R} \times \mathbf{R} \wedge x^2 + y^2 \leq 9 \wedge (x * y \geq 1 \vee y^2 < x) \\ \wedge \sim(x * y \geq 1 \wedge y^2 < x) \}$$

$$R_6 = R_1 \cap \sim R_2 \cap R_3.$$

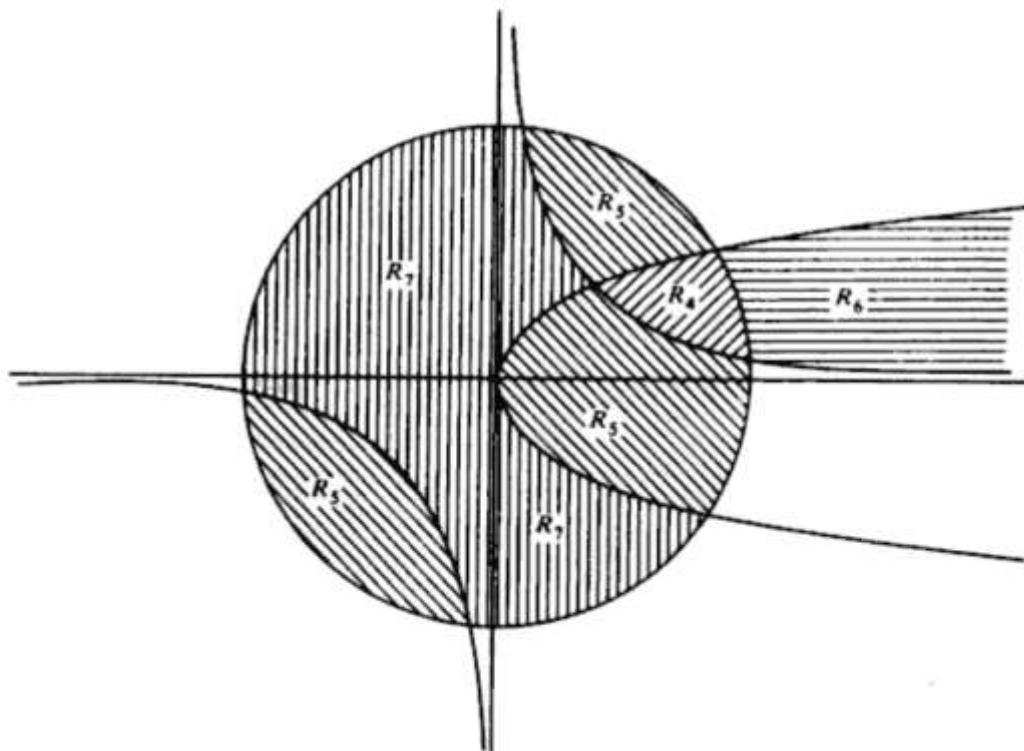
$$= \{ \langle x, y \rangle \mid \langle x, y \rangle \in \mathbf{R} \times \mathbf{R} \wedge x * y \geq 1 \wedge \sim(x^2 + y^2 \leq 9) \wedge y^2 < x \}$$

$$R_7 = \sim(R_1 \cup R_3) \cap R_2$$

$$= \{ \langle x, y \rangle \mid \langle x, y \rangle \in \mathbf{R} \times \mathbf{R} \wedge \sim(x * y \geq 1 \vee y^2 < x) \wedge x^2 + y^2 \leq 9 \}$$

R_4 includes all points lying within the circle and the parabola and above the hyperbola of the first quadrant. R_5 includes all points within the circle which lie either within the parabola or above the hyperbola of the first quadrant, but not both, and all points within the circle and below the hyperbola in the third quadrant. R_6 includes all points lying above the hyperbola and within the parabola in the first quadrant. R_7 includes all points lying within the circle and between the hyperbolic curves but not within the parabola.

These newly defined sets can pictorially be represented as shown in Fig. 2-3.2. The program given in Fig. 2-3.3 reads a number of coordinate points and determines whether these points lie in the sets R_4 to R_7 . Note that the relations R_4 to R_7 are written as predicates P_4 to P_7 in the program.



FUNCTIONS

In this section we study a particular class of relations called functions. We are primarily concerned with discrete functions which transform a finite set into another finite set. There are several such transformations involved in the computer implementation of any program. Computer output can be considered as a function of the input. A compiler transforms a program into a set of machine language instructions (the object program). After introducing the concept of function in general, we discuss unary and binary operations which form a class of functions. Such operations have important applications in the study of algebraic structures in Chaps. 3 and 4. Also discussed is a special class of functions known as hashing functions that are used in organizing files on a computer, along with other techniques associated with such organizations. A PL/I program for the construction of a symbol table is also given.

2-4.1 Definition and Introduction

Definition 2-4.1 Let X and Y be any two sets. A relation f from X to Y is called a *function* if for every $x \in X$ there is a unique $y \in Y$ such that $\langle x, y \rangle \in f$.

Note that the definition of function requires that a relation must satisfy two additional conditions in order to qualify as a function. The first condition is that every $x \in X$ must be related to some $y \in Y$, that is, the domain of f must

UNIT – V

LATTICES AS PARTIALLY ORDERED SETS

A lattice is a partially ordered set, where every pair of elements has both a unique least upper bound and a unique greatest lower bound. This means a lattice is a ordered set that fulfils two additional properties: it's a “join-semilattice” and a “meet-semilattice”. A common example is the power set of a set, ordered by subset inclusion, where the join is the union and the meet is the intersection.

Partially ordered sets:

A set with a relation that is reflexive, antisymmetric, and transitive.

- **Reflexive:** For all $a \in P$, $a \leq a$
- **Antisymmetric:** For all $a, b \in P$, if $a \leq b$ and $b \leq a$, then $a = b$.
- **Transitive:** For all $a, b, c \in P$, if $a \leq b$ and $b \leq c$ then $a \leq c$.

Least Upper Bound (Join):

For any two elements a and b , the join ($a \vee b$) is the smallest element that is greater than or equal to both a and b .

Greatest Lower Bound (Meet):

For any two elements a and b , the join ($a \wedge b$) is the largest element that is less than or equal to both a and b .

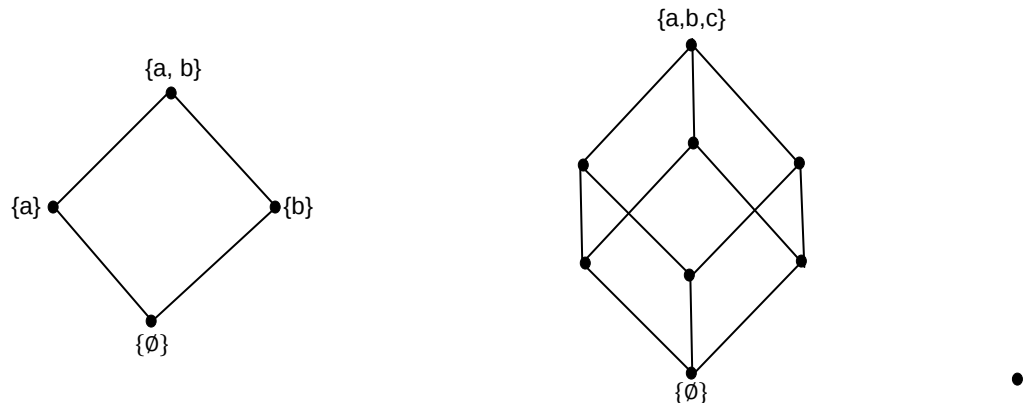
Definition:

A lattice in a partially ordered set $(L, \leq)$ in which every pair of elements $a, b \in L$ has a greatest lower bound and a least upper bound.

Example 1.

Let S be any set and $P(S)$ be its power set. The partially ordered set $(P(S), \leq)$ is a lattice in which the meet and join are the same as the operations $\cap$ and $\cup$ respectively. In

particular, when S has a single element, the corresponding lattice is a chain containing two elements. When S has two and three elements, the diagrams of the corresponding lattices are as shown below.



Example 2.

Is the poset $(\mathbb{Z}^+, 1)$ a lattice?

Solution:

Let a and b be two positive integers. The least upper bound and greatest lower bound of these two integers are the least common multiple and the greatest common divisor of these integers, respectively. It follows that the poset is a lattices.

Example 3.

Let $S = \{a, b, c\}$. Draw the diagram of $(P(S), \subseteq)$.

Solution:

Given $S = \{a, b, c\}$

$$P(S) = \{ \{a\}, \{b\}, \{c\}, \{a, b\}, \{b, c\}, \{a, c\}, \{a, b, c\}, \{\} \}$$

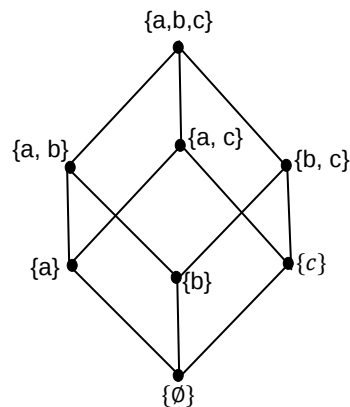
We know that $\{P(S), \subseteq\}$ is a poset.

Since empty set is a subset of every set is $P(S)$, $\{\}$ is a least of $P(S)$.

Similarly, $S = \{a, b, c\}$ contains all elements of $P(S)$. (ie) an element of $P(S)$ is a subset of $\{a, b, c\}$.

Therefore, S is a greatest element is $P(S)$. Since, $(P(S), \subseteq)$ is a lattice.

Hasse Diagram:



PROPERTY 1: Let $(L, \leq)$ be a lattice. For any $a, b, c \in L$

we have $a * a = a$ and $a \oplus a = a$

[Idempotent law]

Proof: Let $a, b, c \in L$, by the definition of GLB of a and b we have

$$a * b \leq a \quad \dots (i)$$

and if $a \leq a$ and $a \leq b$, then

$$a \leq a * b \quad \dots (ii)$$

As $a \leq a$ from (i) and (ii) we have

$$a * a \leq a \text{ and } a \leq a * a$$

By the antisymmetric property it follows that $a = a * a$

Similarly we can prove that $a \oplus a = a$

PROPERTY 2. Show that the operation of meet and join on a lattice are associative.

Solution: To prove : $(a*b) * c = a * (b*c)$

Let $a, b, c \in L$ by the definition we have

and

$$(a*b) * c \leq a*b$$

$$(a*b) * c \leq c$$

By the definition of GLB of a and b , we have $a * b \leq a$ and $a * b \leq b$, so by transitive property of $\leq$ we have

$$(a*b) * c \leq a$$

$$\text{and } (a*b) * c \leq b$$

$$\text{As } (a*b) * c \leq b \text{ and } (a*b) * c \leq c$$

We see that $(a*b) * c$ is lower bound for b and c . From the definition of $b*c$ it follows that $(a * b) * c \leq b*c$

$$\text{As } (a*b) * c \leq a \text{ and } (a*b) * c \leq b * c$$

From the definition of $a * (b*c)$, we have

$$(a*b) * c \leq a * (b*c) \dots (i)$$

$$\text{Now } a * (b*c) \leq a \text{ and } a * (b*c) \leq b * c$$

$$\text{As } b*c \leq b, \text{ by transitivity } a * (b*c) \leq b$$

$$\text{since } a * (b*c) \leq a, \text{ and } a * (b * c) \leq b$$

$$\text{We have } a * (b*c) \leq (a*b)$$

As $a * (b * c) \leq b * c \leq c$

$$a * (b * c) \leq (a * b) * c \dots (ii)$$

From (i) and (ii), by antisymmetric property, it follows that

$$a * (b * c) = (a * b) * c$$

Similarly, we can prove that $a \oplus (b \oplus c) = (a \oplus b) \oplus c$

PROPERTY 3. Show that the operation of meet and join on a lattice are commutative law. i.e., $a * b = b * a$ and $a \oplus b = b \oplus a$

Solution : Given $a, b \in L$ both $a * b$ and $b * a$ are GLB of a and b . By the uniqueness of GLB of a and b , we have $a * b = b * a$. Similarly $a \oplus b = b \oplus a$ holds good.

PROPERTY 4. Absorption law $a * (a \oplus b) = a$ and $a \oplus (a * b) = a$

Solution : Let $a, b \in L$. Then $a \leq a$ and $a \leq a \oplus b$. So $a \leq a * (a \oplus b)$. On the other hand $a * (a \oplus b) \leq a$. By annisymmetric property of $\leq$ we have $a = a * (a \oplus b)$

Similarly we have $a = a \oplus (a * b) \forall a, b \in L$

Theorem 1.

Let $(L, \leq)$ be a lattice in which $*$ and $\oplus$ denotes the operations of meet and join respectively. For any $a, b \in L$.

$$a \leq b \Leftrightarrow a * b = a \Leftrightarrow a \oplus b = b$$

Proof : First let us prove that $a \leq b \Leftrightarrow a * b = a$

Let us assume that $a \leq b$ and also we know that $a \leq a$.

$$\therefore a \leq a * b \dots (1)$$

But, from definition of $a * b$, we have

$$a * b \leq a \dots (2)$$

Hence $a \leq b \Rightarrow a * b = a$ [From (1) and (2)]

Next, assume that $a * b = a$, but it is only possible if $a \leq b$.

That is $a * b = a \Rightarrow a \leq b$

Combining these two results, we get

$$a \leq b \Leftrightarrow a * b = a$$

To show that $a \leq b \Leftrightarrow a \oplus b = b$ in a similar way.

From $a * b = a$, We have

$$b \oplus (a * b) = b \oplus a = a \oplus b$$

But $b \oplus (a * b) = b$

Hence $a \oplus b = b$ follows that $a * b = a$

Theorem 2.

Let $(L, \leq)$ be a lattice. For any $a, b \in L$ the following are equivalent.

(i) $a \leq b$, (ii) $a * b = a$, (iii) $a \oplus b = b$

Proof : At first, consider (i) $\Leftrightarrow$ (ii)

We have $a \leq a$, assume $a \leq b$. Therefore $a \leq a * b$. By the definition of GLB, we have

$$a * b \leq a$$

Hence by antisymmetric property, $a * b = a$

Assume that $a * b = a$, but is only possible if

$$a \leq b \Rightarrow a * b = a \Rightarrow a \leq b.$$

Combining these two results, we have $a \leq b \Leftrightarrow a * b = a$

Similarly, $a \leq b \Leftrightarrow a \oplus b = b$

Alternatively, (ii) $\Leftrightarrow$ (iii) as follows :

Assume $a * b = a$, we have $b \oplus (a * b) = b \oplus a = a \oplus b$, but by absorption $b \oplus (a * b) = b$. Hence $a \oplus b = b$.

By representing similar steps, we can show that $a * b = a$ follows from $a \oplus b = b$.

$$\therefore (ii) \Leftrightarrow (iii)$$

Hence the theorem

Theorem 3.

Let $(L, \leq)$ be a lattice. For any $a, b, c \in L$, the following inequalities, hold :

(1) Distribute Inequalities

$$(i) a \oplus (b * c) \leq (a \oplus b) * (a \oplus c)$$

$$(ii) a * (b \oplus c) \geq (a * b) \oplus (a * c)$$

(2) Modular Inequalities

$$(i) a \leq c \Leftrightarrow a \oplus (b * c) \leq (a \oplus b) * c$$

$$(ii) a \geq c \Leftrightarrow a * (b \oplus c) \geq (a * b) \oplus c$$

Proof :

As (ii) in 1, (ii) in 2 are duals of (i) in 1 and (i) in 2 respectively, it is enough to prove (i) and (i) in 2 only.

Consider (i) in 1.

Let $a, b, c \in L$. As $a \leq a \oplus b$ and $a \leq a \oplus c$

we have $a \leq [(a \oplus b) * (a \oplus c)]$

As $b * c \leq b \leq a \oplus b$ and $b * c \leq c \leq a \oplus c$,

we have $(b * c) \leq (a \oplus b) * (a \oplus c)$.

$(a \oplus b) * (a \oplus c)$ is an upper bound for a and $b * c$ and

hence $a \oplus (b * c) \leq (a \oplus b) * (a \oplus c)$.

Thus (i) in 1 is proved.

The inequality (i) in 2 is special case of (i) in 1.

If $a \leq c$, then $a \oplus c = c$ and from (i) in 1 we obtain

$$a \oplus (b * c) \leq (a \oplus b) * (a \oplus c) = (a \oplus b) * c$$

which is inequality (i) in 2.

Hence the theorem.

Theorem 4.

In a lattice $(L, \leq)$, show that (i) $(a * b) \oplus (c * d) \leq (a \oplus c) * (b \oplus d)$
(ii) $(a * b) \oplus (b * c) \oplus (c * a) \leq (a \oplus b) * (b \oplus c) * (c \oplus a), \forall a, b, c \in L$

Proof : Let $a, b, c \in L$

$$\text{Then } a * b \leq a \text{ (or) } b \leq a \oplus b \quad \dots (1)$$

$$a * b \leq a \leq c \oplus a \quad \dots (2)$$

$$a * b \leq b \leq b \oplus c \quad \dots (3)$$

Using (1), (2) and (3), we get

$$a * b \leq (a \oplus b) * (b \oplus c) * (c \oplus a)$$

Similarly $b * c \leq (a \oplus b) * (b \oplus c) * (c \oplus a)$

$$c * a \leq (a \oplus b) * (b \oplus c) * (c \oplus a)$$

This proves (ii)

We have $a \leq a \oplus c$

$$b \leq b \oplus d$$

$$\therefore (a * b) \leq (a \oplus c) * (b \oplus d)$$

We know that $c \leq a \oplus c \quad \dots (4)$

$$d \leq b \oplus d \quad \dots (5)$$

$$\therefore c * d \leq (a \oplus c) * (b \oplus d)$$

By (4) and (5), $(a * b) \oplus (c * d) \leq (a \oplus c) * (b \oplus d)$. This proves (i)

Theorem 5.

In a lattice $(L, \leq)$, prove that for $a, b, c \in L$

$$(i) (a * b) \oplus (a * c) \leq a * (b \oplus (a * c))$$

$$(ii) (a \oplus b) * (a \oplus c) \geq a \oplus (b * (a \oplus c))$$

Proof : We know that $a * b \leq a, a * c \leq a$

$$\therefore (a * b) \oplus (a * c) \leq a \oplus a = a \quad \dots (1)$$

$$\text{Also } a * b \leq b, a * c \leq a * c$$

$$\Rightarrow (a * b) \oplus (a * c) \leq b \oplus (a * c) \quad \dots (2)$$

$$\text{From (1) and (2), } (a * b) \oplus (a * c) \leq a * (b \oplus (a * c))$$

This proves (i)

$$\text{We know that } a \leq a \oplus b ; a \leq a \oplus c$$

$$\Rightarrow a = a * a \leq (a \oplus b) * (a \oplus c)$$

$$\text{Further } b \leq a \oplus b ; a \oplus c \leq a \oplus c$$

$$\Rightarrow b * (a \oplus c) \leq (a \oplus b) * (a \oplus c)$$

$$\text{By (3) \& (4), } a \oplus (b * (a \oplus c)) \leq (a \oplus b) * (a \oplus c)$$

This proves (ii)

Theorem 6.

In a lattice if $a \leq b \leq c$, show that

$$(i) a \oplus b = b * c \quad (ii) (a * b) \oplus (b * c) = (a \oplus b) * (a \oplus c) = b$$

Proof : Let $a \leq b \leq c$

$$a \leq b \Rightarrow a \oplus b = b, a * b = a$$

$$b \leq c \Rightarrow b \oplus c = c, b * c = b$$

$$a \leq c \Rightarrow a \oplus c = c, a * c = a$$

$$\therefore a \oplus b = b = b * c \quad (i) \text{ follows}$$

$$\text{Now } (a * b) \oplus (b * c) = a \oplus b = b$$

$$(a \oplus b) * (a \oplus c) = b * c = b \quad (ii) \text{ follows}$$

A lattice L is modular if and only if none of its sublattices is isomorphic to the pentagon lattice N_5 .

Proof : Since the pentagon lattice N_5 is not modular lattice. Hence any lattice having a pentagon as a sublattice cannot be modular.

Conversely, let $(L, \leq)$ be any nonmodular lattice we shall prove there is a sublattice which is isomorphic to N_5 .

As $(L, \leq)$ is a nonmodular lattice, then there are elements $a, b, c \in L$ such that

$$a \leq c \text{ and } a \vee (b \wedge c) \neq (a \vee b) \wedge c$$

$$\text{Let } u = b \wedge c$$

$$x = a \vee (b \wedge c)$$

$$y = b$$

$$z = (a \vee b) \wedge c$$

$$v = a \vee b$$

Then all the elements u, x, y, z, v are distinct, we have $u \leq x \leq z \leq v$ and $u \leq y \leq v$

$$(i) u \leq x \wedge y \leq z \wedge y = (a \vee b) \wedge c \wedge b$$

$$= ((a \vee b) \wedge b) \wedge c$$

$$= b \wedge c = u$$

$$\text{Therefore } x \wedge y = z \wedge y = u$$

$$(ii) v \geq z \vee y \geq x \vee y = (a \vee (b \wedge c)) \vee b$$

$$= a \vee ((b \wedge c) \vee b)$$

$$= a \vee b = v$$

$$\text{So that, } x \vee y = z \vee y = v$$