

**DIGITAL NOTES
ON
Discrete Mathematics
B.TECH II YEAR - II
SEM(2021-22)**



DEPARTMENT OF COMPUTER SCIENCE & ENGINEERING

**MALLA REDDY COLLEGE OF ENGINEERING & TECHNOLOGY
(Autonomous Institution – UGC, Govt. of India)**

(Affiliated to JNTUH, Hyderabad, Approved by AICTE - Accredited by NBA & NAAC – ‘A’ Grade - ISO 9001:2015 Certified)
Maisammaguda, Dhulapally (Post Via. Hakimpet), Secunderabad – 500100, Telangana State, INDIA.



MALLA REDDY COLLEGE OF ENGINEERING & TECHNOLOGY

DEPARTMENT OF COMPUTER SCIENCE & ENGINEERING

II Year B.Tech CSE – I Sem

T/P/D C L

3 /-/- 3

(R20A0026) Discrete Mathematics

UNIT-I

Mathematical Logic : Statements and notations, Connectives, Well formed formulas, Truth Tables, tautology, equivalence implication, Normal forms, Quantifiers, universal quantifiers.

Predicates : Predicative logic, Free & Bound variables, Rules of inference, Consistency, proof of contradiction, Automatic Theorem Proving.

UNIT-II

Relations: Properties of Binary Relations, equivalence, transitive closure, compatibility and partial ordering relations, Lattices, Hasse diagram. Functions: Inverse Function Composition of functions, recursive Functions, Lattice and its Properties,

Algebraic structures : Algebraic systems Examples and general properties, Semigroups and monads, groups sub groups' homomorphism, Isomorphism.

UNIT-III

Elementary Combinatorics: Basis of counting, Combinations & Permutations, with repetitions, Constrained repetitions, Binomial Coefficients, Binomial Multinomial theorems, the principles of Inclusion – Exclusion. Pigeon hole principles and its application.

UNIT-IV

Recurrence Relation : Generating Functions, Function of Sequences Calculating Coefficient of generating function, Recurrence relations, Solving recurrence relation by substitution and Generating funds. Characteristics roots solution of In homogeneous Recurrence Relation.

UNIT-V

Graph Theory : Representation of Graph, DFS, BFS, Spanning Trees, planar Graphs. Graph Theory and Applications, Basic Concepts Isomorphism and Sub graphs, Multi graphs and Euler circuits, Hamiltonian graphs, Chromatic Numbers.

TEXT BOOKS:

1. Elements of DISCRETE MATHEMATICS- A computer Oriented Approach- C L Liu, D P Mohapatra. Third Edition, Tata McGraw Hill.
2. Discrete Mathematics for Computer Scientists & Mathematicians, J.L. Mott, A. Kandel, T.P. Baker, PHI.

REFERENCE BOOKS:

1. Discrete Mathematics and its Applications, Kenneth H. Rosen, Fifth Edition.TMH.
2. Discrete Mathematical structures Theory and application-Malik & Sen, Cengage.
3. Discrete Mathematics with Applications, Thomas Koshy, Elsevier.
4. Logic and Discrete Mathematics, Grass Man & Trembley, Pearson Education.



MALLA REDDY COLLEGE OF ENGINEERING & TECHNOLOGY

DEPARTMENT OF COMPUTER SCIENCE & ENGINEERING

INDEX

S. No	Unit	Topic	Page no
1	I	MATHEMATICAL LOGIC	1
2	I	PREDICATES	10
3	I	RULES OF INFERENCE	11
4	II	RELATIONS	22
5	II	FUINCTIONS	31
6	II	ALGEBRAIC STRUCTURES	36
7	III	ELEMENTARY COMBINATORICS	42
8	III	BINOMIAL COEFFICIENTS	49
9	III	PEGION HOLE PRINCIPLE	54
10	IV	RECURRENCE RELATION	57
11	IV	SOLVING RECURRENCE RELATION	61
12	V	GRAPH THEORY	65
13	V	APPLICATIONS OF GRAPH	73
14	V	MULTIGRAPHS & EULER CIRCUITS	76
15	V	CHROMATIC NUMBERS	79



UNIT-I

Mathematical Logic

Statements and notations:

A proposition or statement is a declarative sentence that is either true or false (but not both).

For instance, the following are propositions:

- Paris is in France < (true)
- London is in Denmark < (false)
- $2 < 4$ < (true)
- $4 = 7$ < (false)

However the following are not propositions:

- what is your name? < (this is a question)
- do your homework < (this is a command)
- this sentence is false < (neither true nor false)
- x is an even number < (it depends on what x represents)
- Socrates < (it is not even a sentence)

The truth or falsehood of a proposition is called its truth value.

Connectives:

Connectives are used for making compound propositions. Generally used five connectives are –

- OR ($\vee$)
- AND ($\wedge$)
- Negation/ NOT ($\neg$)
- Implication / if-then ($\rightarrow$)
- If and only if ($\leftrightarrow$).

Well formed formulas (wff):

The strings that produce a proposition when their symbols are interpreted must follow the rules given below, and they are called wffs (well-formed formulas) of the first order predicate logic.

A predicate name followed by a list of variables such as $P(x, y)$, where P is predicate name, and x and y are variables, is called an atomic formula.

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

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

Wffs are constructed using the following rules:

1. *True* and *False* are wffs.
2. Each propositional constant (i.e. specific proposition), and each propositional variable (i.e. a variable representing propositions) are wffs.
3. Each atomic formula (i.e. a specific predicate with variables) is a wff.
4. If A , B , and C are wffs, then so are A , $(A \wedge B)$, $(A \vee B)$, $(A \rightarrow B)$, and $(A \leftrightarrow B)$.
5. If x is a variable (representing objects of the universe of discourse), and A is a wff, then so are $\forall x A$ and $\exists x A$.

For example, "The capital of Virginia is Richmond." is a specific proposition. Hence it is a wff by Rule 2.

Let B be a predicate name representing "being blue" and let x be a variable. Then $B(x)$ is an atomic formula meaning " x is blue". Thus it is a wff by Rule 3. above.

By applying Rule 5. to $B(x)$, $\forall x B(x)$ is a wff and so is $\exists x B(x)$.

Then by applying Rule 4. to them $\forall x B(x) \wedge \exists x B(x)$ is seen to be a wff. Similarly if R is a predicate name representing "being round". Then $R(x)$ is an atomic formula. Hence it is a wff.

By applying Rule 4 to $B(x)$ and $R(x)$, a wff $B(x) \wedge R(x)$ is obtained.

To express the fact that Tom is taller than John, we can use the atomic formula ***taller(Tom, John)***, which is a wff. This wff can also be part of some compound statements such as ***taller(Tom, John) \wedge \neg taller(John, Tom)***, which is also a wff. ***If x is a variable representing people in the world, then $taller(x, Tom)$, $\forall x taller(x, Tom)$, $\exists x taller(x, Tom)$, $\exists x \forall y taller(x, y)$ are all wffs among others. However, $taller(\exists x, John)$ and $taller(Tom \wedge Mary, Jim)$, for example, are NOT wffs.***

Truth Tables:

Logical identity

Logical identity is an operation on one logical value, typically the value of a proposition that produces a value of *true* if its operand is true and a value of *false* if its operand is false.

The truth table for the logical identity operator is as follows:

Logical Identity	
p	p
T	T
F	F

Logical negation

Logical negation is an operation on one logical value, typically the value of a proposition that produces a value of *true* if its operand is false and a value of *false* if its operand is true.

The truth table for NOT p (also written as $\neg p$ or $\sim p$) is as follows:

Logical Negation	
p	$\neg p$
T	F
F	T

Logical conjunction:

Logical conjunction is an operation on two logical values, typically the values of two propositions, that produces a value of *true* if both of its operands are true.

The truth table for p AND q (also written as $p \text{ K } q$, $p \& q$, or $p \text{ q}$) is as follows:

If both p and q are true, then the conjunction $p \text{ K } q$ is true. For all other assignments of logical values to p and to q the conjunction $p \text{ K } q$ is false. It can also be said that if p , then $p \text{ K } q$ is q , otherwise $p \text{ K } q$ is p .

Logical Conjunction		
p	q	$p \wedge q$
T	T	T
T	F	F
F	T	F
F	F	F

Logical disjunction:

Logical disjunction is an operation on two logical values, typically the values of two propositions, that produces a value of *true* if at least one of its operands is true. The truth table for p OR q (also written as $p \vee q$, $p \parallel q$, or $p + q$) is as follows:

Logical Disjunction		
p	q	$p \vee q$
T	T	T
T	F	T
F	T	T
F	F	F

Logical implication:

Logical implication and the material conditional are both associated with an operation on two logical values, typically the values of two propositions, that produces a value of *false* just in the singular case the first operand is true and the second operand is false.

Logical Implication		
p	q	$p \rightarrow q$
T	T	T
T	F	F
F	T	T

F	F	T

The truth table associated with the material conditional if p then q (symbolized as $p \rightarrow q$) and the logical implication p implies q (symbolized as $p \rightarrow q$) is as shown above.

Logical equality:

Logical equality (also known as biconditional) is an operation on two logical values, typically the values of two propositions, that produces a value of *true* if both operands are false or both operands are true. The truth table for p XNOR q (also written as $p \leftrightarrow q$, $p = q$, or $p \equiv q$) is as follows:

Logical Equality		
p	q	$p \equiv q$
T	T	T
T	F	F
F	T	F
F	F	T

Exclusive disjunction:

Exclusive disjunction is an operation on two logical values, typically the values of two propositions, that produces a value of *true* if one but not both of its operands is true. The truth table for p XOR q (also written as $p \oplus q$, or $p \neq q$) is as follows:

Exclusive Disjunction		
p	q	$p \oplus q$
T	T	F
T	F	T
F	T	T
F	F	F

Logical NAND:

The logical NAND is an operation on two logical values, typically the values of two propositions, that produces a value of *false* if both of its operands are true. In other words, it produces a value of *true* if at least one of its operands is false. The truth table for p NAND q (also written as $p \uparrow q$ or $p \mid q$) is as follows:

\

Logical NAND		
<i>P</i>	<i>q</i>	$p \uparrow q$
T	T	F
T	F	T
F	T	T
F	F	T

In the case of logical NAND, it is clearly expressible as a compound of NOT and AND. The negation of a conjunction: $\neg(p \text{ K } q)$, and the disjunction of negations: $(\neg p) \vee (\neg q)$ is same.

Logical NOR

The logical NOR is an operation on two logical values, typically the values of two propositions, that produces a value of *true* if both of its operands are false. In other words, it produces a value of *false* if at least one of its operands is true. $\downarrow$ is also known as the Peirce arrow after its inventor, Charles Sanders Peirce, and is a Sole sufficient operator.

The truth table for **p NOR q** (also written as **p $\downarrow$ q** or **p T q**) is as follows:

Logical NOR		
<i>p</i>	<i>q</i>	$p \downarrow q$
T	T	F
T	F	F
F	T	F
F	F	T

The negation of a disjunction $\neg(p \vee q)$, and the conjunction of negations $(\neg p) \text{ K } (\neg q)$ is same.

Inspection of the tabular derivations for NAND and NOR, under each assignment of logical values to the functional arguments *p* and *q*, produces the identical patterns of functional values for $\neg(p \text{ K } q)$ as for $(\neg p) \vee (\neg q)$, and for $\neg(p \vee q)$ as for $(\neg p) \text{ K } (\neg q)$. Thus

the first and second expressions in each pair are logically equivalent, and may be

substituted for each other in all contexts that pertain solely to their logical values.

This equivalence is one of De Morgan's laws.

The truth value of a compound proposition depends only on the value of its components.

F for false and T for true summarizes the meaning of the connectives in following way:

p	q	$\neg p$	$p \wedge q$	$p \vee q$	$p \oplus q$	$p \rightarrow q$	$p \leftrightarrow q$
T	T	F	T	T	F	T	T
T	F	F	F	T	T	F	F
F	T	T	F	T	T	T	F
F	F	T	F	F	F	T	T

Note that $\vee$ represents a non-exclusive or, i.e., $p \vee q$ is true when any of p , q is true and also when both are true. On the other hand $\oplus$ represents an exclusive or, i.e., $p \oplus q$ is true only when exactly one of p and q is true.

Tautology, Contradiction, Contingency:

A proposition is said to be a tautology if its truth value is T for any assignment of truth values to its components. Example: The proposition $p \vee \neg p$ is a tautology.

A proposition is said to be a contradiction if its truth value is F for any assignment of truth values to its components. Example: The proposition $p \wedge \neg p$ is a contradiction.

A proposition that is neither a tautology nor a contradiction is called a contingency.

p	$\neg p$	$p \vee \neg p$	$p \wedge \neg p$
T	F	T	F
T	F	T	F
F	T	T	F
F	T	T	F

Equivalence Implication:

We say that the statements r and s are logically equivalent if their truth tables are identical.

For example the truth table is:-

p	q	$\neg p \vee q$
T	T	T
T	F	F
F	T	T
F	F	T

shows $\neg p \vee q$ is equivalent to $p \rightarrow q$. It is easily shown that the statements r and s are equivalent if and only if $r \leftrightarrow s$ is a tautology.

Normal forms:

Let $A(P_1, P_2, P_3, \dots, P_n)$ be a statement formula where $P_1, P_2, P_3, \dots, P_n$ are the atomic variables. If A has truth value T for all possible assignments of the truth values to the variables $P_1, P_2, P_3, \dots, P_n$, then A is said to be a tautology. If A has truth value F, then A is said to be identically false or a contradiction.

Disjunctive Normal Forms

A product of the variables and their negations in a formula is called an elementary product. A sum of the variables and their negations is called an elementary sum. That is, a sum of elementary products is called a disjunctive normal form of the given formula.

Example:

$$\begin{aligned} & (A \wedge B \wedge \neg A) \vee (C \wedge \neg B) \vee (A \wedge \neg C) \\ & A \vee (B \wedge C) \end{aligned} \begin{array}{l} (1) \\ (2) \\ (3) \\ (4) \\ (5) \end{array}$$

Principal Disjunctive Normal Form (PDNF)

Let us assume A and B be two statement variables. All possible formulas by using conjunction are given as follows. The total number of formulas for two variables A and B are 22 formulas. They are $A \wedge B$, $A \wedge \neg B$, $\neg A \wedge B$ and $\neg A \wedge \neg B$.

These are called minterms or Boolean conjunctions of A and B . The minterms (2^n terms) are denoted by $M_0, M_1, \dots, M_{2^n-1}$.

A formula equivalent to a given formula consisting of the disjunction of minterms only is called PDNF of the given formula.

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 a given formula.

Example:

$$(A \vee B) \wedge (! A \vee C)$$

$$A \wedge (B \vee C),$$

(1)

(2)

(3)

(4)

Principal Conjunctive Normal Forms (PCNF)

The duals of minterms 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.

For a given formula, an equivalent formula consisting of conjunctions of maxterms only is known as its principal conjunctive normal form. This is also called the product of sums canonical form.

QUANTIFIERS

The variable of predicates is quantified by quantifiers. There are two types of quantifier in **predicate logic** – Universal Quantifier and Existential Quantifier.

Universal Quantifier

Universal quantifier states that the statements within its scope are true for every value of the specific variable. It is denoted by the symbol $\forall$.

$\forall x P(x)$ is read as for every value of x , $P(x)$ is true.

Example – "Man is mortal" can be transformed into the propositional form $\forall x P(x)$ where $P(x)$ is the predicate which denotes x is mortal and the universe of discourse is all men.

Existential Quantifier

Existential quantifier states that the statements within its scope are true for some values of the specific variable. It is denoted by the symbol $\exists$.

$\exists x P(x)$ is read as for some values of x , $P(x)$ is true.

Example – "Some people are dishonest" can be transformed into the propositional form $\exists x P(x)$ where $P(x)$ is the predicate which denotes x is dishonest and the universe of discourse is some people.

Nested Quantifiers

If we use a quantifier that appears within the scope of another quantifier, it is called nested quantifier.

Example

$\exists a \exists b P(x, y)$ where $P(a, b)$ denotes $a + b = 0$

$\exists a \exists b \exists c P(a, b, c)$ where $P(a, b, c)$ denotes $a + (b+c) = (a+b) + c$

Note – $\exists a \exists b P(x, y) \neq \exists a \exists b P(x, y)$

Predicates

Predicative logic:

A predicate or propositional function is a statement containing variables. For instance $-x + 2 = 7$, $-X$ is American, $-x < y$, $-p$ is a prime number are predicates. The truth value of a predicate depends on the value assigned to its variables. For instance if we replace x with 1 in the predicate $-x + 2 = 7$ we obtain $-1 + 2 = 7$, which is false, but if we replace it with 5 we get $-5 + 2 = 7$, which is true.

We represent a predicate by a letter followed by the variables enclosed between parenthesis: $P(x)$, $Q(x, y)$, etc. An example for $P(x)$ is a value of x for which $P(x)$ is true. A counterexample is a value of x for which $P(x)$ is false. So, 5 is an example for $-x + 2 = 7$, while 1 is a counterexample.

Each variable in a predicate is assumed to belong to a universe(or domain) of discourse, for instance in the predicate $-n$ is an odd integer 'n' represents an integer, so the universe of discourse of n is the set of all integers. In $-X$ is American we may assume that X is a human being, so in this case the universe of discourse is the set of all human beings.

Free & Bound variables:

Have a look at the following formula:

$$\neg(\exists x \neg \forall y (\forall z (P(x, y, z) \vee Q(x, y, z)) \wedge \forall y \text{PERSON}(y)))$$

The first occurrence of x is *free*, whereas the second and third occurrences of x are *bound*, namely by the first occurrence of the quantifier $\forall$. The first and second occurrences of the variable y are also

bound, namely by the second occurrence of the quantifier $\forall$.

Informally, the concept of a *bound variable* can be explained as follows: Recall that quantifications are generally of the form:

$$\forall x \phi$$

or

$$\exists x \phi$$

where x may be any variable. Generally, all occurrences of this variable within the quantification are bound. But we have to distinguish two cases. Look at the following formula to see why:

$$\neg x(\text{MAN}(x) \wedge (\forall x \text{WALKS}(x)) \wedge \text{LAFFY}(x))$$

1. x may occur within another, embedded, quantification $\forall y$ or $\exists y$, such as the x in $\text{WALKS}(x)$ in our example. Then we say that it is bound by the quantifier of this embedded quantification (and so on, if there's another embedded quantification over within ϕ).
2. Otherwise, we say that it is bound by the top-level quantifier (like all other occurrences of x in our example).

Here's a full formal simultaneous definition of *free* and *bound*:

1. Any occurrence of any variable is free in any atomic formula.
2. No occurrence of any variable is bound in any atomic formula.
3. If an occurrence of any variable is free in ϕ or in ψ , then that same occurrence is free in $\neg\phi$, $(\phi \rightarrow \psi)$, $(\phi \vee \psi)$, and $(\phi \wedge \psi)$.
4. If an occurrence of any variable is bound in ϕ or in ψ , then that same occurrence is bound in $\neg\phi$, $(\phi \rightarrow \psi)$, $(\phi \vee \psi)$, $(\phi \wedge \psi)$. Moreover, that same occurrence is bound in $\forall y\phi$ and $\exists y\phi$ as well, for any choice of variable y .
5. In any formula of the form $\forall y\phi$ or $\exists y\phi$ (where y can be any variable at all in this case) the occurrence of y that immediately follows the initial quantifier symbol is bound.
6. If an occurrence of a variable x is free in ϕ , then that same occurrence is free in $\forall y\phi$ and $\exists y\phi$, for any variable y distinct from x . On the other hand, all occurrences of x that are free in ϕ , are bound in $\forall x\phi$ and in $\exists x\phi$.

If a formula contains no occurrences of free variables we call it a sentence.

Rules of inference:

The two rules of inference 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 proceeding the actual process of derivation, some important list of implications and equivalences are given in the following tables:

Implications

I1	$P \wedge Q \Rightarrow P$	} Simplification
I2	$PQ \wedge \Rightarrow Q$	
I3	$P \Rightarrow PVQ$	} Addition
I4	$Q \Rightarrow PVQ$	
I5	$\neg P \Rightarrow P \rightarrow Q$	
I6	$Q \Rightarrow P \rightarrow Q$	
I7	$(P \rightarrow Q) \Rightarrow P$	
I8	$(P \rightarrow Q) \Rightarrow Q$	
I9	$P, Q \Rightarrow P \wedge Q$	
I10	$\neg P, PVQ \Rightarrow Q$	(disjunctive syllogism)
I11	$P, P \rightarrow Q \Rightarrow Q$	(modus ponens)
I12	$\neg Q, P \rightarrow Q \Rightarrow \neg P$	(modus tollens)
I13	$P \rightarrow Q, Q \rightarrow R \Rightarrow P \rightarrow R$	(hypothetical syllogism)
I14	$P \vee Q, P \rightarrow Q, Q \rightarrow R \Rightarrow R$	(dilemma)

Equivalences

E1	$\neg \neg P \Leftrightarrow P$	
E2	$P \wedge Q \Leftrightarrow Q \wedge P$	} Commutative laws
E3	$P \vee Q \Leftrightarrow Q \vee P$	
E4	$(P \wedge Q) \wedge R \Leftrightarrow P \wedge (Q \wedge R)$	} Associative laws
E5	$(P \vee Q) \vee R \Leftrightarrow P \vee (Q \vee R)$	
E6	$P \wedge (Q \vee R) \Leftrightarrow (P \wedge Q) \vee (P \wedge R)$	} Distributive laws
E7	$P \vee (Q \wedge R) \Leftrightarrow (P \vee Q) \wedge (P \vee R)$	
E8	$\neg(P \wedge Q) \Leftrightarrow \neg P \vee \neg Q$	
E9	$\neg(P \vee Q) \Leftrightarrow \neg P \wedge \neg Q$	} De Morgan's laws

E10 $P \vee P \leftrightarrow P$
 E11 $P \wedge P \leftrightarrow P$

Example 1. Show that R is logically derived from $P \rightarrow Q$, $Q \rightarrow R$, and P

Solution.	{1}	(1) $P \rightarrow Q$	Rule P
	{2}	(2) P	Rule P
	{1, 2}	(3) Q	Rule (1), (2) and I11
	{4}	(4) $Q \rightarrow R$	Rule P
	{1, 2, 4}	(5) R	Rule (3), (4) and I11.

Example 2. Show that $S \vee R$ tautologically implied by $(P \vee Q) \wedge (P \rightarrow R) \wedge (Q \rightarrow S)$.

Solution .	{1}	(1) $P \vee Q$	Rule P
	{1}	(2) $\neg P \rightarrow Q$	T, (1), E1 and E16
	{3}	(3) $Q \rightarrow S$	P
	{1, 3}	(4) $\neg P \rightarrow S$	T, (2), (3), and I13
	{1, 3}	(5) $\neg S \rightarrow P$	T, (4), E13 and E1
	{6}	(6) $P \rightarrow R$	P
	{1, 3, 6}	(7) $\neg S \rightarrow R$	T, (5), (6), and I13
	{1, 3, 6}	(8) $S \vee R$	T, (7), E16 and E1

Example 3. Show that $\neg Q, P \rightarrow Q \Rightarrow \neg P$

Solution .	{1}	(1) $P \rightarrow Q$	Rule P
	{1}	(2) $\neg P \rightarrow \neg Q$	T, and E 18
	{3}	(3) $\neg Q$	P
	{1, 3}	(4) $\neg P$	T, (2), (3), and I11 .

Example 4 . Prove 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 I12
{4}	(4) $P \vee Q$	P
{1, 2, 4}	(5) Q	T, (3), (4), and I10.
{6}	(6) $Q \rightarrow R$	P
{1, 2, 4, 6}	(7) R	T, (5), (6) and I11
{1, 2, 4, 6}	(8) $R \wedge (P \vee Q)$	T, (4), (7), and I9.

There is 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.

Note. 1. Rule CP follows from the equivalence E10 which states that $(P \wedge R) \rightarrow S \text{ ó } P \rightarrow (R \rightarrow S)$.

2. 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.

3. 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 5 .Show that $R \rightarrow S$ can be derived from the premises $P \rightarrow (Q \rightarrow S)$, $\neg R \vee P$, and Q.

Solution.	{1}	(1) $\neg R \vee P$	P
	{2}	(2) R	P, assumed premise
	{1, 2}	(3) P	T, (1), (2), and I10
	{4}	(4) $P \rightarrow (Q \rightarrow S)$	P
	{1, 2, 4}	(5) $Q \rightarrow S$	T, (3), (4), and I11
	{6}	(6) Q	P
	{1, 2, 4, 6}	(7) S	T, (5), (6), and I11
	{1, 4, 6}	(8) $R \rightarrow S$	CP.

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

Solution.

{1}	(1) $\neg P \vee Q$	P
{2}	(2) P	P, assumed premise
{1, 2}	(3) Q	T, (1), (2) and I11
{4}	(4) $\neg Q \vee R$	P
{1, 2, 4}	(5) R	T, (3), (4) and I11
{6}	(6) $R \rightarrow S$	P
{1, 2, 4, 6}	(7) S	T, (5), (6) and I11
{2, 7}	(8) $P \rightarrow S$	CP

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.

Given premises are: $P \rightarrow Q$, $R \rightarrow \neg Q$ and R conclusion is: $\neg P$

{1}	(1) $P \rightarrow Q$	P
{2}	(2) $R \rightarrow \neg Q$	P
{3}	(3) R	P
{2, 3}	(4) $\neg Q$	T, (2), (3), and I11
{1, 2, 3}	(5) $\neg P$	T, (2), (4) and I12

Consistency of premises:

Consistency

A set of formulas $H_1, H_2, \dots, H_m$ is said to be consistent if their conjunction has the truth value T for some assignment of the truth values to the atomic variables appearing in $H_1, H_2, \dots, H_m$.

Inconsistency

If for every assignment of the truth values to the atomic variables, at least one of the formulas $H_1, H_2, \dots, H_m$ is false, so that their conjunction is identically false, then the formulas

$H_1, H_2, \dots, H_m$ are called inconsistent.

A set of formulas $H_1, H_2, \dots, H_m$ is inconsistent, if their conjunction implies a contradiction, that is $H_1 \wedge H_2 \wedge \dots \wedge H_m \Rightarrow R \wedge \neg R$

Where R is any formula. Note that $R \wedge \neg R$ is a contradiction and it is necessary and sufficient that $H_1, H_2, \dots, H_m$ are inconsistent the formula.

Indirect method of proof

In order to show that a conclusion C follows logically from the premises $H_1, H_2, \dots, H_m$, we assume that C is false and consider $\neg C$ as an additional premise. If the new set of premises is inconsistent, so that they imply a contradiction, then the assumption that $\neg C$ is true does not hold simultaneously with $H_1 \wedge H_2 \wedge \dots \wedge H_m$ being true. Therefore, C is true whenever $H_1 \wedge H_2 \wedge \dots \wedge H_m$ is true.

Thus, C follows logically from the premises $H_1, H_2, \dots, H_m$.

Example 8 Show that $\neg(P \wedge Q)$ follows from $\neg P \wedge \neg Q$.

Solution.

We introduce $\neg(P \wedge Q)$ as an additional premise and show that this additional premise leads to a contradiction.

{1}	(1)	$\neg(P \wedge Q)$	P assumed premise
{1}	(2)	$P \wedge Q$	T, (1) and E1
{1}	(3)	P	T, (2) and I1
{1}	{4}	$\neg P \wedge \neg Q$	P
{4}	(5)	$\neg P$	T, (4) and I1
{1, 4}	(6)	$P \wedge \neg P$	T, (3), (5) and I9

Here (6) $P \wedge \neg P$ is a contradiction. Thus $\{1, 4\}$ viz. $\neg(P \wedge Q)$ and

Example 9 Show that the following premises are inconsistent.

1. If Jack misses many classes through illness, then he fails high school.
2. If Jack fails high school, then he is uneducated.
3. If Jack reads a lot of books, then he is not uneducated.
4. Jack misses many classes through illness and reads a lot of books.

Solution.

P: Jack misses many classes. Q: Jack fails high school.

R: Jack reads a lot of books. S: Jack is uneducated.

The premises are $P \rightarrow Q$, $Q \rightarrow S$, $R \rightarrow \neg S$ and $P \wedge R$

{1}	(1) $P \rightarrow Q$	P
{2}	(2) $Q \rightarrow S$	P
{1, 2}	(3) $P \rightarrow S$	T, (1), (2) and I13
{4}	(4) $R \rightarrow \neg S$	P
{4}	(5) $S \rightarrow \neg R$	T, (4), and E18
{1, 2, 4}	(6) $P \rightarrow \neg R$	T, (3), (5) and I13
{1, 2, 4}	(7) $\neg P \vee \neg R$	T, (6) and E16
{1, 2, 4}	(8) $\neg(P \wedge R)$	T, (7) and E8
{9}	(9) $P \wedge R$	P
{1, 2, 4, 9}	(10) $(P \wedge R) \wedge \neg(P \wedge R)$	T, (8), (9) and I9

The rules above can be summed up in the following table. The "Tautology" column shows how to interpret the notation of a given rule.

Rule of inference	Tautology	Name
$\frac{p}{\therefore \overline{p \vee q}}$	$p \rightarrow (p \vee q)$	
$\frac{p \wedge q}{\therefore \overline{p}}$	$(p \wedge q) \rightarrow p$	Addition Simplification
$\frac{p}{\therefore \overline{p \wedge q}}$	$((p) \wedge (q)) \rightarrow (p \wedge q)$	Conjunction
$\frac{p \quad p \rightarrow q}{\therefore \overline{q}}$	$((p \wedge (p \rightarrow q)) \rightarrow q$	Modus ponens
$\frac{\neg q \quad p \rightarrow q}{\therefore \overline{\neg p}}$	$((\neg q \wedge (p \rightarrow q)) \rightarrow \neg p$	Modus tollens
$\frac{p \rightarrow q \quad q \rightarrow r}{\therefore \overline{p \rightarrow r}}$	$((p \rightarrow q) \wedge (q \rightarrow r)) \rightarrow (p \rightarrow r)$	hypothetical syllogism
$\frac{p \vee q \quad \neg p}{\therefore \overline{q}}$	$((p \vee q) \wedge \neg p) \rightarrow q$	Disjunctive syllogism
$\frac{p \vee q \quad \neg p \vee r}{\therefore \overline{q \vee r}}$	$((p \vee q) \wedge (\neg p \vee r)) \rightarrow (q \vee r)$	Resolution

$$r \rightarrow p$$

Example 1

$$\neg r$$

Let us consider the following assumptions: "If it rains today, then we will not go on a canoe today. If we do not go on a canoe trip today, then we will go on a canoe trip tomorrow. Therefore (Mathematical symbol for "therefore" is $\therefore$), if it rains today, we will go on a canoe trip tomorrow. To make use of the rules of inference in the above table we let p be the proposition "If it rains today", q be "We will not go on a canoe today" and let r be "We will go on a canoe trip tomorrow". Then this argument is of the form:

$$\begin{array}{l} p \rightarrow q \\ q \rightarrow r \\ \therefore \overline{p \rightarrow r} \end{array}$$

Example 2

Let us consider a more complex set of assumptions: "It is not sunny today and it is colder than yesterday". "We will go swimming only if it is sunny", "If we do not go swimming, then we will have a barbecue", and "If we will have a barbecue, then we will be home by sunset" lead to the conclusion "We will be home before sunset." Proof by rules of inference: Let p be the proposition "It is sunny this today", q the proposition "It is colder than yesterday", r the proposition "We will go swimming", s the proposition "We will have a barbecue", and t the proposition "We will be home by sunset". Then the hypotheses become

and $s \rightarrow t$. Using our intuition we conjecture that the conclusion might be t . Using the Rules of Inference table we can proof the conjecture easily:

Step	Reason
1. $\neg p \wedge q$	Hypothesis
2. $\neg p$	Simplification using Step 1
3.	Hypothesis
4.	Modus tollens using Step 2 and 3
5.	Hypothesis
6. s	Modus ponens using Step 4 and 5
7.	Hypothesis

8. t Modus ponens using Step 6 and
 $s \rightarrow t$ 7

Proof of contradiction:

The "Proof by Contradiction" is also known as reductio ad absurdum, which is probably Latin for "reduce it to something absurd".

Here's the idea:

1. Assume that a given proposition is untrue.
2. Based on that assumption reach two conclusions that contradict each other.

This is based on a classical formal logic construction known as Modus Tollens: If P implies Q and Q is false, then P is false. In this case, Q is a proposition of the form (R and not R) which is always false. P is the negation of the fact that we are trying to prove and if the negation is not true then the original proposition must have been true. If computers are not "not stupid" then they are stupid. (I hear that "stupid computer!" phrase a lot around here.)

Example:

Lets prove that there is no largest prime number (this is the idea of Euclid's original proof). Prime numbers are integers with no exact integer divisors except 1 and themselves.

1. To prove: "There is no largest prime number" by contradiction.
2. Assume: There is a largest prime number, call it p.
3. Consider the number N that is one larger than the product of all of the primes smaller than or equal to p. $N = 1 * 2 * 3 * 5 * 7 * 11 * \dots * p + 1$. Is it prime?
4. N is at least as big as p+1 and so is larger than p and so, by Step 2, cannot be prime.

5. On the other hand, N has no prime factors between 1 and p because they would all leave a remainder of 1. It has no prime factors larger than p because Step 2 says that there are no primes larger than p. So N has no prime factors and therefore must itself be prime (see note below). We have reached a contradiction (N is not prime by Step 4, and N is prime by Step 5) and therefore our original assumption that there is a largest prime must be false.

Note: The conclusion in Step 5 makes implicit use of one other important theorem: The Fundamental Theorem of Arithmetic: Every integer can be uniquely represented as the product of primes. So if N had a composite (i.e. non-prime) factor, that factor would itself have prime factors which would also be factors of N.

Automatic Theorem Proving:

Automatic Theorem Proving (ATP) deals with the development of computer programs that show that some statement (the *conjecture*) is a *logical consequence* of a set of statements (the *axioms* and *hypotheses*). ATP systems are used in a wide variety of domains.

The language in which the conjecture, hypotheses, and axioms (generically known as *formulae*) are written is a logic, often classical 1st order logic, but possibly a non-classical logic and possibly a higher order logic. These languages allow a precise formal statement of the necessary information, which can then be manipulated by an ATP system. This formality is the underlying strength of ATP: there is no ambiguity in the statement of the problem, as is often the case when using a natural language such as English.

ATP systems are enormously powerful computer programs, capable of solving immensely difficult problems. Because of this extreme capability, their application and operation sometimes needs to be guided by an expert in the domain of application, in order to solve problems in a reasonable amount of time. Thus ATP systems, despite the name, are often used by domain experts in an interactive way. The interaction may be at a very detailed level, where the user guides the inferences made by the system, or at a much higher level where the user determines intermediate lemmas to be proved on the way to the proof of a conjecture. There is often a synergetic relationship between ATP system users and the systems themselves:

- The system needs a precise description of the problem written in some logical form,
- the user is forced to think carefully about the problem in order to produce an appropriate formulation and hence acquires a deeper understanding of the problem,
- the system attempts to solve the problem, if successful the proof is a useful output,
- if unsuccessful the user can provide guidance, or try to prove some intermediate result, or examine the formulae to ensure that the problem is correctly described,
- and so the process iterates.

ATP is thus a technology very suited to situations where a clear thinking domain expert can interact with a powerful tool, to solve interesting and deep problems. There are many ATP systems readily available for use.

UNIT II Relations

Introduction

The elements of a set may be related to one another. For example, in the set of natural numbers there is the less than' relation between the elements. The elements of one set may also be related to the elements another set.

Binary Relation

A binary relation between two sets A and B is a rule R which decides, for any elements, whether a is in relation R to b. If so, then we write $a R b$. If a is not in relation R to b, then $a \not R b$.

We can also consider $a R b$ as the ordered pair (a, b) in which case we can define a binary relation from A to B as a subset of $A \times B$. This subset is denoted by the relation R.

In general, any set of ordered pairs defines a binary relation.

For example, the relation of father to his child is $F = \{(a, b) / a \text{ is the father of } b\}$ In this relation F, the first member is the name of the father and the second is the name of the child.

The definition of relation permits any set of ordered pairs to define a relation.

For example, the set S given by

$$S = \{(1, 2), (3, a), (b, a), (b, \text{Joe})\}$$

Definition

The domain D of a binary relation S is the set of all first elements of the ordered pairs in the relation.(i.e) $D(S) = \{a / \exists b \text{ for which } (a, b) \in S\}$

The range R of a binary relation S is the set of all second elements of the ordered pairs in the relation. (i.e) $R(S) = \{b / \exists a \text{ for which } (a, b) \in S\}$

For example

For the relation $S = \{(1, 2), (3, a), (b, a), (b, \text{Joe})\}$

$D(S) = \{1, 3, b, b\}$ and

$R(S) = \{2, a, a, \text{Joe}\}$

Let X and Y be any two sets. A subset of the Cartesian product $X \times Y$ defines a relation, say C. For any such relation C, we have $D(C) \subseteq X$ and $R(C) \subseteq Y$, and the relation C is said to from X to Y. If $Y = X$, then C is said to be a relation form X to X. In such case, c is called a relation in X. Thus any relation in X is a subset of $X \times X$. The set $X \times X$ is called a *universal relation* in X, while the empty set which is also a subset of $X \times X$ is called a *void relation* in X.

For example: Let L denote the relation —less than or equal to< and D denote the relation —divides< where $x D y$ means — x divides y < . Both L and D are defined on the set $\{1, 2, 3, 4\}$

$L = \{(1, 1), (1, 2), (1, 3), (1, 4), (2, 2), (2, 3), (2, 4), (3, 3), (3, 4), (4, 4)\}$

$D = \{(1, 1), (1, 2), (1, 3), (1, 4), (2, 2), (2, 4), (3, 3), (4, 4)\}$

$L \subset D = \{(1, 1), (1, 2), (1, 3), (1, 4), (2, 2), (2, 4), (3, 3), (4, 4)\} = D$

Properties of Binary Relations:

Definition: A binary relation R in a set X is **reflexive if, for every $x \in X$, $x R x$, That is $(x, x) \in R$** , or R is reflexive in X $\hat{=} (x) (x \in X \Rightarrow x R x)$.

For example:-

- The relation $\leq$ is reflexive in the set of real numbers.
- The set inclusion is reflexive in the family of all subsets of a universal set.
- The relation equality of set is also reflexive.
- The relation is parallel in the set lines in a plane.
- The relation of similarity in the set of triangles in a plane is reflexive.

Definition: A relation R in a set X is symmetric if for every x and y in X, whenever $x R y$, then $y R x$. (i.e) R is symmetric in X $\hat{=} (x) (y) (x \in X \wedge y \in X \wedge x R y \Rightarrow y R x)$

For example:-

- The relation equality of set is symmetric.
- The relation of similarity in the set of triangles in a plane is symmetric.
- The relation of being a sister is not symmetric in the set of all people.
- However, in the set females it is symmetric.

Definition: A relation R in a set X is **whenever $x R y$ and $y R z$, then $x R z$. (i.e) transitive if, for every x, y, and z are in X, R is transitive in X $\hat{=} (x) (y) (z) (x \in X \wedge y \in X \wedge z \in X \wedge x R y \wedge y R z \Rightarrow x R z)$**

For example:-

- The relations $<$, $\leq$, $>$, $\geq$ and $=$ are transitive in the set of real numbers
- The relations $\subset$, $\subseteq$, $\supset$, $\supseteq$ and equality are also transitive in the family of sets.
- The relation of similarity in the set of triangles in a plane is transitive.

Definition: A relation R in a set X is **irreflexive** if, for every $x \in X$, $(x, x) \notin X$.

For example:-

- < The relation < is irreflexive in the set of all real numbers.
- < The relation proper inclusion is irreflexive in the set of all nonempty subsets of a universal set.
- Let $X = \{1, 2, 3\}$ and $S = \{(1, 1), (1, 2), (3, 2), (2, 3), (3, 3)\}$ is neither irreflexive nor reflexive.

Definition: A relation R in a set x is **anti symmetric** if , for every x and y in X,
whenever $x R y$ and $y R x$, Then $x = y$.
Symbolically, $(x) (y) (x \in X \wedge y \in X \wedge x R y \wedge y R x \Rightarrow x = y)$

For example

- The relations $\leq$, $\geq$ and $=$ are anti symmetric
- The relation $\supset$ is anti symmetric in set of subsets.
- The relation —divides< is anti symmetric in set of real numbers.
- Consider the relation —is a son of< on the male children in a family. Evidently the relation is not symmetric, transitive and reflexive.
- The relation — is a divisor of — is reflexive and transitive but not symmetric on the set of natural numbers.
- Consider the set H of all human beings. Let r be a relation — is married to —
R is symmetric.
- Let I be the set of integers. R on I is defined as $a R b$ if $a - b$ is an even number. R is an reflexive, symmetric and transitive

Equivalence Relation:

Definition: A relation R in a set A is called an **equivalence** relation if

- $a R a$ for every i.e. R is reflexive
- $a R b \Rightarrow b R a$ for every $a, b \in A$ i.e. R is symmetric
- $a R b$ and $b R c \Rightarrow a R c$ for every $a, b, c \in A$, i.e. R is transitive.

For example

- < The relation equality of numbers on set of real numbers.
- < The relation being parallel on a set of lines in a plane.

Problem1: Let

R in T as $R = \{(a, b) / (a, b) \in T \text{ and } a \text{ is similar to } b\}$

We have to show that relation R is an Equivalence relation

Solution :

- < A triangle a is similar to itself. $a R a$
- < If the triangle a is similar to the triangle b , then triangle b is similar to the triangle a then $a R b \Rightarrow b R a$
- < If a is similar to b and b is similar to c , then a is similar to c (i.e) $a R b$ and $b R c \Rightarrow a R c$.

Hence R is an equivalence relation.

Problem 2: Let $x = \{1, 2, 3, \dots, 7\}$ and $R = \{(x, y) / x - y \text{ is divisible by } 3\}$ Show that R is an equivalence relation.

Solution: For any $a \in X$, $a - a$ is divisible by 3, Hence $a R a$, R is reflexive

For any $a, b \in X$, if $a - b$ is divisible by 3, then $b - a$ is also divisible by 3, R is symmetric.

For any $a, b, c \in X$, if $a R b$ and $b R c$, then $a - b$ is divisible by 3 and $b - c$ is divisible by 3. So that $(a - b) + (b - c)$ is also divisible by 3, hence $a - c$ is also divisible by 3. Thus R is transitive.

Hence R is equivalence.

Problem3 . Let Z be the set of all integers. Let m be a fixed integer. Two integers a and b are said to be congruent modulo m if and only if m divides $a - b$, in which case we write $a \equiv b \pmod{m}$. This relation is called the relation of congruence modulo m and we can show that is an equivalence relation.

Solution :

- < $a - a = 0$ and m divides $a - a$ (i.e) $a R a$, $(a, a) \in R$, R is reflexive .
- < $a R b \Rightarrow m \text{ divides } a - b$

$m \text{ divides } b - a \Rightarrow b \equiv a \pmod{m} \Rightarrow b R a$ that is R is symmetric.

- $a R b$ and $b R c \Rightarrow a \equiv b \pmod{m}$ and $b \equiv c \pmod{m}$ $\Rightarrow m \text{ divides } a - b \text{ and } m \text{ divides } b - c$
 $\Rightarrow a - b = km$ and $b - c = lm$ for some $k, l \in \mathbb{Z}$
 $\Rightarrow (a - b) + (b - c) = km + lm$
 $\Rightarrow a - c = (k + l)m$

- ☐ $a \equiv c \pmod{m}$
- ☐ $a R c$
- ☐ R is transitive

Hence the congruence relation is an equivalence relation.

Equivalence Classes:

Let R be an equivalence relation on a set A . For any $a \in A$, the equivalence class generated by a is the set of all elements $b \in A$ such that $a R b$ and is denoted $[a]$. It is also called the R – equivalence class and denoted by $a \in A$. i.e., $[a] = \{b \in A / b R a\}$

Let Z be the set of integer and R be the relation called —congruence modulo 3 defined by $R = \{(x, y) / x \in Z \wedge y \in Z \wedge (x-y) \text{ is divisible by } 3\}$

Then the equivalence classes are

$$[0] = \{\dots -6, -3, 0, 3, 6, \dots\}$$

$$[1] = \{\dots, -5, -2, 1, 4, 7, \dots\}$$

$$[2] = \{\dots, -4, -1, 2, 5, 8, \dots\}$$

Composition of binary relations:

Definition: Let R be a relation from X to Y and S be a relation from Y to Z . Then the relation $R \circ S$ is given by $R \circ S = \{(x, z) / x \in X \wedge z \in Z \wedge \exists y \in Y \text{ such that } (x, y) \in R \wedge (y, z) \in S\}$ is called the composite relation of R and S .

The operation of obtaining $R \circ S$ is called the **composition of relations**.

Example: Let $R = \{(1, 2), (3, 4), (2, 2)\}$ and $S = \{(4, 2), (2, 5), (3, 1), (1, 3)\}$

Then $R \circ S = \{(1, 5), (3, 2), (2, 5)\}$ and $S \circ R = \{(4, 2), (3, 2), (1, 4)\}$

It is to be noted that $R \circ S \neq S \circ R$.

Also $R \circ (S \circ T) = (R \circ S) \circ T = R \circ S \circ T$

Note: We write $R \circ R$ as R^2 ; $R \circ R \circ R$ as R^3 and so on.

Definition

Let R be a relation from X to Y , a relation $\check{R}$ from Y to X is called the converse of R , where the ordered pairs of $\check{R}$ are obtained by interchanging the numbers in each of the ordered pairs of R . This means for $x \in X$ and $y \in Y$, that $x R y \iff y \check{R} x$.

Then the relation $\check{R}$ is given by $\check{R} = \{(x, y) / (y, x) \in R\}$ is called the converse of R Example:

Let $R = \{(1, 2), (3, 4), (2, 2)\}$

Then $\check{R} = \{(2, 1), (4, 3), (2, 2)\}$

Note: If R is an equivalence relation, then $\check{R}$ is also an equivalence relation.

Definition Let X be any finite set and R be a relation in X . The relation $R^+ = R \cup R^2 \cup R^3 \dots$ in X is called the *transitive closure* of R in X

Example: Let $R = \{(a, b), (b, c), (c, a)\}$.

Now $R^2 = R \circ R = \{(a, c), (b, a), (c, b)\}$

$R^3 = R^2 \circ R = \{(a, a), (b, b), (c, c)\}$

$R^4 = R^3 \circ R = \{(a, b), (b, c), (c, a)\} = R$

$R^5 = R^3 \circ R^2 = R^2$ and so on.

Thus, $R^+ = R \cup R^2 \cup R^3 \cup R^4 \cup \dots$

$= R \cup R^2 \cup R^3$.

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

We see that R^+ is a transitive relation containing R . In fact, it is the smallest transitive relation containing R .

Partial Ordering Relations:

Definition

A binary relation R in a set P is called **partial order relation** or *partial ordering* in P iff R is reflexive, anti symmetric, and transitive.

A partial order relation is denoted by the symbol $\preceq$. If $\preceq$ is a partial ordering on P , then the ordered pair $(P, \preceq)$ is called a *partially ordered set* or a *poset*.

< Let R be the set of real numbers. The relation —less than or equal to $\leq$ or $\geq$, is a partial ordering on R .

< Let X be a set and $r(X)$ be its power set. The relation subset, $\subseteq$ on X is partial ordering.

< Let S_n be the set of divisors of n . The relation D means —divides $\leq$ on S_n , is partial ordering on S_n .

In a partially ordered set $(P, \preceq)$, an element $y \in P$ is said to cover an element $x \in P$ if $x \preceq y$ and if there does not exist any element $z \in P$ such that $x \preceq z$ and $z \preceq y$; that is, y covers $x \iff (x \preceq y \implies (x \preceq z \preceq y \implies x = z \vee z = y))$

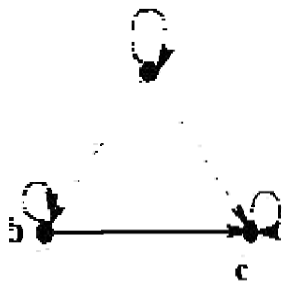
A partial order relation $\preceq$ on a set P can be represented by means of a diagram known as a Hasse diagram or partial order set diagram of $(P, \preceq)$. In such a diagram, each element is represented by a small circle or a dot. The circle for $x \in P$ is drawn below the circle for $y \in P$ if $x \preceq y$, and a line is drawn between x and y if y covers x .

If $x \preceq y$ but y does not cover x , then x and y are not connected directly by a single line. However, they are connected through one or more elements of P .

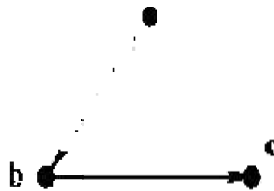
Hasse Diagram:

A Hasse diagram is a digraph for a poset which does not have loops and arcs implied by the transitivity.

Example 10: For the relation $\{< a, a >, < a, b >, < a, c >, < b, b >, < b, c >, < c, c >\}$ on set $\{a, b, c\}$, the Hasse diagram has the arcs $\{< a, b >, < b, c >\}$ as shown below

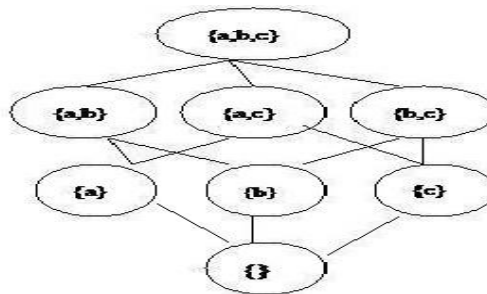


Digraph for Partial Order



Hasse Diagram

Ex: Let A be a given finite set and $r(A)$ its power set. Let $\hat{I}$ be the subset relation on the elements of $r(A)$. Draw Hasse diagram of $(r(A), \hat{I})$ for $A = \{a, b, c\}$



Lattice and its Properties:

Introduction:

A lattice is 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.

The glb of a subset, $\{a, b\} \subseteq L$ will be denoted by $a * b$ and the lub by $a \wedge b$.

.

Usually, for any pair $a, b \in L$, $\text{GLB } \{a, b\} = a * b$, is called the meet or product and $\text{LUB}\{a, b\} = a \wedge b$, is called the join or sum of a and b .

Example1 Consider a non-empty set S and let $P(S)$ be its power set. The relation $\subseteq$ —contained in— is a partial ordering on $P(S)$. For any two subsets $A, B \in P(S)$

$\text{GLB } \{A, B\}$ and $\text{LUB } \{A, B\}$ are evidently $A \cap B$ and $A \cup B$ respectively.

Example2 Let I^+ be the set of positive integers, and D denote the relation of —division— in I^+ such that for any $a, b \in I^+$, $a D b$ iff a divides b . Then (I^+, D) is a lattice in which

the join of a and b is given by the least common multiple (LCM) of a and b , that is,

$a \vee b = \text{LCM of } a \text{ and } b$, and the meet of a and b , that is, $a \wedge b$ is the greatest common divisor (GCD) of a and b .

A lattice can be conveniently represented by a diagram.

For example, let S_n be the set of all divisors of n , where n is a positive integer. Let D denote the relation —division— such that for any $a, b \in S_n$, $a D b$ iff a divides b .

Then (S_n, D) is a lattice with $a \wedge b = \text{gcd}(a, b)$ and $a \vee b = \text{lcm}(a, b)$.

Take $n=6$. Then $S_6 = \{1, 2, 3, 6\}$. It can be represented by a diagram in

Fig(1). Take $n=8$. Then $S_8 = \{1, 2, 4, 8\}$

Two lattices can have the same diagram. For example if $S = \{1, 2, 3\}$ then $(p(s), \subseteq)$ and (S_6, D) Have the same diagram viz. fig(1), but the nodes are differently labeled.

We observe that for any partial ordering relation $\leq$ on a set S the converse relation $\geq$ is also partial ordering relation on S . If $(S, \leq)$ is a lattice With meet $a \wedge b$ and join $a \vee b$, then $(S, \geq)$ is the also a lattice with meet $a \vee b$ and join $a \wedge b$ i.e., the GLB and LUB get interchanged. Thus we have **the principle of duality of lattice as follows.**

Any statement about lattices involving the operations $\wedge$ and $\vee$ and the relations $\leq$ and $\geq$ remains true if $\wedge$, $\vee$, $\geq$ and $\leq$ are replaced by $\vee$, $\wedge$, $\leq$ and $\geq$ respectively.

The operation $\wedge$ and $\vee$ are called duals of each other as are the relations $\leq$ and $\geq$. Also, the lattice $(L, \leq)$ and $(L, \geq)$ are called the duals of each other.

Properties of lattices:

Let $(L, \leq)$ be a lattice with the binary operations $*$ and $\wedge$ then for any $a, b, c \in L$,

$$\leq \quad a * a = a \quad \quad \quad a \wedge a = a \quad \quad \quad (\text{Idempotent})$$

$$\leq \quad a * b = b * a, \quad \quad \quad a \wedge b = b \wedge a \quad \quad \quad (\text{Commutative})$$

$$\leq \quad (a * b) * c = a * (b * c), \quad (a \wedge b) \wedge c = a \wedge (b \wedge c)$$

O (Associative) $\leq$

$$a * (a \wedge b) = a, \quad a \wedge (a * b) = a \quad (\text{absorption})$$

For any $a \in L$, $a \in a$, $a \in \text{LUB } \{a, b\} \Rightarrow a \in a * (a \wedge b)$. On the other hand, $\text{GLB } \{a, a \wedge b\} \in a$ i.e., $(a \wedge b) \wedge a$, hence $a * (a \wedge b) = a$

Theorem 1

Let $(L, \leq)$ be a lattice with the binary operations $*$ and $\wedge$ denote the operations of meet and join respectively For any $a, b \in L$,

$$a \leq b \iff a * b = a \iff a \wedge b = b$$

Proof

Suppose that $a \leq b$. we know that $a \in a$, $a \in \text{GLB } \{a, b\}$, i.e., $a \in a * b$.

But from the definition of $a * b$, we get $a * b \in a$.

Hence $a \in b \Rightarrow a * b = a$ (1)

Now we assume that $a * b = a$; but is possible only if $a \leq b$,

that is $a * b = a \Rightarrow a \leq b$ (2)

From (1) and (2), we get $a \leq b \iff a * b = a$.

Suppose $a * b = a$.

then $b \wedge (a * b) = b \wedge a = a \wedge b$ (3)

but $b \wedge (a * b) = b$ (by iv)..... (4)

Hence $a \wedge b = b$, from (3) $\Rightarrow$ (4)

Suppose $a \wedge b = b$, i.e., $\text{LUB } \{a, b\} = b$, this is possible only if $a \leq b$, thus(3) $\Rightarrow$ (1)

(1) $\Rightarrow$ (2) $\Rightarrow$ (3) $\Rightarrow$ (4). Hence these are equivalent.

Let us assume $a * b = a$.

Now $(a * b) \wedge b = a \wedge b$

We know that by absorption law , $(a * b) \wedge b = b$

so that $a \wedge b = b$, therefore $a * b = a \wedge b = b$ (5)

similarly, we can prove $a \wedge b = b \wedge a * b = a$ (6)

From (5) and (6), we get

$$a * b = a \iff a \wedge b = b$$

Hence the theorem.

Theorem2 For any $a, b, c \in L$, where $(L, \leq)$ is a lattice. b

$$\leq c \iff \{ a * b \leq a * c \text{ and } \\ \{ a \wedge b \leq a \wedge c$$

Proof Suppose $b \leq c$. we have proved that $b \leq a \wedge b * c = b$ (1)

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

$$= a * (b * c) \text{ (by Idempotent)}$$

$$= a * b \quad \text{(by (1))}$$

Thus $(a * b) * (a * c) = a * b$ which $\Rightarrow (a * b) \leq (a * c)$

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

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

$$= a \wedge c$$

which $\Rightarrow (a \wedge b) \leq (a \wedge c)$

note: These properties are known as **isotonicity**.

Functions

Introduction

A function is a special type of relation. It may be considered as a relation in which each element of the domain belongs to only one ordered pair in the relation. Thus a function from A to B is a subset of $A \times B$ having the property that for each $a \in A$, there is one and only one $b \in B$ such that $(a, b) \in G$.

Definition

Let A and B be any two sets. A relation f from A to B is called a function if for every $a \in A$ there is a unique $b \in B$ such that $(a, b) \in f$.

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 $a \in A$ must be related to some $b \in B$, (i.e) the domain of f must be A and not merely subset of A. The second requirement of uniqueness can be expressed as $(a, b) \in f \wedge (b, c) \in f \Rightarrow b = c$

Intuitively, a function from a set A to a set B is a rule which assigns to every element of A, a unique element of B. **If $a \in A$, then the unique element of B assigned to a under f is denoted by f**

(a). The usual notation for a function f from A to B is $f: A \rightarrow B$ defined by $a \mapsto f(a)$ where $a \in A$, $f(a)$ is called the image of a under f and a is called pre image of f(a).

- < Let $X = Y = \mathbb{R}$ and $f(x) = x^2 + 2$. $D_f = \mathbb{R}$ and $R_f \subseteq \mathbb{R}$.
- < Let X be the set of all statements in logic and let $Y = \{\text{True}, \text{False}\}$. A mapping $f: X \rightarrow Y$ is a function.
- < A program written in high level language is mapped into a machine language by a compiler. Similarly, the output from a compiler is a function of its input.
- < Let $X = Y = \mathbb{R}$ and $f(x) = x^2$ is a function from $X \rightarrow Y$, and $g(x^2) = x$ is not a function from $X \rightarrow Y$.

A mapping $f: A \rightarrow B$ is called one-to-one (injective or 1-1) if distinct elements of A are mapped into distinct elements of B. (i.e) f is one-to-one if

$$a_1 \neq a_2 \Rightarrow f(a_1) \neq f(a_2) \text{ or equivalently } f(a_1) = f(a_2) \Rightarrow a_1 = a_2$$

For example, $f: \mathbb{N} \rightarrow \mathbb{N}$ given by $f(x) = x$ is 1-1 where N is the set of a natural numbers.

A mapping $f: A \rightarrow B$ is called onto (surjective) if for every $b \in B$ there is an $a \in A$ such that $f(a) = b$. i.e. if every element of B has a pre-image in A. Otherwise it is called into.

For example, $f: \mathbb{Z} \rightarrow \mathbb{Z}$ given by $f(x) = x + 1$ is an onto mapping. A mapping is both 1-1 and onto is called bijective

For example $f: \mathbb{R} \rightarrow \mathbb{R}$ given by $f(x) = x + 1$ is bijective.

Definition: A mapping $f: \mathbb{R} \rightarrow b$ is called a **constant mapping** if, for all $a \in A$, $f(a) = b$, a fixed element.

For example $f: \mathbb{Z} \rightarrow \mathbb{Z}$ given by $f(x) = 0$, for all $x \in \mathbb{Z}$ is a constant mapping.

Definition

A mapping $f: A \rightarrow A$ is called the **identity mapping of A** if $f(a) = a$, for all $a \in A$. Usually it is denoted by IA or simply I .

Composition of functions:

If $f: A \rightarrow B$ and $g: B \rightarrow C$ are two functions, then the composition of functions f and g , denoted by $g \circ f$, is the function is given by $g \circ f: A \rightarrow C$ and is given by

$g \circ f = \{(a, c) / a \in A \wedge c \in C \wedge \exists b \in B : f(a) = b \wedge g(b) = c\}$ and $(g \circ f)(a) = (g(f(a)))$

Example 1: Consider the sets $A = \{1, 2, 3\}$, $B = \{a, b\}$ and $C = \{x, y\}$. Let $f: A \rightarrow B$ be defined by $f(1) = a$; $f(2) = b$ and $f(3) = b$ and Let $g: B \rightarrow C$ be defined by $g(a) = x$ and $g(b) = y$

(i.e) $f = \{(1, a), (2, b), (3, b)\}$ and $g = \{(a, x), (b, y)\}$. Then $g \circ f: A \rightarrow C$ is defined by

$$(g \circ f)(1) = g(f(1)) = g(a) = x$$

$$(g \circ f)(2) = g(f(2)) = g(b) = y$$

$$(g \circ f)(3) = g(f(3)) = g(b) = y$$

$$\text{i.e., } g \circ f = \{(1, x), (2, y), (3, y)\}$$

If $f: A \rightarrow A$ and $g: A \rightarrow A$, where $A = \{1, 2, 3\}$, are given by

$$f = \{(1, 2), (2, 3), (3, 1)\} \text{ and } g = \{(1, 3), (2, 2), (3, 1)\}$$

Then $g \circ f = \{(1, 2), (2, 1), (3, 3)\}$, $f \circ g = \{(1, 1), (2, 3), (3, 2)\}$

$$f \circ f = \{(1, 3), (2, 1), (3, 2)\} \text{ and } g \circ g = \{(1, 1), (2, 2), (3, 3)\}$$

Example 2: Let $f(x) = x+2$, $g(x) = x - 2$ and $h(x) = 3x$ for $x \in \mathbb{R}$, where $\mathbb{R}$ is the set of real numbers.

$$\text{Then } f \circ f = \{(x, x+4) / x \in \mathbb{R}\}$$

$$f \circ g = \{(x, x) / x \in \mathbb{R}\}$$

$$g \circ f = \{(x, x) / x \in \mathbb{R}\}$$

$$g \circ g = \{(x, x-4) / x \in \mathbb{R}\}$$

$$h \circ g = \{(x, 3x-6) / x \in \mathbb{R}\}$$

$$h \circ f = \{(x, 3x+6) / x \in \mathbb{R}\}$$

Inverse functions:

Let $f: A \rightarrow B$ be a one-to-one and onto mapping. Then, its inverse, denoted by f^{-1} is given by $f^{-1} = \{(b, a) / (a, b) \in f\}$ Clearly $f^{-1}: B \rightarrow A$ is one-to-one and onto.

Also we observe that $f \circ f^{-1} = IB$ and $f^{-1} \circ f = IA$.
If f^{-1} exists then f is called invertible.

For example: Let $f: \mathbb{R} \rightarrow \mathbb{R}$ be defined by $f(x) = x + 2$
Then $f^{-1}: \mathbb{R} \rightarrow \mathbb{R}$ is defined by $f^{-1}(x) = x - 2$

Theorem: Let $f: X \rightarrow Y$ and $g: Y \rightarrow Z$ be two one to one and onto functions. Then $g \circ f$ is also one to one and onto function.

Proof

Let $f: X \rightarrow Y$ and $g: Y \rightarrow Z$ be two one to one and onto functions. Let $x_1, x_2 \in X$

- < $g \circ f(x_1) = g \circ f(x_2)$,
- < $g(f(x_1)) = g(f(x_2))$,
- < $g(x_1) = g(x_2)$ since $[f \text{ is } 1-1]$

$x_1 = x_2$ since $[g \text{ is } 1-1]$
so that $g \circ f$ is 1-1.

By the definition of composition, $g \circ f: X \rightarrow Z$ is a function.

We have to prove that every element of $z \in Z$ an image element for some $x \in X$ under $g \circ f$.

Since g is onto $\exists y \in Y$: $g(y) = z$ and f is onto from X to Y ,
 $\exists x \in X$: $f(x) = y$.

Now, $g \circ f(x) = g(f(x))$
 $= g(y)$ [since $f(x) = y$]
 $= z$ [since $g(y) = z$] which shows that $g \circ f$ is onto.

Theorem $(g \circ f)^{-1} = f^{-1} \circ g^{-1}$ (i.e) the inverse of a composite function can be expressed in terms of the composition of the inverses in the reverse order.

Proof. $f: A \rightarrow B$ is one to one and onto. $g: B \rightarrow C$ is one to one and onto.

$g \circ f: A \rightarrow C$ is also one to one and onto. $\exists (g \circ f)^{-1}$:

$C \rightarrow A$ is one to one and onto.

Let $a \in A$, then there exists an element $b \in B$ such that $f(a) = b$ $\exists a = f^{-1}(b)$

(c). Now $b \in B$ $\exists$ there exists an element $c \in C$ such that $g(b) = c$ $\exists b = g^{-1}(c)$.

Then $(g \circ f)(a) = g[f(a)] = g(b) = c$ $\exists a = (g \circ f)^{-1}(c)$ (1)

$(f^{-1} \circ g^{-1})(c) = f^{-1}(g^{-1}(c)) = f^{-1}(b) = a$ $\exists a = (f^{-1} \circ g^{-1})(c)$

)(2) Combining (1) and (2), we have $(g \circ f)^{-1} = f^{-1} \circ g^{-1}$

Theorem: If $f: A \rightarrow B$ is an invertible mapping ,
then $f \circ f^{-1} = I_B$ and $f^{-1} \circ f = I_A$

Proof: f is invertible, then f^{-1} is defined by $f(a) = b \iff f^{-1}(b) = a$ where $a \in A$ and $b \in B$.

Now we have to prove that $f \circ f^{-1} = I_B$

. Let $b \in B$ and $f^{-1}(b) = a, a \in A$

then $f \circ f^{-1}(b) = f(f^{-1}(b))$
 $= f(a) = b$

therefore $f \circ f^{-1}(b) = b \iff b \in B \Rightarrow f \circ f^{-1} = I_B$
Now $f^{-1} \circ f(a) = f^{-1}(f(a)) = f^{-1}(b) = a$
therefore $f^{-1} \circ f(a) = a \iff a \in A \Rightarrow f^{-1} \circ f = I_A$.
Hence the theorem.

Recursive Functions:

The term "recursive function" is often used informally to describe any function that is defined with recursion. There are several formal counterparts to this informal definition, many of which only differ in trivial respects.

Kleene (1952) defines a "partial recursive function" of nonnegative integers to be any function f that is defined by a noncontradictory system of equations whose left and right sides are composed from

(1) function symbols (for example, f, g etc.), (2) variables for nonnegative integers (for example, x, y, z , etc.), (3) the constant 0, and (4) the successor function $S(x) = x + 1$.

For example,

$$f(x, 0) = 0 \quad (1)$$

$$f(x, S(y)) = g(f(x, y), x) \quad (2)$$

$$g(x, 0) = x \quad (3)$$

$$g(x, S(y)) = S(g(x, y)) \quad (4)$$

defines $f(x, y)$ to be the function x that computes the product of x and y

Note that the equations might not uniquely determine the value of f for every possible input, and in that sense the definition is "partial." If the system of equations determines the value of f for every input, then the definition is said to be "total." When the term "recursive function" is used alone, it is usually implicit that "total recursive function" is intended. Note that some authors use the term "general recursive function" to mean partial recursive function, although others use it to mean "total recursive function."

The set of functions that can be defined recursively in this manner is known to be equivalent to the set of functions computed by Turing machines and by the lambda calculus.

Algebraic systems:

An algebraic system, loosely speaking, is a set, together with some operations on the set. Before formally defining what an algebraic system is, let us recall that a n -ary operation (or operator) on a set A is a function whose domain is A^n and whose range is a subset of A . Here, n is a non-negative integer. When $n=0$, the operation is usually called a nullary operation, or a constant, since one element of A is singled out to be the (sole) value of this operation. A finitary operation on A is just an n -ary operation for some non-negative integer n .

Definition. An *algebraic system* is an ordered pair (A, O) , where A is a set, called the underlying set of the algebraic system, and O is a set, called the operator set, of finitary operations on A .

We usually write A , instead of (A, O) , for brevity.

A prototypical example of an algebraic system is a group, which consists of the underlying set G , and a set O consisting of three operators: a constant e called the multiplicative identity, a unary operator called the multiplicative inverse, and a binary operator called the multiplication.

For a more comprehensive listing of examples, please see this entry.

Remarks.

- < An algebraic system is also called algebra for short. Some authors require that A be non-empty. Note that A is automatically non-empty if O contains constants. A *finite algebra* is an algebra whose underlying set is finite.
- < By definition, all operators in an algebraic system are finitary. If we allow O to contain infinitary operations, we have an *infinitary algebraic system*. Other generalizations are possible. For example, if the operations are allowed to be multivalued, the algebra is said to be a *multialgebra*. If the operations are not everywhere defined, we get a *partial algebra*. Finally, if more than one underlying set is involved, then the algebra is said to be *many-sorted*.

The study of algebraic systems is called the theory of universal algebra. The first important thing in studying algebraic system is to compare systems that are of the same "type". Two algebras are said to have the same *type* if there is a one-to-one correspondence between their operator sets such that an n -ary operator in one algebra is mapped to an n -ary operator in the other algebra.

Examples:

Some recurring universes: $\mathbb{N}$ =natural numbers; $\mathbb{Z}$ =integers; $\mathbb{Q}$ =rational numbers; $\mathbb{R}$ =real numbers; $\mathbb{C}$ =complex numbers.

$\mathbb{N}$ is a pointed unary system, and under addition and multiplication, is both the standard interpretation of Peano arithmetic and a commutative semiring.

Boolean algebras are at once semigroups, lattices, and rings. They would even be abelian groups if the identity and inverse elements were identical instead of complements.

Group-like structures

- < Nonzero $\mathbb{N}$ under addition (+) is a magma.
- < $\mathbb{N}$ under addition is a magma with an identity.
- < **$\mathbb{Z}$ under subtraction (−) is a quasigroup.**
- < Nonzero $\mathbb{Q}$ under division ($\div$) is a quasigroup. $a^{-1} * b$, and $y * a = b$ if
- < Every group is a loop, because $a * x = b$ if and only if $x = a^{-1} * b$ and only if $y = b * a^{-1}$.
- < 2x2 matrices(of non-zero determinant) with matrix multiplication form a group.
- < $\mathbb{Z}$ under addition (+) is an abelian group.
- < Nonzero $\mathbb{Q}$ under multiplication ($\times$) is an abelian group.
- < Every cyclic group G is abelian, because if x, y are in G , then $xy = yx$. In particular, $\mathbb{Z}$ is an abelian group under addition, as is the integers modulo n $\mathbb{Z}/n\mathbb{Z}$.
- < A monoid is a category with a single object, in which case the composition of morphisms and the identity morphism interpret monoid multiplication and identity element, respectively.
- < The Boolean algebra 2 is a boundary algebra.

General Properties:

Property of Closure

If we take two *real numbers* and multiply them together, we get another real number. (The real numbers are all the rational numbers and all the irrational numbers.) Because this is always true, we say that the real numbers are "closed under the operation of multiplication": there is no way to escape the set. When you combine any two elements of the set, the result is also included in the set.

Real numbers are also closed under addition and subtraction. They are not closed under the square root operation, because the square root of -1 is not a real number.

Inverse

The inverse of something is that thing turned inside out or upside down. The inverse of an operation undoes the operation: division undoes multiplication.

A number's *additive inverse* is another number that you can add to the original number to get the additive identity. For example, the additive inverse of 67 is -67, because $67 + -67 = 0$, the additive identity.

Similarly, if the product of two numbers is the *multiplicative identity*, the numbers are *multiplicative inverses*. Since $6 * 1/6 = 1$ (the multiplicative identity), the multiplicative inverse of 6 is $1/6$.

Zero does not have a multiplicative inverse, since no matter what you multiply it by, the answer is always 0, not 1.

Equality

The equals sign in an equation is like a scale: both sides, left and right, must be the same in order for the scale to stay in balance and the equation to be true.

The *addition property of equality* says that if $a = b$, then $a + c = b + c$: if you add the same number to (or subtract the same number from) both sides of an equation, the equation continues to be true.

The *multiplication property of equality* says that if $a = b$, then $a * c = b * c$: if you multiply (or divide) by the same number on both sides of an equation, the equation continues to be true.

The *reflexive property of equality* just says that $a = a$: anything is congruent to itself: the equals sign is like a mirror, and the image it "reflects" is the same as the original.

The *symmetric property of equality* says that if $a = b$, then $b = a$.

The *transitive property of equality* says that if $a = b$ and $b = c$, then $a = c$.

Semi groups and monoids:

In the previous section, we have seen several algebraic system with binary operations. Here we consider an algebraic system consisting of a set and an associative binary operation on the set and then the algebraic system which possess an associative property with an identity element. These algebraic systems are called semigroups and monoids.

Semi group

Let S be a nonempty set and let $*$ be a binary operation on S . The algebraic system $(S, *)$ is called a semi-group if $*$ is associative

if $a * (b * c) = (a * b) * c$ for all $a, b, c \in S$.

Example The $\mathbb{N}$ of natural numbers is a semi-group under the operation of usual addition of numbers.

Monoids

Let M be a nonempty set with a binary operation $*$ defined on it. Then $(M, *)$ is called a monoid if

- $*$ is associative

(i.e) $a * (b * c) = (a * b) * c$ for all $a, b, c \in M$ and
there exists an element e in M such that

$a * e = e * a = a$ for all $a \in M$

e is called the identity element in $(M, *)$.

It is easy to prove that the identity element is unique. From the definition it follows that $(M, *)$ is a semigroup with identity.

Example1 Let S be a nonempty set and $\mathcal{P}(S)$ be its power set. The algebras $(\mathcal{P}(S), \cup)$ and $(\mathcal{P}(S), \cap)$ are monoids with the identities $\emptyset$ and S respectively.

Example2 Let $\mathbb{N}$ be the set of natural numbers, then $(\mathbb{N}, +)$, $(\mathbb{N}, \times)$ are monoids with the identities 0 and 1 respectively.

Groups Sub Groups:

Recalling that an algebraic system $(S, *)$ is a semigroup if the binary operation $*$ is associative. If there exists an identity element $e \in S$, then $(S, *)$ is monoid. A further condition is imposed on the elements of the monoid, i.e., the existence of an inverse for each element of S then the algebraic system is called a group.

Definition

Let G be a nonempty set, with a binary operation $*$ defined on it. Then the algebraic system $(G, *)$ is called a group if

- $*$ is associative i.e. $a * (b * c) = (a * b) * c$ for all $a, b, c \in G$.
- there exists an element e in G such that $a * e = e * a = a$ for all $a \in G$
- for each $a \in G$ there is an element denoted by a^{-1} in G such that $a * a^{-1} = a^{-1} * a = e$, a^{-1} is called the inverse of a .

From the definition it follows that $(G, *)$ is a monoid in which each element has an inverse w.r.t. $*$ in G .

A group $(G, *)$ in which $*$ is commutative is called an abelian group or a commutative group. If $*$ is not commutative then $(G, *)$ is called a non-abelian group or non-commutative group.

The order of a group $(G, *)$ is the number of elements of G , when G is finite and is denoted by $o(G)$ or $|G|$

Examples 1. $(\mathbb{Z}_5, +)$ is an abelian group of order 5.

2. $G = \{1, -1, i, -i\}$ is an abelian group with the binary operation x is defined as $1 \times 1 = 1, -1 \times -1 = 1, i \times i = -1, -i \times -i = 1, \dots$

Homomorphism of semigroups and monoids

Semigroup homomorphism.

Let $(S, *)$ and (T, D) be any two semigroups. A mapping $g: S \rightarrow T$ such that any two elements $a, b \in S$, $g(a * b) = g(a) D g(b)$ is called a semigroup homomorphism.

Monoid homomorphism

Let $(M, *, e_M)$ and (T, D, e_T) be any two monoids. A mapping $g: M \rightarrow T$ such that any two elements $a, b \in M$,

$$\begin{aligned} g(a * b) &= g(a) D g(b) \\ \text{and } g(e_M) &= e_T \end{aligned}$$

is called a monoid homomorphism.

Theorem 1 Let $(s, *)$, (T, D) and (V, Δ) be semigroups. A mapping $g: S \rightarrow T$ and $h: T \rightarrow V$ be semigroup homomorphisms. Then $(hog): S \rightarrow V$ is a semigroup homomorphism from $(S, *)$ to (V, Δ) .

Proof. Let $a, b \in S$. Then

$$\begin{aligned} (h \circ g)(a * b) &= h(g(a * b)) \\ &= h(g(a) D g(b)) \\ &= h(g(a)) \Delta h(g(b)) \\ &= (h \circ g)(a) \Delta (h \circ g)(b) \end{aligned}$$

Theorem 2 Let $(s, *)$ be a given semigroup. There exists a homomorphism $g: S \rightarrow SS$, where $(SS, \circ)$ is a semigroup of function from S to S under the operation of composition.

Proof For any element $a \in S$, let $g(a) = f_a$ where $f_a \in SS$ and f_a is defined by

$$f_a(b) = a * b \quad \text{for any } a, b \in S$$

$$g(a * b) = f_{a * b}$$

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

$$\text{where } = f_a(f_b(c)) = (f_a \circ f_b)(c).$$

Therefore, $g(a * b) = f a * b = f a \circ f b = g(a) \circ g(b)$, this shows that $g: S \rightarrow SS$ is a homomorphism.

Theorem 3 For any commutative monoid $(M, *)$, the set of idempotent elements of M forms a submonoid.

Proof. Let S be the set of idempotent elements of M .

Since the identity element $e \in M$ is idempotent, $e \in S$.

Let $a, b \in S$, so that $a * a = a$ and $b * b = b$

Now $(a * b) * (a * b) = (a * b) * (b * a)$ [$(M, *)$ is a commutative monoid]

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

$$= a * b * a$$

$$= a * a * b$$

$$= a * b$$

Hence $a * b \in S$ and $(S, *)$ is a submonoid.

Isomorphism:

-1

In abstract algebra, an isomorphism is a bijective map f such that both f and its inverse f^{-1} are homomorphisms, i.e., *structure-preserving* mappings. In the more general setting of category theory, an **isomorphism** is a morphism $f: X \rightarrow Y$ in a category for which there exists an

"inverse" $f^{-1}: Y \rightarrow X$, with the property that both f

$$f \circ f^{-1} = \text{id}_Y \text{ and } f^{-1} \circ f = \text{id}_X.$$

UNIT-III

Elementary Combinatorics

Basis of counting:

If X is a set, let us use $|X|$ to denote the number of elements in X .

Two Basic Counting Principles

Two elementary principles act as —building blocks< for all counting problems. The first principle says that the whole is the sum of its parts; it is at once immediate and elementary.

Sum Rule: The principle of disjunctive counting :

If a set X is the union of disjoint nonempty subsets $S_1, \dots, S_n$, then $|X| = |S_1| + |S_2| + \dots + |S_n|$.

We emphasize that the subsets $S_1, S_2, \dots, S_n$ must have no elements in common. Moreover, since $X = S_1 \cup S_2 \cup \dots \cup S_n$, each element of X is in exactly one of the subsets S_i . In other words, $S_1, S_2, \dots, S_n$ is a partition of X .

If the subsets $S_1, S_2, \dots, S_n$ were allowed to overlap, then a more profound principle will be needed--the principle of inclusion and exclusion.

Frequently, instead of asking for the number of elements in a set *per se*, some problems ask for how many ways a certain event can happen.

The difference is largely in semantics, for if A is an event, we can let X be the set of ways that A can happen and count the number of elements in X . Nevertheless, let us state the sum rule for counting events.

If $E_1, \dots, E_n$ are mutually exclusive events, and E_1 can happen e_1 ways, E_2 happen e_2 ways, ..., E_n can happen e_n ways, E_1 or E_2 or ... or E_n can happen $e_1 + e_2 + \dots + e_n$ ways.

Again we emphasize that mutually exclusive events E_1 and E_2 mean that E_1 or E_2 can happen but both cannot happen simultaneously.

The sum rule can also be formulated in terms of choices: If an object can be selected from a reservoir in e_1 ways and an object can be selected from a separate reservoir in e_2 ways and an object can be selected from a separate reservoir in e_2 ways, then the selection of one object from either one reservoir or the other can be made in $e_1 + e_2$ ways.

Product Rule: The principle of sequencing counting

If $S_1, \dots, S_n$ are nonempty sets, then the number of elements in the Cartesian product $S_1 \times S_2 \times \dots \times S_n$ is the product $\prod_{i=1}^n |S_i|$. That is,

$$|S_1 \times S_2 \times \dots \times S_n| = \prod_{i=1}^n |S_i|.$$

Observe that there are 5 branches in the first stage corresponding to the 5 elements of S_1 and to each of these branches there are 3 branches in the second stage corresponding to the 3 elements of S_2 giving a total of 15 branches altogether. Moreover, the Cartesian product $S_1 \times S_2$ can be partitioned as $(a_1 \times S_2) \cup (a_2 \times S_2) \cup (a_3 \times S_2) \cup (a_4 \times S_2) \cup (a_5 \times S_2)$, where $(a_i \times S_2) = \{(a_i, b_1), (a_i, b_2), (a_i, b_3)\}$. Thus, for example, $(a_3 \times S_2)$ corresponds to the third branch in the first stage followed by each of the 3 branches in the second stage.

More generally, if $a_1, \dots, a_n$ are the n distinct elements of S_1 and $b_1, \dots, b_m$ are the m distinct elements of S_2 , then $S_1 \times S_2 = \bigcup_{i=1}^n (a_i \times S_2)$.

For if x is an arbitrary element of $S_1 \times S_2$, then $x = (a, b)$ where $a \in S_1$ and $b \in S_2$. Thus, $a = a_i$ for some i and $b = b_j$ for some j . Thus, $x = (a_i, b_j) \in (a_i \times S_2)$ and therefore $x \in \bigcup_{i=1}^n (a_i \times S_2)$.

Conversely, if $x \in \bigcup_{i=1}^n (a_i \times S_2)$, then $x \in (a_i \times S_2)$ for some i and thus $x = (a_i, b_j)$ where b_j is some element of S_2 . Therefore, $x \in S_1 \times S_2$.

Next observe that $(a_i \times S_2)$ and $(a_j \times S_2)$ are disjoint if $i \neq j$ since if

$x \in (a_i \times S_2) \cap (a_j \times S_2)$ then $x = (a_i, b_k)$ for some k and $x = (a_j, b_l)$ for some l .

But then $(a_i, b_k) = (a_j, b_l)$ implies that $a_i = a_j$ and $b_k = b_l$. But since $i \neq j$, $a_i \neq a_j$.

Thus, we conclude that $S_1 \times S_2$ is the disjoint union of the sets $(a_i \times S_2)$. Furthermore $|a_i \times S_2| = |S_2|$ since there is obviously a one-to-one correspondence between the sets $a_i \times S_2$ and S_2 , namely, $(a_i, b_j) \rightarrow b_j$.

Then by the sum rule $|S_1 \times S_2| = \sum_{i=1}^n |a_i \times S_2|$

$$7. (n \text{ summands}) |S_2| + |S_2| + \dots + |S_2|$$

$$8. n |S_2|$$

$$9. nm.$$

Therefore, we have proven the product rule for two sets. The general rule follows by mathematical induction.

We can reformulate the product rule in terms of events. If events $E_1, E_2, \dots, E_n$ can happen $e_1, e_2, \dots, e_n$ ways, respectively, then the sequence of events E_1 first,

followed by $E_2, \dots$, followed by E_n can happen $e_1 e_2 \dots e_n$ ways.

In terms of choices, the product rule is stated thus: If a first object can be chosen e_1 ways, **a second e_2 ways , ..., and an n th object can be made in $e_1 e_2 \dots e_n$ ways.**

Combinations & Permutations:

Definition.

A combination of n objects taken r at a time (called an r -combination of n objects) is an unordered selection of r of the objects.

A permutation of n objects taken r at a time (also called an r -permutation of n objects) is an ordered selection or arrangement of r of the objects.

Note that we are simply defining the terms r -combinations and r -permutations here and have not mentioned anything about the properties of the n objects.

For example, these definitions say nothing about whether or not a given element may appear more than once in the list of n objects.

In other words, it may be that the n objects do not constitute a set in the normal usage of the word.

SOLVED PROBLEMS

Example 1. Suppose that the 5 objects from which selections are to be made are: a, a, a, b, c . then the 3-combinations of these 5 objects are : aaa, aab, aac, abc . The permutations are:

$aaa, aab, aba, baa, aac, aca, caa,$
 $abc, acb, bac, bca, cab, cba.$

Neither do these definitions say anything about any rules governing the selection of the r -objects: on one extreme, objects could be chosen where all repetition is forbidden, or on the other extreme, each object may be chosen up to t times, or then again may be some rule of selection between these extremes; for instance, the rule that would allow a given object to be repeated up to a certain specified number of times.

We will use expressions like $\{3 . a , 2 . b , 5 . c\}$ to indicate either

(1) that we have $3 + 2 + 5 = 10$ objects including $3a$'s , $2b$'s and $5c$'s, or (2) that we have 3 objects a, b, c , where selections are constrained by the conditions that a can be selected at most three times, b can be selected at most twice, and c can be chosen up to five times.

The numbers 3, 2 and 5 in this example will be called repetition numbers.

Example 2 The 3-combinations of $\{3 . a, 2 . b, 5 . c\}$ are:

$aaa, aab, aac, abb,$
 $abc, ccc, ccb, cca,$
 $cbb.$

Example 3. **The 3-combinations of $\{3 . a, 2 . b, 2 . c , 1 . d\}$ are:**

**$aaa, aab, aac, aad, bba, bbc, bbd,$
 $cca, ccb, ccd, abc, abd, acd, bcd.$**

In order to include the case where there is no limit on the number of times an object can be repeated in a selection (except that imposed by the size of the selection) we use the **symbol ∞ as a repetition number to mean that an object can occur an infinite number of times.**

Example 4. The 3-combinations of $\{\infty. a, 2.b, \infty.c\}$ are the same as in Example 2 even though a and c can be repeated an infinite number of times. This is because, in 3-combinations, 3 is the limit on the number of objects to be chosen.

If we are considering **selections where each object has ∞ as its repetition number then** we designate such selections as selections with unlimited repetitions. In particular, a selection of r objects in this case will be called r-combinations with unlimited repetitions and any ordered arrangement of these r objects will be an r-permutation with unlimited repetitions.

Example5 The combinations of a ,b, c, d with unlimited repetitions are the 3-combinations of $\{\infty . a , \infty. b, \infty. c, \infty. d\}$. These are 20 such 3-combinations, namely:

aaa, aab, aac, aad,
 bbb, bba, bbc, bbd,
 ccc, cca, ccb, ccd,
 ddd, dda, ddb, ddc,
 abc, abd, acd, bcd.

2-combinations with Unlimited Repetitions

aa
 ab
 ac
 ad
 bb
 bc
 bd
 cc
 cd
 dd
 10

2-permutations with Unlimited Repetitions

Aa
 ab, ba
 ac, ca
 ad, da
 Bb
 bc, cb
 bd, db
 Cc
 cd, dc
 Dd
 16

Moreover, there are $4^3 = 64$ of 3-permutations with unlimited repetitions since the first position can be filled 4 ways (with a, b, c, or d), the second position can be filled 4 ways, and likewise for the third position.

The 2-permutations of $\{\infty. a, \infty. b, \infty. c, \infty. d\}$ do not present such a formidable list and so we tabulate them in the following table.

We list some more examples just for concreteness. We might, for example, consider selections of $\{\infty.a, \infty. b, \infty. c\}$ where b can be chosen only even number of times. Thus, 5-combinations with these repetition numbers and this constraint would be those 5-combinations with unlimited repetitions and where b is chosen 0, 2, or 4 times.

Example6 The 3-combinations of $\{\infty .a, \infty .b, 1 .c, 1 .d\}$ where b can be chosen only an even number of times are the 3-combinations of a, b, c, d where a can be chosen up 3 times, b can be chosen 0 or 2 times, and c and d can be chosen at most once. The 3-combinations subject to these constraints are:

aaa, aac, aad, bbc, bbd, acd.

As another example, we might be interested in, selections of $\{\infty.a, 3.b, 1.c\}$ where a can be chosen a prime number of times. Thus, the 8-combinations subject to these constraints would be all those 8-combinations where a can be chosen 2, 3, 5, or 7 times, b can chosen up to 3 times, and c can be chosen at most once.

There are, as we have said, an infinite variety of constraints one could place on selections. You can just let your imagination go free in conjuring up different constraints on the selection, would constitute an r-combination according to our definition. Moreover, any arrangement of these r objects would constitute an r-permutation.

While there may be an infinite variety of constraints, we are primarily interested in two major types: one we have already described—combinations and permutations with unlimited repetitions, the other we now describe.

If the repetition numbers are all 1, then selections of r objects are called r-combinations without repetitions and arrangements of the r objects are r-permutations without repetitions. We remind you that r-combinations without repetitions are just subsets of the n elements containing exactly r elements. Moreover, we shall often drop the repetition number 1 when considering r-combinations without repetitions. For example, when considering r-combinations of $\{a, b, c, d\}$ we will mean that each repetition number is 1 unless otherwise designated, and, of course, we mean that in a given selection an element need not be chosen at all, but, if it is chosen, then in this selection this element cannot be chosen again.

Example7. Suppose selections are to be made from the four objects a, b, c, d.

2-combinations without Repetitions	2-Permutations without Repetitions
ab	ab, ba
ac	ac, ca
ad	ad, da
bc	bc, cb

bd	bd, db
cd	cd, dc
6	12

There are six 2-combinations without repetitions and to each there are two 2-permutations giving a total of twelve 2-permutations without repetitions.

Note that total number of 2-combinations with unlimited repetitions in Example 5 included six 2-combinations without repetitions of Example.7 and as well 4 other 2-combinations where repetitions actually occur. Likewise, the sixteen 2-permutations with unlimited repetitions included the twelve 2-permutations without repetitions.

3-combinations without Repetitions	3-Permutations without Repetitions
abc	abc, acb, bac, bca, cab, cba
abd	abd, adb, bad, bda, dab, dba
acd	acd, adc, cad, cda, dac, dca
bcd	bcd, bdc, cbd, cdb, dbc, dcb
4	24

Note that to each of the 3-combinations without repetitions there are 6 possible 3- permutations without repetitions. Momentarily, we will show that this observation can be generalized.

Combinations And Permutations With Repetitions:

General formulas for enumerating combinations and permutations will now be presented. At this time, we will only list formulas for combinations and permutations without repetitions or with unlimited repetitions. We will wait until later to use generating functions to give general techniques for enumerating combinations where other rules govern the selections.

Let $P(n, r)$ denote the number of r -permutations of n elements without repetitions.

Theorem 5.3.1.(Enumerating r -permutations without repetitions).

$$P(n, r) = n(n-1)..... (n - r + 1) = n! / (n-r)!$$

Proof. Since there are n distinct objects, the first position of an r -permutation may be filled in n ways. This done, the second position can be filled in $n-1$ ways since no repetitions are allowed and there are $n - 1$ objects left to choose from. The third can be filled in $n-2$ ways. By applying the product rule, we conduct that

$$P(n, r) = n(n-1)(n-2)..... (n - r + 1).$$

From the definition of factorials, it follows that

$$P(n, r) = n! / (n-r)!$$

When $r = n$, this formula becomes

$$P(n, n) = n! / 0! = n!$$

When we explicit reference to r is not made, we assume that all the objects are to be arranged; thus we talk about the permutations of n objects we mean the case $r=n$. Corollary 1. There are $n!$ permutations of n distinct objects.

Example 1.

There are $3! = 6$ permutations of $\{a, b, c\}$.

There are $4! = 24$ permutations of $\{a, b, c, d\}$. The number of 2-permutations $\{a, b, c, d, e\}$ is $P(5, 2) = 5! / (5 - 2)! = 5 \times 4 = 20$.

The number of 5-letter words using the letters a, b, c, d , and e at most once is $P(5, 5) = 120$.

Example 2 There are $P(10, 4) = 5,040$ 4-digit numbers that contain no repeatd digits since each such number is just an arrangement of four of the digits $0, 1, 2, 3, \dots, 9$ (leading zeroes are allowed). There are $P(26, 3) P(10, 4)$ license plates formed by 3 distinct letters followed by 4 distinct digits.

Example3. In how many ways can 7 women and 3 men be arranged in a row if the 3 men must always stand next to each other?

There are $3!$ ways of arranging the 3 men. Since the 3 men always stand next to each other, we treat them as a single entity, which we denote by X . Then if $W_1, W_2, \dots, W_7$ represents the women, we next are interested in the number of ways of arranging $\{X, W_1, W_2, W_3, \dots, W_7\}$. There are $8!$ permutations these 8 objects. Hence there are $(3!)(8!)$ permutations altogether. (of course, if there has to be a prescribed order of an arrangement on the 3 men then there are only $8!$ total permutations).

Example4. In how many ways can the letters of the English alphabet be arranged so that there are exactly 5 letters between the letters a and b ?

There are $P(24, 5)$ ways to arrange the 5 letters between a and b , 2 ways to place a and b , and then $20!$ ways to arrange any 7-letter word treated as one unit along with the remaining 19 letters. The total is $P(24, 5)(20!)(2)$.

permutations for the objects are being arranged in a line. If instead of arranging objects in a line, we arrange them in a circle, then the number of permutations decreases.

Example 5. In how many ways can 5 children arrange themselves in a ring?

Solution. Here, the 5 children are not assigned to particular places but are only arranged relative to one another. Thus, the arrangements (see Figure 2-3) are considered the same if the children are in the same order clockwise. Hence, the position of child C_1 is immaterial and it is only the position of the 4 other children relative to C_1 that counts. Therefore, keeping C_1 fixed in position, there are $4!$ arrangements of the remaining children.

Binomial Coefficients: In mathematics, the binomial coefficient $\binom{n}{k}$ is the coefficient of the x^k term in the polynomial expansion of the binomial power $(1 + x)^n$.

In combinatorics, $\binom{n}{k}$ is interpreted as the number of k -element subsets (the k -combinations) of an n -element set, that is the number of ways that k things can be "chosen" from a set of n things.

Hence, $\binom{n}{k}$ is often read as " n choose k " and is called the choose function of n and k . The notation $\binom{n}{k}$ was introduced by Andreas von Ettingshausen in 182, although the numbers were already known centuries before that (see Pascal's triangle). Alternative notations include $C(n, k)$,

nC_k , C_k^n , C_k^n , C_n^k , in all of which the C stands for combinations or choices.

For natural numbers (taken to include 0) n and k , the binomial coefficient $\binom{n}{k}$ can be defined as the coefficient of the monomial X^k in the expansion of $(1 + X)^n$. The same coefficient also

occurs (if $k \leq n$) in the binomial formula

$$(x + y)^n = \sum_{k=0}^n \binom{n}{k} x^{n-k} y^k$$

(valid for any elements x, y of a commutative ring), which explains the name "binomial coefficient".

Another occurrence of this number is in combinatorics, where it gives the number of ways, disregarding order, that k objects can be chosen from among n objects; more formally, the number of k -element subsets (or k -combinations) of an n -element set. This number can be seen to be equal to the one of the first definition, independently of any of the formulas below to compute

it: if in each of the n factors of the power $(1 + X)^n$ one temporarily labels the term X with an index i (running from 1 to n), then each subset of k indices gives after expansion a contribution

X^k , and the coefficient of that monomial in the result will be the number of such subsets. This shows in particular that $\binom{n}{k}$ is a natural number for any natural numbers n and k . There are many other combinatorial interpretations of binomial coefficients (counting problems for which the answer is given by a binomial coefficient expression), for instance the number of words formed of n bits (digits 0 or 1) whose sum is k , but most of these are easily seen to be equivalent to counting k -combinations.

Several methods exist to compute the value of $\binom{n}{k}$ without actually expanding a binomial power or counting k -combinations.

Binomial Multinomial theorems:

Binomial theorem:

In elementary algebra, the binomial theorem describes the algebraic expansion of powers of a binomial. According to the theorem, it is possible to expand the power $(x + y)^n$ into a sum

involving terms of the form $a x^b y^c$, where the coefficient of each term is a positive integer, and the sum of the exponents of x and y in each term is n . For example,

$$(x + y)^4 = x^4 + 4x^3y + 6x^2y^2 + 4xy^3 + y^4.$$

The coefficients appearing in the binomial expansion are known as binomial coefficients. They are the same as the entries of Pascal's triangle, and can be determined by a simple formula

involving factorials. These numbers also arise in combinatorics, where the coefficient of $x^{n-k} y^k$ is equal to the number of different combinations of k elements that can be chosen from an n -element set.

According to the theorem, it is possible to expand any power of $x + y$ into a sum of the form

$$(x + y)^n = \binom{n}{0} x^n y^0 + \binom{n}{1} x^{n-1} y^1 + \binom{n}{2} x^{n-2} y^2 + \binom{n}{3} x^{n-3} y^3 + \dots \\ \dots + \binom{n}{n-1} x^1 y^{n-1} + \binom{n}{n} x^0 y^n,$$

where $\binom{n}{k}$ denotes the corresponding binomial coefficient. Using summation notation, the formula above can be written

$$(x + y)^n = \sum_{k=0}^n \binom{n}{k} x^{n-k} y^k.$$

This formula is sometimes referred to as the **Binomial Formula** or the **Binomial Identity**.

A variant of the binomial formula is obtained by substituting 1 for x and x for y , so that it involves only a single variable. In this form, the formula reads

$$(1 + x)^n = \binom{n}{0} x^0 + \binom{n}{1} x^1 + \binom{n}{2} x^2 + \dots + \binom{n}{n-1} x^{n-1} + \binom{n}{n} x^n,$$

or equivalently

$$(1 + x)^n = \sum_{k=0}^n \binom{n}{k} x^k.$$

Multinomial theorem:

In mathematics, the **multinomial theorem** says how to write a power of a sum in terms of powers of the terms in that sum. It is the generalization of the binomial theorem to polynomials.

For any positive integer m and any nonnegative integer n , the multinomial formula tells us how a polynomial expands when raised to an arbitrary power:

$$(x_1 + x_2 + \cdots + x_m)^n = \sum_{k_1+k_2+\cdots+k_m=n} \binom{n}{k_1, k_2, \dots, k_m} x_1^{k_1} x_2^{k_2} \cdots x_m^{k_m}.$$

The summation is taken over all sequences of nonnegative integer indices k_1 through k_m such the sum of all k_i is n . That is, for each term in the expansion, the exponents must add up to n .

Also, as with the binomial theorem, quantities of the form x that appear are taken to equal 1 (even when x equals zero). Alternatively, this can be written concisely using multiindices as

where $\alpha = (\alpha_1, \alpha_2, \dots, \alpha_m)$ and $x^\alpha = x_1^{\alpha_1} x_2^{\alpha_2} \cdots x_m^{\alpha_m}$.

Example

$$(x_1 + \cdots + x_m)^n = \sum_{|\alpha|=n} \binom{n}{\alpha} x^\alpha$$

$$(a + b + c)^3 = a^3 + b^3 + c^3 + 3a^2b + 3a^2c + 3b^2a + 3b^2c + 3c^2a + 3c^2b + 6abc.$$

We could have calculated each coefficient by first expanding

$$(a + b + c)^2 = a^2 + b^2 + c^2 + 2ab + 2bc + 2ac, \text{ then self-multiplying it again to get } (a + b + c)^3$$

(and then if we were raising it to higher powers, we'd multiply it by itself even some more).

However this process is slow, and can be avoided by using the multinomial theorem. The multinomial theorem "solves" this process by giving us the closed form for any coefficient we might want. It is possible to "read off" the multinomial coefficients from the terms by using the multinomial coefficient formula. For example:

$$\begin{array}{lcl} \begin{matrix} 2 & 0 & 1 \\ a & b & c \end{matrix} & \text{has the coefficient} & \binom{3}{2, 0, 1} = \frac{3!}{2! \cdot 0! \cdot 1!} = \frac{6}{2 \cdot 1 \cdot 1} = 3 \\ \begin{matrix} 1 & 1 & 1 \\ a & b & c \end{matrix} & \text{has the coefficient} & \binom{3}{1, 1, 1} = \frac{3!}{1! \cdot 1! \cdot 1!} = \frac{6}{1 \cdot 1 \cdot 1} = 6 \end{array}$$

We could have also had a 'd' variable, or even more variables—hence the *multinomial* theorem.

The principles of Inclusion – Exclusion:

Let $|A|$ denote the cardinality of set A , then it follows immediately that

$$|A \cup B| = |A| + |B| - |A \cap B|, \quad (1)$$

where $\cup$ denotes union, and $\cap$ denotes intersection. The more general statement

$$\left| \bigcup_{i=1}^N E_i \right| \leq \sum_{i=1}^N |E_i|, \quad (2)$$

also holds, and is known as Boole's inequality.

This formula can be generalized in the following beautiful manner. Let $\mathcal{A} = \{A_i\}_{i=1}^p$ be a p -system of consisting of sets $A_1, \dots, A_p$, then

$$|A_1 \cup A_2 \cup \dots \cup A_p| = \sum_{1 \leq i \leq p} |A_i| - \sum_{1 \leq i_1 < i_2 \leq p} |A_{i_1} \cap A_{i_2}| + \sum_{1 \leq i_1 < i_2 < i_3 \leq p} |A_{i_1} \cap A_{i_2} \cap A_{i_3}| - \dots + (-1)^{p-1} |A_1 \cap A_2 \cap \dots \cap A_p|$$

where the sums are taken over k -subsets of $\mathcal{A}$. This formula holds for infinite sets S as well as finite sets.

The principle of inclusion-exclusion was used by Nicholas Bernoulli to solve the recontres problem of finding the number of derangements.

For example, for the three subsets $A_1 = \{2, 3, 7, 9, 10\}$, $A_2 = \{1, 2, 3, 9\}$, and $A_3 = \{2, 4, 9, 10\}$ of $S = \{1, 2, \dots, 10\}$, the following table summarizes the terms appearing the sum.

#	term	set	length
1	A_1	{2, 3, 7, 9, 10}	5
	A_2	{1, 2, 3, 9}	4
	A_3	{2, 4, 9, 10}	4
2	$A_1 \cap A_2$	{2, 3, 9}	3
	$A_1 \cap A_3$	{2, 9, 10}	3
	$A_2 \cap A_3$	{2, 9}	2
3		{2, 9}	2

$$(5 + 4 + 4) - (3 + 3 + 2) + 2 = 7$$

$$A_1 \cap A_2 \cap A_3 = \{ \}$$

$|A_1 \cup A_2 \cup A_3|$ is therefore equal to, corresponding to the seven elements $A_1 \cup A_2 \cup A_3 = \{1, 2, 3, 4, 7, 9, 10\}$.

Pigeon hole principles and its application:

The statement of the *Pigeonhole Principle*:

If m pigeons are put into m pigeonholes, there is an empty hole iff there's a hole with more than one pigeon.

If $n > m$ pigeons are put into m pigeonholes, there's a hole with more than one pigeon.

Example:

Consider a chess board with two of the diagonally opposite corners removed. Is it possible to cover the board with pieces of domino whose size is exactly two board squares?

Solution

No, it's not possible. Two diagonally opposite squares on a chess board are of the same color. Therefore, when these are removed, the number of squares of one color exceeds by 2 the number of squares of another color. However, every piece of domino covers exactly two squares and these are of different colors. Every placement of domino pieces establishes a 1-1 correspondence between the set of white squares and the set of black squares. If the two sets have different number of elements, then, by the Pigeonhole Principle, no 1-1 correspondence between the two sets is possible.

Generalizations of the pigeonhole principle

A generalized version of this principle states that, if n discrete objects are to be allocated to m containers, then at least one container must hold no fewer than $\lceil n/m \rceil$ objects, where $\lceil x \rceil$ is the ceiling function, denoting the smallest integer larger than or equal to x . Similarly, at least one container must hold no more than $\lfloor n/m \rfloor$ objects, where $\lfloor x \rfloor$ is the floor function, denoting the largest integer smaller than or equal to x .

A probabilistic generalization of the pigeonhole principle states that if n pigeons are randomly put into m pigeonholes with uniform probability $1/m$, then at least one pigeonhole will hold more than one pigeon with probability

$$1 - \frac{(m)_n}{m^n},$$

where $(m)_n$ is the falling factorial $m(m-1)(m-2)\dots(m-n+1)$. For $n=0$ and for $n=1$ (and $m > 0$), that probability is zero; in other words, if there is just one pigeon, there cannot be a conflict. For $n > m$ (more pigeons than pigeonholes) it is one, in which case it coincides with the ordinary pigeonhole principle. But even if the number of pigeons does not exceed the number of pigeonholes ($n \leq m$), due to the random nature of the assignment of pigeons to pigeonholes there is often a substantial chance that clashes will occur. For example, if 2 pigeons are randomly assigned to 4 pigeonholes, there is a 25% chance that at least one pigeonhole will hold more than one pigeon; for 5 pigeons and 10 holes, that probability is 69.76%; and for 10 pigeons and 20 holes it is about 93.45%. If the number of holes stays fixed, there is always a greater probability of a pair when you add more pigeons. This problem is treated at much greater length at birthday paradox.

A further probabilistic generalisation is that when a real-valued random variable X has a finite mean $E(X)$, then the probability is nonzero that X is greater than or equal to $E(X)$, and similarly the probability is nonzero that X is less than or equal to $E(X)$. To see that this implies the standard pigeonhole principle, take any fixed arrangement of n pigeons into m holes and let X be the number of pigeons in a hole chosen uniformly at random. The mean of X is n/m , so if there are more pigeons than holes the mean is greater than one. Therefore, X is sometimes at least 2.

Applications:

The pigeonhole principle arises in computer science. For example, collisions are inevitable in a hash table because the number of possible keys exceeds the number of indices in the array. No hashing algorithm, no matter how clever, can avoid these collisions. This principle also proves that any general-purpose lossless compression algorithm that makes at least one input file smaller will make some other input file larger. (Otherwise, two files would be compressed to the same smaller file and restoring them would be ambiguous.)

A notable problem in mathematical analysis is, for a fixed irrational number a , to show that the set $\{[na]: n \text{ is an integer}\}$ of fractional parts is dense in $[0, 1]$. After a moment's thought, one

finds that it is not easy to explicitly find integers n_1, n_2 such that $|n_1a - n_2a| < \epsilon$, where $\epsilon > 0$ is a name small positive number and a is some arbitrary irrational number. But if one takes M such that $1/M < \epsilon$, by the pigeonhole principle there must be $n_1, n_2 \in \{1, 2, \dots, M+1\}$ such that n_1a and n_2a are in the same integer subdivision of size $1/M$ (there are only M such subdivisions between consecutive integers). In particular, we can find n_1, n_2 such that n_1a is in $(p + k/M, p + (k+1)/M)$, and n_2a is in $(q + k/M, q + (k+1)/M)$, for some p, q integers and $k \in \{0, 1, \dots, M-1\}$. We can then easily verify that $(n_2 - n_1)a$ is in $(q - p - 1/M, q - p + 1/M)$. This implies that $[na] < 1/M < \epsilon$, where $n = n_2 - n_1$ or $n = n_1 - n_2$. This shows that 0 is a limit point of $\{[na]\}$. We can then use this fact to prove the case for p in $(0, 1]$: find n such that $[na] < 1/M < \epsilon$; then if $p \in (0, 1/M]$, we are done. Otherwise $p \in (1/M, (k+1)/M]$, and by setting $n = \sup\{n : [na] < 1/M\}$, we have $|(k+1)na - p| < 1/M < \epsilon$.

UNIT-IV

Recurrence Relation

Generating Functions:

In mathematics, a **generating function** is a formal power series in one indeterminate, whose coefficients encode information about a sequence of numbers a_n that is indexed by the natural numbers. Generating functions were first introduced by Abraham de Moivre in 1730, in order to solve the general linear recurrence problem. One can generalize to formal power series in more than one indeterminate, to encode information about arrays of numbers indexed by several natural numbers.

Generating functions are not functions in the formal sense of a mapping from a domain to a codomain; the name is merely traditional, and they are sometimes more correctly called generating series.

Ordinary generating function

The *ordinary generating function* of a sequence a_n is

$$G(a_n; x) = \sum_{n=0}^{\infty} a_n x^n.$$

When the term *generating function* is used without qualification, it is usually taken to mean an ordinary generating function.

If a_n is the probability mass function of a discrete random variable, then its ordinary generating function is called a probability-generating function.

The ordinary generating function can be generalized to arrays with multiple indices. For example, the ordinary generating function of a two-dimensional array $a_{m,n}$ (where n and m are natural numbers) is

$$G(a_{m,n}; x, y) = \sum_{m,n=0}^{\infty} a_{m,n} x^m y^n.$$

Example:

$$G(n^2; x) = \sum_{n=0}^{\infty} n^2 x^n = \frac{x(x+1)}{(1-x)^3}$$

Exponential generating function

The *exponential generating function* of a sequence a_n is

$$\text{EG}(a_n; x) = \sum_{n=0}^{\infty} a_n \frac{x^n}{n!}.$$

Example:

$$\text{EG}(n^2; x) = \sum_{n=0}^{\infty} \frac{n^2 x^n}{n!} = x(x+1)e^x$$

Function of Sequences:

Generating functions giving the first few powers of the nonnegative integers are given in the following table.

n^p	$f(x)$	series
1	$\frac{x}{1-x}$	$x + x^2 + x^3 + \dots$
n	$\frac{x}{(1-x)^2}$	$x + 2x^2 + 3x^3 + 4x^4 + \dots$
n^2	$\frac{x(x+1)}{(1-x)^3}$	$x + 4x^2 + 9x^3 + 16x^4 + \dots$
n^3	$\frac{x(x^2+4x+1)}{(1-x)^4}$	$x + 8x^2 + 27x^3 + \dots$
n^4	$\frac{x(x+1)(x^2+10x+1)}{(1-x)^5}$	$x + 16x^2 + 81x^3 + \dots$

There are many beautiful generating functions for special functions in number theory. A few particularly nice examples are

$$f(x) = \frac{1}{(x)_{\infty}} \quad (2)$$

$$\sum_{n=0}^{\infty} P(n) x^n \quad (3)$$

$$= 1 + x + 2x^2 + 3x^3 + \dots \quad (4)$$

for the partition function P , where $(q)_\infty$ is a q -Pochhammer symbol, and

$$\frac{x}{1-x-x^2} \quad (5)$$

$$\sum_{n=0}^{\infty} F_n x^n \quad (6)$$

$$= x + x^2 + 2x^3 + 3x^4 + \dots \quad (7)$$

for the Fibonacci numbers F_n .

Generating functions are very useful in combinatorial enumeration problems. For example, the subset sum problem, which asks the number of ways C_m to select out of given integers such that their sum equals s , can be solved using generating functions.

Calculating Coefficient of generating function:

By using the following polynomial expansions, we can calculate the coefficient of a generating function.

Polynomial Expansions:

$$1) \frac{1}{1-x} = 1 + x + x^2 + x^3 + \dots$$

$$2) \frac{1}{1-x^2} = 1 + x^2 + x^4 + \dots$$

$$3) (1-x)^n = 1 - C(n,1)x + C(n,2)x^2 - \dots + (-1)^r C(n,r)x^r + \dots + (-1)^n C(n,n)x^n$$

$$4) (1+x)^n = 1 + C(n,1)x + C(n,2)x^2 + \dots + C(n,r)x^r + \dots + C(n,n)x^n$$

$$5) \frac{1}{(1-x)^n} = 1 + C(n-1,1)x + C(n-1,2)x^2 + \dots + C(n-1,r)x^r + \dots$$

$$6) \text{ If } h(x) = f(x)g(x), \text{ where } f(x) = a_0 + a_1x + a_2x^2 + \dots \text{ and } g(x) = b_0 + b_1x + b_2x^2 + \dots, \text{ then}$$

$$h(x) = a_0b_0 + (a_0b_1 + a_1b_0)x + (a_0b_2 + a_1b_1 + a_2b_0)x^2 + \dots + (a_0b_r + a_1b_{r-1} + \dots + a_rb_0)x^r + \dots$$

Example

Find the coefficient of x^{16} in $(x^2 + x^3 + x^4 + \dots)^5$

x^{16} in $x^{16}(1-x)^{-5}$ [i.e., the x^6 term in $(1-x)^{-5}$ is

multiplied by x^{10} to become the x^{16} term in $x^{10}(1-x)^{-5}$]

To simplify the expression, we extract x^2 from each polynomial factor and then apply identity (2).

$$\begin{aligned}(x^2 + x^3 + x^4 + \dots)^5 &= [x^2(1 + x + x^2 + \dots)]^5 \\ &= x^{10}(1 + x + x^2 + \dots)^5 \\ &= x^{10} \frac{1}{(1-x)^5}\end{aligned}$$

Thus the coefficient of x^{16} in $(x^2 + x^3 + x^4 + \dots)^5$ is the coefficient of x^6 in $x^{10}(1-x)^{-5}$ multiplied by x^{10} [i.e., the x^6 term in $x^{10}(1-x)^{-5}$]

$$\frac{1}{(1-x)^5} = 1 + C(1+n-1, 1)x + C(2+n-1, 2)x^2 + \dots + C(r+n-1, r)x^r + \dots$$

From expansion (5) we see that the coefficient of x^6 in $(1-x)^{-5}$ is $C(6+5-1, 6)$

More generally, the coefficient of x^r in

x^r in $x^{10}(1-x)^{-5}$ equals the coefficient of x^{r-10} in $(1-x)^{-5}$, namely, $C((r-10)+5-1, (r-10))$.

Recurrence relations:

Introduction : A recurrence relation is a formula that relates for any integer $n \geq 1$, the n -th term of a sequence $A = \{a_n\}_{n=0}^{\infty}$ to one or more of the terms $a_0, a_1, \dots, a_{n-1}$. Example. If S_n denotes the sum of the first n positive integers, then

10. $S_n = n + S_{n-1}$. Similarly if d is a real number, then the n th term of an arithmetic progression with common difference d satisfies the relation

11. $a_n = a_{n-1} + d$. Likewise if p_n denotes the n th term of a geometric progression with common ratio r , then

$p_n = r p_{n-1}$. We list other examples

as: $a_n - 3a_{n-1} + 2a_{n-2} = 0$.

$a_n - 3a_{n-1} + 2a_{n-2} = n^2 + 1$.

$a_n - (n-1)a_{n-1} - (n-1)a_{n-2} = 0$.

$a_n - 9a_{n-1} + 26a_{n-2} - 24a_{n-3} = 5n$.

$a_n - 3(a_{n-1})^2 + 2a_{n-2} = n$.

$a_n = a_0 a_{n-1} + a_1 a_{n-2} + \dots + a_{n-1} a_0$.

$a_{2n} + (a_{n-1})^2 = -1$.

Definition. Suppose n and k are nonnegative integers. A recurrence relation of the form $c_0(n)a_n + c_1(n)a_{n-1} + \dots + c_k(n)a_{n-k} = f(n)$ for $n \geq k$, where $c_0(n), c_1(n), \dots, c_k(n)$, and $f(n)$ are functions of n is said to be a linear recurrence relation. If $c_0(n)$ and $c_k(n)$ are not identically zero, then it is said to be a linear recurrence relation *degree* k . If $c_0(n), c_1(n), \dots, c_k(n)$ are constants, then the recurrence relation is known as a linear relation with constant coefficients. If $f(n)$ is identically zero, then the recurrence relation is said to be homogeneous; otherwise, it is inhomogeneous.

Thus, all the examples above are linear recurrence relations except (8), (9), and (10); the relation (8), for instance, is not linear because of the squared term.

The relations in (3), (4), (5), and (7) are linear with constant coefficients.

Relations (1), (2), and (3) have degree 1; (4), (5), and (6) have degree 2; (7) has degree

3. Relations (3), (4), and (6) are homogeneous.

There are no general techniques that will enable one to solve all recurrence relations. There are, nevertheless, techniques that will enable us to solve linear recurrence relations with constant coefficients.

SOLVING RECURRENCE RELATIONS BY SUSTITUTION AND GENERATING FUNCTIONS

We shall consider four methods of solving recurrence relations in this and the next two sections:

1. Substitution (also called iteration),
2. Generating functions,
3. Characteristics roots,

In the substitution method the recurrence relation is used repeatedly to solve for a general expression for a_n in terms of n . We desire that this expression involve no other terms of the sequence except those given by boundary conditions.

The mechanics of this method are best described in terms of examples. We used this method in Example 5.3.4. Let us also illustrate the method in the following examples.

Example

Solve the recurrence relation $a_n = a_{n-1} + f(n)$ for $n \geq 1$ by substitution

$$a_1 = a_0 + f(1)$$

$$a_2 = a_1 + f(2) = a_0 + f(1) + f(2)$$

$$a_3 = a_2 + f(3) = a_0 + f(1) + f(2) + f(3)$$

·
·
·

$$\begin{aligned}
 a_n &= a_0 + f(1) + f(2) + \dots + f(n) \\
 &= a_0 + \sum_{k=1}^n f(k)
 \end{aligned}$$

Thus, a_n is just the sum of the $f(k)$'s plus a_0 .

More generally, if c is a constant then we can solve $a_n = c a_{n-1} + f(n)$ for $n \geq 1$ in the same way:

$$a_1 = c a_0 + f(1)$$

$$\begin{aligned}
 a_2 &= c a_1 + f(2) = c(c a_0 + f(1)) + f(2) \\
 &= c^2 a_0 + c f(1) + f(2)
 \end{aligned}$$

$$\begin{aligned}
 a_3 &= c a_2 + f(3) = c(c^2 a_0 + c f(1) + f(2)) + f(3) \\
 &= c^3 a_0 + c^2 f(1) + c f(2) + f(3)
 \end{aligned}$$

.

.

.

$$\begin{aligned}
 a_n &= c a_{n-1} + f(n) = c(c^{n-1} a_0 + c^{n-2} f(1) + \dots + c^{n-2} f(n-1)) + f(n) \\
 &= c^n a_0 + c^{n-1} f(1) + c^{n-2} f(2) + \dots + c f(n-1) + f(n)
 \end{aligned}$$

Or

$$a_n = c^n a_0 + \sum_{k=1}^n c^{n-k} f(k)$$

Solution of Linear Inhomogeneous Recurrence Relations:

The equation $a_n + p_1 a_{n-1} + p_2 a_{n-2} = f(n)$, where p_1 and p_2 are constant, and $f(n)$ is not identically 0, is called a second-order linear inhomogeneous recurrence relation (or difference equation) with constant coefficients. The homogeneous case, which we've looked at already, occurs when

$f(n) \equiv 0$. The inhomogeneous case occurs more frequently. The homogeneous case is so important largely because it gives us the key to solving the inhomogeneous equation. If you've studied linear differential equations with constant coefficients, you'll see the parallel. We will call the

difference obtained by setting the right-hand side equal to 0, the associated homogeneous equation. We know how to solve this. Say that V is a solution. Now suppose that $g(x)$ is any particular solution of the inhomogeneous equation. (That is, it solves the equation, but does not necessarily match the initial data.) Then $U = V + g(x)$ is a solution to the inhomogeneous equation, which you can see simply by substituting U into the equation. On the other hand, every solution U of the inhomogeneous equation is of the form $U = V + g(x)$ where V is a solution of the homogeneous equation, and $g(x)$ is a particular solution of the inhomogeneous equation. The proof of this is straightforward. If we have two solutions to the inhomogeneous equation, say U_1 and U_2 , then their difference $U_1 - U_2 = V$ is a solution to the homogeneous equation, which you can check by substitution. But then $U_1 = V + U_2$, and we can set $U_2 = g(x)$, since by assumption, U_2 is a particular solution. This leads to the following theorem: the general solution to the inhomogeneous equation is the general solution to the associated homogeneous equation, plus any particular solution to the inhomogeneous equation. This gives the following procedure for solving the inhomogeneous equation:

4. Solve the associated homogeneous equation by the method we've learned. This will involve variable (or undetermined) coefficients.
5. Guess a particular solution to the inhomogeneous equation. It is because of the guess that I've called this a procedure, not an algorithm. For simple right-hand sides, we can say how to compute a particular solution, and in these cases, the procedure merits the name algorithm.
6. The general solution to the inhomogeneous equation is the sum of the answers from the two steps above.
7. Use the initial data to solve for the undetermined coefficients from step 1.

To solve the equation $y'' - 6y' + 8y = 3$. Let's suppose that we are also given the initial data $y(0) = 3$, $y'(0) = 3$. The associated homogeneous equation is $y'' - 6y' + 8y = 0$, so the characteristic equation is $\lambda^2 - 6\lambda + 8 = 0$, which has roots $\lambda = 2$ and $\lambda = 4$. Thus, the general solution to the associated homogeneous equation is $y_h = 12e^{2x} + 24e^{4x}$. When the right-hand side is a polynomial, as in this case, there will always be a particular solution that is a polynomial.

Usually, a polynomial of the same degree will work, so we'll guess in this case that there is a constant C that solves the homogeneous equation. If that is so, then $y = C$, $y' = 0$, $y'' = 0$, and substituting into the equation gives $C - 6C + 8C = 3$, and we find that $C = 1$. Now, the general solution to the inhomogeneous equations is $y = 12e^{2x} + 24e^{4x} + 1$. Reassuringly, this is the answer given in the back of the book. Our initial data lead to the equations $1 + 12 + 24 = 3$ and $2 + 4 + 1 = 3$, whose solution is $1 = 3$, $2 = -1$. Finally, the solution to the inhomogeneous equation, with the initial condition given, is $y = 3 \cdot 12e^{2x} - 1 \cdot 24e^{4x} + 1$. Sometimes, a polynomial of the same degree as the right-hand side doesn't work. This happens when the characteristic equation has 1 as a root. If our equation had been $y'' - 6y' + 5y = 3$, when we guessed that the particular solution was a constant, we'd have arrived at the equation $C - 6C + 5C = 3$, or $0 = 3$. The way to deal with this is to increase the degree of the polynomial. Instead of assuming that the solution is constant, we'll assume that it's linear. In fact, we'll guess that it is of the form

$g_n = nC$. Then we have $nC - 6C - 1C + 5C - 2C = 3$, which simplifies to $6C - 10C = 3$ so that $C = -3/4$. Thus, $g_n = -3n/4$. This won't be enough if 1 is a root of multiplicity 2, that is, if $-1/2$ is a factor of the characteristic polynomial. Then there is a particular solution of the form $g_n = Cn^2$. For second-order equations, you never have to go past this. If the right-hand side is a polynomial of degree greater than 0, then the process works just the same, except that you start with a polynomial of the same degree, increase the degree by 1, if necessary, and then once more, if need be. For example, if the right-hand side were $f_n = 2n - 1$, we would start by guessing a particular solution $g_n = C_1n + C_2$. If it turned out that 1 was a characteristic root, we would amend our guess to $g_n = C_1n^2 + C_2n + C_3$. If 1 is a double root, this will fail also, but $g_n = C_1n^3 + C_2n^2 + C_3n + C_4$ will work in this case.

Another case where there is a simple way of guessing a particular solution is when the right-hand side is an exponential, say $f_n = r^n$. In that case, we guess that a particular solution is just a constant multiple of r^n , say $g_n = Cn$. Again, we gave trouble when 1 is a characteristic root. We then guess that $g_n = knCn$, which will fail only if 1 is a double root. In that case we must use $g_n = kn^2Cn$, which is as far as we ever have to go in the second-order case. These same ideas extend to higher-order recurrence relations, but we usually solve them numerically, rather than exactly. A third-order linear difference equation with constant coefficients leads to a cubic characteristic polynomial. There is a formula for the roots of a cubic, but it's very complicated.

For fourth-degree polynomials, there's also a formula, but it's even worse. For fifth and higher degrees, no such formula exists. Even for the third-order case, the exact solution of a simple-looking inhomogeneous linear recurrence relation with constant coefficients can take pages to write down. The coefficients will be complicated expressions involving square roots and cube roots. For most, if not all, purposes, a simpler answer with numerical coefficients is better, even though they must in the nature of things, be approximate.

The procedure I've suggested may strike you as silly. After all, we've already solved the characteristic equation, so we know whether 1 is a characteristic root, and what its multiplicity is. Why not start with a polynomial of the correct degree? This is all well and good, while you're taking the course, and remember the procedure in detail. However, if you have to use this procedure some years from now, you probably won't remember all the details. Then the method I've suggested will be valuable. Alternatively, you can start with a general polynomial of the maximum possible degree. This leads to a lot of extra work if you're solving by hand, but it's the approach I prefer for computer solution.

UNIT V

Graph Theory

Representation of Graphs:

There are two different sequential representations of a graph. They are

- Adjacency Matrix representation

- Path Matrix representation

Adjacency Matrix Representation

Suppose G is a simple directed graph with m nodes, and suppose the nodes of G have been ordered and are called $v_1, v_2, \dots, v_m$. Then the adjacency matrix $A = (a_{ij})$ of the graph G is the $m \times m$ matrix defined as follows:

$a_{ij} = 1$ if v_i is adjacent to v_j , that is, if there is an edge (v_i, v_j)
 $a_{ij} = 0$ otherwise

Suppose G is an undirected graph. Then the adjacency matrix A of G will be a symmetric matrix, i.e., one in which $a_{ij} = a_{ji}$; for every i and j .

Drawbacks

12. It may be difficult to insert and delete nodes in G .

13. If the number of edges is $O(m)$ or $O(m \log^2 m)$, then the matrix A will be sparse, hence a great deal of space will be wasted.

Path Matrix Representation

Let G be a simple directed graph with m nodes, $v_1, v_2, \dots, v_m$. The path matrix of G is the m -square matrix $P = (p_{ij})$ defined as follows:

$p_{ij} = 1$ if there is a path from v_i to v_j
 $p_{ij} = 0$ otherwise

Graphs and Multigraphs

A graph G consists of two things:

1. A set V of elements called nodes (or points or vertices)

2. A set E of edges such that each edge e in E is identified with a unique

Sometimes we indicate the parts of a graph by writing $G = (V, E)$.

Suppose $e = [u, v]$. Then the nodes u and v are called the endpoints of e , and u and v are said to be adjacent nodes or neighbors. The degree of a node u , written $\deg(u)$, is the number of edges containing u . If $\deg(u) = 0$ — that is, if u does not belong to any edge— then u is called an isolated node.

Path and Cycle

A path P of length n from a node u to a node v is defined as a sequence of $n + 1$ nodes. $P = (v_0, v_1, v_2, \dots, v_n)$ such that $u = v_0$; v_{i-1} is adjacent to v_i for $i = 1, 2, \dots, n$ and $v_n = v$.

Types of Path

1. Simple Path
2. Cycle Path

(i) Simple Path

Simple path is a path in which first and last vertex are different ($V_0 \neq V_n$)

(ii) Cycle Path

Cycle path is a path in which first and last vertex are same ($V_0 = V_n$). It is also called as Closed path.

Connected Graph

A graph G is said to be connected if there is a path between any two of its nodes.

Complete Graph

A graph G is said to be complete if every node u in G is adjacent to every other node v in G .

Tree

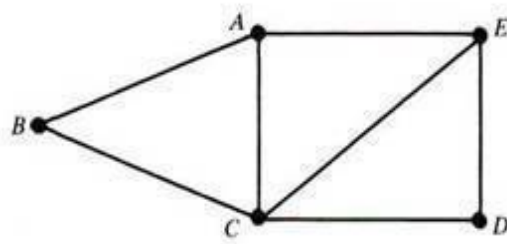
A connected graph T without any cycles is called a tree graph or free tree or, simply, a tree.

Labeled or Weighted Graph

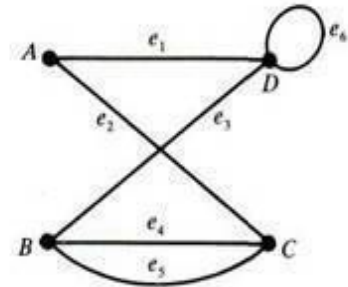
If the weight is assigned to each edge of the graph then it is called as Weighted or Labeled graph.

The definition of a graph may be generalized by permitting the following:

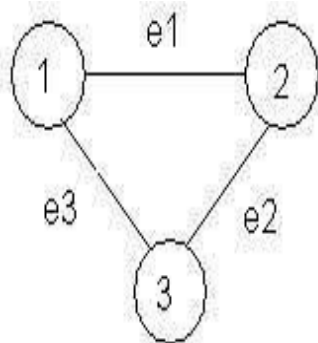
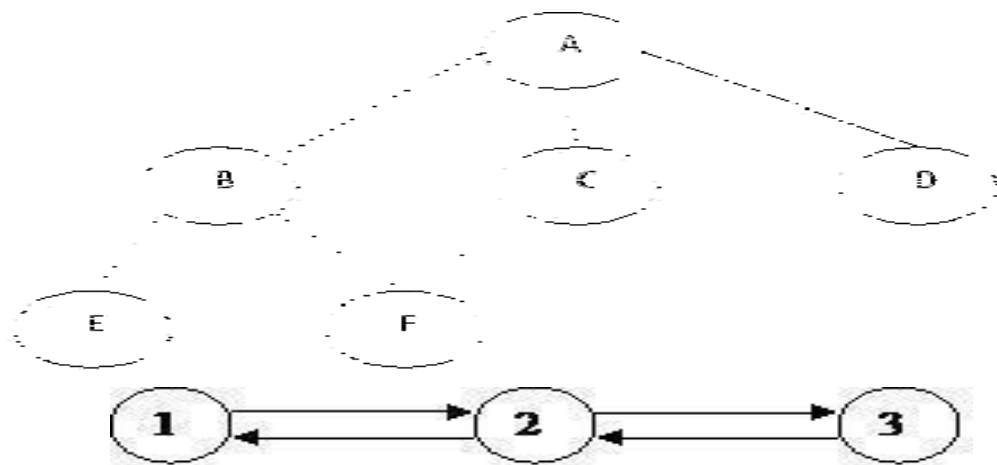
- **Multiple edges:** Distinct edges e and e' are called multiple edges if they connect the same endpoints, that is, if $e = [u, v]$ and $e' = [u, v]$.
- **Loops:** An edge e is called a loop if it has identical endpoints, that is, if $e = [u, u]$.
- **Finite Graph:** A multigraph M is said to be finite if it has a finite number of nodes and a finite number of edges.



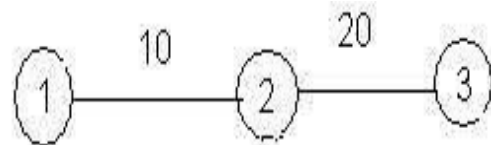
(a) Graph.



(b) Multigraph.



(a)



(b)

weighted or Labeled Graph

Directed Graphs

A directed graph G , also called a digraph or graph is the same as a multigraph except that each edge e in G is assigned a direction, or in other words, each edge e is identified with an ordered pair (u, v) of nodes in G .

Outdegree and Indegree

Indegree : The indegree of a vertex is the number of edges for which v is head

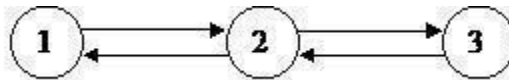
Example:

Indegree of 1 = 1

Indegree of 2 = 2

Outdegree : The outdegree of a node or vertex is the number of edges for which v is tail.

Example



Outdegree of 1 = 1

Outdegree of 2 = 2

Simple Directed Graph

A directed graph G is said to be simple if G has no parallel edges. A simple graph G may have loops, but it cannot have more than one loop at a given node.

Graph Traversal

The breadth first search (BFS) and the depth first search (DFS) are the two algorithms used for traversing and searching a node in a graph. They can also be used to find out whether a node is reachable from a given node or not.

Depth First Search (DFS)

The aim of DFS algorithm is to traverse the graph in such a way that it tries to go far from the root node. Stack is used in the implementation of the depth first search. Let's see how depth first search works with respect to the following graph:

As stated before, in DFS, nodes are visited by going through the depth of the tree from the starting node. If we do the depth first traversal of the above graph and print the visited node, it will be -A B E F C D< . DFS visits the root node and then its children nodes until it reaches the end node, i.e. E and F nodes, then moves up to the parent nodes.

Algorithmic Steps

6. **Step 1:** Push the root node in the Stack.
7. **Step 2:** Loop until stack is empty.
8. **Step 3:** Peek the node of the stack.
9. **Step 4:** If the node has unvisited child nodes, get the unvisited child node, mark it as traversed and push it on stack.
10. **Step 5:** If the node does not have any unvisited child nodes, pop the node from the stack.

Based upon the above steps, the following Java code shows the implementation of the DFS algorithm:

```
public void dfs()
{
    //DFS uses Stack data structure

    Stack s=new Stack();
    s.push(this.rootNode);
    rootNode.visited=true;
    printNode(rootNode);
    while(!s.isEmpty())
    {
        Node n=(Node)s.peek();

        Node child=getUnvisitedChildNode(n);
        if(child!=null)
        {
            child.visited=true; printNode(child);
            s.push(child);

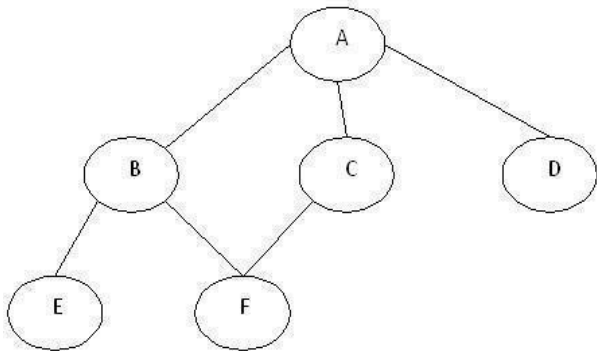
        }
        else
        {
            s.pop();

        }
    }
}
```

```
//Clear visited property of
nodes clearNodes();
}
```

Breadth First Search (BFS)

This is a very different approach for traversing the graph nodes. The aim of BFS algorithm is to traverse the graph as close as possible to the root node. Queue is used in the implementation of the breadth first search. Let's see how BFS traversal works with respect to the following graph:



If we do the breadth first traversal of the above graph and print the visited node as the output, it will print the following output. -A B C D E F. The BFS visits the nodes level by level, so it will start with level 0 which is the root node, and then it moves to the next levels which are B, C and D, then the last levels which are E and F.

Algorithmic Steps

1. **Step 1:** Push the root node in the Queue.
2. **Step 2:** Loop until the queue is empty.
3. **Step 3:** Remove the node from the Queue.
4. **Step 4:** If the removed node has unvisited child nodes, mark them as visited and insert the unvisited children in the queue.

Based upon the above steps, the following Java code shows the implementation of the BFS algorithm:

```
public void bfs()
{
    //BFS uses Queue data structure

    Queue q=new LinkedList();
    q.add(this.rootNode);
    printNode(this.rootNode);
    rootNode.visited=true;
    while(!q.isEmpty())
    {
        Node n=(Node)q.remove();
        Node child=null;
        while((child=getUnvisitedChildNode(n))!=null)
```

```

        { child.visited=true; printNode(child); q.add(child);

            }
    }

    //Clear visited property of
    nodes clearNodes();
}

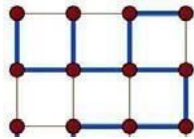
```

Spanning Trees:

In the mathematical field of graph theory, a spanning tree T of a connected, undirected graph G is a tree composed of all the vertices and some (or perhaps all) of the edges of G . Informally, a spanning tree of G is a selection of edges of G that form a tree *spanning* every vertex. That is, every vertex lies in the tree, but no cycles (or loops) are formed. On the other hand, every bridge of G must belong to T .

A spanning tree of a connected graph G can also be defined as a maximal set of edges of G that contains no cycle, or as a minimal set of edges that connect all vertices.

Example:



A spanning tree (blue heavy edges) of a grid graph.

Spanning forests

A spanning forest is a type of subgraph that generalises the concept of a spanning tree. However, there are two definitions in common use. One is that a spanning forest is a subgraph that consists of a spanning tree in each connected component of a graph. (Equivalently, it is a maximal cycle-free subgraph.) This definition is common in computer science and optimisation. It is also the definition used when discussing minimum spanning forests, the generalization to disconnected graphs of minimum spanning trees. Another definition, common in graph theory, is that a spanning forest is any subgraph that is both a forest (contains no cycles) and spanning (includes every vertex).

Counting spanning trees

The number $t(G)$ of spanning trees of a connected graph is an important invariant. In some cases, it is easy to calculate $t(G)$ directly. It is also widely used in data structures in different computer languages. For example, if G is itself a tree, then $t(G)=1$, while if G is the cycle graph C_n with n vertices, then $t(G)=n$. For any graph G , the number $t(G)$ can be calculated using Kirchhoff's matrix-tree theorem (follow the link for an explicit example using the theorem).

Cayley's formula is a formula for the number of spanning trees in the complete graph K_n with n vertices. The formula states that $t(K_n) = n^{n-2}$.

Another way of stating Cayley's formula is that Cayley's formula is a formula for the number of spanning trees in the complete graph K_n with n vertices. The formula states that $t(K_n) = n^{n-2}$.

Another way of stating Cayley's formula is that there are exactly n^{n-2} labelled trees with n vertices. Cayley's formula can be proved using Kirchhoff's matrix-tree theorem or via the Prüfer code.

If G is the complete bipartite graph $K_{p,q}$, then $t(G) = p^{q-1} q^{p-1}$, while if G is the n -dimensional hypercube graph Q_n , $t(G) = 2^{2^n - n - 1} \prod_{k=2}^n k^{\binom{n}{k}}$. These formulae are also consequences of the matrix-tree theorem.

If G is a multigraph and e is an edge of G , then the number $t(G)$ of spanning trees of G satisfies the *deletion-contraction recurrence* $t(G) = t(G-e) + t(G/e)$, where $G-e$ is the multigraph obtained by deleting e and G/e is the contraction of G by e , where multiple edges arising from this contraction are not deleted.

Uniform spanning trees

A spanning tree chosen randomly from among all the spanning trees with equal probability is called a uniform spanning tree (UST). This model has been extensively researched in probability and mathematical physics.

Algorithms

The classic spanning tree algorithm, depth-first search (DFS), is due to Robert Tarjan. Another important algorithm is based on breadth-first search (BFS).

Planar Graphs:

In graph theory, a planar graph is a graph that can be embedded in the plane, i.e., it can be drawn on the plane in such a way that its edges intersect only at their endpoints.

A planar graph already drawn in the plane without edge intersections is called a plane graph or planar embedding of the graph. A plane graph can be defined as a planar graph with a mapping from every node to a point in 2D space, and from every edge to a plane curve, such that the extreme points of each curve are the points mapped from its end nodes, and all curves are disjoint except on their extreme points. Plane graphs can be encoded by combinatorial maps.

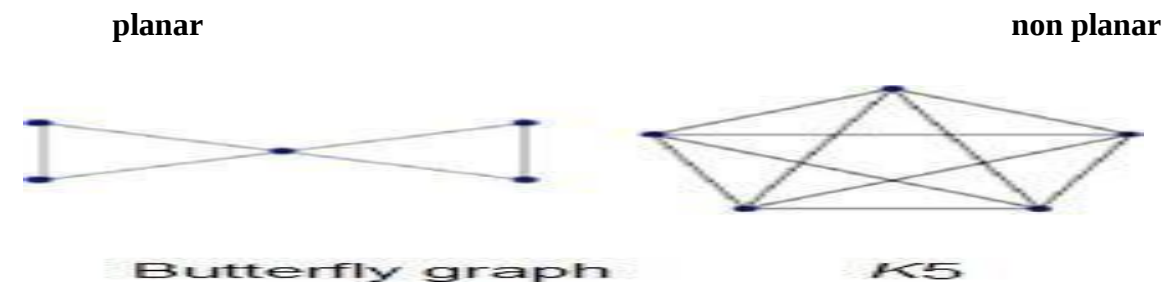
It is easily seen that a graph that can be drawn on the plane can be drawn on the sphere as well, and vice versa.

The equivalence class of topologically equivalent drawings on the sphere is called a planar map. Although a plane graph has an external or unbounded face, none of the faces of a planar map have a particular status.

Applications

- Telecommunications – e.g. spanning trees
- Vehicle routing – e.g. planning routes on roads without underpasses
- VLSI – e.g. laying out circuits on computer chip.
- The puzzle game Planarity requires the player to "untangle" a planar graph so that none of its edges intersect.

Example graphs



Graph Theory and Applications:

Graphs are among the most ubiquitous models of both natural and human-made structures. They can be used to model many types of relations and process dynamics in physical, biological and social systems. Many problems of practical interest can be represented by graphs.

In computer science, graphs are used to represent networks of communication, data organization, computational devices, the flow of computation, etc. One practical example: The link structure of a website could be represented by a directed graph. The vertices are the web pages available at the website and a directed edge from page *A* to page *B* exists if and only if *A* contains a link to *B*. A similar approach can be taken to problems in travel, biology, computer chip design, and many other fields. The development of algorithms to handle graphs is therefore of major interest in computer science. There, the transformation of graphs is often formalized and represented by graph rewrite systems. They are either directly used or properties of the rewrite systems (e.g. confluence) are studied. Complementary to graph transformation systems focussing on rule- based in-memory manipulation of graphs are graph databases geared towards transaction-safe, persistent storing and querying of graph-structured data.

Graph-theoretic methods, in various forms, have proven particularly useful in linguistics, since natural language often lends itself well to discrete structure. Traditionally, syntax and compositional semantics follow tree-based structures, whose expressive power lies in the Principle of Compositionality, modeled in a hierarchical graph. Within lexical semantics, especially as applied to computers, modeling word meaning is easier when a given word is understood in terms of related words; semantic networks are therefore important in computational linguistics. Still other methods in phonology (e.g. Optimality Theory, which uses lattice graphs) and morphology (e.g. finite-state morphology, using finite-state transducers) are common in the analysis of language as a graph. Indeed, the usefulness of this area of mathematics to linguistics has borne organizations such as TextGraphs, as well as various 'Net' projects, such as WordNet, VerbNet, and others.

Graph theory is also used to study molecules in chemistry and physics. In condensed matter physics, the three dimensional structure of complicated simulated atomic structures can be studied quantitatively by gathering statistics on graph-theoretic properties related to the topology of the atoms. For example, Franzblau's shortest-path (SP) rings. In chemistry a graph makes a natural model for a molecule, where vertices represent atoms and edges bonds. This approach is especially used in computer processing of molecular structures, ranging from chemical editors to database searching. In statistical physics, graphs can represent local connections between interacting parts of a system, as well as the dynamics of a physical process on such systems.

Graph theory is also widely used in sociology as a way, for example, to measure actors' prestige or to explore diffusion mechanisms, notably through the use of social network analysis software. Likewise, graph theory is useful in biology and conservation efforts where a vertex can represent regions where certain species exist (or habitats) and the edges represent migration paths, or movement between the regions. This information is important when looking at breeding patterns or tracking the spread of disease, parasites or how changes to the movement can affect other species.

In mathematics, graphs are useful in geometry and certain parts of topology, e.g. Knot Theory. Algebraic graph theory has close links with group theory.

A graph structure can be extended by assigning a weight to each edge of the graph. Graphs with weights, or weighted graphs, are used to represent structures in which pairwise connections have some numerical values. For example if a graph represents a road network, the weights could represent the length of each road.

Basic Concepts Isomorphism:

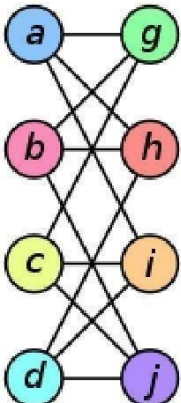
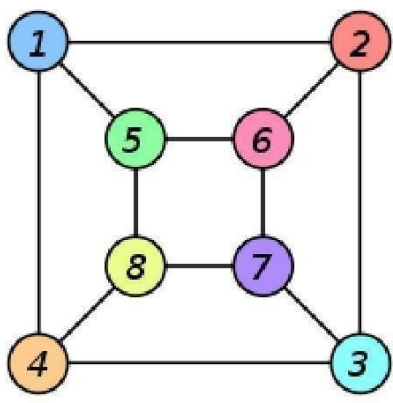
Let G_1 and G_2 be two graphs and let f be a function from the vertex set of G_1 to the vertex set of G_2 . Suppose that f is one-to-one and onto & $f(v)$ is adjacent to $f(w)$ in G_2 if and only if v is adjacent to w in G_1 .

Then we say that the function f is an isomorphism and that the two graphs G_1 and G_2 are isomorphic. So two graphs G_1 and G_2 are isomorphic if there is a one-to-one correspondence between vertices of G_1 and those of G_2 with the property that if two vertices of G_1 are adjacent then so are their images in G_2 . If two graphs are isomorphic then as far as we are concerned they are the same graph though the location of the vertices may be different. To show you how the program can be used to explore isomorphism draw the graph in figure 4 with the program (first get the null graph on four vertices and then use the right mouse to add edges).

Save this graph as Graph 1 (you need to click Graph then Save). Now get the circuit graph with 4 vertices. It looks like figure 5, and we shall call it C(4).

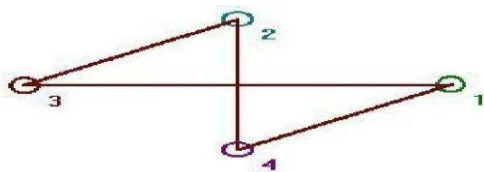
Example:

The two graphs shown below are isomorphic, despite their different looking drawings.

Graph G	Graph H	An isomorphism between G and H
		$f(a) = 1$ $f(b) = 6$ $f(c) = 8$ $f(d) = 3$ $f(g) = 5$ $f(h) = 2$ $f(i) = 4$ $f(j) = 7$

Subgraphs:

A subgraph of a graph G is a graph whose vertex set is a subset of that of G , and whose adjacency relation is a subset of that of G restricted to this subset. In the other direction, a supergraph of a graph G is a graph of which G is a subgraph. We say a graph G contains another graph H if some subgraph of G is H or is isomorphic to H .

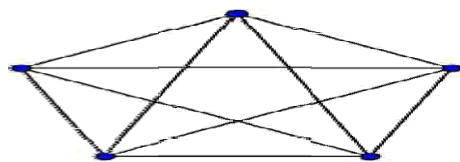


A subgraph H is a spanning subgraph, or factor, of a graph G if it has the same vertex set as G . We say H spans G .

A subgraph H of a graph G is said to be induced if, for any pair of vertices x and y of H , xy is an edge of H if and only if xy is an edge of G . In other words, H is an induced subgraph of G if it has all the edges that appear in G over the same vertex set. If the vertex set of H is the subset S of $V(G)$, then H can be written as $G[S]$ and is said to be induced by S .

A graph that does *not* contain H as an induced subgraph is said to be H -free.

A universal graph in a class K of graphs is a simple graph in which every element in K can be embedded as a subgraph.



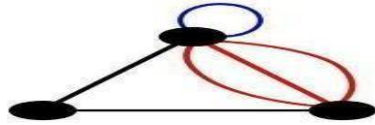
K_5 , a complete graph. If a subgraph looks like this, the vertices in that subgraph form a clique of size 5.

Multi graphs:

In mathematics, a multigraph or pseudograph is a graph which is permitted to have multiple edges, (also called "parallel edges"), that is, edges that have the same end nodes. Thus two vertices may be connected by more than one edge. Formally, a multigraph G is an ordered pair $G := (V, E)$ with

- V a set of *vertices* or *nodes*,
- E a multiset of unordered pairs of vertices, called *edges* or *lines*.

Multigraphs might be used to model the possible flight connections offered by an airline. In this case the multigraph would be a directed graph with pairs of directed parallel edges connecting cities to show that it is possible to fly both *to* and *from* these locations.



A multigraph with multiple edges (red) and a loop (blue). Not all authors allow multigraphs to have loops.

Euler circuits:

In graph theory, an Eulerian trail is a trail in a graph which visits every edge exactly once. Similarly, an Eulerian circuit is an Eulerian trail which starts and ends on the same vertex. They were first discussed by Leonhard Euler while solving the famous Seven Bridges of Königsberg problem in 1736. Mathematically the problem can be stated like this:

Given the graph on the right, is it possible to construct a path (or a cycle, i.e. a path starting and ending on the same vertex) which visits each edge exactly once?

Euler proved that a necessary condition for the existence of Eulerian circuits is that all vertices in the graph have an even degree, and stated without proof that connected graphs with all vertices of even degree have an Eulerian circuit. The first complete proof of this latter claim was published in 1873 by Carl Hierholzer.

The term Eulerian graph has two common meanings in graph theory. One meaning is a graph with an Eulerian circuit, and the other is a graph with every vertex of even degree. These definitions coincide for connected graphs.

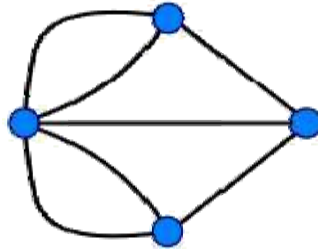
For the existence of Eulerian trails it is necessary that no more than two vertices have an odd degree; this means the Königsberg graph is *not* Eulerian. If there are no vertices of odd degree, all Eulerian trails are circuits. If there are exactly two vertices of odd degree, all Eulerian trails start at one of them and end at the other. Sometimes a graph that has an Eulerian trail but not an Eulerian circuit is called semi-Eulerian.

An Eulerian trail, Eulerian trail or Euler walk in an undirected graph is a path that uses each edge exactly once. If such a path exists, the graph is called traversable or semi-eulerian.

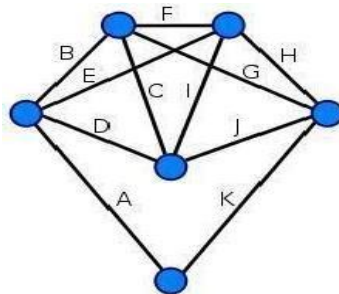
An Eulerian cycle, Eulerian circuit or Euler tour in an undirected graph is a cycle that uses each edge exactly once. If such a cycle exists, the graph is called unicursal. While such graphs are Eulerian graphs, not every Eulerian graph possesses an Eulerian cycle.

For directed graphs path has to be replaced with directed path and cycle with directed cycle.

The definition and properties of Eulerian trails, cycles and graphs are valid for multigraphs as well.



This graph is not Eulerian, therefore, a solution does not exist.



Every vertex of this graph has an even degree, therefore this is an Eulerian graph. Following the edges in alphabetical order gives an Eulerian circuit/cycle.

Hamiltonian graphs:

In the mathematical field of graph theory, a Hamiltonian path (or traceable path) is a path in an undirected graph which visits each vertex exactly once. A Hamiltonian cycle (or Hamiltonian circuit) is a cycle in an undirected graph which visits each vertex exactly once and also returns to the starting vertex. Determining whether such paths and cycles exist in graphs is the Hamiltonian path problem which is NP-complete.

Hamiltonian paths and cycles are named after William Rowan Hamilton who invented the Icosian game, now also known as *Hamilton's puzzle*, which involves finding a Hamiltonian cycle in the edge graph of the dodecahedron. Hamilton solved this problem using the Icosian Calculus, an algebraic structure based on roots of unity with many similarities to the quaternions (also invented by Hamilton). This solution does not generalize to arbitrary graphs.

A *Hamiltonian path* or *traceable path* is a path that visits each vertex exactly once. A graph that contains a Hamiltonian path is called a **traceable graph**. A graph is **Hamilton-connected** if for every pair of vertices there is a Hamiltonian path between the two vertices.

A *Hamiltonian cycle*, *Hamiltonian circuit*, *vertex tour* or *graph cycle* is a cycle that visits each vertex exactly once (except the vertex which is both the start and end, and so is visited twice). A graph that contains a Hamiltonian cycle is called a **Hamiltonian graph**.

Similar notions may be defined for *directed graphs*, where each edge (arc) of a path or cycle can only be traced in a single direction (i.e., the vertices are connected with arrows and the edges traced "tail-to-head").

A **Hamiltonian decomposition** is an edge decomposition of a graph into Hamiltonian circuits.

Examples

- a complete graph with more than two vertices is Hamiltonian
- every cycle graph is Hamiltonian
- every tournament has an odd number of Hamiltonian paths
- every platonic solid, considered as a graph, is Hamiltonian

Chromatic Numbers:

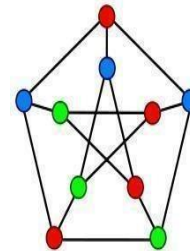
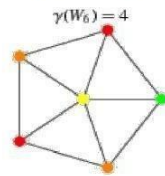
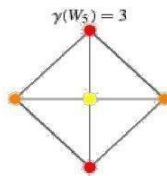
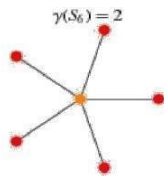
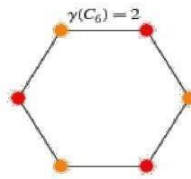
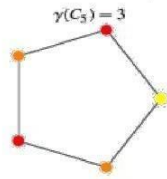
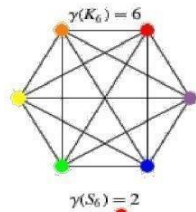
In graph theory, graph coloring is a special case of graph labeling; it is an assignment of labels traditionally called "colors" to elements of a graph subject to certain constraints. In its simplest form, it is a way of coloring the vertices of a graph such that no two adjacent vertices share the same color; this is called a vertex coloring. Similarly, an edge coloring assigns a color to each edge so that no two adjacent edges share the same color, and a face coloring of a planar graph assigns a color to each face or region so that no two faces that share a boundary have the same color.

Vertex coloring is the starting point of the subject, and other coloring problems can be transformed into a vertex version. For example, an edge coloring of a graph is just a vertex coloring of its line graph, and a face coloring of a planar graph is just a vertex coloring of its planar dual. However, non-vertex coloring problems are often stated and studied *as is*. That is partly for perspective, and partly because some problems are best studied in non-vertex form, as for instance is edge coloring.

The convention of using colors originates from coloring the countries of a map, where each face is literally colored. This was generalized to coloring the faces of a graph embedded in the plane. By planar duality it became coloring the vertices, and in this form it generalizes to all graphs. In mathematical and computer representations it is typical to use the first few positive or nonnegative integers as the "colors". In general one can use any finite set as the "color set". The nature of the coloring problem depends on the number of colors but not on what they are.

Graph coloring enjoys many practical applications as well as theoretical challenges. Beside the classical types of problems, different limitations can also be set on the graph, or on the way a color is assigned, or even on the color itself. It has even reached popularity with the general public in the form of the popular number puzzle Sudoku. Graph coloring is still a very active field of research.

	$\gamma(G)$
	n
	$\begin{cases} 3 & \text{for } n \text{ odd} \\ 2 & \text{for } n \text{ even} \end{cases}$
$S_n \quad n > 1$	
$W_n \quad n > 2$	$\begin{cases} 3 & \text{for } n \text{ odd} \\ 4 & \text{for } n \text{ even} \end{cases}$



A proper vertex coloring of the Petersen graph with 3 colors, the minimum number possible.

Vertex coloring

When used without any qualification, a coloring of a graph is almost always a *proper vertex coloring*, namely a labelling of the graph's vertices with colors such that no two vertices sharing the same edge have the same color. Since a vertex with a loop could never be properly colored, it is understood that graphs in this context are loopless.

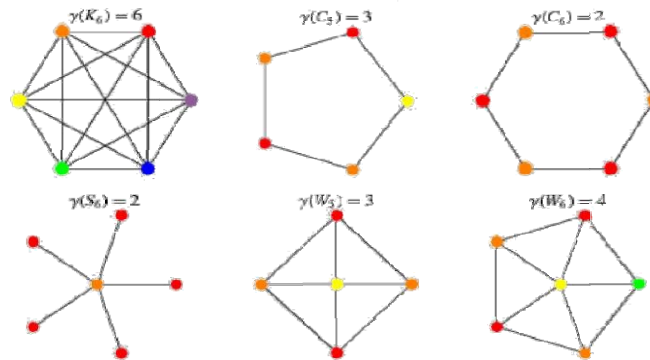
The terminology of using *colors* for vertex labels goes back to map coloring. Labels like *red* and *blue* are only used when the number of colors is small, and normally it is understood that the labels are drawn from the integers $\{1, 2, 3, \dots\}$.

A coloring using at most k colors is called a (proper) k -coloring. The smallest number of colors needed to color a graph G is called its chromatic number, $\chi(G)$. A graph that can be assigned a (proper) k -coloring is k -colorable, and it is k -chromatic if its chromatic number is exactly k . A subset of vertices assigned to the same color is called a *color class*, every such class forms an independent set. Thus, a k -coloring is the same as a partition of the vertex set into k independent sets, and the terms *k-partite* and *k-colorable* have the same meaning.

This graph can be 3-colored in 12 different ways.

The following table gives the chromatic number for familiar classes of graphs.

	$\gamma(G)$
n	
	$\begin{cases} 3 & \text{for } n \text{ odd} \\ 2 & \text{for } n \text{ even} \end{cases}$
$S_n, n \geq 1$	
$W_n, n \geq 2$	$\begin{cases} 3 & \text{for } n \text{ odd} \\ 4 & \text{for } n \text{ even} \end{cases}$



graph G complete graph K_n cycle graph $C_n, n \geq 1$

star graph $S_n, n \geq 1$

wheel graph $W_n, n \geq 3$

wheel graph $W_n, n \geq 3$

wheel graph $W_n, n \geq 3$

