

Hall Ticket Number:

|  |  |  |  |  |  |  |  |  |  |
|--|--|--|--|--|--|--|--|--|--|
|  |  |  |  |  |  |  |  |  |  |
|--|--|--|--|--|--|--|--|--|--|

## IV/IV B. Tech (Regular/Supplementary) DEGREE EXAMINATION

March, 2019

Electronics &amp; Instrumentation Engineering

Eighth Semester

VLSI Design

Time: Three Hours

Maximum : 60 Marks

Answer Question No.1 compulsorily.

(1X12 = 12 Marks)

Answer ONE question from each unit.

(4X12=48 Marks)

(1X12=12 Marks)

1 Answer all questions

- What is photolithography in VLSI design.
- Illustrate latch up condition in CMOS circuits? How to prevent it?
- Implement OR gate using pass transistor design.
- What is propagation delay.
- Define sheet resistance.
- What are stick diagrams.
- Explain switch logic.
- List some architectural issues in VLSI design
- Distinguish between PAL & PLA.
- Distinguish function and procedures.
- What is clocked CMOS logic.
- Give the general structure of a VHDL program.

## UNIT I

- Explain in detail about NMOS fabrication technology. 6M
  - Describe the CMOS inverter & derive its dc characteristics. 6M

(OR)

- Explain in detail about the  $I_{ds}$  and  $V_{ds}$  relationship 6M
  - Explain the alternative forms of pull-ups for an inverter 6M

## UNIT II

- Describe three sources of wiring capacitances. Explain the effect of wiring capacitances on the performance of a VLSI circuit. 6M
  - Draw a layout of 2 input CMOS AND gate indicating all the regions and layers. 6M

(OR)

- Explain about various design rules. 6M
  - Explain how scaling factors influence the design of a device. 6M

## UNIT III

- Discuss architectural issues in design process. 6M
  - Explain about design of an ALU subsystem. 6M

(OR)

- Compare switch logic and gate logic. 6M
  - Give the subsystem design considerations of a four-bit multiplier. 6M

## UNIT IV

- Design a full adder using PLA. 6M
  - Explain about ASIC'S 6M

(OR)

- Explain about structural design elements in VHDL. 6M
  - Write a program to implement a 4-bit adder in data flow model. 6M

Scheme of Evaluation

(2)

1.  $12 \times 1 = 12M$

2 (a) Six steps -  $6 \times 1 = 6M$

(b) circuit diagram -  $2M$

characteristic -  $2M$

Derivation -  $2M$

3. (a) Exp. for  $T_{sd}$  -  $2M$

$I_{ds}$  vs  $V_{ds}$  for non saturated region -  $3M$

$I_{ds}$  vs  $V_{ds}$  for saturated region -  $1M$

(b) Any two forms of pull up -  $2 \times 3M$

4 (a) Fringing fields -  $2M$

Interlayer capacitance -  $2M$

peripheral capacitance -  $2M$

(b) CMOS methodology -  $1M$

CMOS AND gate circuit diagram -  $3M$

Layout diagram -  $2M$

5 (a) Any six rules -  $6 \times 1 = 6M$

(b) Any six scaling factors -  $6 \times 1 = 6M$

6 (a) Any five issues -  $5 \times 1 = 5M$

(b) Block diagram -  $3M$

Explanation -  $3M$

7 (a) Any three Comparisons -  $3 \times 2 = 6M$

(b) Design Considerations -  $3M$

Design of multiplier -  $3M$

8 (a) Truth table of Full adder -  $2M$

Sum and carry expressions -  $1M$

Realisation using PLA -  $3M$

(b) Classification of ASICs -  $1M$

Description of each -  $5 \times 1 = 5M$

9 (a) Description of Structural Design -  $2M$

Syntax of VHDL Component declaration -  $2M$

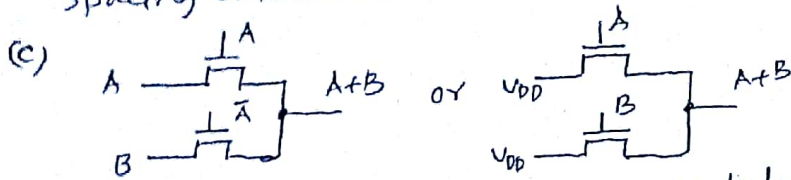
Syntax of VHDL Component statement -  $2M$

(b) Package declaration -  $1M$

Entity declaration -  $2M$

Architecture body -  $3M$

1. (a) photolithography is the technology to create a pattern on the silicon wafer using an ultraviolet (UV) ray of light.
- (b) In CMOS circuits, the parasitic NPN and PNP bipolar junction transistor (BJT) form a thyristor of PNPN structure. When the thyristor is triggered by noise/glitch, the BJTs become on, and it leads to a formation of a stable circuit between the power and ground lines, known as CMOS latch up. Prevention techniques are (i) Use of p+ and n+ guard rings around nmos and pmos connected to ground and VDD respectively. (ii) Keeping sufficient spacing between nmos and pmos transistors.



- (d) The delay unit  $\tau$  is defined as the product of sheet resistance,  $R_s$  and standard unit gate capacitance,  $C_g$  i.e.,  $\tau = R_s \square C_g$
- (e) Sheet resistance,  $R_s = \rho/t$  ohm per square where  $\rho$  = resistivity,  $t$  = thickness
- (f) A stick diagram is a simple way of representing the layout by using thick lines with their interconnections.
- (g) Switch logic is the one which uses pass transistors & transmission gates.
- (h) Requirement of a logical and systematic approach.

(i) PAL

1. Programmable AND, fixed OR
2. Cheaper than PLA
3. Moderate flexible
4. It can implement SOP with limited number of terms

PLA

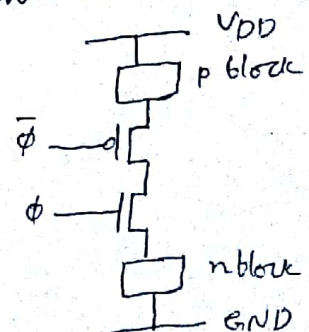
- (1) Both AND, OR are programmable.
- (2) Costlier than PAL.
- (3) Extremely flexible.
- (4) It can implement SOP with any number of terms.

(j) Functions :- Functions are used to describe frequently used sequential algorithms that return a single value. This value is returned to the calling program using a return statement.

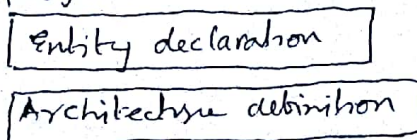
Procedures :- These are used to partition large behavioral descriptions. A procedure can return zero or more values using parameters of mode out and inout.

(k) The arrangement of a clocked CMOS logic is shown below.

$\phi$  is the clock signal. The logic in clocked CMOS is evaluated only during the ON period of the clock.

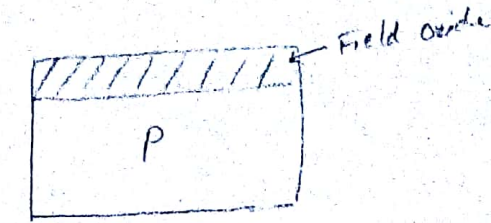


(L) VHDL program file structure is shown below.



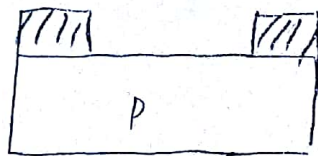
## 2(a) Fabrication of n-MOSFET -

4



(a) Field oxide

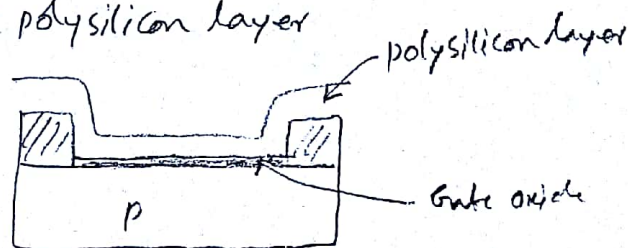
→ open a window in the field oxide



→ Subject to gate oxidation

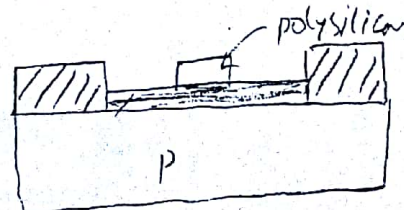


→ deposit polysilicon layer

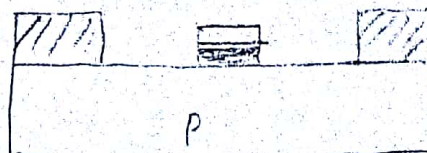


By chemical vapour deposition (CVD), a layer of polysilicon is deposited on the top of silicon

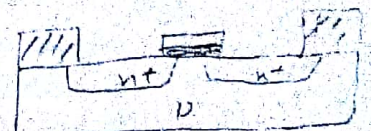
→ Retain the polysilicon only over the gate and remove it from the rest of the region by etching process



→ Open windows for source and drain. A quick dip in hydrofluoric acid will remove thin oxide



→ Realise source and drain



2(b)

CMOS Inverter:-

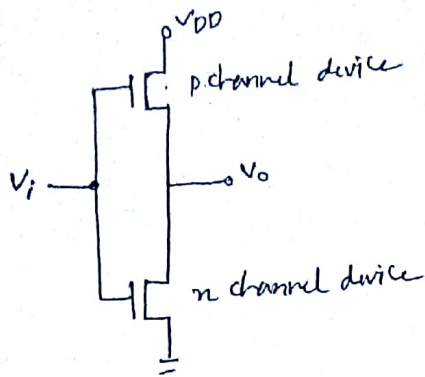
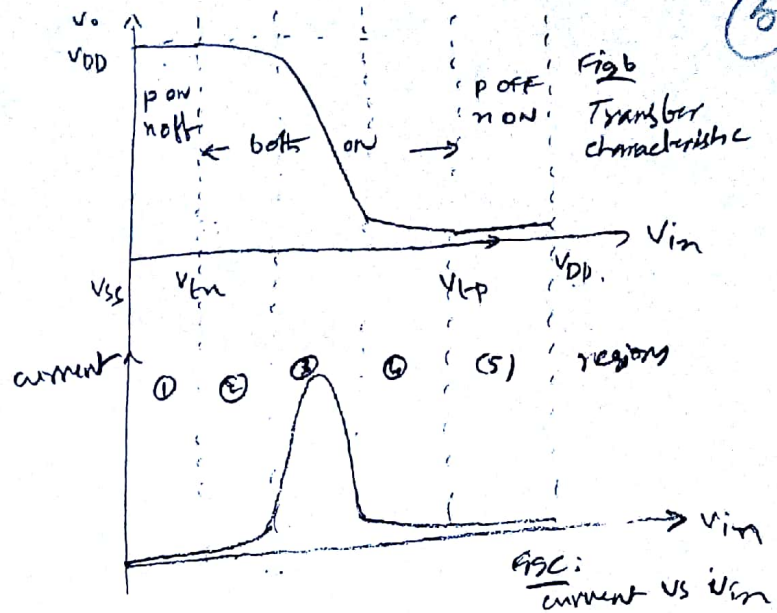


Fig: CMOS Inverter



The current voltage relationships for the MOS transistor are

$$I_{ds} = K \frac{W}{L} \left[ (V_{gs} - V_t) V_{ds} - \frac{V_{ds}^2}{2} \right]; \text{ in the resistive region}$$

$$I_{ds} = K \frac{W}{L} \frac{(V_{gs} - V_t)^2}{2}; \text{ in the saturation region}$$

In both cases the factor  $K$  is a technology dependent parameter such that  $K = \frac{\epsilon_{ins} \epsilon_0 \mu}{D}$

The factor  $\frac{W}{L}$  is contributed by geometry.

It is common practice to write  $\beta = K \frac{W}{L}$

$\beta$  may be applied to both nmos and pmos transistor as follows:

$$\beta_n = \frac{\epsilon_{ins} \epsilon_0 \mu_n}{D} \frac{W_n}{L_n}$$

$$\beta_p = \frac{\epsilon_{ins} \epsilon_0 \mu_p}{D} \frac{W_p}{L_p}$$

where  $W_n$  and  $L_n$ ,  $W_p$  and  $L_p$  are the n and p transistor dimensions respectively

CMOS inverter has five distinct regions of operation which are shown in the above figure. ①

In region 1 for which  $V_{in} = \text{logic } 0$ , the p transistor is fully turned on and the n transistor is fully turned off. Thus no current flows through the inverter and the output is directly connected to  $V_{DD}$ .

In region 5,  $V_{in} = \text{logic } 1$ , n transistor is fully on while the p-transistor is fully off. Again, no current flows and a good logic '0' appears at the output.

In region 2 the input voltage has increased to a level which just exceeds the threshold voltage of the n-transistor. The n transistor conducts and has a large voltage between source and drain, so it is in saturation. The p transistor is also conducting but with only a small voltage across it, it operates in the unsaturated resistive region. A small current now flows through the inverter from  $V_{DD}$  to  $V_{ss}$ .

Region 4 is similar to region 2 but with the roles of p and n transistors reversed. However, the current magnitudes in regions 2 and 4 are small and most of the energy consumed in switching from one state to the other is due to the larger current which flows in region 3.

Region 3 is the region in which the inverter exhibits gain and in which both transistors are in saturation.

The currents in each device must be same since the transistors are in series.

$$I_{dsp} = -I_{dsn}$$

$$\frac{\beta_p}{2} (V_{in} - V_{DD} - V_{tp})^2 = -\frac{\beta_n}{2} (V_{in} - V_{tn})^2$$

$$V_{in} - V_{DD} - V_{tp} = -\sqrt{\frac{\beta_n}{\beta_p}} (V_{in} - V_{tn})$$

$$V_{in} - V_{DD} - V_{tp} = -\sqrt{\frac{\beta_n}{\beta_p}} V_{in} + \sqrt{\frac{\beta_n}{\beta_p}} V_{tn}$$

$$\left(1 + \sqrt{\frac{\beta_n}{\beta_p}}\right) V_{in} = V_{DD} + V_{tp} + V_{tn} \sqrt{\frac{\beta_n}{\beta_p}}$$

$$V_{in} = \frac{V_{DD} + V_{tp} + V_{tn} \sqrt{\frac{\beta_n}{\beta_p}}}{1 + \sqrt{\beta_n/\beta_p}}$$

(7)

In region 3, since both transistors are in saturation, they act as current sources so that the equivalent circuit in this region is two current sources in series between  $V_{DD}$  and  $V_{SS}$  with the output voltage coming from their common point.

The region is inherently unstable in consequence and the changover from one logic level to the other is rapid.

If  $\beta_n = \beta_p$  and if  $V_{tn} = -V_{tp}$ , then from the above equation

$$V_{in} = 0.5 V_{DD}$$

This implies that the changover between logic levels is symmetrically disposed about the point at which

$$V_{in} = V_{out} = 0.5 V_{DD}$$

3(a)

8

Drain to source current  $I_{DS}$  versus voltage  $V_{DS}$  relationship

Gate voltage induces charge in the channel region between source and drain. The electric field created by the voltage  $V_{DS}$ , moves this charge from source to drain.

$$\text{Current from drain to source, } I_{DS} = \frac{\text{charge induced in channel } (Q_c)}{\text{Electron transit time } (\tau)} \rightarrow (1)$$

$$\text{Transit time, } \tau_{sd} = \frac{\text{length of the channel } (L)}{\text{velocity } (v)}$$

$$\text{but velocity, } v = \mu E_{ds}$$

where  $\mu$  = electron or hole mobility (surface)  
 $E_{ds}$  = electric field (drain to source)

$$\text{Now } E_{ds} = \frac{V_{ds}}{L}$$

$$\therefore v = \mu \frac{V_{ds}}{L}$$

$$\text{Thus, } \tau_{sd} = \frac{L}{\mu \frac{V_{ds}}{L}} = \frac{L^2}{\mu V_{ds}} \rightarrow (2)$$

The non saturated region:-

Charge induced in channel due to gate voltage is due to the voltage difference between the gate and the channel,  $V_{gs}$  (assuming substrate connected to source).

The voltage along the channel varies linearly with distance  $x$  from the source due to IR drop in the channel and assuming that the device is not saturated then the average value is  $\frac{V_{ds}}{2}$ .

Furthermore, the effective gate voltage  $V_g = V_{gs} - V_t$  where  $V_t$  = threshold voltage.

$$\text{Note that charge/unit area} = E_g \epsilon_{ins} \epsilon_0$$

where  $E_g$  = average electric field gate to channel

$\epsilon_{ins}$  = relative permittivity of insulation between gate and channel.

$\epsilon_0$  = permittivity of free space

Thus induced charge,

$$Q_c = E_g \epsilon_{ins} \epsilon_0 W L \quad \text{where } W = \text{width of the channel}$$

$L = \text{length of the channel}$

$$\text{Now } E_g = \frac{[(V_{gs} - V_t) - \frac{V_{ds}}{2}]}{D}$$

where  $D$  = oxide thickness

$$\text{Thus } Q_c = W L \epsilon_{ins} \epsilon_0 \frac{[(V_{gs} - V_t) - \frac{V_{ds}}{2}]}{D} \rightarrow (3)$$

using equations (2) and (3), eq(1) can be written as

$$I_{DS} = \frac{W L \epsilon_{ins} \epsilon_0 [(V_{gs} - V_t) - \frac{V_{ds}}{2}]}{D} \bigg/ \frac{L^2}{\mu V_{ds}}$$

(9)

$$I_{ds} = \frac{\epsilon_{ins} \epsilon_0 \mu}{D} \frac{W}{L} \left[ (V_{gs} - V_t) - \frac{V_{ds}}{2} \right] V_{ds}$$

$$\text{Let } K = \frac{\epsilon_{ins} \epsilon_0 \mu}{D}$$

$$\therefore I_{ds} = K \frac{W}{L} \left[ (V_{gs} - V_t) V_{ds} - \frac{V_{ds}^2}{2} \right] \rightarrow (4)$$

$$= \beta \left[ (V_{gs} - V_t) V_{ds} - \frac{V_{ds}^2}{2} \right] \text{ where } \beta = K \frac{W}{L}$$

The gate/channel capacitance is given by  $C_g = \frac{\epsilon_{ins} \epsilon_0 WL}{D}$

$\therefore$  The above relation shown be shown as

$$I_{ds} = \frac{C_g \mu}{L^2} \left[ (V_{gs} - V_t) V_{ds} - \frac{V_{ds}^2}{2} \right] \rightarrow (4a)$$

The Saturated region:-

Saturation begins when  $V_{ds} = V_{gs} - V_t$  since at this point the IR drop in the channel equals the effective gate to channel voltage at the drain and we may assume that the current remains fairly constant as  $V_{ds}$  increases further. Thus

$$\begin{aligned} I_{ds} &= K \frac{W}{L} \left[ (V_{gs} - V_t) V_{ds} - \frac{V_{ds}^2}{2} \right] \\ &= K \frac{W}{L} \left[ (V_{gs} - V_t)^2 - \frac{(V_{gs} - V_t)^2}{2} \right] \\ &= K \frac{W}{L} \frac{(V_{gs} - V_t)^2}{2} \rightarrow (5) \end{aligned}$$

or we may write

$$I_{ds} = \frac{\beta}{2} (V_{gs} - V_t)^2$$

or

$$I_{ds} = \frac{C_g \mu}{2 L^2} (V_{gs} - V_t)^2$$

The expressions derived for  $I_{ds}$  hold for both enhancement and depletion mode devices.

3(b)

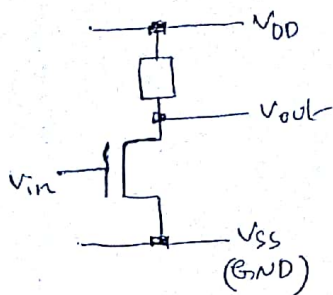
10

## Alternate forms of pull up :-

There are four possible arrangements

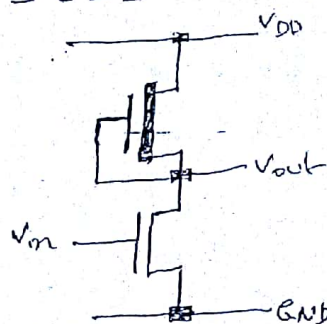
- (i) Load resistance  $R_L$
- (ii) nmos depletion mode transistor pull up
- (iii) nmos enhancement mode pull up
- (iv) Complementary transistor pull up (cmos)

(i) Load resistance  $R_L$  :-



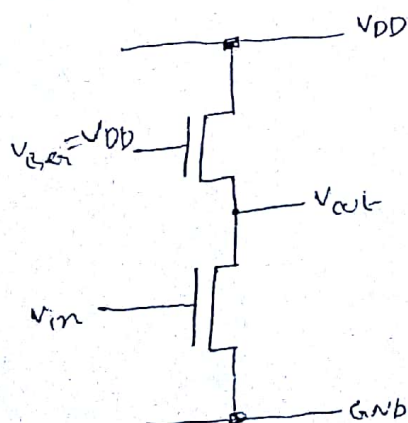
This arrangement is not often used because of the large space requirement of resistors produced in a silicon substrate.

(ii) nmos depletion mode transistor pull up :-



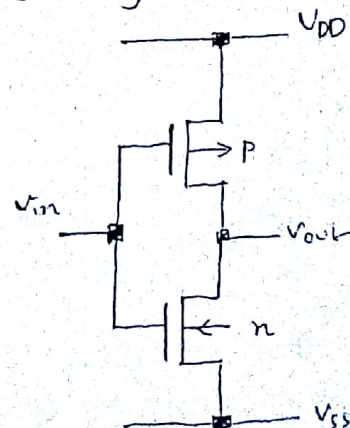
- Dissipation is high since rail to rail current flows when  $V_{in} = \text{logical } 1$ .
- Switching of output from 1 to 0 begins when  $V_{in}$  exceeds  $V_t$  of pull down device.
- When switching the output from 1 to 0, the pull up device is non-saturated initially and this presents lower resistance through which to charge capacitive loads.

(iii) nmos enhancement mode pull up :-



- Dissipation is high since current flows when  $V_{in} = \text{logical } 1$
- $V_{out}$  can never reach  $V_{DD}$  (logical 1) if  $V_{gate} = V_{DD}$  as is normally the case

(iv) Complementary transistor pull up (cmos)



- No current flow either for logical 0 or for logical 1 inputs
- Full logical 1 and 0 levels are presented at the output.
- For devices of similar dimensions the p-channel is slower than the n-channel device.

## (a) Wiring Capacitances:-

(11)

Earlier, the area capacitances associated with the layers to substrate and from gate to channel are considered.

However, there are other significant sources of capacitance which contribute to the overall wiring capacitance.

Three such sources are

(i) Fringing Fields:- The fringe field is the peripheral magnetic field outside of the magnetic core.

Capacitance due to bringing field effects can be a major component of the overall capacitance of interconnect wires.

For fine line metallization, the value of bringing field capacitance ( $C_{ff}$ ) can be of the same order as that of the area capacitance.

Thus,  $C_{ff}$  should be taken into account if accurate prediction of performance is needed.

$$C_{ff} = \epsilon_{SiO_2} \epsilon_0 l \left[ \frac{\pi}{\ln \left\{ 1 + \frac{2d}{t} \left( 1 + \sqrt{1 + \frac{t}{d}} \right) \right\}} - \frac{t}{4d} \right]$$

Where  $l$  = wire length

$t$  = thickness of wire

$d$  = wire to substrate separation

Then, total wire capacitance

$$C_w = C_{area} + C_{ff}$$

## (ii) Interlayer Capacitances:-

Quite obviously the parallel plate effects are present between one layer and another.

For example, for a given area, metal to polysilicon capacitance must be higher than metal to substrate.

The reason for not taking such effects into account for simple calculations is that the effects occur only when layers cross or when one layer underlies another, and in consequence interlayer capacitance is highly dependent on layout. However, for regular structures it is readily calculated and contributes significantly to the accuracy of circuit modelling and delay calculations.

## Peripheral Capacitance..

(12)

The source and drain n-diffusion regions (n-active regions for orbit process) form junctions with the p-substrate or p-well at well defined and uniform depths;

Similarly for p-diffusion (p-active) regions in n-substrates or n-wells

For diffusion regions, each diode thus formed has associated with it a peripheral (side-wall) capacitance in picofarads per unit length, which, in total, can be considerably greater than the area capacitance of the diffusion region to substrate;

The smaller the source to drain area, the greater becomes the relative value of the peripheral capacitance.

For orbit processes, the n-active and p-active regions are formed by impurity implant at the surface of the silicon and thus, having negligible depths, they have negligible peripheral capacitance.

However, for n- and p-regions formed by a diffusion process, the peripheral capacitance is important and becomes particularly so as we shrink the device dimensions.

In order to calculate the total diffusion capacitance, we must add the contributions of area and their peripheral components.

$$C_{\text{total}} = C_{\text{area}} + C_{\text{periph}}$$

Table : Typical values of for diffusion capacitances

| Diffusion Capacitance             | Typical values                                |                                                |                                                |
|-----------------------------------|-----------------------------------------------|------------------------------------------------|------------------------------------------------|
|                                   | 5 $\mu\text{m}$                               | 2 $\mu\text{m}$                                | 1.2 $\mu\text{m}$                              |
| Area C ( $C_{\text{area}}$ )      | $1.0 \times 10^{-4} \text{ pF}/\mu\text{m}^2$ | $1.75 \times 10^{-4} \text{ pF}/\mu\text{m}^2$ | $3.75 \times 10^{-4} \text{ pF}/\mu\text{m}^2$ |
| Periphery ( $C_{\text{periph}}$ ) | $8.0 \times 10^{-4} \text{ pF}/\mu\text{m}$   | negligible*                                    | negligible*                                    |

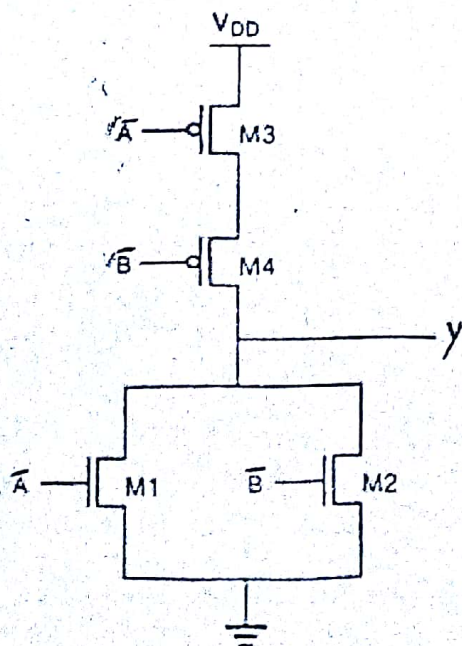
\* Assuming implanted regions of negligible depths.

40)

## CMOS Design methodology :-

It has three steps

1. Given the Boolean expression, take its complement.
2. Design pull down network (PDN) by realizing
  - AND terms using series connected nMOSFETs
  - OR terms using parallel connected nMOSFETs.
3. Design pull up network just reverse (& dual) of the PDN.
  - AND terms using parallel connected pMOSFETs
  - OR terms using series connected pMOSFETs.



13

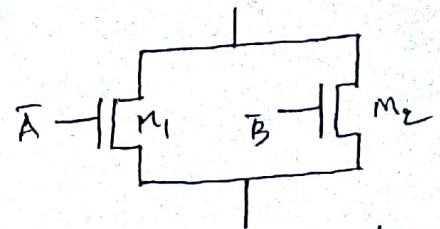
## Design of two input- AND gate using CMOS Technology

If A and B are inputs and Y is the output,  $Y = AB$   
As per CMOS methodology

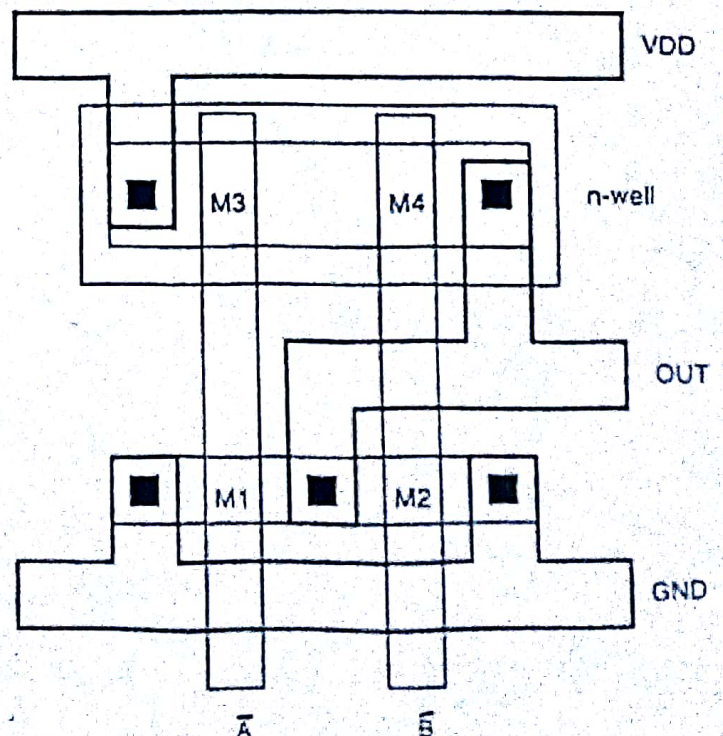
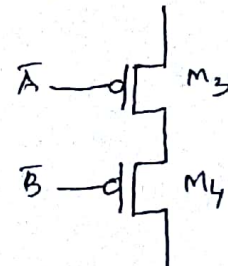
1. Take Complement of Y.

$$\begin{aligned}\bar{Y} &= \overline{AB} \\ &= \bar{A} + \bar{B}\end{aligned}$$

2. Design the Pull down Network. There is only one OR term. Use parallel connected nMOSFETs



3. Design the pull up network. There is only one OR term. Use series connected pMOSFETs.



# $\lambda$ based Design Rules:-

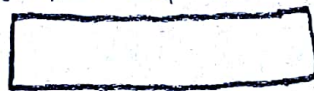
14

Features of  $\lambda$ -based Design rules:

- $\lambda$  is the size of a minimum feature
- All the dimensions are specified in integer multiple of  $\lambda$
- Specifying  $\lambda$  particularises the scalable rules
- parasitics are generally not specified in  $\lambda$  units
- These rules specify geometry of masks, which will provide reasonable yields.

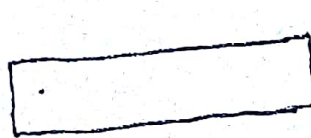
Guidelines for using  $\lambda$ -based Design rules:

Diffusion not lower than  $2\lambda$



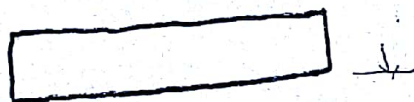
; minimum width of diffusion is  $2\lambda$

poly not lower than  $2\lambda$



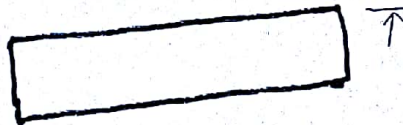
; minimum width of poly is  $2\lambda$

Diffusion and diffusion not lower than  $3\lambda$



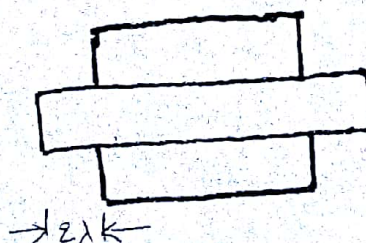
$3\lambda$

; minimum distance between two diffusion layers  $3\lambda$

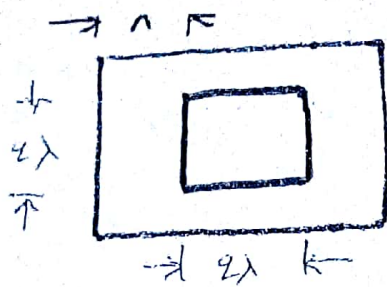


$3\lambda$

poly cross diffusion



It is necessary for the poly to completely cross active, otherwise the transistor that has been created crossing of diffusion and poly, will be shorted by diffused path of source and drain.



Contact cut on Metal

(15)

Contact window will be of  $2\lambda$  by  $2\lambda$  that is minimum feature size while metal deposition is of  $4\lambda$  by  $4\lambda$  for reliable contact.

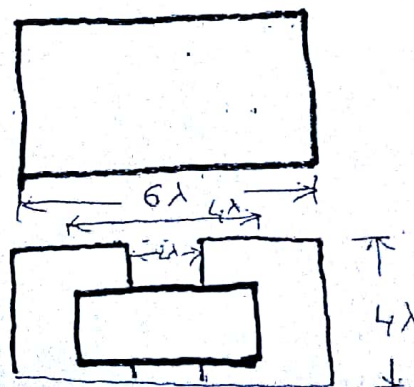
In Metal:

Two metal wires have  $3\lambda$  distance between them to overcome capacitance coupling and high frequency coupling. Metal wires width can be as large as possible to decrease resistance.



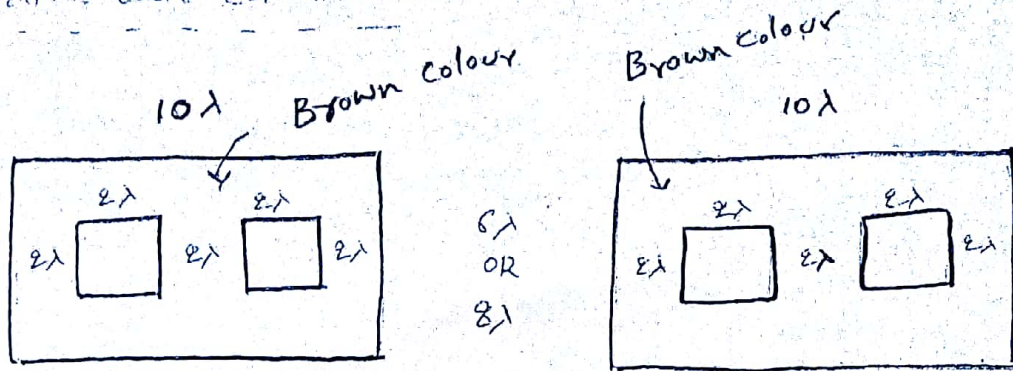
Buttering Contact:

Buttering Contact is used to make poly and silicon contact. Window's original width is  $4\lambda$ , but on overlapping width is  $2\lambda$ . So actual contact area is  $6\lambda$  by  $4\lambda$ .



The distance between two wells

16



The distance between two wells depends on the well potential as shown above.

The reason for  $8\lambda$  is that if both wells are at same high potential then the depletion region between them may touch each other causing punch-through.

The reason for  $6\lambda$  is that if both wells are at different potentials then depletion region of one well will be smaller, so both depletion regions will not touch each other. So  $6\lambda$  will be good enough.

The active region has length  $10\lambda$  which is distributed over the followings

- $2\lambda$  for source diffusion
- $2\lambda$  for drain diffusion
- $2\lambda$  for channel length
- $2\lambda$  for source side encroachment
- $2\lambda$  for drain side encroachment

5(b)

Scaling factors for device parameters --

In this section, simple derivations and calculations reveal the effects of scaling.

(1) Gate Area  $A_g$  :-

$$A_g = L W$$

where  $L$  and  $W$  are the channel length and width respectively.

Both are scaled by  $\frac{1}{\alpha}$ .

Thus  $A_g$  is scaled by  $\frac{1}{\alpha^2}$ .

(2) Gate Capacitance per unit area  $C_0$  &  $C_{ox}$ 

$$C_0 = \frac{\epsilon_{ox}}{D}$$

where  $\epsilon_{ox}$  = permittivity of the gate oxide (thin oxide)  $[= \epsilon_{ins} \epsilon_0]$   
and  $D$  = gate oxide thickness.

Here  $D$  is scaled by  $\frac{1}{\beta}$ .

Thus  $C_0$  is scaled by  $\frac{1}{(1/\beta)} = \beta$ .

(3) Gate Capacitance  $C_g$  :-

$$C_g = C_0 L W$$

Thus  $C_g$  is scaled by  $\beta \frac{1}{\alpha^2} = \frac{\beta}{\alpha^2}$ .

(4) Parasitic Capacitance  $C_x$  :-

$C_x$  is proportional to  $\frac{A_x}{d}$

where  $d$  = depletion width around source & drain

$d$  is scaled by  $\frac{1}{\alpha}$

$A_x$  = area of the depletion region around source & drain.

$A_x$  is scaled by  $\frac{1}{\alpha^2}$

Thus  $C_x$  is scaled by  $\frac{1}{\alpha^2} \cdot \frac{1}{(1/\alpha)} = \frac{1}{\alpha}$

(5) Carrier density in channel  $Q_{on}$  :-

$$Q_{on} = C_0 V_{gs}$$

where  $Q_{on}$  is the average charge per unit area in the channel in the on state. Note that  $C_0$  is scaled by  $\beta$  and  $V_{gs}$  is scaled by  $\frac{1}{\beta}$ .

Thus  $Q_{on}$  is scaled by 1.

(6) Channel Resistance  $R_{on}$  :-

$$R_{on} = \frac{L}{W} \frac{1}{Q_{on} \mu}$$

where  $\mu$  is the carrier mobility in the channel and is assumed constant.

Thus  $R_{on}$  is scaled by  $\frac{1}{\alpha} \frac{1}{1/\alpha} = 1$

(7) Gate Delay  $T_d$  :-

$T_d$  is proportional to  $R_{on} C_g$

Thus  $T_d$  is scaled by  $1 \frac{\beta}{\alpha^2} = \frac{\beta}{\alpha^2}$

(8) Maximum operating frequency  $f_o$ .

$$f_o = \frac{W}{L} \frac{\mu C_o V_{DD}}{C_g}$$

&  $f_o$  is inversely proportional to delay  $T_d$ .

Thus  $f_o$  is scaled by  $\frac{\alpha^2}{\beta}$

(9) Saturation current  $I_{DSS}$  :-

$$I_{DSS} = \frac{C_o \mu}{2} \frac{W}{L} (V_{GS} - V_t)^2$$

Noting that both  $V_{GS}$  and  $V_t$  are scaled by  $\frac{1}{\beta}$ , we have

$I_{DSS}$  scaled by  $\beta \left(\frac{1}{\beta}\right)^2 = \frac{1}{\beta}$

(10) Current density  $J$  :-

$$J = \frac{I_{DSS}}{A}$$

where  $A$  is the cross-sectional area of the channel in the 'on' state which is scaled by  $\frac{1}{\alpha^2}$ .

So,  $J$  is scaled by  $\frac{1/\beta}{1/\alpha^2} = \frac{\alpha^2}{\beta}$

(11) Switching Energy per gate  $E_g$  :-

$$E_g = \frac{1}{2} C_g (V_{DD})^2$$

So,  $E_g$  is scaled by  $\frac{\beta}{\alpha^2} \frac{1}{\beta^2} = \frac{1}{\alpha^2 \beta}$

## 6(a) Architectural issues :-

In the design of a VLSI system, a logical and systematic approach is essential.

The design time tend to rise exponentially with increased complexity.

### Guidelines in designing a VLSI system

1. Define the requirements (properly and carefully).
2. Partition the overall architecture into appropriate subsystems
3. Consider Communication paths carefully in order to develop sensible interrelationships between subsystems.
4. Draw a block plan of how the system is to map onto the Silicon
5. Aim for regular structures so that design is largely a matter of replication.
6. Draw suitable (stick & symbolic) diagrams of the least-cells of the subsystems.
7. Convert each cell to a layout
8. Carefully and thoroughly carry out a design rule check on each cell.
9. Simulate the performance of each cell/subsystem.

The whole design process will be greatly assisted if considerable care is taken with

- (1) the partitioning of the system so that there are clean or clear subsystems with a minimum interdependence and complexity of interconnection between them.
- (2) the design simplification within subsystems so that architectures are adopted which allow the exploitation of a cellular design concept. This allows the system to be composed of relatively few standard cells which are replicated to form highly regular structures.

6(b)

(20)

Design of an ALU Subsystem

The heart of the ALU is a 4 bit adder circuit. A 4 bit adder must take the sum of two 4 bit numbers.

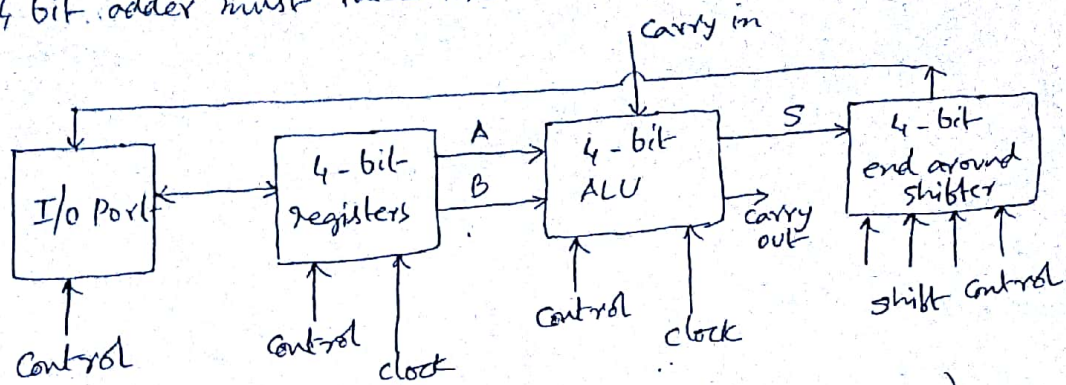


Fig 8.1: 4 bit data path for processor (block diagram)

The shifter circuit accepts and shifts a 4 bit parallel sum from the ALU.

The sum is to be stored in parallel at the output of the adder from where it may be fed through the shifter and back to the register array.

Thus, a single 4-bit data bus is needed from the adder to the shifter and another 4-bit bus is required from the shifted output back to the register array (since the shifter is merely a switch array with no storage capability).

As far as the input to the adder is concerned, the two 4 bit parallel numbers to be added are to be presented in parallel on two 4 bit buses.

Assume clock phase  $\phi$ , as being the phase during which signals are fed along buses to the adder input and during which their sum is stored at the adder output.

Thus clock signals are required by the ALU as shown.

The shifter is unclocked but must be connected to four shift control lines. It is also necessary to provide a 'carry out' signal from the adder and, in the general case, to provide both a possible 'carry in' signal, as indicated in the above fig.

7(a)

## Switch logic and gate logic

(21)

Features of switch logic are

- It is based on pass transistors & on transmission gates
- It is better for small arrays
- It takes no static current from the supply rails.

The power dissipation is low.

- It is similar to logic arrays based on relay contacts
- Logic levels propagated through pass transistors are degraded by threshold voltage effects.

- 
- pass transistors have significant disadvantage in that load presented to the inputs depend on the load connected to the output, and output rise/fall times are a strong function of input rise/fall times. This makes timing analysis difficult. As a result pass transistors tend to be only used in localized environments where the input drivers and output loads are well understood.

- pass transistor logic often uses buffer transistors.

Gate logic is based on the general arrangement typified by the inverter circuits (the inverter being the simplest gate).

The drawbacks of switch logic can be overcome with gate logic but requires more transistors.

- Lower unit cost
- Higher reliability
- Lower power dissipation, lower weight and lower volume.
- Better performance - particularly in terms of speed power product.
- Enhanced repeatability
- Reduced design/development periods.

If the subsystems are regular and therefore composed of relatively few actual logic cell circuits, then the designer can concentrate on the main problem of VLSI design - the routing of interconnections and communication paths.

### Array multiplier

The design of an unsigned array multiplier is based on the simple pen and paper method of multiplication - taking one bit at a time from the multiplier, the partial products are generated by multiplying each bit to the multiplicand.

These partial products are then written in rows, by shifting each row by one bit position to the right.

These rows are then added column-wise. This multiplier has a very regular architecture and is most suitable for VLSI implementation.

The following fig illustrates an unsigned array multiplier.

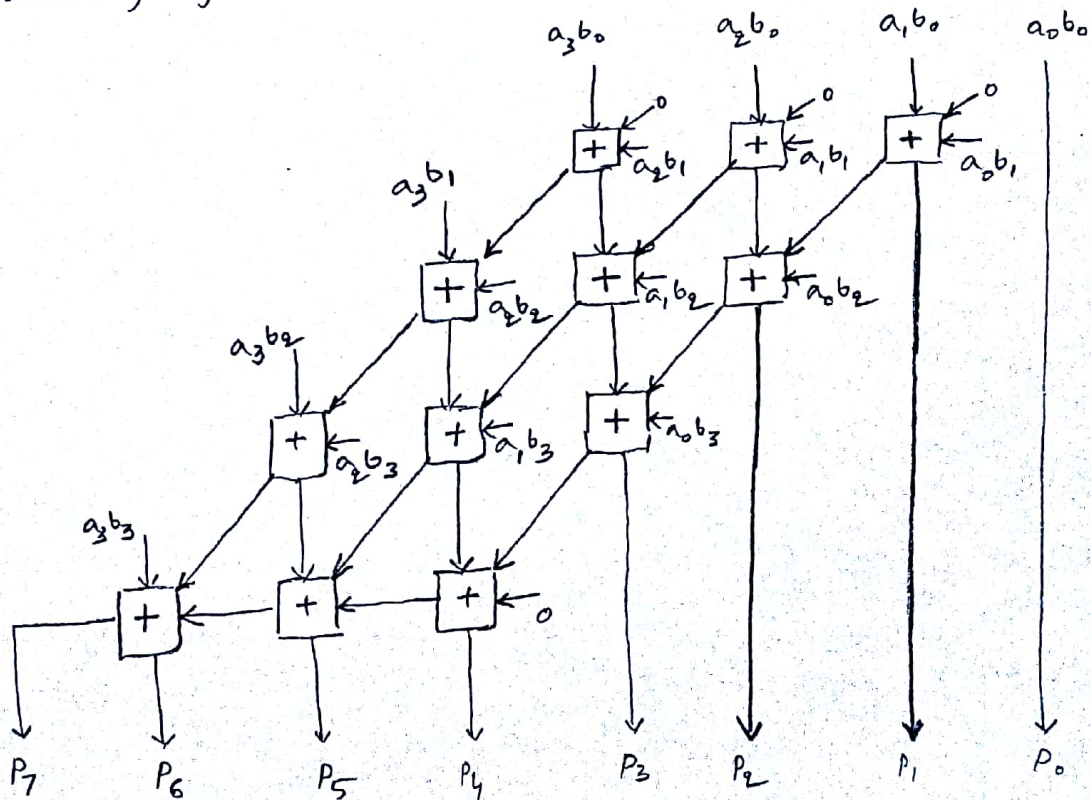


Fig: Unsigned array multiplier

# Full Adder using PLA

Truth table for Full Adder

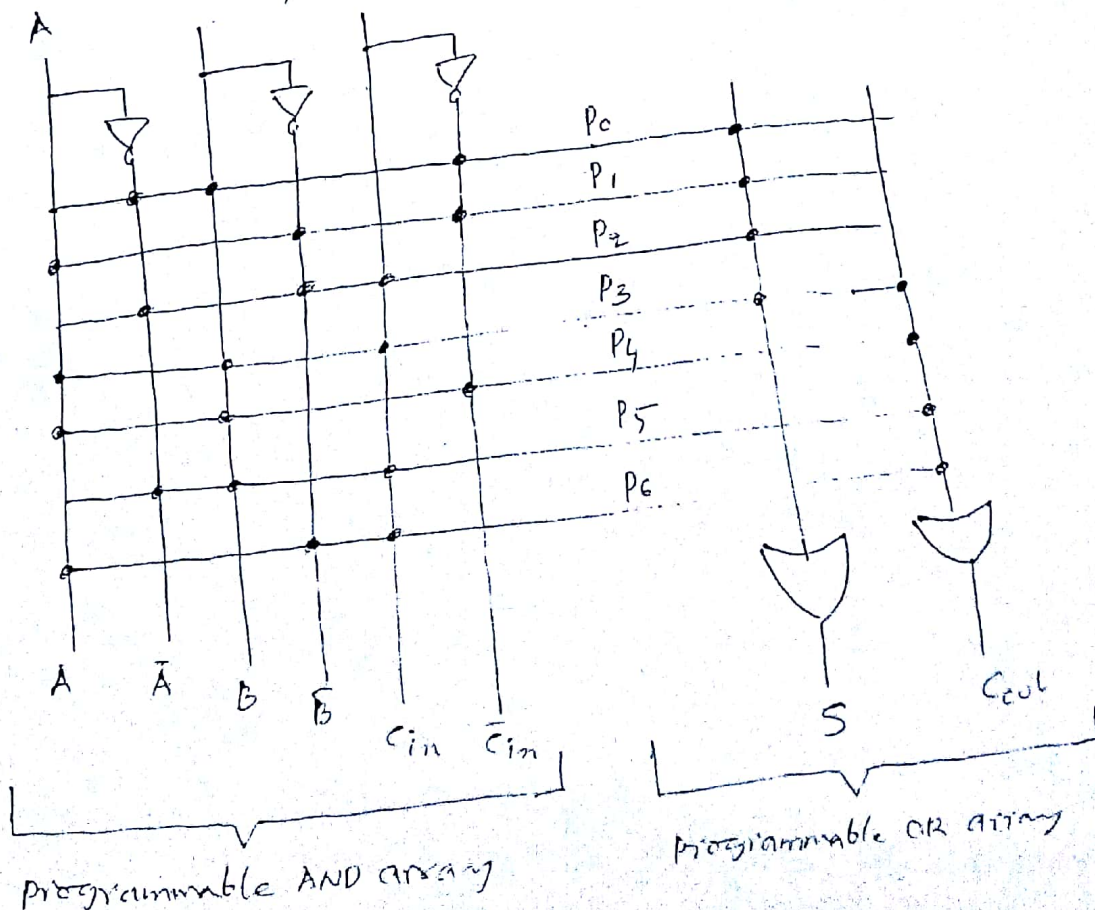
| Inputs |   |                 | Outputs |       |
|--------|---|-----------------|---------|-------|
| A      | B | C <sub>in</sub> | S       | Carry |
| 0      | 0 | 0               | 0       | 0     |
| 0      | 1 | 0               | 1       | 0     |
| 1      | 0 | 0               | 1       | 0     |
| 1      | 1 | 0               | 0       | 1     |
| 0      | 0 | 1               | 1       | 0     |
| 0      | 1 | 1               | 0       | 1     |
| 1      | 0 | 1               | 0       | 1     |
| 1      | 1 | 1               | 1       | 1     |

S = Sum  
C = Carry

From the truth table

$$S = \overline{A} \overline{B} C_{in} + A \overline{B} \overline{C}_{in} + \overline{A} B C_{in} + A B \overline{C}_{in}$$

$$C_{out} = A B \overline{C}_{in} + \overline{A} B C_{in} + A \overline{B} C_{in} + A B C_{in}$$



8(b)

Full form of ASIC is Application Specific Integrated Circuit.

Types of ASICs:-

- (i) Full custom ASICs
- (ii) Standard cell based ASICs
- (iii) Gate array based ASICs
- (iv) Programmable logic devices
- (v) Field Programmable Gate Arrays

In a full custom ASIC, an engineer designs some or all of the logic cells, circuits or layout specifically for one ASIC.

Standard cell based ASIC uses predesignated logic cells (AND gates, OR gates, multiplexers and flip flops for example) known as standard cells.

In a gate array based ASIC, the transistors are predefined on the silicon wafer. Both cell based and gate array ASICs use predefined cells, but there is a difference - we can change the transistor sizes in a standard cell to optimise speed and performance, but the device sizes in a gate array are fixed.

Programmable logic devices (PLDs) are standard ICs that are available in standard configurations from a catalog of parts and are sold in very high volume to many different customers. These are subdivided in ROM, Programmable array logic (PAL), and Programmable logic array (PLA).

A step above the PLD in complexity is the field Programmable gate array (FPGA). FPGA is an IC which can be hardware programmed to implement various logic functions.

9(a)

## Structural Design elements

The body of an architecture is a series of concurrent statements.

In VHDL, each concurrent statement executes simultaneously with the other concurrent statements in the same architecture body.

Thus, in a VHDL architecture body, if the last statement updates a signal that is used by the first statement, then the simulator will go back to that first statement and update its results according to the signal that just changed.

VHDL has several different concurrent statements, as well as a mechanism for bundling a set of sequential statements to operate as a single concurrent statement.

The most basic of VHDL's concurrent statements is the component statement, whose syntax is shown below

Syntax of a VHDL Component statement

label : component-name port-map (signal 1, signal 2, ..., signal n);

label : component-name port-map (port 1  $\Rightarrow$  signal 1, port 2  $\Rightarrow$  signal 2, ..., port n  $\Rightarrow$  signal n);

Component declaration is shown below

Syntax of a VHDL Component declaration

Component component-name

port (signal-names : mode signal-type;

signal-names : mode signal-type;

⋮

signal-names : mode signal-type);

end Component;

A complex digital system is implemented in a hierarchical order. The whole circuit is partitioned into blocks, and the blocks are then partitioned into sub-blocks. This process is continued to a level where the leaf level or the last level is merely the basic logic gates. Therefore, any hardware description language supports this hierarchical design concept. In the VHDL, this hierarchical design concept is addressed by the structural modeling style. In this style of modeling, a bigger digital block is implemented by a set of smaller blocks or basic logic gates. The blocks or gates that are used to design the bigger block are called components. The components are connected using signals to implement the bigger block.

9(b)

# 4 bit Adder using Dataflow modeling

26

library IEEE;

use IEEE.std\_logic\_1164.all;

use IEEE.std\_logic\_arith.all;

use IEEE.std\_logic\_unsigned.all;

entity adder4bit is

port (A, B : in std\_logic\_vector (3 downto 0);

cin : in std\_logic;

cout : out std\_logic;

sum : out std\_logic\_vector (3 downto 0));

generic (n : integer := 4);

end adder4bit;

architecture Dataflow of adder4bit is

signal result : std\_logic\_vector (4 downto 0);

signal cin\_s : std\_logic\_vector (3 downto 0);

begin

cin\_s <= (others => '0');

result <= ('0' & A) + ('0' & B) + (cin\_s & cin);

cout <= result(4);

sum <= result(3 downto 0);

end Dataflow;

Scheme & Solutions

Prepared by

B. V. Kumaraswamy

(B. V. KUMARASWAMY)

EIE Dept.

9441503437

Ch. He 27/3/2019

Signature of the HoD.

H. Srinivas

External Examiner