

Architectures de Processeurs pour la Mobilité

S'ARMer, s'ATOMiser,
FUSIONner...

Du mobile au HPC
Le prochain « détournement » ?

« Détourne-toi des préceptes de ceux qui spéculent sur le monde mais
dont les raisons ne sont pas confirmées par l'expérience. »
de Léonard de Vinci

Centre Blaise Pascal

« catalyser » l'informatique scientifique



Nasa X29

- Cellule de F5
- Moteur de F18
- Servos de F16
- Études
 - Flèche inversée
 - Incidence $>50^\circ$
 - « *Fly-By-Wire* »

Maison de la Simulation ET centre d'essais

Centre Blaise Pascal

Une Plate-Forme Expérimentale

Avec des plateaux techniques:

- Cluster de 8 à 50 nœuds
- Nœuds de 2 à 12 cœurs
- Stations GPU de la GeForce 8400 à la Tesla M2090
- Intégration logicielle (versions de distribution)
- Exploration : expérience/démonstrateur/prototype
- Intégration matérielle (Sparc, PowerPC, ... et ARM)
- Plateau COMOD : « Compute On My Own Device »

Pour étudiants/enseignants/chercheurs/ingénieurs

L'informatique Mobile & HPC

Le « Grand Détournement »

- Avril 2010, café « Aramis » sur les GPU
 - **Détournement** des puces graphiques pour le calcul
- Novembre 2011, JRES à Toulouse
 - **Détournement** de stockage local pour le distribuer
- Juin 2012, journée « Aramis » sur la virtualisation
 - **Conclusion** : « émergence des puces ARM en HPC »
- Avril 2013 :
 - Quelle pertinence à me présenter à vous ?
 - Puces de portables vers HPC, prochain détournement ?
 - **Et l'histoire, un perpétuel recommencement ?**

Le « grand détournement » : d'un GPU à CPU faible TDP

- **Pour le GPU :**
 - « Gros » multiplicateur de matrice
 - Multiplication d'unités simples (plus spécialisées) ~ 1000
- **Pour le CPU à faible TDP**
 - Consommation basse
 - Fréquence adaptative
 - Multiplication de « cœurs » (généraliste & spécialisé)
- **Exploitable comme :**
 - Serveur élémentaire ?
 - Socle pour Sidus ?

L'actualité récente, à la une « Le Monde Informatique »

- **5 mars 2013** : *Les serveurs ARM de HP bientôt équipés de puces Texas Instruments*
 - HP prévoit de commercialiser ces serveurs à base de puces ARM au deuxième trimestre de cette année.
- **2 avril 2013** : *Dell travaille sur des prototypes de superordinateurs ARM*
 - « Les processeurs ARM peuvent permettre de faire des économies budgétaires par FLOP et par rack, et certaines institutions n'auront pas de craintes à utiliser un supercalculateur construit autour de processeurs ARM »
- **8 avril 2013** : *Intel dévoile Avoton, sa prochaine puce pour serveur Atom*
 - la dernière puce Atom pour serveur d'Intel, Avoton, optimise les performances et l'efficacité énergétique.
- **9 avril 2013** : *HP lève le voile sur les serveurs du projet Moonshot*
 - Le constructeur affirme ainsi qu'un système Moonshot permet d'économiser 89% d'énergie, 80% d'espace et de réduire ses coûts de 77%



Deux justifications

- La convergence ?

- Même socle technique
 - Poste/Station de travail
 - Machine virtuelle locale
 - Nœud de Cluster
- **Solution** : Sidus
- Quid architectures ?
 - ARM
 - Intel Atom
 - AMD Fusion

Côté « utilisateur »

- La rationalisation ?

- Ratio prix CPU/nœud
- Multiplication cœurs
 - S2 : 2, S4 : 8
- Recours Accélérateur
 - TPM, GPU
- Gestion énergie
- Paradigme du service :
 - 1 service
 - 1 cœur

Côté « constructeur »

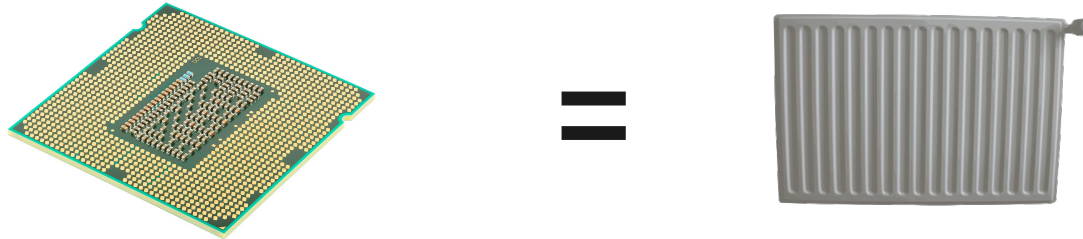
SIDUS : déduplication extrême d'environnement de travail

- *Single Instance Distributing Universal System*
 - Boot PXE
 - Exploitation des ressources locales, stockage distant
- Universalité :
 - Pour des nœuds de cluster
 - Pour des postes de travail (salle libre service)
 - Pour des machines « on demand »
 - (Pour l'exploration de machines compromises)
- Socle nécessaire :
 - DHCP/PXE/TFTP/NFS

Un facteur discriminant : la TDP

Thermal Design Power

- **TDP** : puissance à fournir pour évacuer la chaleur
 - Puissance à dissiper : $TDP \sim C V^2 f$
 - C : Capacité ($\sim 1/\text{Résistance}$), V : tension, f : fréquence
- Pour un physicien : 100W sur 4 cm²



- Refroidir (ou ne pas « trop » consommer) :
 - Une Nécessité...
 - Sinon...



Comparaison entre architectures

Les points communs

- Des processeurs « RISC »
 - Eh oui ! Intel depuis le P6 (1995), AMD depuis le K5
 - Nombre de registres étendus
 - Mécanismes de *pipeline*
- Des fréquences « comparables » : jusqu'à 1.8 GHz
- Des nombres de cœurs « comparables » : de 2 à 8
- Pour la majorité, « *System On Chip* » : CPU+GPU
 - NorthBridge souvent obsolète (au sens propre ;-)
 - SouthBridge (PCI, Sata, IDE, Ethernet, USB, Bios, ...)

Architecture ARM

- Une vieille idée, trop en avance sur son temps...

- Acorn Computer, 1987 : Acorn Risc Machine

- Des Architectures « v » :

- La « 7 » la plus répandue, la « 8 » apporte le 64 bits

- Des familles « Cortex »

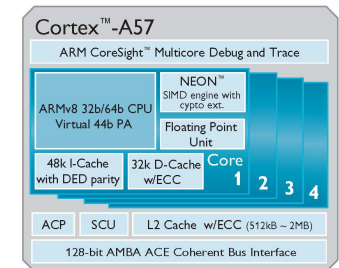
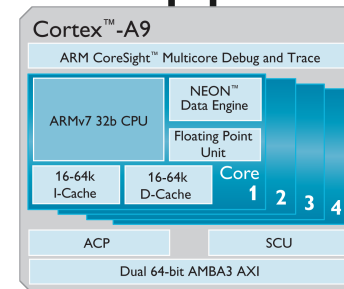
- La « A » la plus répandue (A9)

- Des constructeurs : 32 19/9/4 Asie/USA/Europe

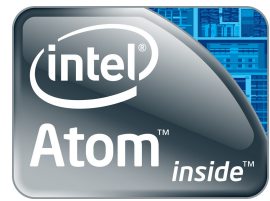
- Samsung (Exynos, Apple A6), Qualcomm, TI

- Des OS : 6 dont Linux (via Android/Ubuntu/Debian)

- Association au « Mali » pour le GPU



Architecture Intel Atom



- Architecture RISC (Bonnell)
- Des «familles» Z, E, D, N et S (après M et A)
 - D : Desktop, N : Netbook, E : Embedded, S : Server
 - Z : les autres (tablettes, etc...)
- Des constructeurs : bien, comment dire... un seul !
- Des OS : tous ceux des x86 (et x86_64)
- Des circuits graphiques, euh, comment dire... :(
 - Association au Poulsbo dans les premiers
- Une saturation du marché...

Architecture AMD Fusion



- Rachat ATI par AMD en juillet 2006
- Fusion : une nouvelle idée ?
 - APU : Accelerated Power Unit
 - SoC : Radeon HD + North Bridge
- Des « séries » A, E2, Z
 - A : APU (en fait Desktop), E : Embedded, Z : ?
- Des constructeurs : GlobalFoundries, vendu !!!
- Projection : FPU du GPU comme FPU primaire
 - Déjà le PileDriver partagé...

Comparaison entre architectures

Les différences

« Famille » Serveur

- 32/64 depuis 2006
- L3/L2/L1: xM/256k/32k
- Pas de SoC !

« Famille » ARMv7

- 64 uniquement ARMv8
- L3/L2/L1 : 0/64k/64k
- PowerVR, GeForce ULP

« Famille » ATOM

- 64 ? Ca dépend !
- L3/L2/L1 : 0/512k/32k
- Poulsbo & co...

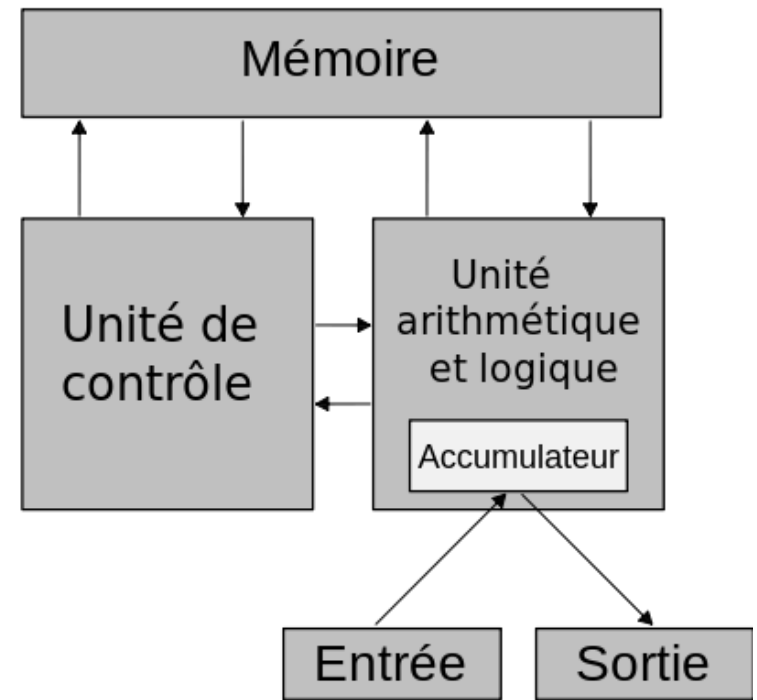
« Famille » Fusion

- 32/64 depuis le début !
- L3/L2/L1 : 0/512k/32k
- Vrai GPU : Radeon HD

Processeurs embarqués

Le retour de Von Neumann

- Le CPU, « vraiment » UC
 - Processeurs spécialisés
 - Contrôleur mémoire
 - Graphique : GPU
 - Entrée/Sortie : NB intégré
 - Une architecture simplifiée
 - Pas de VT-x/VT-d
 - Pas de cache L3
- Un « System On Chip »
 - SouthBridge intégré parfois



L'informatique embarquée

Architecture de test



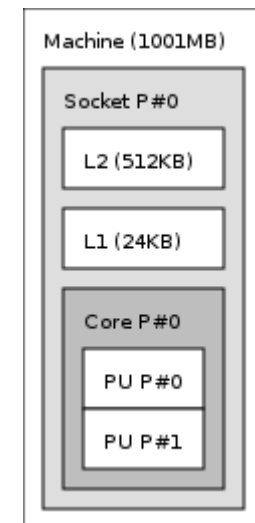
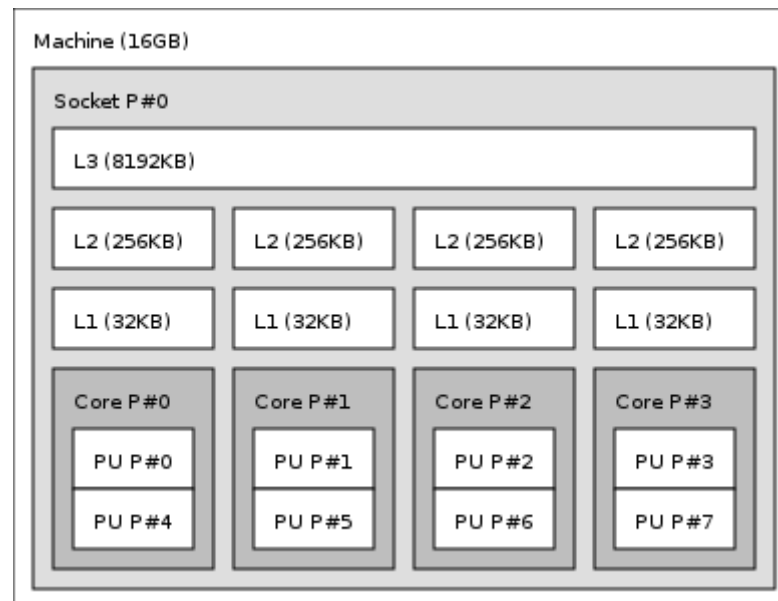
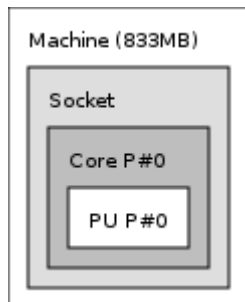
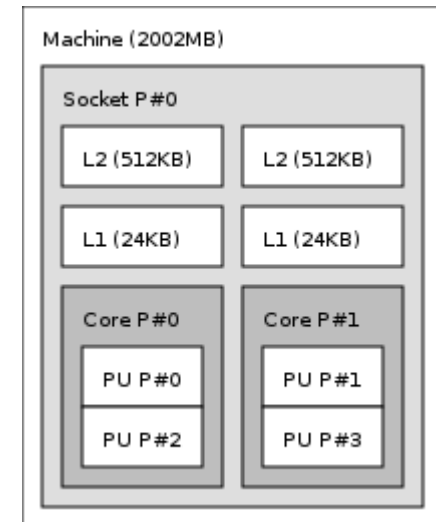
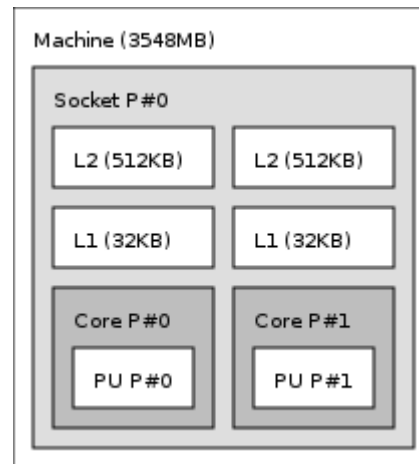
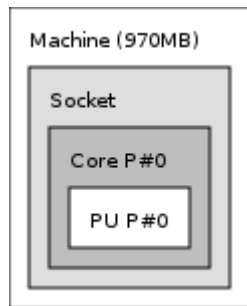
CBP / HPC : Dryden / NASA

Quelles architectures de test ?

- **Pour les ARM** : Tablette tactile & Smartphone
 - Asus TF700 : Tegra 3 avec Cortex A9
 - Samsung Galaxy S3 : Exynos 4412 avec Cortex A9
- **Pour les Intel Atom** : Tablette tactile, Netbook
 - Dell Mini 1010 : Atom Z530 avec un Pouslbo
 - Clevo : Atom N550 (une « vraie » bonne surprise)
 - Acer Iconia W510 : Atom z2760 & Windows 8 :'-()
- **Pour le Fusion** : Portable « jetable »
 - Asus x55u : Fusion E2-1800 (avec Radeon HD7340)

L'informatique embarquée

Un « petit » hwloc-ls



Comment les comparer ?

Même OS ? Même Kernel ?

- **Même Kernel ?**

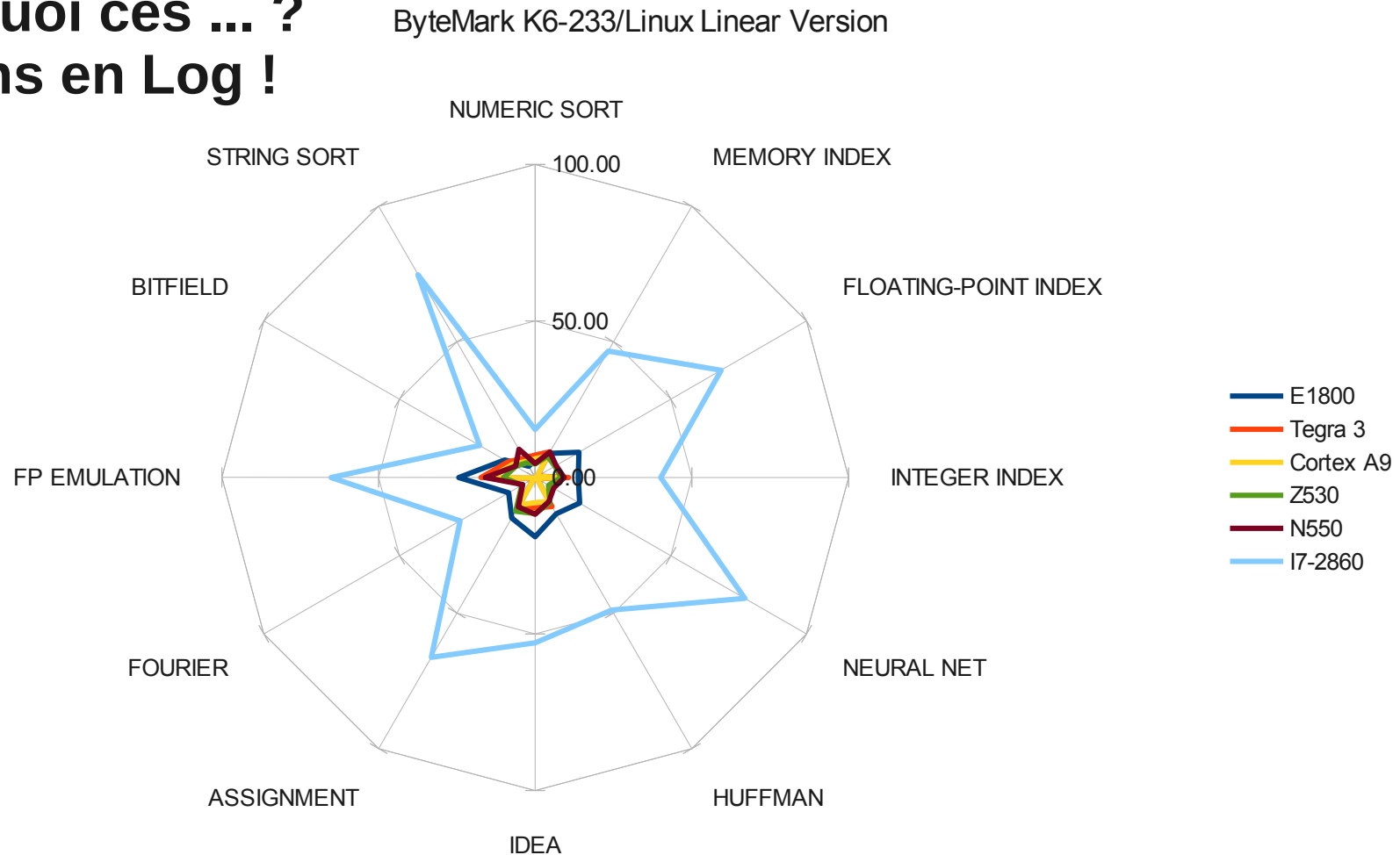
- **ARM** : tous sur Android 4 ICS ou JB
- **Atom** : Iconia irréductible à Windows 8
- **Fusion** : un vrai chemin de croix dans le BIOS !

- **Même OS ?**

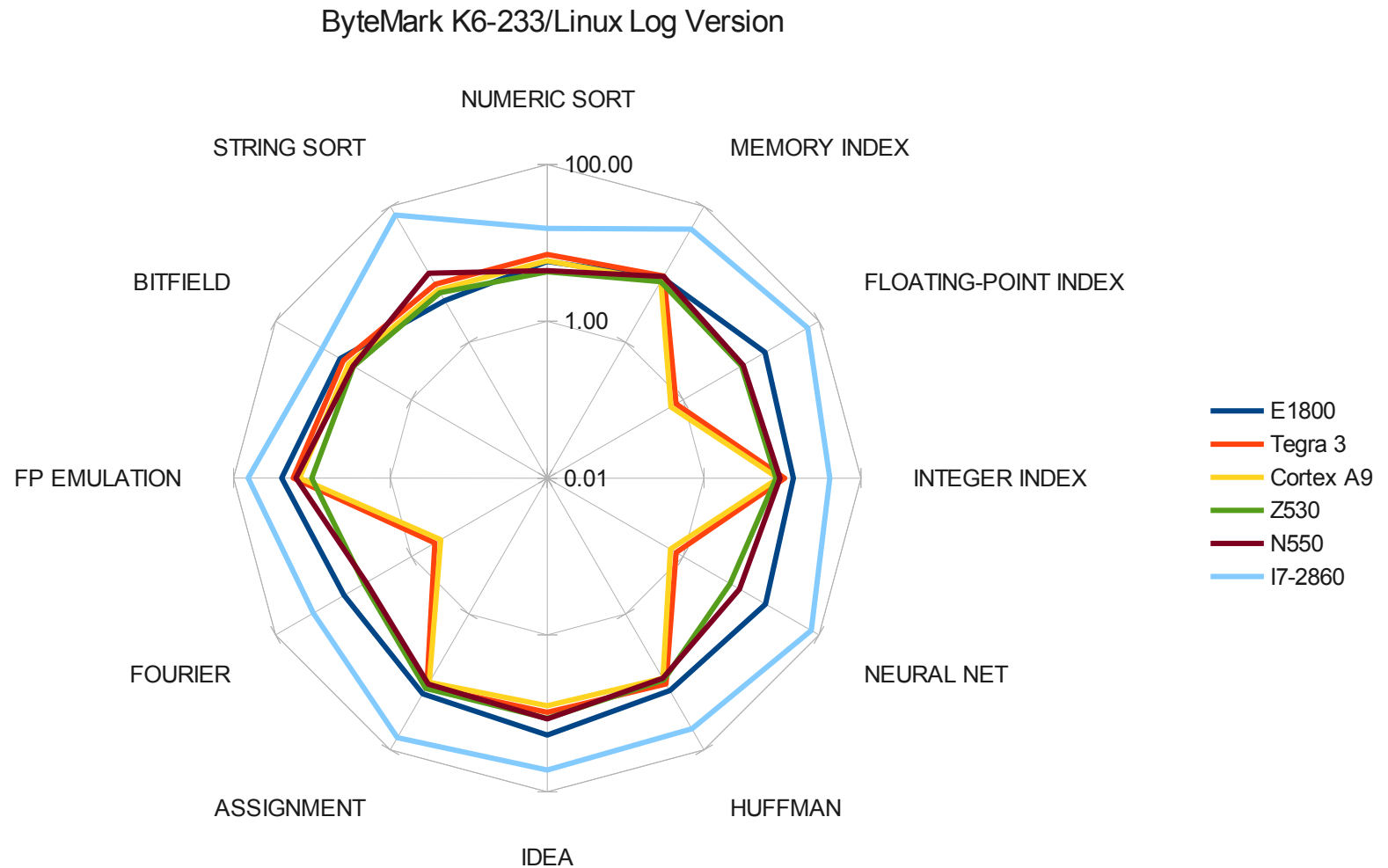
- **ARM** : CyanogenMod 10.1 avec chroot Debian
- **Atom** : variable !
 - *No problemo* sur Clevo M1110 & N550 (Sidus 32 ou 64)
 - *No problemo* sur Dell Mini1010 & Z530 (Sidus 32)
 - Impossible sur Acer Iconia W510 & Z2760
- **Fusion** : un peu galère (UEFI & Secure Boot)

Un « vieux » test des nineties ByteMark devenu NBench

**C'est quoi ces ... ?
Passons en Log !**

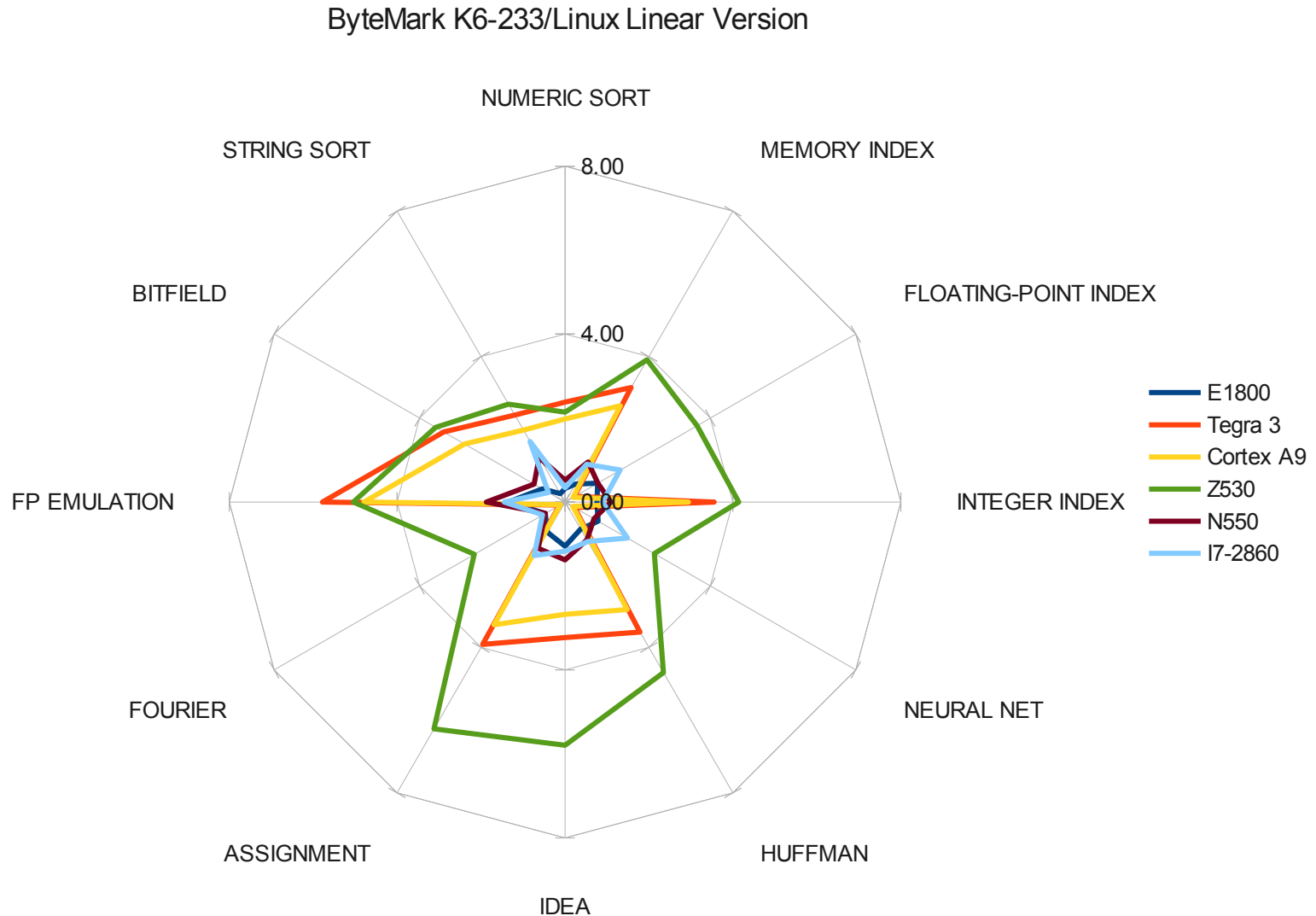


Nbench pour les 6 architectures : Passage en log pour le désastre



Nbench rapporté au TDP

Des éléments s'affinent...



Quelles conclusions « Nbench » ?

- Pour les architectures ARM
 - Un comportement honorable pour les tests entiers
 - Des performances indigentes en flottant
- Pour les architectures Atom
 - La TDP « plombe » les séries N 64 bits
 - La série Z530, malgré son âge, est la plus équilibrée
- Pour les architectures Fusion
 - Une performance moyenne pour un portable
 - La TDP anéantit son classement

Un test universel pour le //isme

Le « Monte Carlo Pi »

- Pertinent ? Pi est irrationnel, et pourtant
- Facile : exemple très simple à implémenter
- Très « optimisable » par les compilateurs
- Facilement parallélisable : pThreads, OpenMP, MPI
- Implémentation GPU (Cuda & OpenCL)
- Quel intérêt ?
 - Test : mix entre calculs entier/flottant/test/boucle
 - Capacité à exploiter les cœurs périphériques

Un test universel pour le //isme

Le « Monte Carlo Pi »

Cœur de Calcul

```

#define znew ((z=36969*(z&65535)+(z>>16))<<16)
#define wnew ((w=18000*(w&65535)+(w>>16))&65535)
#define MWC (znew+wnew)
#define MWCfp MWC * 2.328306435454494e-10f
  
```

```

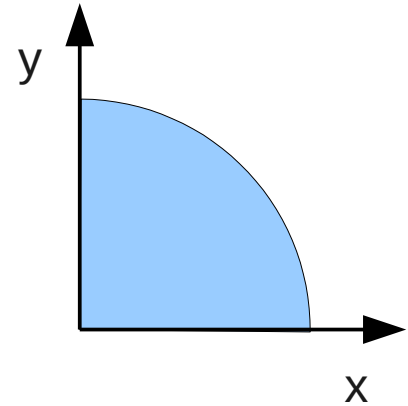
LENGTH MainLoopGlobal(LENGTH iterations,unsigned
int seed_w,unsigned int seed_z)
{
    unsigned int z=seed_z;
    unsigned int w=seed_w;
    unsigned long total=0;
    for (LENGTH i=0;i<iterations;i++) {
        float x=MWCfp ;
        float y=MWCfp ;
        // Matching test
        int inside=((x*x+y*y) < 1.0f) ? 1:0;
        total+=inside;
    }
    return(total);
}
  
```

Programme principal

```

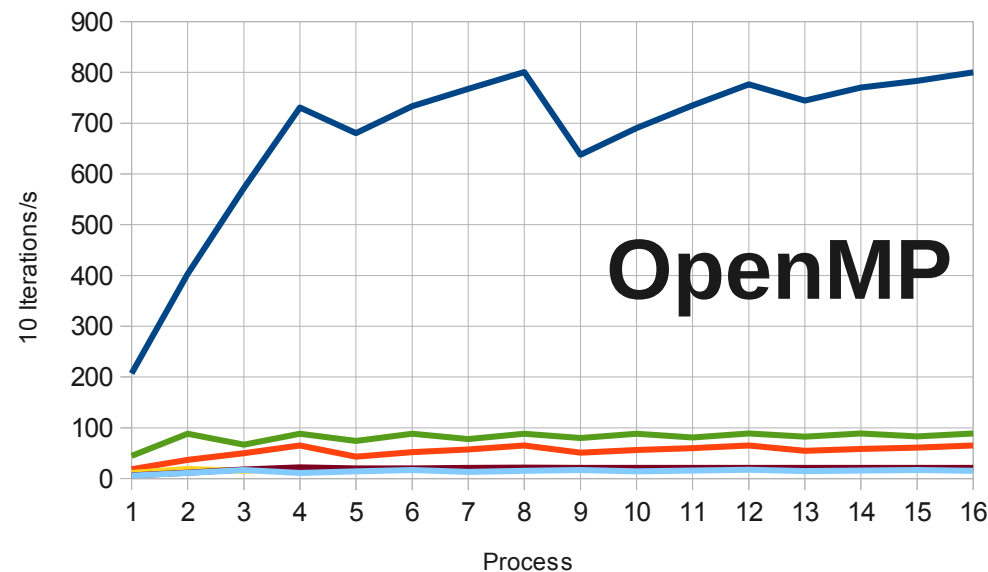
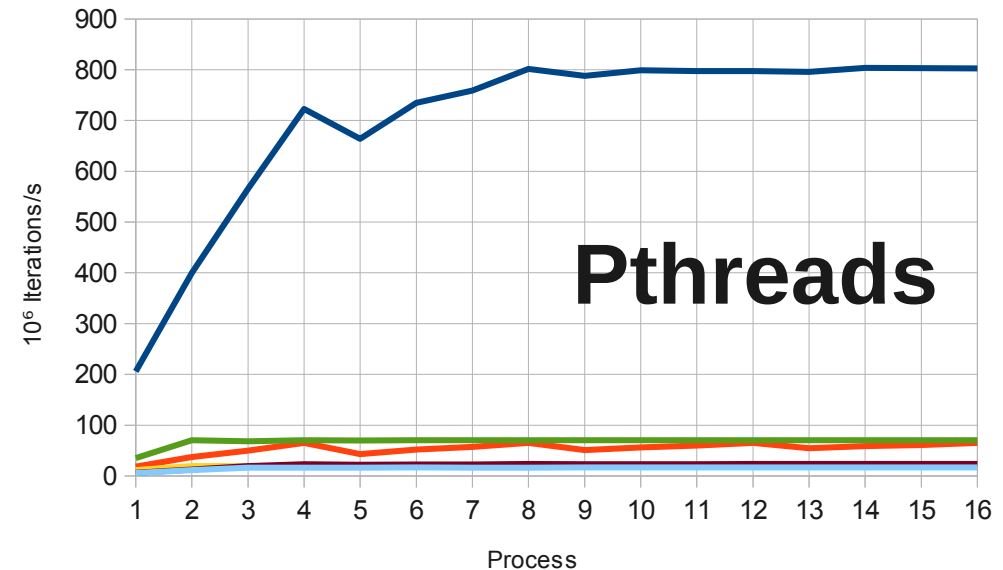
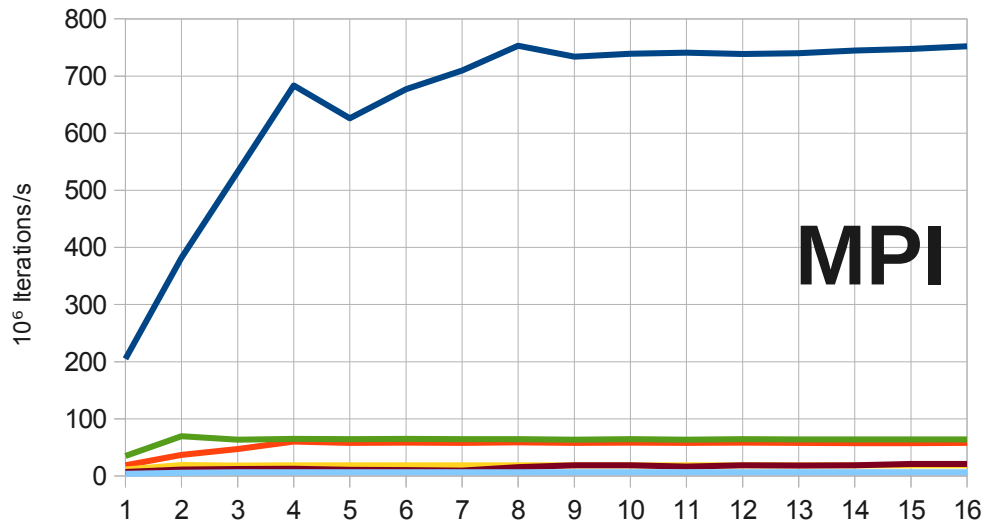
#include <math.h>
#include <stdio.h>
#include <stdlib.h>

int main(int argc, char *argv[]) {
    unsigned int seed_w=10,seed_z=10;
    LENGTH iterations=ITERATIONS;
    if (argc > 1) {
        iterations=(LENGTH)atol(argv[1]);
    }
    else {
        printf("\n\tPi : Estimate Pi with Monte Carlo exploration\n\n\t\t#1 :
number of iterations (default 1 billion)\n\n");
    }
    float pi=(float)MainLoopGlobal(iterations,seed_w,seed_z)/
(float)iterations*4;
    printf("\tPi=%f with error %f and %lu iterations\n\n",pi,
        fabs(pi-4*atan(1))/pi,(unsigned long)iterations);
}
  
```



Comparaison PiBench //isme

Des résultats... à analyser...

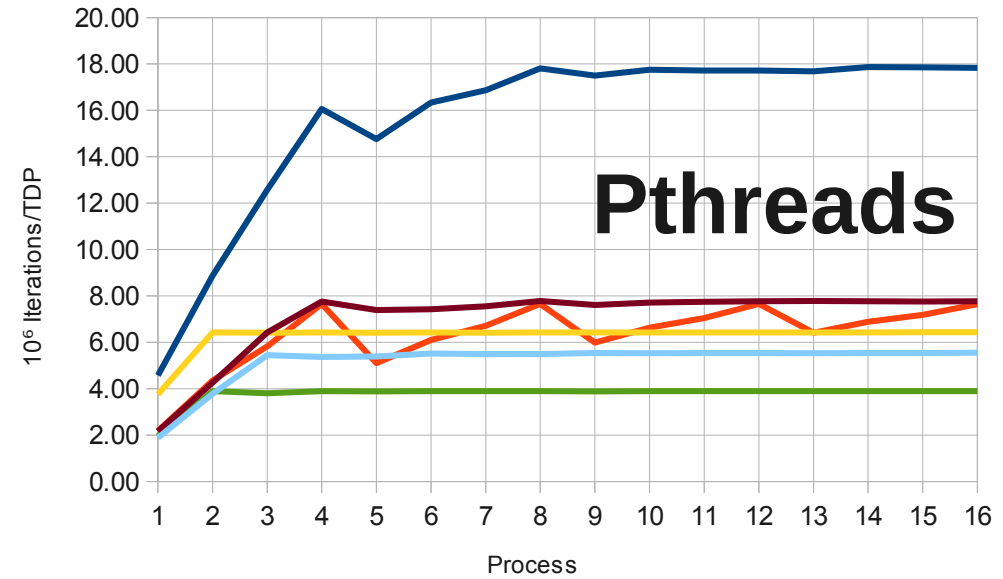
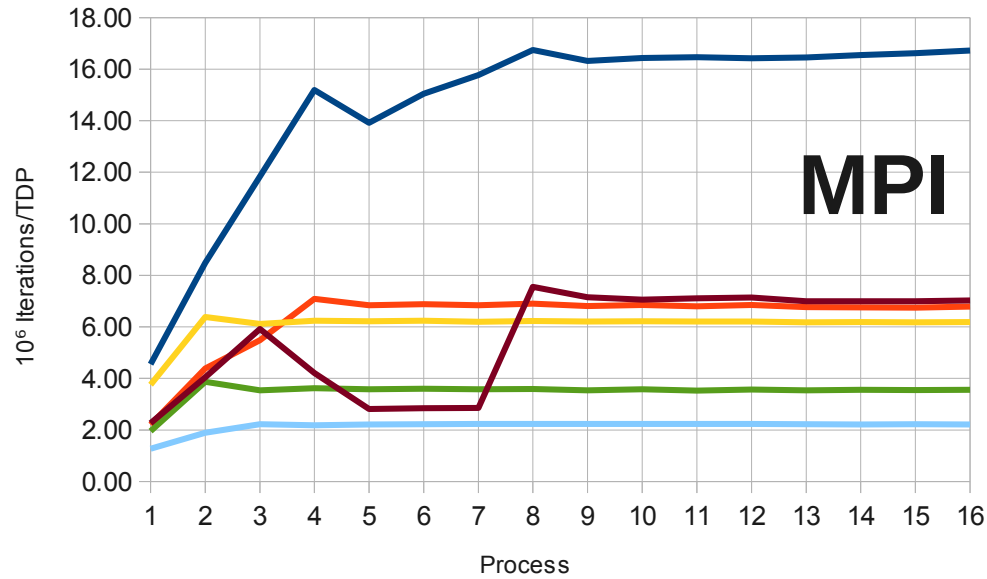


i7-2860
 N550
 Z530
 E2-1800
 Tegra3
 Cortex A9

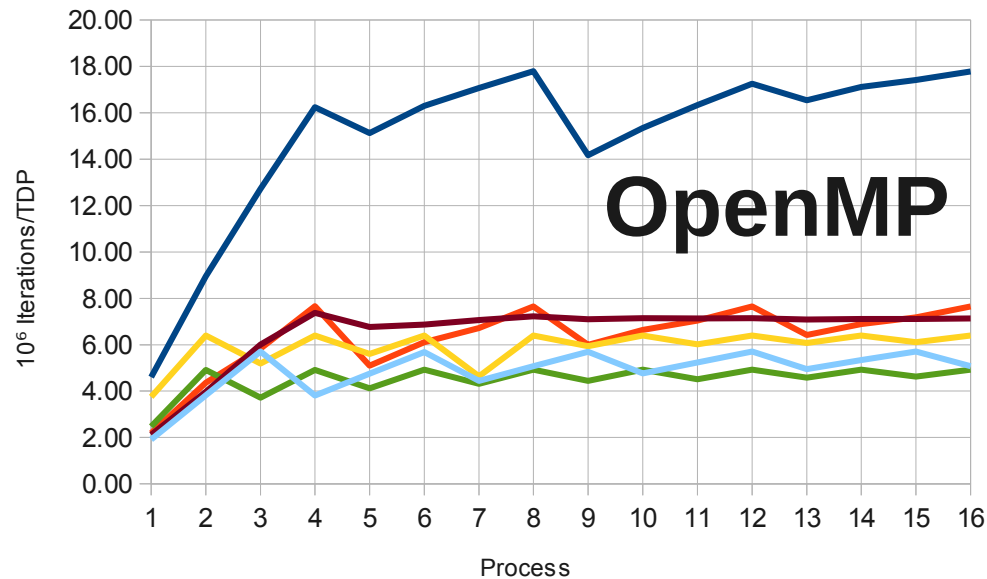
Intel i7
 Atom 64
 Atom 32
 Fusion
 Tegra 3
 Exynos 4412

Comparaison PiBench //isme

Retour à la consommation



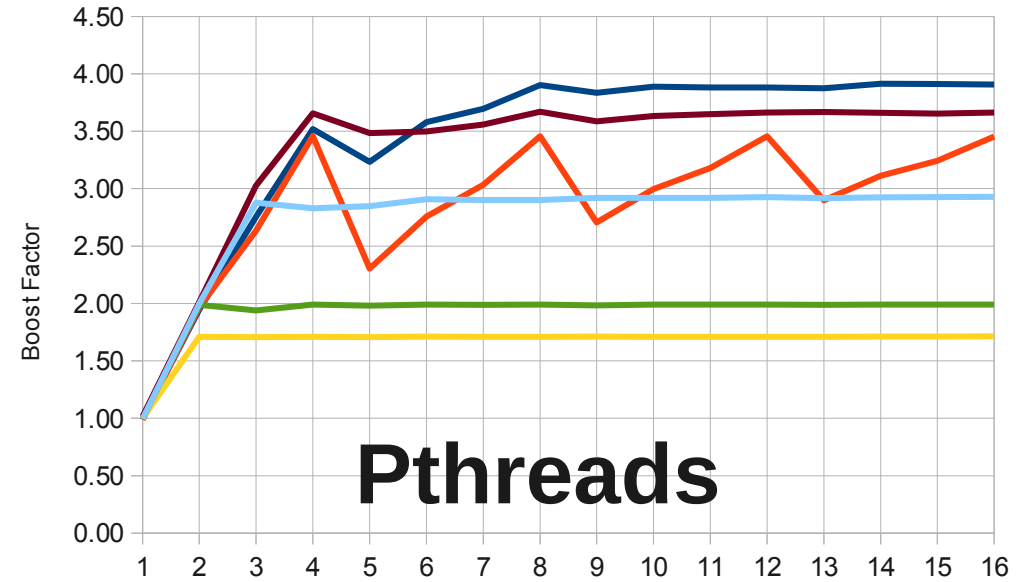
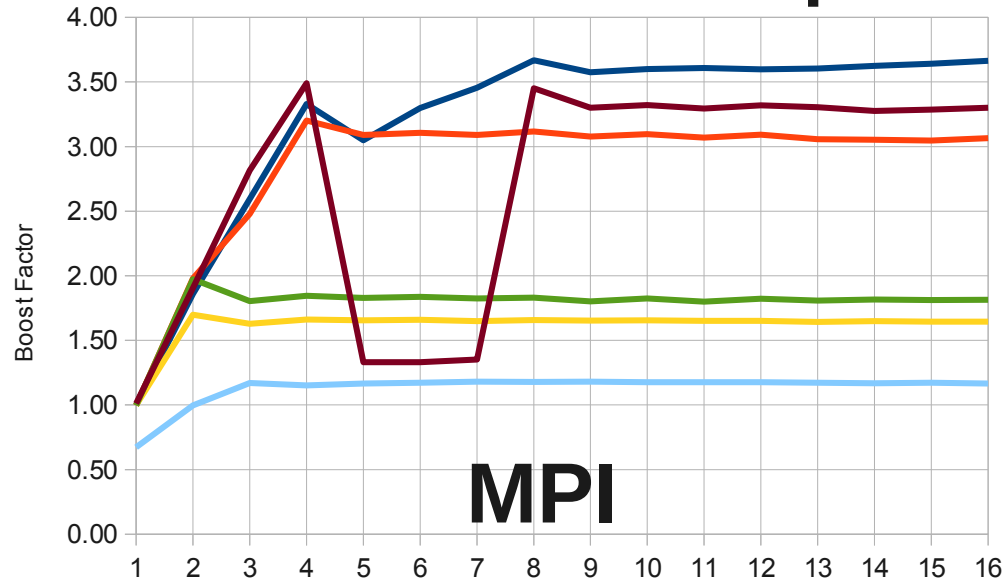
■ i7-2860
 ■ N550
 ■ Z530
 ■ E2-1800
 ■ Tegra3
 ■ Cortex A9



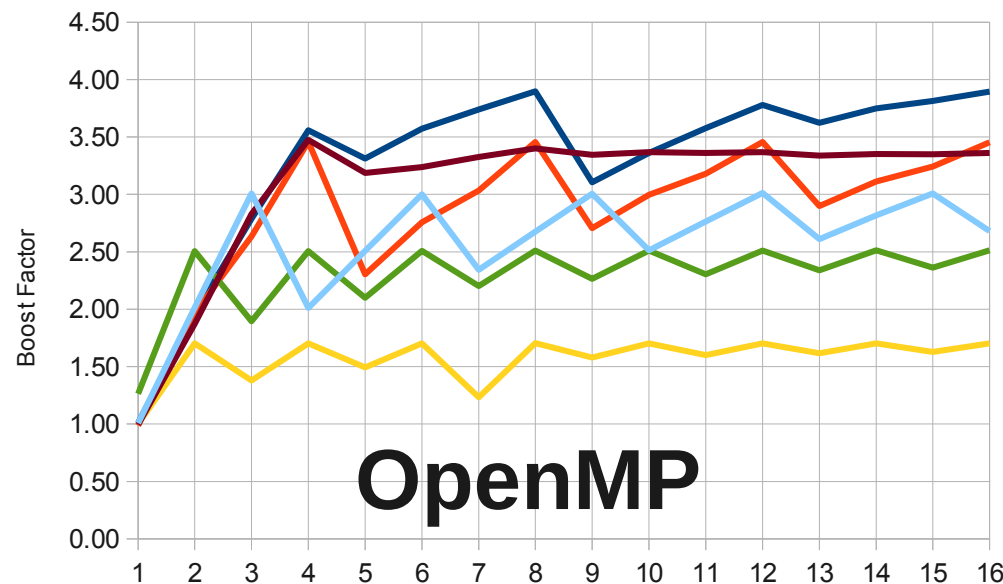
■ Intel i7
 ■ Atom 64
 ■ Atom 32
 ■ Fusion
 ■ Tegra 3
 ■ Exynos 4412

Comparaison PiBench //isme

Des comportements étonnants ?



■ i7-2860
 ■ N550
 ■ Z530
 ■ E2-1800
 ■ Tegra3
 ■ Cortex A9

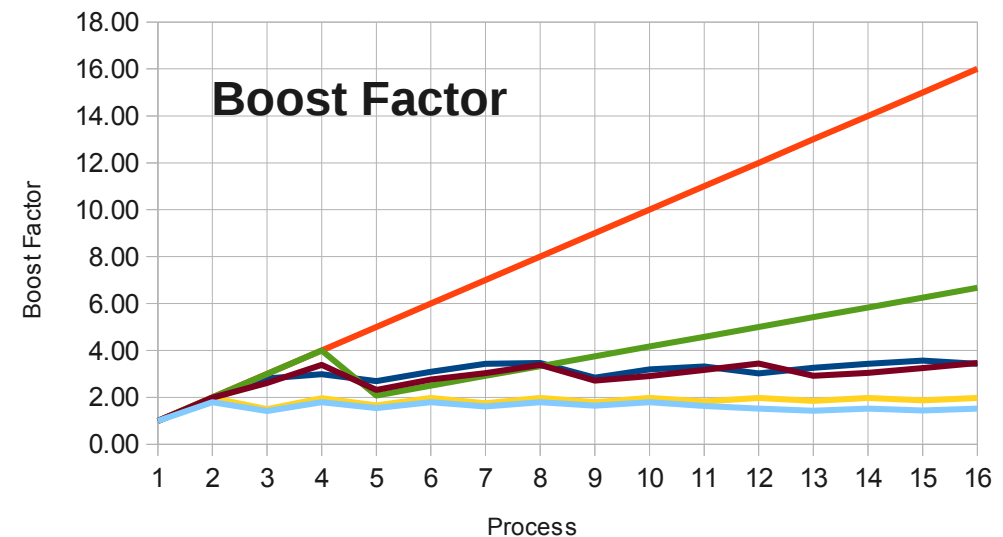
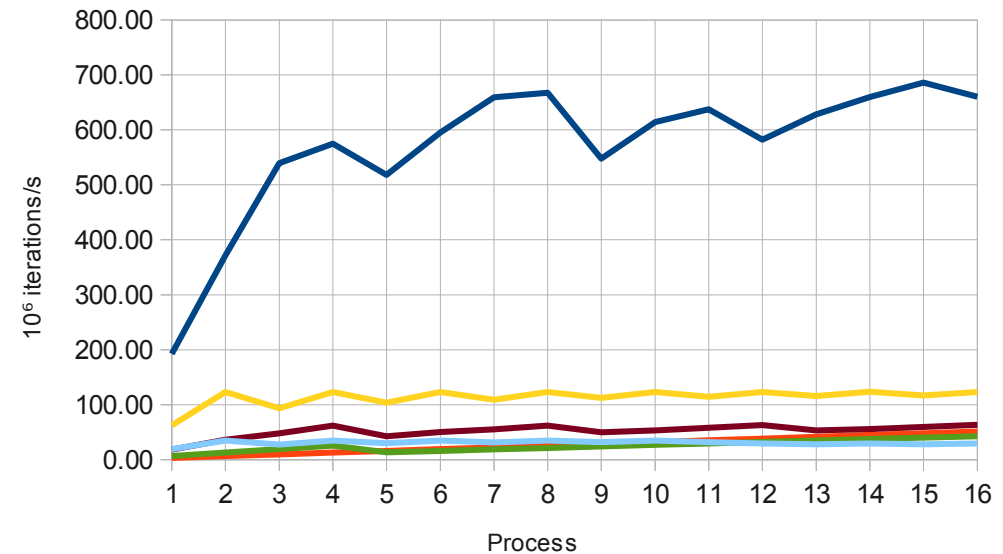
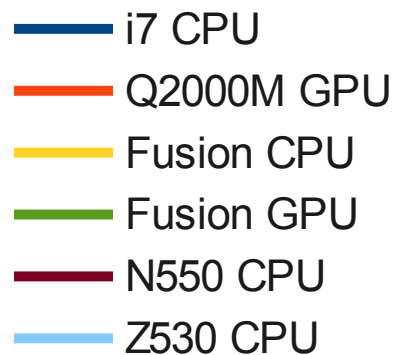


■ Intel i7
 ■ Atom 64
 ■ Atom 32
 ■ Fusion
 ■ Tegra 3
 ■ Exynos 4412

Et en OpenCL ?

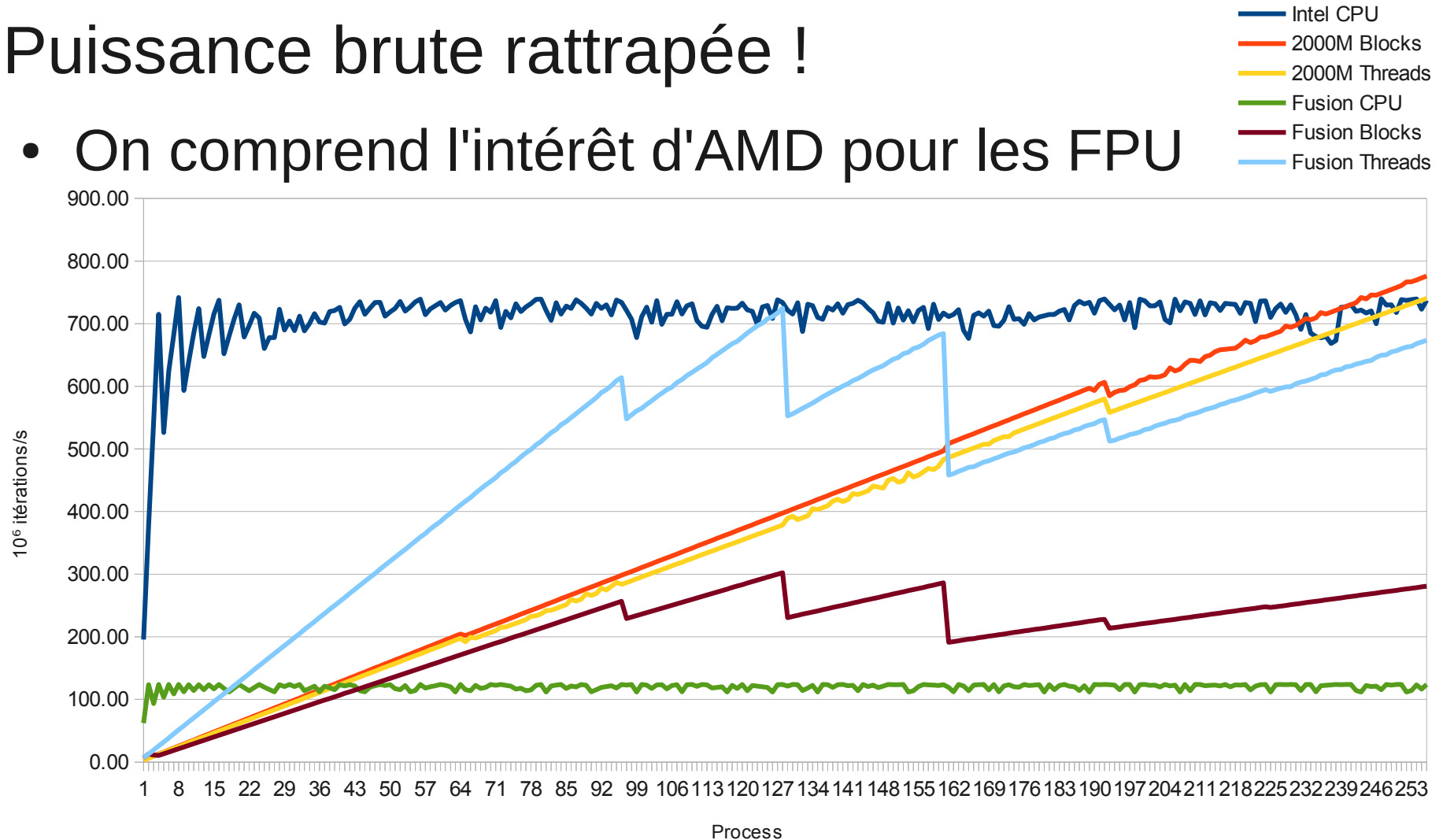
Comparaison CPU/GPU

- Pour les CPU ?
 - Comportement conforme
 - I7 pas terrible
- Pour les GPU :
 - Bien linéaires...



OpenCL : et pour de nombreux threads ?

- Puissance brute rattrapée !
 - On comprend l'intérêt d'AMD pour les FPU



La puissance « brute » est à chercher dans les GPU !

Le Futur ? Retour au problème...

- **La convergence ?** (un même socle pour HPC)
 - Le poste de travail évolue : de - en - de postes de travail
 - Les BYOD « doivent » permettre de travailler !
 - Les applications le doivent, mais les OS plus encore...
- **La rationalisation ?** (des CPU universels)
 - De vieilles recettes recyclées (APU dans CPU)
 - Serveur élémentaire ? Pas pour tout de suite...
- **Pour nous, ITRF BAP E :**
 - Investissement sur de nouveaux langages : OpenCL
 - Émergence de nouveaux métiers : pilote d'essais

Iconographie

- <http://digital-lifestyles.info/2008/03/03/intel-atom-proc>
- Eric Gaba – Wikimedia Commons user: Sting
- <http://www.lemondeinformatique.fr/actualites/lire-intel>
- <http://www.lemondeinformatique.fr/actualites/lire-hp-l>
- Photos : ACER, ASUS, Dell, Samsung,
- Logos : AMD, Intel, ARM