

De l'hôtel à projets au centre d'essais

Un hôtel, trois missions

Conférences

- organisation de *Workshops*, séminaires
- noeud CECAM en France
- Computational Physics* aux Houches

Formations

- internationales : Erasmus Mundus
- continues : toutes disciplines
- initiales : de L3 à M2

Projets

- Hébergement "physique"
- Gestion de projet : forge
- "Paillasse numériques"

Un centre d'essais, trois quêtes

Scalabilité

- traitement : du cluster au GPU
- transport : du série à l'OmniPath
- stockage : du zRam au GlusterFS

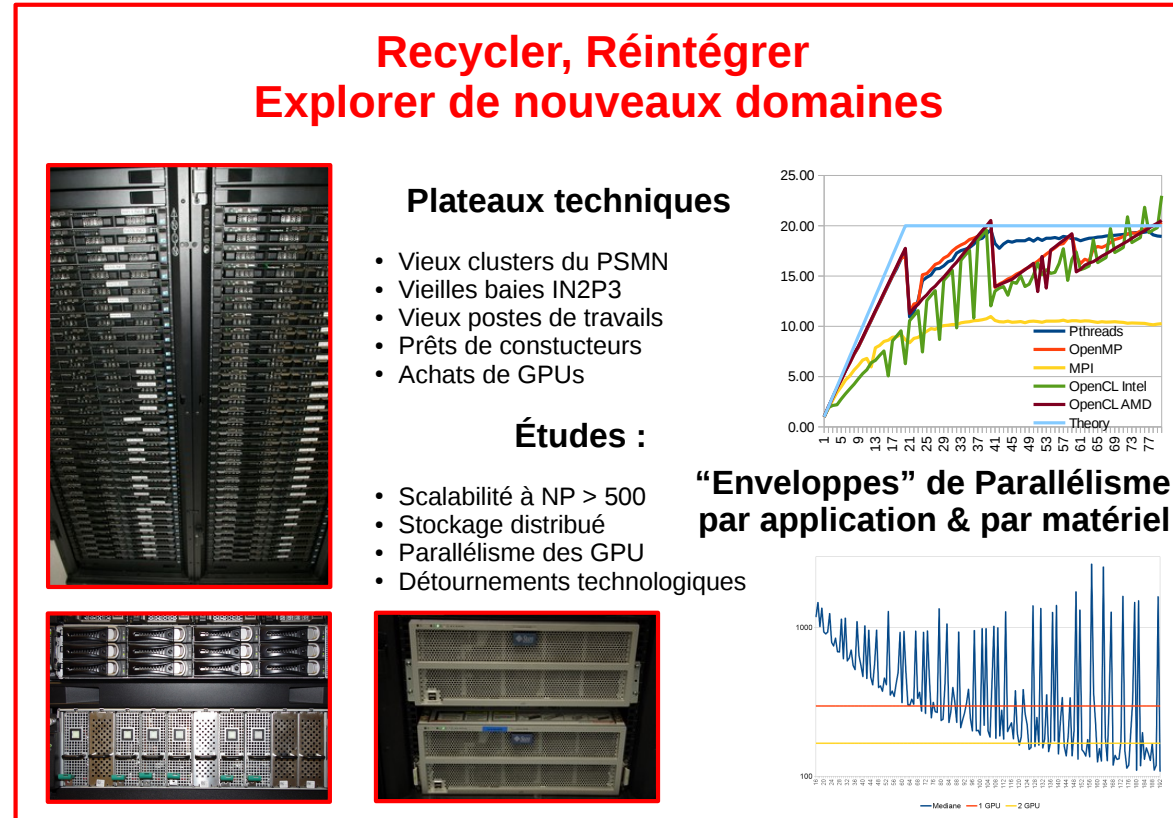
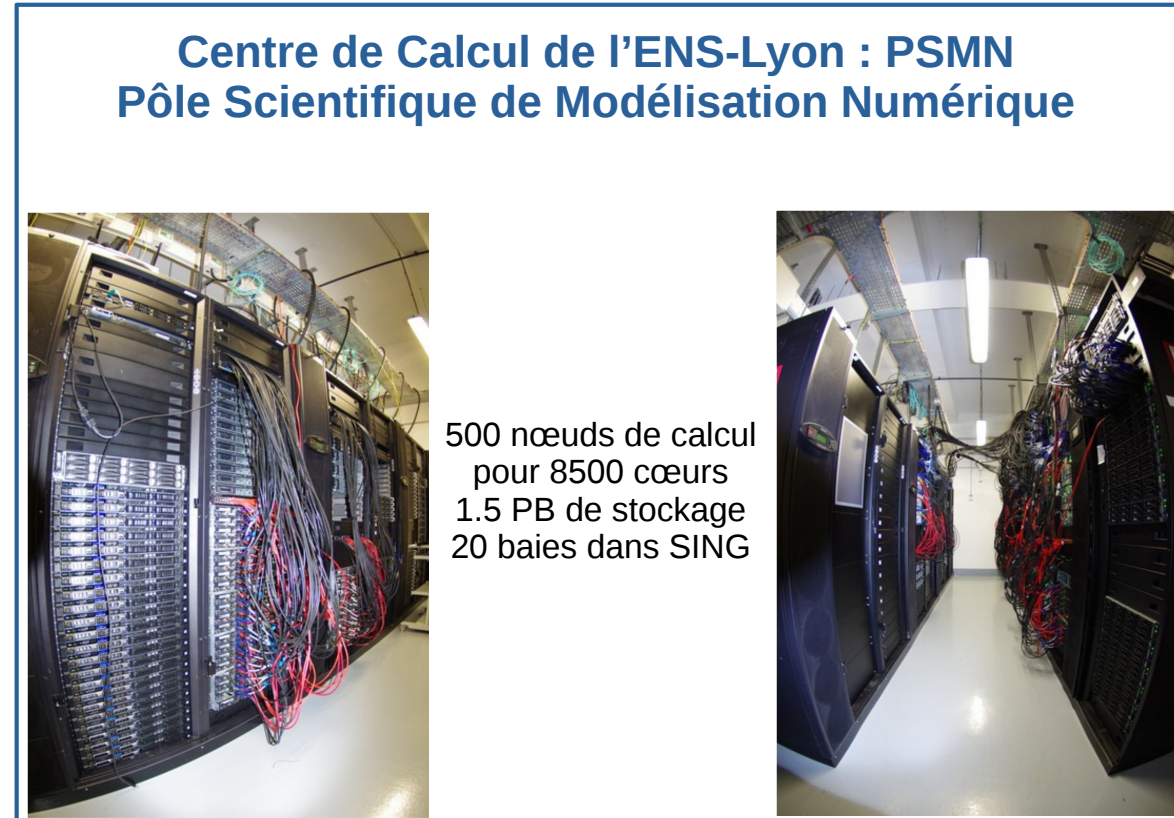
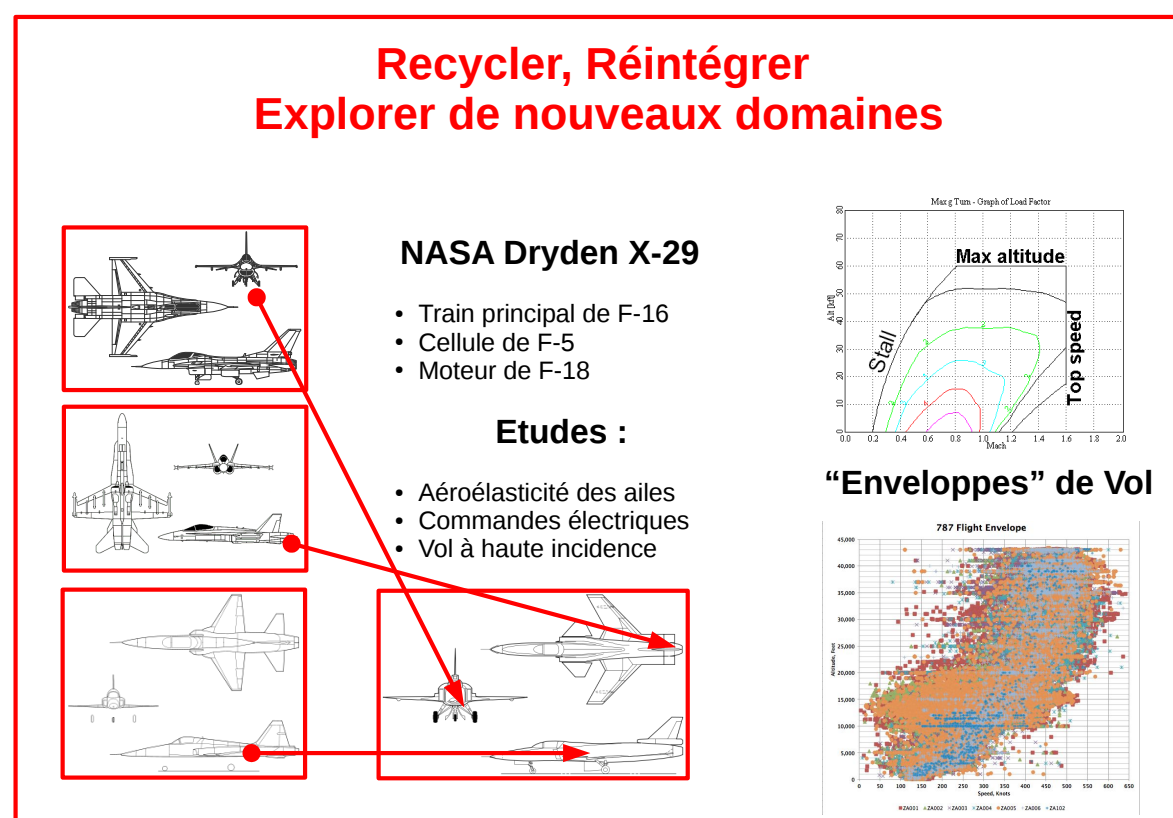
Reproductibilité

- dans l'espace : systèmes différents
- dans le temps : répétabilité
- avec définition de "systèmes socle"

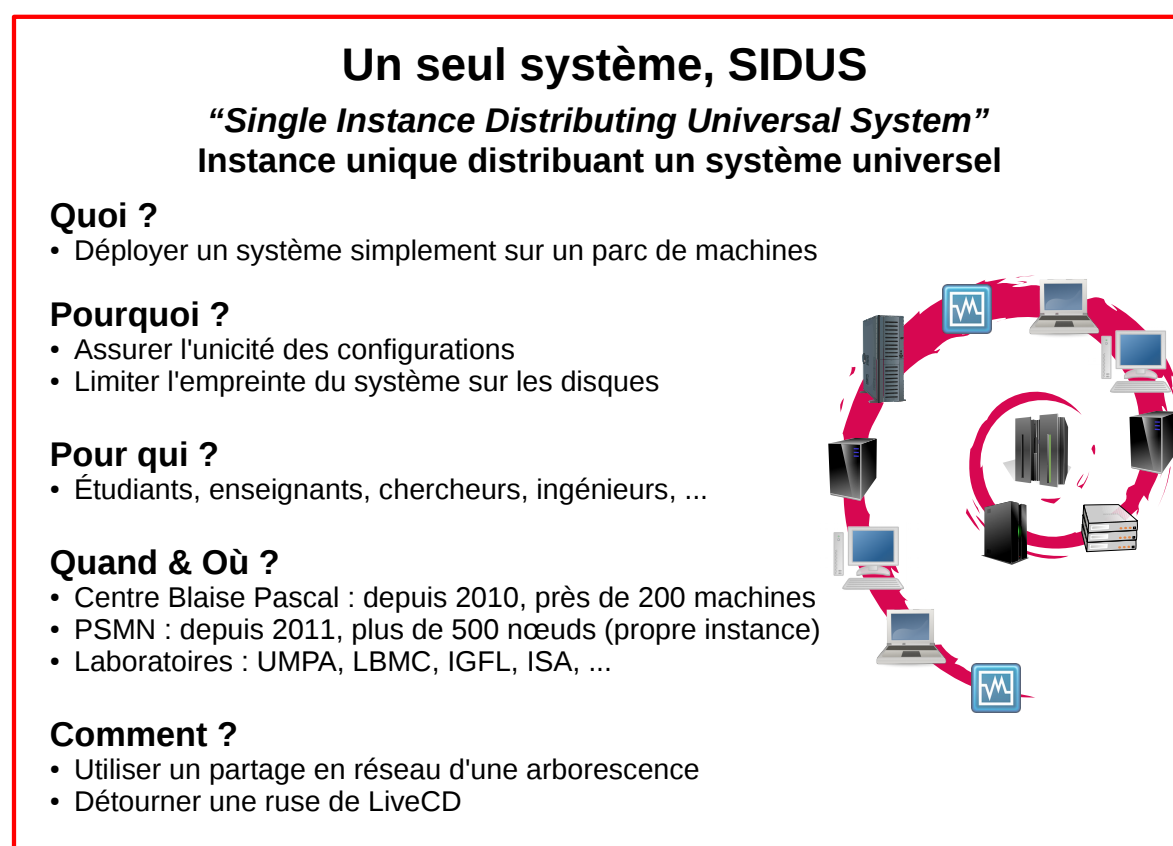
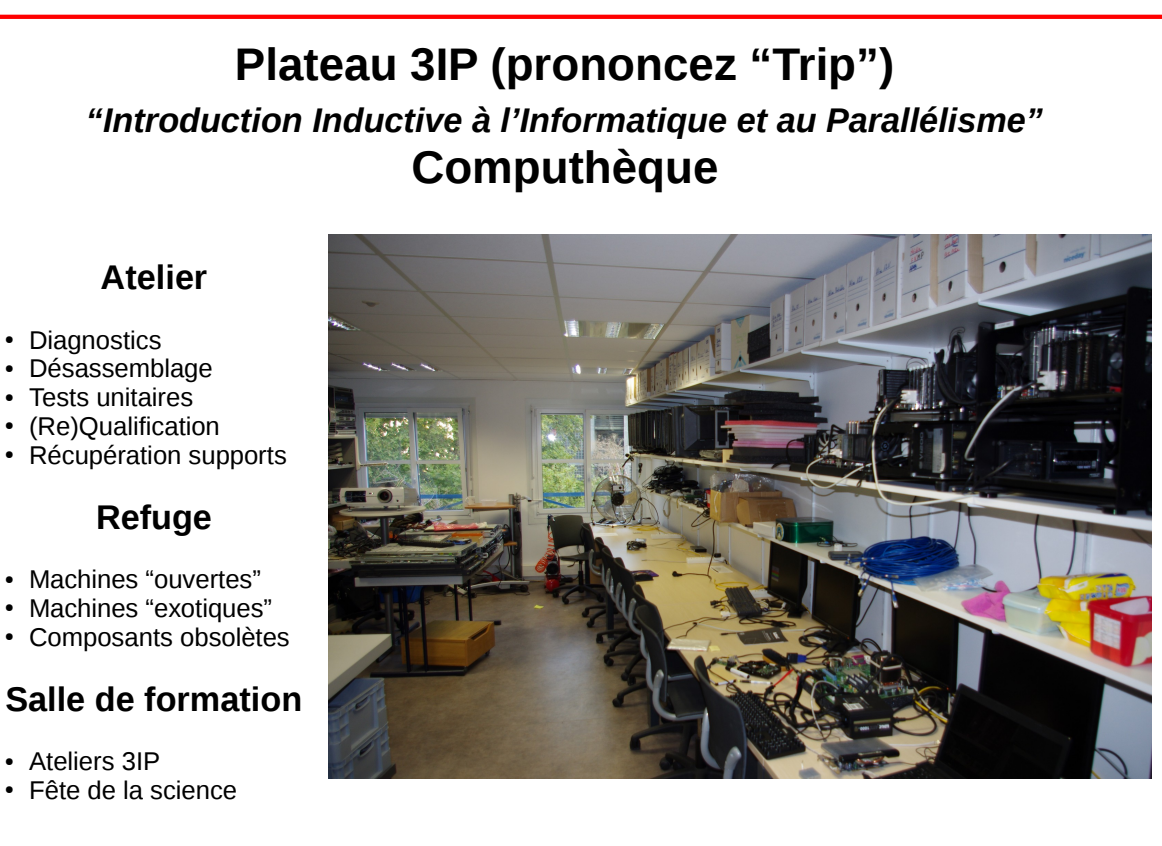
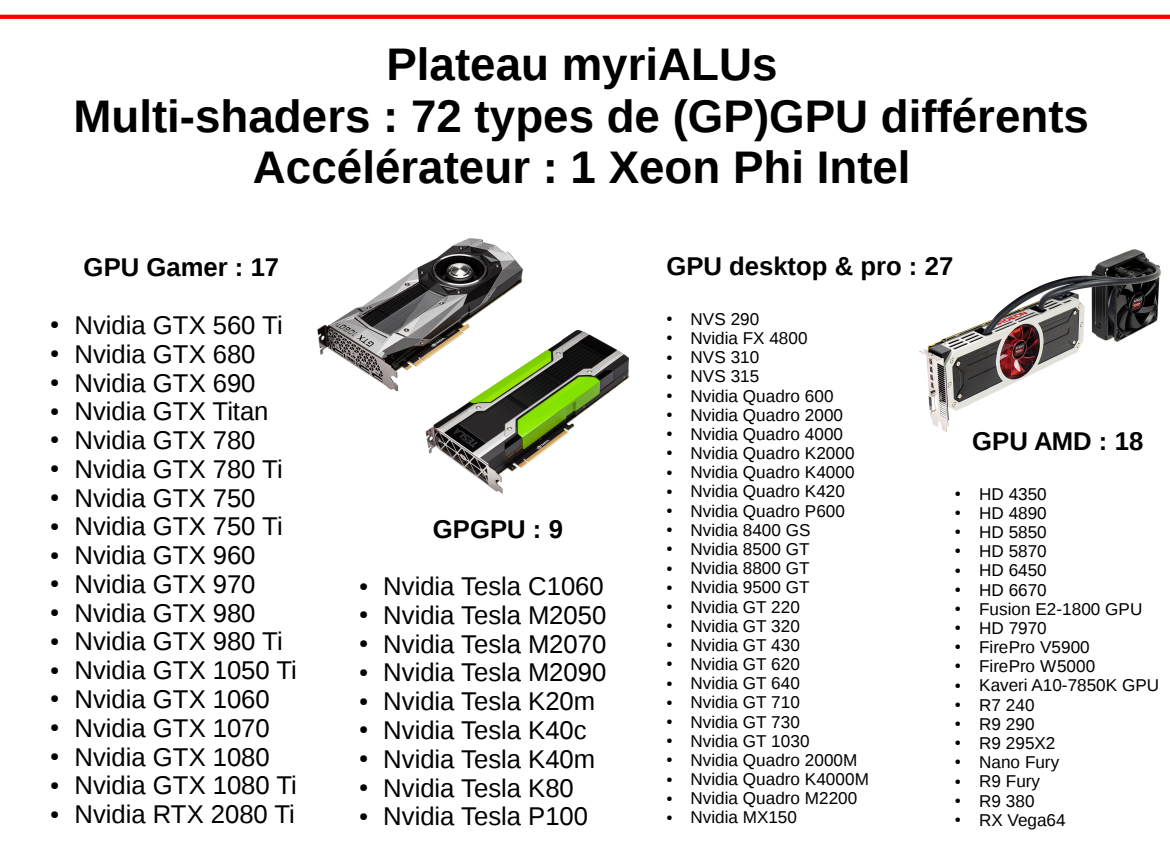
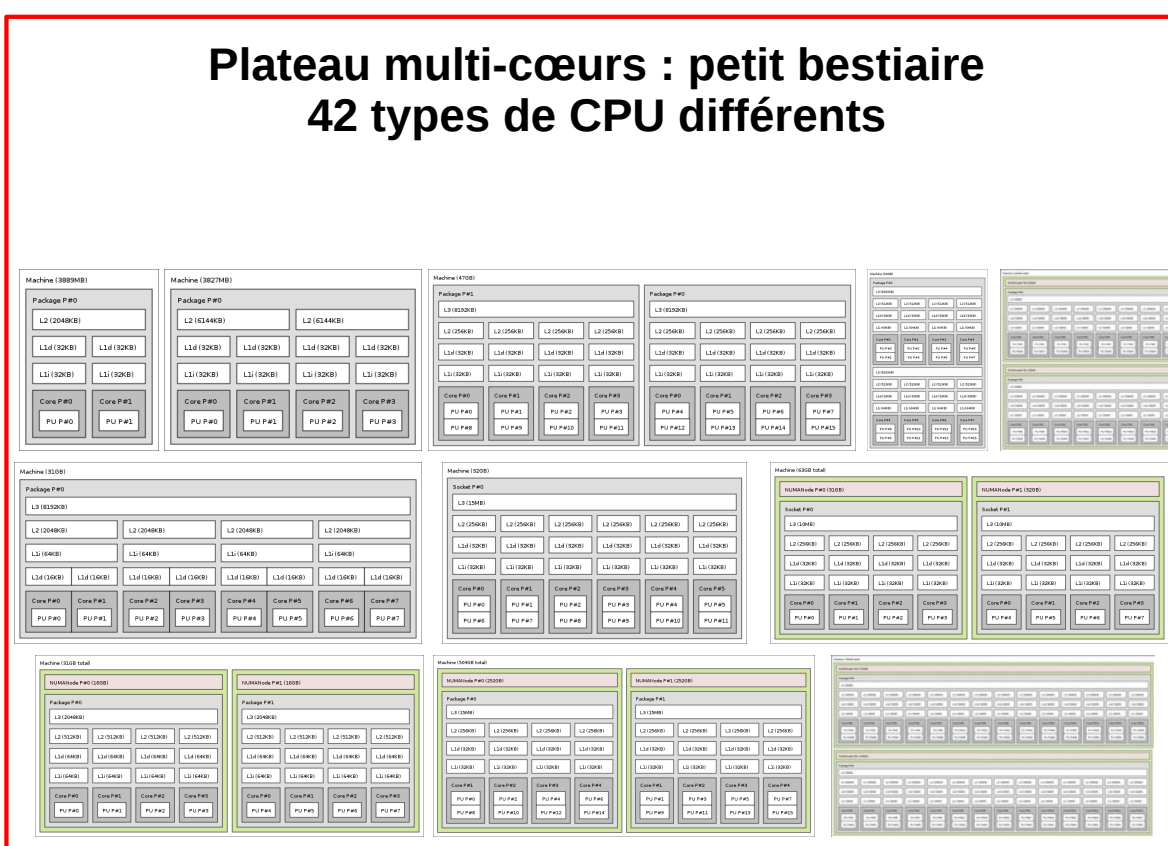
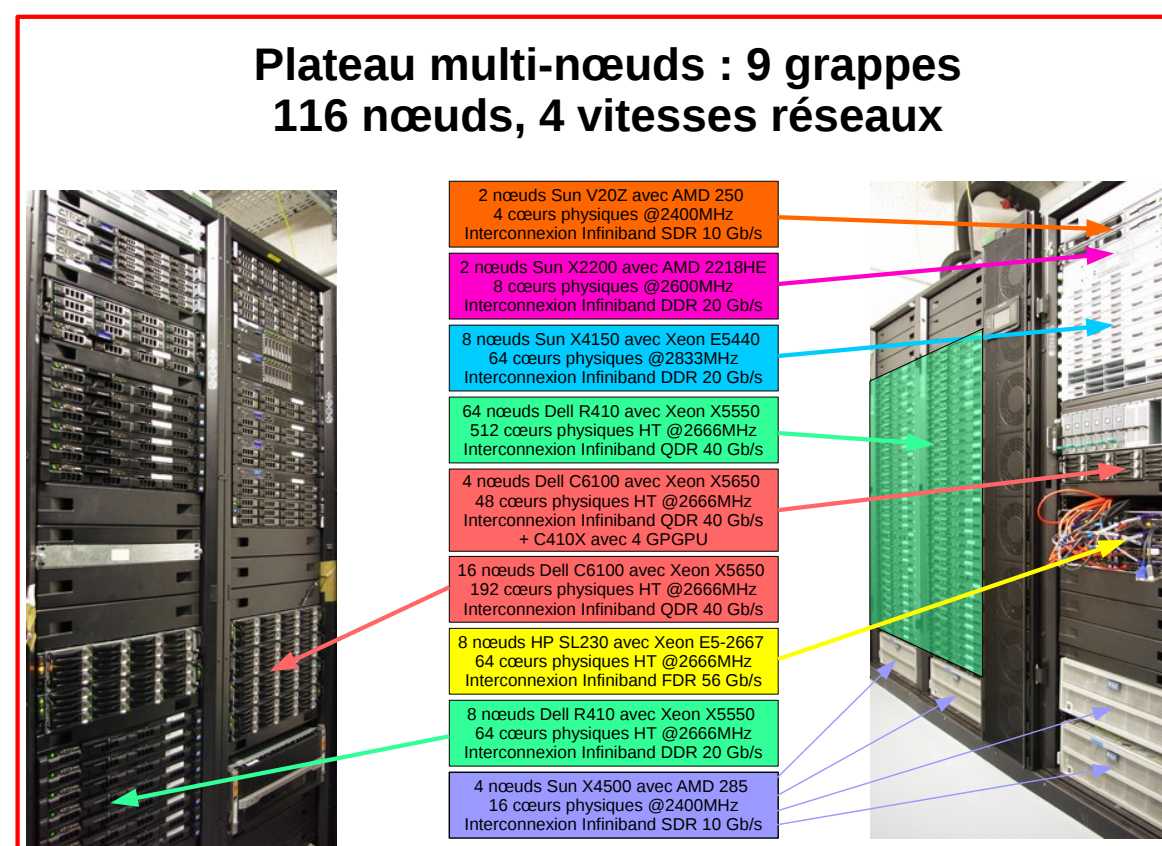
Simplicité

- de déploiement d'OS : SIDUS
- d'accès au stockage
- de traitement : HPDA@TheEdge

Quoi ? Petite analogie aéronautique



Comment ? Des plateaux techniques & un unique système : SIDUS

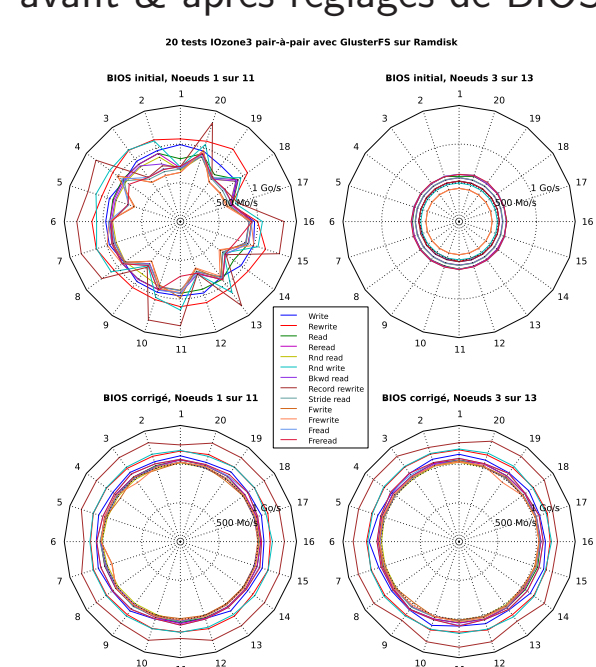


Pour Quoi ? Quelques exemples d'études

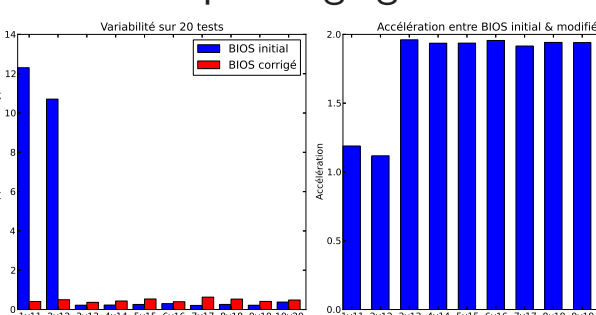
Étude : GlusterFS comme scratch en HPC

- L'objectif** : évaluer les performances de GlusterFS
- Le contexte** : Equip@Meso, quel scratch ?
- Le banc d'essai** : 20 noeuds, 2 réseaux, 1 OS, 1 test
 - 20 HP SL230 bi-SandyBridge avec 64 Go RAM
 - Réseaux Gigabit Ethernet & InfiniBand FDR
 - GlusterFS en IP over IB
 - Socle logiciel SIDUS
- Les expériences** : 20 tests, 10 clients sur 10 serveurs
 - Utilisation comme Ramdisk de TMPFS ou de BRD
 - Test de système de fichiers IOzone3
 - 10 couples client/serveur
- Les résultats** : faible performance ou forte variabilité
 - performance autour de 1 Go/s de transfert
 - forte variabilité, bonne performance sur 2 couples
 - faible variabilité, faible performance sur 8 couples
- Une interrogation** : quelles sources de variabilité ?
 - Noeuds identiques, interconnexions identiques
 - Même OS au bit près par SIDUS
 - Seule différence possible : le BIOS
- La conclusion** : le BIOS responsable !
 - réglages d'usine déléterent en performance (C-States)
 - mixage de réglages catastrophique en reproductibilité

IOzone3 sur 2 client/serveur avant & après réglages de BIOS



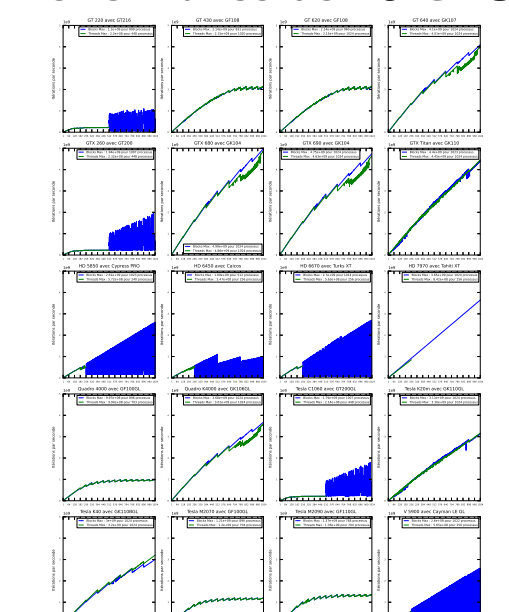
Variabilité & Performance pour les 10 client/serveur avant & après réglages de BIOS



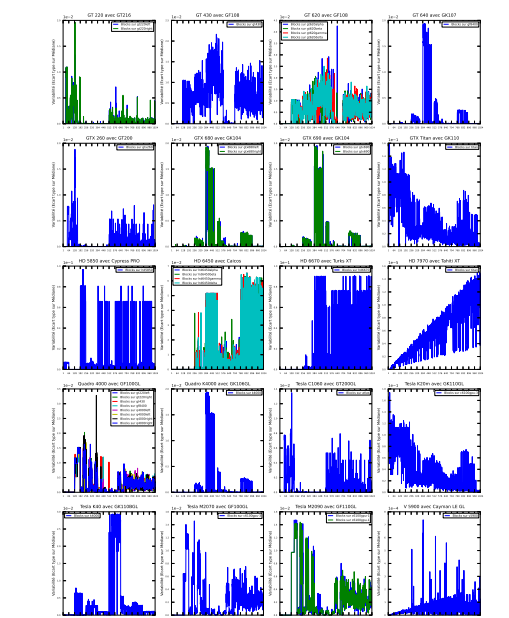
Étude : quel parallélisme massif des GPU ?

- L'objectif** : évaluer le comportement en parallélisme des GPU
- Le contexte** : contexte *many-cores*, comment comparer ?
 - Deux constructeurs "sérieux" : Nvidia & AMD
 - Deux langages : CUDA pour Nvidia, OpenCL pour tous
 - Deux parallélismes : *Blocks/WorkItems* & *Threads*
- Le banc d'essai** : 26 stations/noeuds, 20 GPU différents, 1 OS
 - Dell C6100, Precision T7500, T5600, T7600, Optiplex 745
 - De 1 à 2 GPU différents dans chaque hôte
 - 20 GPU de référence différente : Nvidia, AMD
 - Socle logiciel SIDUS
- Les expériences** : 1 OS, 1 langage, 1 API, 1 test
 - Langage OpenCL appelé en Python/Numpy
 - Programme test : Pi par Monte Carlo
 - Cœur de calcul comprenant un cycle de 32 opérations :
 - Parallélisation par distribution des itérations
- Une interrogation** : quelles sources de variabilité ?
 - Noeuds ou stations identiques
 - Même OS au bit près par SIDUS
- Les conclusions** : test simple mais instructif !
 - Performance : 1024 tâches faibles pour tester GPU récents
 - Comportement linéaire en parallélisme peu respecté
 - Comportement comparable pour des générations identiques
 - Variabilité : une signature aussi pertinente que la performance

Performance de 20 GPU



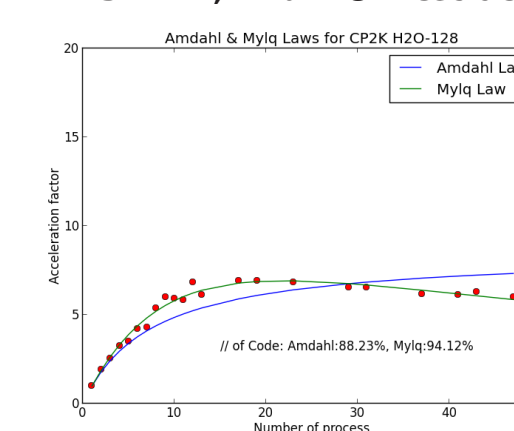
Variabilité de 20 GPU



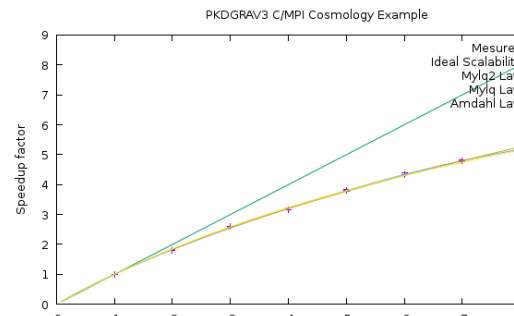
Étude : la loi d'Amdahl représentative ?

- L'objectif** : évaluer la pertinence de la loi d'Amdahl
- Le contexte** : codes hybrides OpenMP/MPI
 - 3 applications multi-noeuds, multi-cœurs
 - Deux parallélismes : MPI ou OpenMP/MPI
 - Loi d'Amdahl : $T_n = s + p/n$
 - Loi de Mylq : $T_n = s + p/n + cn + c_2n^2$
 - s, p, c, c_2 : séquentielle, parallèle, linéaire, quadratique
- Le banc d'essai** : 3 clusters, 1 OS
 - 48/64 Dell R410 avec 384/512 cœurs physiques
 - 8 Dell C6320 avec 128 cœurs physiques
 - Socle logiciel SIDUS
- Les expériences** : 1 OS, 3 clusters, 3 tests
 - chimie quantique : CP2K, cube 128 H2O
 - astrophysique : PKDGRAV3, 256³ particules
 - programme "test" : Pi Monte Carlo avec 10¹² itérations
- Une interrogation** : quelle scalabilité pour ces codes ?
 - CP2K : plus de gain au delà de 16 noeuds
 - PKDGRAV3 : tassement au delà de 6 noeuds
 - PiMC : baisse au delà de $N/P = 400$
- Les conclusions** :
 - loi d'Amdahl peu représentative des applications
 - loi de Mylq préférable, au moins au 1er ordre
 - scalabilité temporelle mitigée pour codes métier

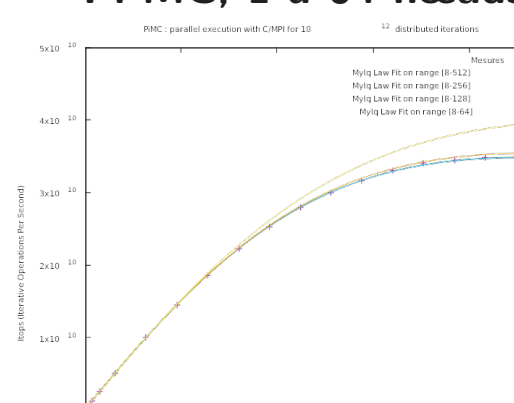
CP2K, 1 à 48 noeuds



PKDGRAV3, 1 à 8 noeuds



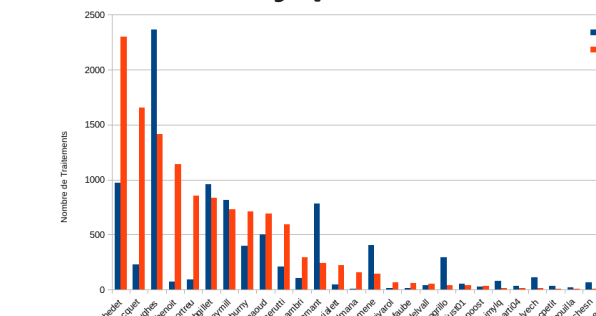
Pi MC, 1 à 64 noeuds



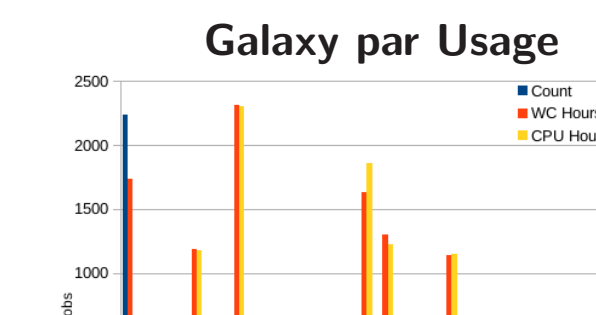
Prototypage : Portail Galaxy local

- L'objectif** : intégration d'un portail Galaxy au centre de calcul
- Le contexte** :
 - explosion du volume de données des biologistes
 - appropriation du centre de calcul difficile
 - standard émergent de portail de traitement
- Le banc d'essai** : entre réel & virtuel
 - 1 machine virtuelle pour le portail
 - 1 cluster de 8 noeuds pour les traitements
- Les étapes** :
 - version 1.0 2015Q2
 - version 2.0 2017Q3 : IB, volume, PostGreSQL
 - version 3.0 2019Q1 : stockage distribué, Shibboleth
- Les conclusions** :
 - développement en cycle court parfait
 - usages très différents entre laboratoires : pas de panacée...
 - intégration sur infrastructure "centre de calcul" difficile
 - métrologie fine indispensable pour évaluation des besoins

Galaxy par Utilisateur



Galaxy par Usage



Applications ≠, empreintes système ≠

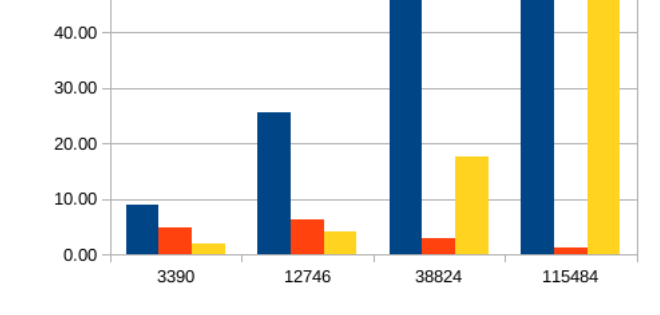
Étude : Repeat* = "crash applications" ?

- L'objectif** : quelle infrastructure pour traiter les données ?
- Le contexte** :
 - workflow trop lourd sur une station de travail
 - portage sur PSMN entraînant une coupure du service
 - investigations poussées sur équipements du CBP
- Le banc d'essai** : déploiement sur plateformes différents
 - déploiement sur station de travail avec SSD
 - déploiement sur noeud de cluster avec GlusterFS
 - déploiement sur serveur avec Ramdisk
- Les étapes** :
 - octobre 2013 : station + SSD
 - septembre 2015 : noeud de calcul + GlusterFS
 - octobre 2015 : serveur "musclé" avec un Ramdisk
- Les conclusions** : une infrastructure nécessaire
 - des traitements générant des millions de petits fichiers
 - un SSD "mort" à la fin de la première campagne
 - un GlusterFS inadapté (comme tout système partagé)
 - une machine à dédier pour le traitement

Exécution d'une même tâche



Stockage GlusterFS ou Ramdisk

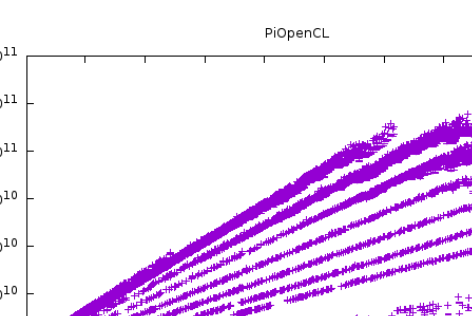


Plus le calcul avance, plus c'est lent sur le GlusterFS

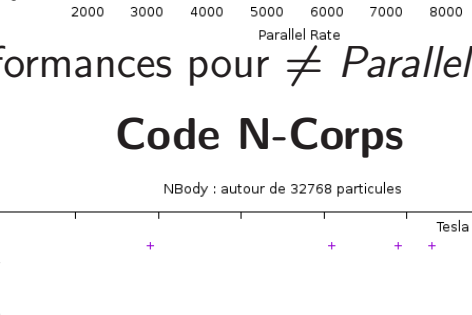
Caractérisation continue : tous les (GP)GPU

- L'objectif** : quel "régime de parallélisme" (PR) optimal ?
- Le contexte** :
 - une nouvelle génération de GPU tous les 2 ans
 - un nombre de cœurs passant de 240 à 4096 en 8 ans
 - une fréquence allant de 500 MHz à 1500 MHz
 - des unités de traitement : *cuda-cores* ou *shader-processors*
- Le banc d'essai** : 4 codes de tests
 - xGEMM* : le "classique" de BLAS (toutes implémentations)
 - PiXPU.py ou PiOpenCL : ALU-bound, gros grain
 - NBody.py : MemoryBound, grain fin
 - Splutter.py : MemoryBound, fonctions atomiques
 - xTSRV* : un assemblage de BLAS (toutes implémentations)
- Les bizarreries** :
 - optimums de performance : PR = 16x nombre de shaders
 - chutes de performance sur Nvidia : PR = nombre premier
- Les conclusions** :
 - Nvidia GTX intéressantes mais uniquement pour FP32
 - choix du régime de parallélisme crucial
 - AMD intéressantes mais pénibles à faire fonctionner

Code Pi Monte Carlo



Performances pour ≠ Parallel Rates



Code N-Corps

