

NOV - 2016

Total No. of Questions : 10]

SEAT No. :

P1756

[5058]-396

[Total No. of Pages : 2

T.E.(Computer Engineering)
PRINCIPLES OF CONCURRENT AND DISTRIBUTED
PROGRAMMING
(2012 Course)(Semester-II) (End Semster)(310249)

Time : 2½ Hours]

[Max. Marks : 70

Instructions to the candidates:

- 1) *Answer Q1 or Q2, Q3 or Q4, Q5 or Q6, Q 7 or Q 8, and Q 9 or Q10.*
- 2) *Neat diagrams must be drawn wherever necessary.*
- 3) *Figures to the right indicate full marks.*
- 4) *Assume suitable data if necessary.*

- Q1)** a) Write short note on basic user interface elements in GUI model. [6]
b) Explain shared-state concurrent model. [4]

OR

- Q2)** a) What are the dependency relationships between two tasks? [6]
b) What is the structure of a YACC file? [4]

- Q3)** a) Write short note on shared memory. [6]
b) Explain various types of parallelism. [4]

OR

- Q4)** a) Write a note on Handler's classification. [6]
b) Define speed up with respect to parallel algorithms. [4]

- Q5)** a) What is DCE? Explain it along with its components. [10]
b) Explain the following terms with respect to operating system: [8]
i) System image.
ii) Autonomy
iii) Fault Tolerance Capability.

OR

P.T.O.

- Q6)** a) What are various models used in distributed computing environment?[10]
b) Why is scalability an important feature in the design of a Distributed OS? Discuss the guiding principles for design scalable distributed system.[8]

- Q7)** a) What is Virtualization? Explain types of virtualization. [8]
b) Explain the common approaches to virtual computer systems. [8]

OR

- Q8)** a) What is Xen domain? Also explain hypervisor. [8]
b) What are the hardware-related issues that should be considered while specifying the physical systems that will host the virtual machines? [8]
Q9) a) Explain global memory in CUDA. [8]
b) How CPUs and GPUs are different? [4]
c) Write short notes on: [4]

CUDA threads

CUDA blocks

OR

- Q10)** a) Explain threads in CUDA. Also explain problem decomposition. [8]
b) Explain texture memory in CUDA. [4]
c) Write short notes on: [4]
 • CUDA grids
 • CUDA Kernels

