

TITAN



TITAN: Built for Science

The first thousand nodes of Titan are scheduled for installation in late 2011, but the Oak Ridge Leadership Computing Facility (OLCF) began preparing for the arrival of its next leadership resource long before the hardware was purchased. Titan's novel architecture alone is no high-performance computing (HPC) game-changer without applications capable of utilizing its innovative computing environment.

In 2009 the OLCF began compiling a list of candidate applications that were to be the vanguards of Titan—the first codes that would be adapted to take full advantage of its mixed architecture. This list was gleaned from research done for the 2009 OLCF report *Preparing for Exascale: OLCF Application Requirements and Strategy* as well as from responses from current and former INCITE awardees.

Initially 50 applications were considered, but this list was eventually pared down to a set of six critical codes from various domain sciences:

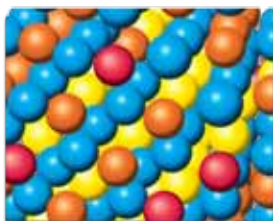
S3D, used by a team led by Jacqueline Chen of Sandia National Laboratories, is a direct numerical simulation code that models combustion. In 2009, a team led by Chen used Jaguar to create the world's first fully resolved simulation of small lifted autoigniting hydrocarbon jet flames, allowing for representation of some of the fine-grained physics relevant to stabilization in a direct-injection diesel engine.

WL-LSMS calculates the interactions between electrons and atoms in magnetic materials—such as those found in computer hard disks and the permanent magnets found in electric motors. It uses two methods. The first is locally self-consistent multiple scattering,

Preparing for Exascale: Six Critical Codes

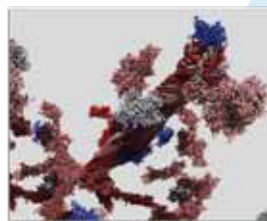
WL-LSMS

Role of material disorder, statistics, and fluctuations in nanoscale materials and systems.



LAMMPS

A simulation model investigating the properties of lignocellulose.

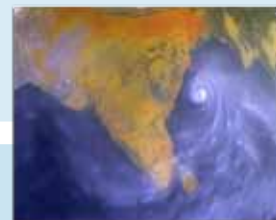


S3D

Combustion simulations to enable the next generation of diesel/bio fuels to burn more efficiently.

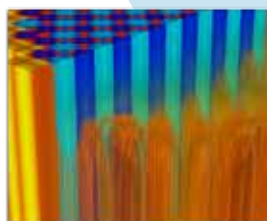
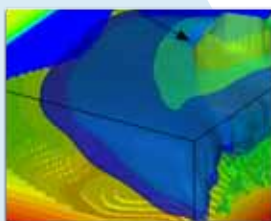
CAM-SE

Answers questions about specific climate change adaptation and mitigation scenarios.



PFLOTRAN

Stability and viability of large-scale CO₂ sequestration; predictive containment groundwater transport.



Denovo

High-fidelity radiation transport calculations that can be used in a variety of nuclear energy and technology applications.

which describes the journeys of scattered electrons at the lowest possible temperature by applying density functional theory to solve the Dirac equation, a relativistic wave equation for electron behavior. The second is the Monte Carlo Wang-Landau method, which guides calculations to explore system behavior at all temperatures, not just absolute zero. The two methods were combined in 2008 by a team of researchers from Oak Ridge National Laboratory (ORNL) to calculate magnetic materials at a finite temperature without adjustable parameters. The combined code was one of the first codes to break the petascale barrier on Jaguar.

PFLOTRAN, developed by Peter C. Lichtner of Los Alamos National Laboratory, simulates groundwater contamination flows.

Developed by ORNL's Tom Evans, Denovo allows fully consistent multi-step approaches to high-fidelity nuclear reactor simulations.

CAM-SE represents two models that work in conjunction to simulate global atmospheric conditions. CAM (Community Atmosphere Model) is a global atmosphere model for weather and climate research. HOMME, the High Order Method Modeling Environment, is an atmospheric dynamical core, which solves fluid and thermodynamic equations on resolved scales. In order for HOMME to be a useful tool for atmospheric scientists, it is necessary to couple this core to physics packages—such as CAM—regularly employed by the climate modeling community.

LAMMPS, the Large-scale Atomic/Molecular Massively Parallel Simulator, was developed by a group of researchers at SNL. It is a classical molecular dynamics code that can be used to model atoms or, more generically, as a parallel particle simulator at the atomic, meso, or continuum scales.

Once the applications were chosen, teams of experts were assembled to work on each. These teams included one liaison from the OLCF's Scientific Computing Group who is familiar with the code; a number of people from the applications' development teams; one or more representatives from hardware developer Cray; and one or more individuals from NVIDIA, which will be supplying the graphics processing units (GPUs) to be used in Titan.

"The main goal is to get these six codes ready so that when Titan hits the floor, researchers can effectively use the system," said Bronson Messer, an astrophysicist at ORNL. Using a single problem from varying domains, the teams have worked to identify parts of each code that are capable of being accelerated via GPUs.

Guiding principles

Before work began on these six codes, the development teams acknowledged some fundamental principles intended to optimize their adaptations. First, because these applications are current INCITE codes, they are under constant development and will continue to be after Titan is in production—any changes made must be flexible. Second, and perhaps most important, these applications

are used by research teams the world over on various platforms.

"We had to make sure we made changes to the codes that won't just die on the vine," said Messer. "We had to ensure that our changes at the very least do no harm while they are running on other, non-GPU platforms." The teams discovered that some of their modifications not only made the codes functional on hybrid systems, but actually helped performance on non-GPU architectures. A prime example is Denovo, which has experienced a twofold increase in performance speed on traditional central processing unit (CPU)-structured systems since being adapted.

"We've made changes that we're sure are going to remain within the 'production trunk' of all these codes—there won't be one version that runs on Titan and another version for traditional architectures," said Messer. It's essential for developers to be able to check a code out of a repository that can be compiled on any architecture; otherwise, the work done by these teams won't survive time.

The development teams have learned plenty from their work thus far. First, application adaptation for GPU architectures changes data structures (the way information is stored and organized in a computer so that it can be used efficiently), a fact that is creating the most difficult work for the teams. GPUs have to be fed information correctly. Like voracious animals—they can't be sated with little "bites" of information, but require huge amounts of data at once. Developers have carefully examined the codes line by line to identify large chunks of data that can be "fed" to these ravenous GPUs.

Second, the teams have learned how to marshal data for the GPUs. This requires the developers to know a lot about the hardware of the GPU to code effectively.

These two discoveries are the key to enabling strong scaling of applications at the exascale. "What we've discovered at the petascale is that research teams have run out of weak scaling," said Messer, referring to the label associated with increasing the size of the problem to be solved. Users have reached the point where they can make quantitative predictions with their codes and don't have to increase the problem size, or they've simply reached the limit of where the code base can take them.

Ultimately, said Messer, all of the applications running on Titan will be able to exploit the GPUs to achieve a level of simulation virtually unthinkable just a few years ago. Simply put, standing up and operating Titan will be America's first step toward the exascale and will cement its reputation as the world leader in supercomputing. With Titan, America can continue to solve the world's greatest scientific challenges one simulation at a time.