

GUJARAT TECHNOLOGICAL UNIVERSITY
M. E. - SEMESTER – I • EXAMINATION – WINTER • 2014

Subject code: 3715504**Date: 08-01-2015****Subject Name: Multicore & GPU Based Programming****Time: 02:30 pm - 05:00 pm****Total Marks: 70****Instructions:**

1. Attempt all questions.
2. Make suitable assumptions wherever necessary.
3. Figures to the right indicate full marks.

- Q.1** (a) Explain Amdahl's law in detail. **07**
(b) A desktop with old single core processor was running a pattern matching application. The application was spending 20% of its time on I/O and rest of the time on pattern matching using parallelization techniques. To make the overall execution faster by 4 times, how much faster CPU shall be used to replace current old generation CPU? **07**

- Q.2** (a) Write in detail at least four difference between collective vs. point-to-point communication calls in MPI. **07**
(b) Write one MPI program using point-to-point communication call where all other processes except root is sending "Hello World" message to root. After collecting "Hello World" message from all other processes, root process will display that message on terminal. **07**

OR

- (b) Describe with block diagram - detail of Multicore architecture. Explain in detail on how it is different from SMP System architecture. **07**
- Q.3** (a) Discuss the differences in how a GPU and a vector processor might execute the following code: **07**
 sum = 0.0;
 for (i = 0; i < n; i++) {
 y[i] += a * x[i] ;
 sum += z[i] * z[i] ;
 }
(b) Explain loop-carried dependency with appropriate example. Explain how we can modify a code snippet to avoid loop carried dependency. **07**

OR

- Q.3** (a) Explain why the performance of hardware multithreaded processing core might degrade if it had large caches while running many threads. **07**
(b) Write down definition of "Semaphore" and "Conditional Variable" and their performance implications. Write down difference between these two **07**

- Q.4** (a) Explain with block diagram, Nvidia Geforce GTX 280 architecture. Explain each component of the block diagram separately. **07**
(b) Write a CUDA program (including CUDA kernel) to find Transpose of a 4 X 4 Matrix. **07**

OR

- Q.4 (a)** Define the difference between coalesced and non-coalesced memory read. Explain with diagram and program snippet, which one is better for performance and why. **07**
- (b)** Write a CUDA program (including CUDA kernel) to find sum of column elements of a 4 X 4 Matrix and put resultant sum values into a 4 cell vector. **07**
- Q.5 (a)** Explain the Memory Hierarchy of OpenCL Programming Paradigm with appropriate block diagram. **07**
- (b)** Design a flow chart for complete OpenCL Program including kernel execution. **07**
- OR**
- Q.5 (a)** OpenCL vs CUDA Which one is better? why? defend your answer with strong justification. **07**
- (b)** The DOT function is the most complex of the BLAS level 1. Its formula is as **07**

$$\text{sum} += x[i] * y[i];$$

Write a complete OpenCL Program to perform DOT operation.
