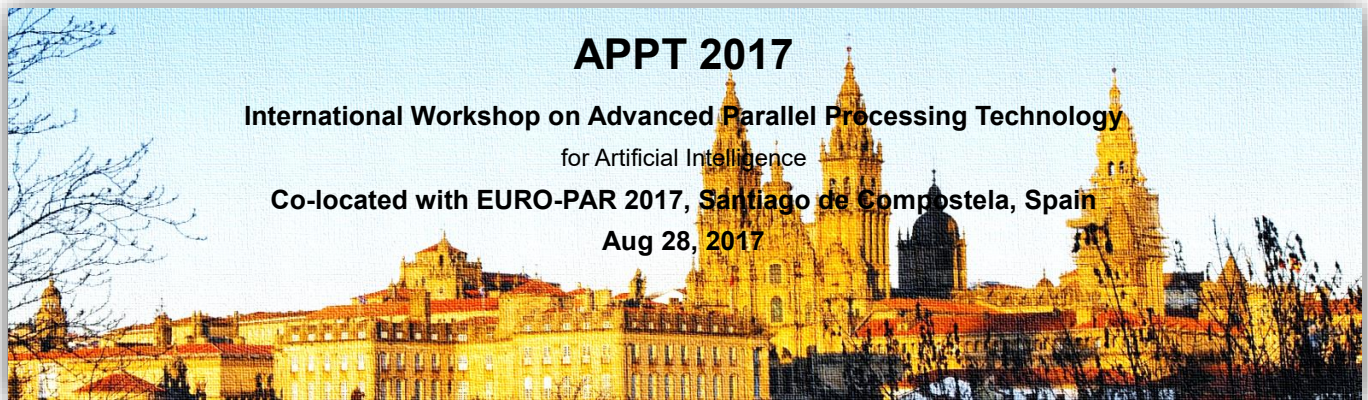




APPT 2017

International Workshop on Advanced Parallel Processing Technology

for Artificial Intelligence

Co-located with EURO-PAR 2017, Santiago de Compostela, Spain

Aug 28, 2017

Conference website: <http://appt2017.tcarch.org>

APPT is a biennial forum on parallel and distributed processing. Its scope covers all aspects of parallel and distributed computing technologies, including architectures, software systems and tools, algorithms, and applications. Recently, artificial intelligence (AI) is changing our daily life ranging from communication, social network, online searching, to navigation and health care. AI techniques such as computer vision are among the most compute-intensive and data-intensive techniques. The demand for its computation power is ever increasing, which drives the society, the academia, and the industry to investigate new parallel processing techniques. In recent years, multi-core processors, GPU, and FPGA are widely adopted to accelerate the speed of AI algorithms, and OpenMP, OpenCL, Spark, and Tensorflow are continually being improved to ease and speed up the development of AI applications. The computing models, methodologies, computer architectures, computer software, programming models, and analysis frameworks are deeply challenged by AI algorithms and applications. We invite original submissions, which advance the parallel processing techniques for AI to the forum. APPT will present and discuss recent advancement on parallel processing techniques with a special interest in AI, including many sub-areas (but not limited to) as follows,

- (a) Machine Learning
- (b) Computer Vision and Pattern Recognition
- (c) Natural Language Processing
- (d) Knowledge Representation and Reasoning
- (e) Planning and Scheduling
- (f) Autonomous Agents and Multi-agent Systems

Submissions characterized by improving AI using parallel processing are encouraged. The scope of APPT 2017 includes, but not limited to following topics,

★ Parallel Architectures

- Special-purpose architectures
- Accelerator architectures for AI applications
- Optimization on memory hierarchy
- Network and interconnect architectures
- Interconnect architectures for multicore
- Parallel I/O and storage systems
- Power-efficient architectures

★ Parallel Algorithms

- AI algorithms and applications
- Massively parallel applications
- Fault-tolerance of algorithms
- Communication & synchronization protocols
- Network algorithms, scheduling and load balancing

★ Parallel Software

- Parallel programming languages, models, compilers, environments, and tools
- Multicore runtime systems
- Resource management
- High availability, reliability and fault tolerance
- Scalability issues in parallel software and systems
- Performance modeling, analysis and tuning tools

★ Distributed and cloud computing

- Resource provisioning, allocation and scheduling
- Managing, debugging, and diagnosing
- Big data and cloud processing for AI applications
- Cloud infrastructure technologies
- Distributed caching , and query processing

Important dates:

- | | |
|-----------------------------|----------------|
| 1. Paper Submission: | April 30, 2017 |
| 2. Notification to Authors: | June 10, 2017 |
| 3. Camera-ready submission: | July 10, 2017 |