

2nd International Conference on **Computer Graphics & Animation**

September 21-22, 2015 San Antonio, USA

GPUMLib framework: Using the GPU to empower machine learning research

Noel Lopes

Polytechnic of Guarda, Portugal

The amount of information being produced by humans is continuously increasing, to the point that we are generating, capturing and sharing an unprecedented volume of data from which useful and valuable information can be extracted. However, obtaining the information represents only a fraction of the time and effort needed to analyze it. Hence, we need scalable fast Machine Learning (ML) tools that can cope with large amounts of data in a realistic time frame. As problems become increasingly challenging and demanding, they become, in many cases, intractable by traditional CPU architectures. Accordingly, novel real-world ML applications will most likely demand tools that take advantage of new high-throughput parallel architectures. In this context, today GPUs (Graphics Processing Units) can be used as inexpensive highly-parallel programmable devices, providing remarkable performance gains as compared to the CPU (it is not uncommon to obtain speed-ups of one or two orders of magnitude). However, mapping algorithms to the GPU is not an easy task. To mitigate this effort we are in the process of building an open source GPU Machine Learning Library – GPUMLib to help ML researchers and practitioners' worldwide. This presentation focuses on the challenges of implementing GPU ML algorithms using CUDA. Moreover, it presents an overview of GPUMLib algorithms and tools and highlights its main benefits.

Biography

Noel Lopes is Professor at the Polytechnic of Guarda, Portugal and a Researcher at CISUC – University of Coimbra, Portugal. Currently, he is focused on extracting information from large repositories and streams of data, using supervised, unsupervised and semi-supervised machine learning algorithms. Accordingly, a line of research being pursued consists of developing parallel Graphics Processing Unit (GPU) implementations of machine learning algorithms with the objective of decreasing substantially the time required to execute them, providing the means to study larger datasets.

noel@ipg.pt

Notes: