

PERFORMANCE ANALYSIS OF THE SUPERCOMPUTER BASED ON RASPBERRY PI NODES

Pham Van Hai^{1*}, Ho Khanh Lam², Le Thanh Hai³

Abstract: *In this paper, a new Raspberry Pi supercomputer cluster architecture is proposed. Generally, to gain speed at petaflops and exaflops, typical modern supercomputers based on 2009-2018 computing technologies must consume between 6 MW and 20 MW of electrical power, almost all of which is converted into heat, requiring high cost for cooling technology and Cooling Towers. The management of heat density has remained a key issue for most centralized supercomputers. In our proposed architecture, supercomputers with highly energy-efficient mobile ARM processors are a new choice as it enables them to address performance, power, and cost issues. With ARM's recent introduction of its energy-efficient 64-bit CPUs targeting servers, Raspberry Pi cluster module-based supercomputing is now within reach. But how is the performance of supercomputers-based mobile multicore processors? Obtained experimental results reported on the proposed approach indicate the lower electrical power and higher performance in comparison with the previous approaches.*

Keywords: Mobile Arm processor; Raspberry Pi supercomputer cluster; Performance analysis.

1. INTRODUCTION

A modern supercomputer usually consumes between 4 and 6 megawatts-enough electricity to supply something like 5000 homes. The power cost and cool the system can be significant, e.g. 4 MW at \$0.10/kWh is \$400 an hour or about \$3.5 million per year. To gain speed at petaflops and exaflops, most modern supercomputers based on multicore Intel CPU and manycore Nvidia coprocessors technologies of 2009-2018 must consume between 6 MW and 24 MW of electrical power, almost all of which is converted into heat, requiring the high cost of cooling technology and Cooling Towers, dissipate the huge amount of CO₂ in the environment. For example, the Tianhe-2 with a peak performance of 33.86 petaflops consumes 24 MW of power, which is enough to power for 24,000 U.S. homes on average for a whole month [1].

The Green500 list ranks computers from the TOP500 list of supercomputers in terms of energy efficiency, typically measured as LINPACK FLOPS per watt. FLOPS per watt are a common measuring unit. There are other energy efficiency measures, such as SWaP (space, wattage, and performance), a Sun Microsystems metric for data centers, incorporating energy and space. SwaP is a simple calculation that allows any customer to run the metric and use the results to compare computer systems from different vendors [2]. All the data used to populate the SWaP metric are publicly available. Space: Space, a server occupies can be measured by the rack unit height of the system; information found on a company's website. Wattage: Metrics on the server's power consumption can be obtained using data from actual benchmark runs or vendor site planning guides, recorded in watts. Estimated system power consumption from calculators or datasheets can be also used to assess wattage. Performance: is measured by any appropriate benchmark. We can use other green computing metrics [3] to evaluate energy consumption and energy dissipation in the environment by data centers, which are installed supercomputer systems.

Green computing represents a responsible way to address the issue of global warming so that supercomputer architects will have new choices of energy-efficient key technology.

Compute efficiency and energy efficiency are more than ever major concerns for future Exascale computing systems. Arm low-power processors dominate the mobile world of smartphones, tablets, and embedded IoT devices. With data centers consuming ever more power, the idea of using highly energy-efficient Arm chips in servers is enticing, especially for energy-hungry High-Performance Computing (HPC) configurations. The new ARM cortex is more energy-efficient than a new Intel Atom processor but only in terms of GHz/ W). However, it is only a simple 32 bit RISC, whereas Intel has powerful 64 bit CISC architecture. Intel x86 processors can deliver up to 3.6 GHz while consuming up to 130 W or at the low end of 1.8 GHz at 40 W. The ARM line of chips has been reported to deliver 1 GHz at 700 mW (down by x50 in terms of GHz/W) and can reach up to 2 GHz; while still consuming less than a watt (down by x 75 in terms of GHz/W), so their power-saving effect seems to be substantial.

The direction will be given in this paper to address the existing problems of the methods introduced earlier:

- Low energy consumption by employing Arm chips or mobile chips;
- To archive computing power equivalent to devices that consume a lot of power, the article proposes a solution for pairing Raspberry-based structure and distribution of work to achieve the highest computing performance.
- From the use of low power consumption devices, compact, high mobility will help in solving the problem of saving space to place equipment and auxiliary resources (cooling system,...) to the system. The system can be operated

2. RELATED WORK

2.1. Supercomputer architecture based on Raspberry Pi 3 Model B+ nodes

2.1.1. New projects on Raspberry Pi in the World

The human appetite for computation has grown even faster than the processing power that Moore's law predicted. We need even more powerful processors just to keep up with modern applications, such as interactive multimedia communications. To deal with such obstacles, we need these powerful processors which use much less energy than those we have been accustomed to. In other words, we need mobile supercomputers that provide massive computational performance from the power in a battery. These supercomputers will make our devices much easier to use. They will perform real-time speech recognition, video transmission and analysis, and high bandwidth communication.

As a result, they can work without considering where the next electrical outlet will be. A mobile supercomputer employs natural I/O interfaces to the mobile user. For example, an input could come through a continuous real-time speech processing component. Los Alamos National Laboratory (LANL) which is home to several of the world's most powerful supercomputers, including Trinity, worked with Australian BitScope Designs to create its new Pi-powered supercomputer from 750

individual mini-computers. The device is based on five rack-mount BitScope Cluster Modules [5]. Each one has 150 Raspberry Pi 3 nodes (144 active nodes and 6 spare nodes) networked together (that's 750 total Pis). Each Raspberry Pi 3 has a Broadcom BCM2837 system-on-a-chip (SoC) with four 64-bit CPU cores clocked at 1.2 GHz. They're ARM Cortex-A53 reference cores, which are the same thing found in many budget smartphones running Qualcomm and MediaTek SoCs. This adds up to 3,000 available CPU cores for the full system, but it uses only a fraction of the power needed for a computer like Trinity. LANL estimates that the system will need less than 5 W/node, in typical operation, you need only 6 kW to run 1000 nodes including network fabric and airflow, just 1,000 watts at idle, and 2,000 watts during typical usage. LANL intends to boost this setup to 40,000 cores next year, according to the Raspberry Pi Foundation. That increase would mean a cluster of around 10,000 Raspberry Pi boards. Oracle's new Supercomputer has 1,060 Raspberry Pi 3 B+ boards (4,240 available CPU cores for the full system) [6]. The Pi 3 B+ has a 1.4 GHz 64-bit quad-core Broadcom BCM2837 Arm Cortex A53-architecture processor compared with the Raspberry Pi 3 Model B's 1.2 GHz CPU. It doesn't seem like much of a change, but it does amount to a 15 percent performance improvement. MPI, an acronym for Message Passing Interface, a library specification for parallel computing architecture makes communication of information between various nodes and clusters possible.

Today, MPI is the most common protocol used in high-performance computing (HPC) and Raspberry Pi supercomputers of LANL and Oracle. The University of Bristol and the University of Leicester will, in total, run more than 12,000 Arm-based cores, hosted by HPE (Hewlett Packard Enterprise) Apollo 70 HPC systems. The clusters at each university will be identical, consisting of 64 HPE Apollo 70 systems, each equipped with two 32 core Cavium ThunderX2 (64-bit Arm®v8 processors), 128 GB of memory composed of 16 DDR4 DIMMs with Mellanox InfiniBand interconnects. The operating system is SUSE Linux Enterprise Server for HPC. Each cluster is expected to occupy two computer racks and consume a total of approximately 30KW of power. By introducing new processors like Arm to the HPC ecosystem, June 18, 2018 – HPE announced its collaboration with Sandia National Laboratories and the U.S. Department of Energy (DOE) to deliver the world's largest Arm supercomputer. Seattle-based supercomputer maker Cray is already drafting ARM chips for its existing machines and has just announced that its new Cray Blade will also be based on 64-bit Armv8-A architecture. Cray Blade will use the Cavium ThunderX2 processors and will feature a full software setting, including the Cray Linux and Cray programming environments, and ARM-optimized compilers, libraries, and tools for running supercomputing workloads. As part of the Vanguard program, Astra, the new Arm-based system, will be used by the National Nuclear Security Administration (NNSA) to run advanced modeling and simulating workloads for addressed areas such as national security, energy, and science. Within the past 18 months, Fujitsu chose ARM cores when it started building its Post-K supercomputer in Japan. In the UK, supercomputer Islamabad - which will be used, among other purposes for large-scale simulations - incorporates more than 10,000 64-bit ARM processors at its core.

2.1.2. Our design: Supercomputer Cluster Module

We also design and build the Supercomputer Cluster Module based on Raspberry Pi 3 B+, including:

- 24 Raspberry Pi 3 Model B+ boards (1 master, 23 data nodes).
- 24 SD cards (Samsung Evo+) 32 GB for each Pi board. The choice was made based on microSD Card Benchmarks [7]
- DC Power Adaptors 5V/2.5A for Pi boards.
- Networking Cables RJ-45. 2 TP-Link Switch 10/100/1000 Mbps Ethernet 16 ports.
- Our system runs on Raspbian OS.
- The cluster uses MPICH, high performance and widely portable implementation of the **Message Passing Interface (MPI)** standard (use: mpich-3.1 and mpi4py-1.3.1).
- Total cost of Cluster Module: 1,550 USD (~ 36,200,000 VND).



Figure 1. Raspberry Pi 3 B+ cluster 24 nodes.

2.2. Performance analysis of supercomputers based on Raspberry Pi 3 Model B+

Performance evaluation of a system is a very important step in building any complex system. Performance evaluation techniques can be classified into two categories:

- Performance modeling: Modeling is usually carried out before the prototype is built. By modeling the system, several design decisions can be drawn so that unexpected problems can be avoided while prototyping. Performance models can be further classified into simulation-based models and analytical models.

- Performance measurement: Measurement can be performed after the prototype is built. By measuring the performance of the prototype, we can validate the models used in the design process and provide feedback for the improvement of future designs.

3. PROPOSED APPROACH

Analytical modeling techniques for performance analysis of multi-core and many-core processor computing systems, and Multi-Cluster Architecture based queuing networks [10-12], and Timed Petri nets [13-16] are widely used.

For performance analysis of parallel computer system, there is Amdahl's law: analyze whether a program merits parallelization, Gustafson-Barsis's law: evaluate the performance of a parallel program, Karp-Flatt metric: decide whether the principle barrier to speedup is due to inherently sequential code or parallel overhead, and Isoefficiency metric: evaluate the scalability of a parallel program executing on a parallel computer.

The common speedup for multiprocessor parallel computing systems states that of a p -size program's execution time with parallelization overhead $k(n,p)$:

$$\text{Speedup}(p, n) \leq \frac{T_{\text{seq}}(p) + T_{\text{par}}(p)}{T_{\text{seq}}(p) + \frac{T_{\text{par}}(p)}{n} + k(p, n)} \quad (1)$$

Where:

- $T(p, 1) = T_{\text{seq}}(p) + T_{\text{par}}(p)$ - p -size program's execution time in a sequence (one processor) computer;

- $T(p, 1) = T_{\text{seq}}(p) + T_{\text{par}}(p)$ - p -size program's execution time in a parallel (multi- processor) computer.

The performance measurement for multiprocessors, a modern version of Amdahl's law for n physical processors (cores) in multiprocessor parallel computing systems, states that if parallel fraction f of a p -size program's execution time with no scheduling parallelization overhead $k(n,p)$, while the remaining fraction, $1-f$, is sequential, the speedup on n physical processor (n -cores) system is governed by [17-19]:

$$\text{Speedup}(f, n) = \frac{1}{(1 - f) + \frac{f}{n}} \quad (2)$$

Besides Speedup measures by using Amdahl's law for the performance analysis of multiprocessor computers, authors of paper [20] present performance analysis for the Single Board Computer Clusters (SBCs). The metrics have been accompanied by increases in energy efficiency (GFLOPS/W) and value for money (GFLOPS/\$) for three different SBC clusters composed of Raspberry Pi 3 Model B, Raspberry Pi 3 Model B+, and Odroid C2 nodes respectively. A 16-node SBC cluster can achieve up to 60GFLOPS, running at 80W.

Table 1. Raspberry Pi 32/64-bit Benchmarks and Stress Tests [22].

Summary	Introduction	Whetstone Benchmark
Dhrystone Benchmark	Linpack 100 Benchmark	Livermore Loops Benchmark
FFT Benchmarks	BusSpeed Benchmark	Memspeed Benchmark
NeonSpeed Benchmark	MultiThreading Benchmarks	MP-Whetstone Benchmark
MP-Dhrystone Benchmark	MP NEON Linpack Benchmarks	MP-BusSpeed Benchmark
MP-BusSpeed Disassembly	MP-RandMem Benchmark	MP-MFLOPS Benchmarks

MP-MFLOPS Disassembly	MP-MFLOPS Sumchecks	OpenMP-MFLOPS Benchmarks
OpenMP-MemSpeed Benchmarks	Stress Testing Benchmarks	Integer Stressing Benchmark
Single Precision Stress Benchmarks	Double Precision Stress Benchmark	High-Performance Linpack
DriveSpeed Benchmark	USB 3 and 2 Benchmarks	Drive Write/Reboot/Read Tests
LAN and Wifi Benchmarks	Java Whetstone Benchmark	JavaDraw Benchmark
OpenGL Benchmark	Stress Tests	HP Linpack Stress Test
Integer Stress Test	Single Precision FPU Stress Test	Double Precision FPU Stress Test
OpenGl + 3 x Livermore Loops	Input/Output Stress Test	CPU + Main SD + USB + LAN Test

Many performance benchmark types can be run for systems based on Raspberry Pi modules. There are 32-bit and 64-bit benchmarks and stress tests on the appropriate range of Raspberry Pi computers, up to model 3 B+/Pi 4B (table 1). Thermal Management, Whetstone, Dhrystone, SysBench CPU, Linpack, Python GPIO, SysBench RAM, Ethernet, Power Draw, Wi-Fi, and Thermal Throttling Benchmarks [21] show Pi 3 B+ is a worthy upgrade over the predecessor designs.

3.1. Propose performance measures of interconnection network topologies

3.1.1. Use topology parameters to compare interconnection topologies

Performance metrics most commonly used in comparing interconnect network topologies of multiprocessor computers are an average and a maximum number of hops and bisection bandwidth (W) and an average number of hops and network diameter. If the network is bisected into two partitions, we can define bisection-width (B) of a network - is the minimum number of links to be removed to disconnect the network into two halves of equal size, it identifies a potential bottleneck of a network and is implied on its internal bandwidth. A low B can slow down many collective communication operations and thus can severely limit the performance of applications. However, achieving a high B may require a non-constant network degree. The average number of hops (H) and diameter (D) of a network - the largest distance between any pair of nodes between nodes/switches in the network serve as measures of network latency, even though they are only a partial indication of the actual message latency. To support efficient communication between any pair of nodes, the D should be minimized and B should be increased, or the rate D/B should be minimized. The maximum total message latency through the diameter includes the link (wire) latency and it increases proportionally with the number of nodes.

There are the following topology parameters of networks:

- k_i - Number of nodes in each dimension I , for an asymmetric topology: $k = k_i$. asymmetric = an asymmetric;
- n - Number of dimensions;

- **Total number of nodes (N):**
nD-Mesh: $N = k^n$, nD-Torus: $N = k^n$, hypercube (n-cube): $N = 2^n$, full connected: N .
- **Node degree (d)** - The maximum number of edges (or links) connected to a node.
nD-Mesh: $2n$, nD-Torus: $2n$, hypercube: $\log_2 N = n$, full connected: $N - 1$.
- **Number of links (L):**
2D-Mesh: $2N^{1/2}(N^{1/2} - 1)$, 3D-Mesh: $N^{2/3}(N^{2/3} - 1)$ for N is odd, and nN for N is even. nD-Torus: nN , hypercube: nN , full connected: $N(N-1)/2$.
- **Diameter (D)**- nD-Mesh: $n(\sqrt[n]{N} - 1)$, nD-Torus: $(n/2)(N^{1/2})$, hypercube: $2\log_2 N = 2\log_2 2^n = 2n$, full connected: $D = 1$.
- **Average number of hops (H):**
2D-Mesh: $(N^{1/2} - 1)$, 3D-Mesh: $(N^{2/3} - 2)$, nD-Torus: $(n/4)N^{1/n}$, hypercube: $\log_2 N = n$, full connected: 1.
- **Bisection Width (B):**
nD-Mesh: $N^{n-1/n}$; nD-Torus: $nN^{n-1/n}$; hypercube: $(N/2)^2$ if N even, $N^2/4$ if N odd; full connected: $(N/2)^2$ if N even, $N^2/4$ if N odd.
- **Bisection bandwidth (W)** = bisection width $\times$ channel bandwidth
The channel bandwidths in our design are 10 Gb/s links between 10 Gb/s switch nodes.
- **Cost (C)** is defined as the product of the node degree (d) and diameter (i.e.).
$$C = d \times D$$

Table 2. Parameters of different network topologies with 48-port 10 GbE switches.

Topology parameters	2D mesh	3D mesh	2D torus	3D torus	Hypercube (n-cube)	Full connected (FCN)
Performance:						
Bisection width (B)	8	16	16	48	32	1024
Average number of	7	14	4	3	3	1
Diameter (D)	14	9	8	6	6	1
Bisection bandwidth	80	40	160	480	320	10240
Cost:						
Total number of nodes	64	64	64	64	64	64
Node degree (d)	4	6	4	6	3	63
Number of links (L)	112	192	128	192	384	2016
Number of 48-port Ethernet switches	64	64	64	64	64	64
Number of Computer	60	60	60	60	60	60
Cost (C)	56	54	32	36	18	63
W/C	1.4	0.7	5	13.3	17.8	0.02

Note*: every switch connects 40 pi nodes, 8 ports for the interconnection of switches.

From table 2, nD torus, and hypercube topologies have the best performance/cost than other ones, but the hypercube has a big number of links, so we choose nD torus for our raspberry Pi supercomputer.

Table 3. Topology parameters of nD torus with 48-port 10 GbE switches.

Topology parameters	2D tori 4x4	3D tori 4x4x4	2D tori 6x6	3D tori 6x6x6	2D tori 8x8	3D tori 8x8x8	2D tori 10x10	3D tori 10x10x10	2D tori 12x12	3D tori 12x12x12	2D tori 14x14	3D tori 14x14x14
Performance												
B	8	32	12	72	16	128	20	200	24	432	28	588
H	2	3	3	4.5	4	6	5	7.5	6	9	7	10.5
D	4	6	6	9	8	12	10	15	12	18	14	21
W Gb/s	80	320	120	720	160	1280	200	2000	240	4320	280	5880
Cost:												
k	4	4	6	6	8	8	10	10	12	12	14	14
N (switches)	16	64	36	216	64	512	100	1000	144	1728	196	2744
d	4	6	4	6	4	6	4	6	4	6	4	6
L	32	192	72	452	128	1536	200	3000	288	1296	392	8232
Cost (C)	16	36	24	54	32	72	40	90	48	108	56	122
Pi nodes (P)	640	60	1440	8640	60	20480	4000	40000	5760	69120	7840	109760

The nD torus has a disadvantage: the complexity of wiring when the network size is big as the number of links is big, like a 3D torus. But there are important performance benefits of nD torus networks that fit in the design of HPC based Arm low-power efficiency and energy efficiency multiprocessors:

- Better fairness, higher speed, and lower latency because the high that affected bandwidth and every link in the folded nD torus network are very short-almost as short as the nearest-neighbor links in a simple grid interconnect-therefore low-latency so data have more options to travel from one node to another which greatly increases speed.

- Lower energy consumption since data tend to travel through fewer hops (value H is small), the energy consumption tends to be lower.

Figure 2 shows the bisection bandwidth of 2D torus is linearly increased over the number of nodes/switches, but the 3D torus networks have the bisection bandwidth is increased higher than 2D torus with the network size.

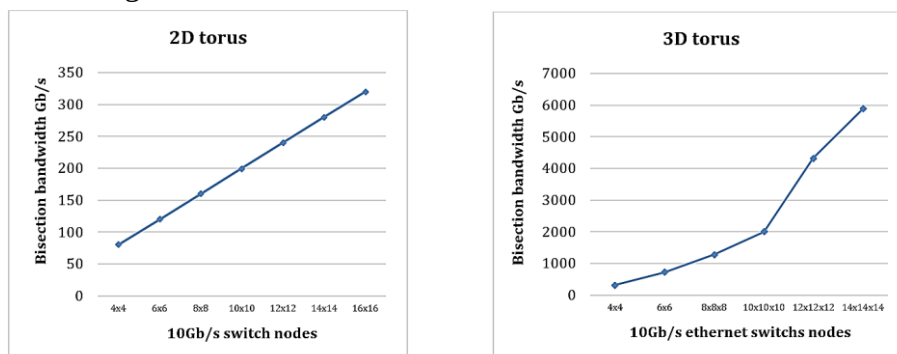


Figure 2. Bisection bandwidth of 2D and 3D torus interconnections of 10 Gb/s ethernet switches.

3.1.2. Amdahl's law expansion to compare interconnection topologies

To define how the speedup is dependent on the interconnection network topology

latency, we propose to use the speedup formula (2) of Amdahl's law with adding the rate of latency, D/B , since the diameter and bisection-width effect on the communication overhead:

$$\text{Speedup}(f, n, D, B) = \frac{1}{(1-f) + \frac{f}{n} + \frac{D}{B}} \quad (3)$$

Amdahl's law could be applied only to the cases in which the problem size is fixed. In practice, as more computing resources are available, they tend to get used to larger problems (larger datasets), and the time spent in the parallelizable part often grows much faster than the inherently serial work. According to Amdahl's law, the speedup is limited by the serial part $(1-f)$ of the program. For example, if $f = 80\%$ of the program can be parallelized, the theoretical maximum speedup using parallel computing would be 20 times. So we let the size $f = 80\%$ and is constant, then:

$$\text{Speedup}(f, n, D, B) = \frac{1}{0.2 + \frac{0.8}{n} + \frac{D}{B}} \quad (4)$$

The values of $n = P, D, B$ used from table 4 to define Speedup in formula (4) for the 2D torus and 3D torus. The results are exhibited in figure 3, that 3D torus is better than 2D torus in the speedup and the total number of Raspberry Pi nodes.

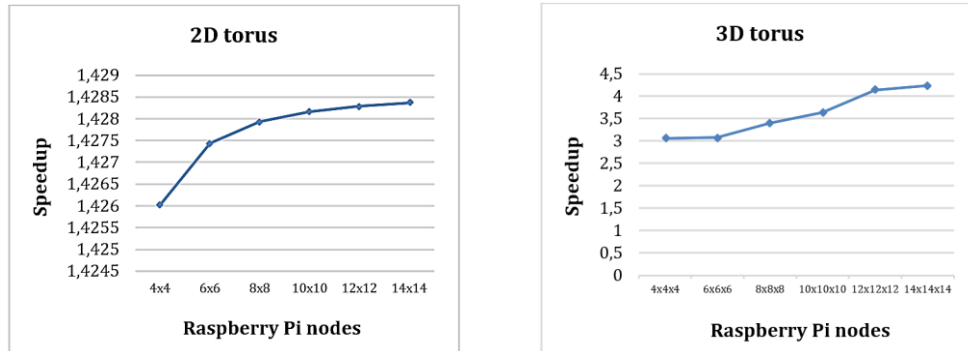


Figure 3. The speedup of nD torus independent of n, D, B .

3.1.3. The power consumption and performance (GFlops) of Raspberry Pi node

From the Benchmarks [20, 23], we can define the power consumption of Raspberry Pi nodes in full load state (stress –CPU 4). The average power usage of the single board raspberry Pi 3 Model B+ (SBR3B+) is 5.1 W - 5.66 W. From the Linpack test [24], the speed of SBR3B+ is 224.89 MIPS for SP (Single-precision floating-point), and 209.23 MIPS for DP (double precision floating point). The speed/power of SBR3B+ node is 39.73 MIPS/W (224.89 MIPS/5.66 W for SP) or 36.97 MIPS/W (209.23 MIPS/5.66 W for DP). So the 3D torus $8 \times 8 \times 8$ has a power consumption of about 116 KW only of 20480 Pi 3 model B+ nodes (116 KW = 5.66 W x 20480). This power consumption of Raspberry Pi Nodes is much smaller than that of the interconnection networks of nodes/switches.

4. CONCLUSIONS

In this work, we have proposed using topology parameters to compare interconnection cluster topologies for the design of Raspberry Pi supercomputers.

From the obtained results, we proposed to use nD torus topology, since their performance/cost parameters are the best. We also proposed the simple expansion of Amdahl's law formula with the rate of D/B as the factor that affected the communication overhead of the interconnection network topology to evaluate the speedup and compare different network topology types.

For future work, the supercomputer cluster module of 24-Raspberry Pi 3 Model B+ could be applied to Raspberry Pi 4 8GB cluster models in 3D torus interconnection network of 10GBASE-T switches to create a higher performance (TFlops/watt) we expect.

REFERENCES

- [1]. Victor Tangermann. *"The Eight most powerful supercomputers in the world"*. September 28th, 2017.
- [2]. Greenhill, David. *"SWaP Space Watts and Power"* (PDF). US EPA Energystar. Retrieved 14 November 2013.
- [3]. Girish Kumar Patnaik et. al. *"Green Computing Metrics, Methods and Models"*. International Journal of Engineering Research & Technology (IJERT). ISSN: 2278-0181. Vol. 3 Issue 3, March – 2014.
- [4]. *"Mont-Blanc. European Modular and Power-Efficient HPC Processor"*. Copyright 2011 - 2020 © All Rights Reserved
- [5]. *"Scalable clusters make HPC R&D easy as Raspberry Pi"*. Bitscope.com/cluster
- [6]. Gerald Venza. *"Building the world's largest Raspberry Pi cluster"*.
- [7]. www.pidramble.com/wiki/benchmarks/microsd-cards.
- [8]. Nikhil Jain et.al. *"Predicting the Performance Impact of Different Fat-Tree Configurations"*. Lawrence Livermore National Laboratory.
- [9]. Tomohiro Inoue, Fujitsu Limited. *"The 6D Mesh/Torus Interconnect of K Computer"*.
- [10]. G. Bolch, S. Greiner, H. de Meer, and K. S. Trivedi, *"Queueing Networks and Markov Chains"*, John Wiley, 2nd edition, 2006.
- [11]. Norhazlina Hamid, Robert John Walters, Gary Brian Wills. *"An Analytical Model of Multi-Core Multi-Cluster Architecture (MCMCA)"*. Open Journal of Cloud Computing (OJCC) Volume 2, Issue 1, 2015. ISSN 2199-1987.
- [12]. Xiaoyue Pan. *"Performance Modeling of Multicore Systems"*. ISSN 1651-6214 ISBN 978-91-554-9451-3.
- [13]. Murata, T.: *"Petri nets: properties, analysis, and applications"*, Proceedings of IEEE, 77 (4), 1989, 541-580.
- [14]. Falko Bause, Pieter S. Kritzinger, *"Stochastic Petri Nets"*. Bause and Kritzinger, 2002.
- [15]. M. Ajmone Marsan, Gianfranco Balbo, Gianni Conte, Susanna Donatelli, Giuliana Franceschinis, *"Modelling with generalized stochastic Petri nets"*. Università degli Studi di Torino.
- [16]. Viktor Mashkov & Jiri Barilla & Pavel Similar. *"Applying Petri Nets to Modeling of Many-Core Processor Self-Testing when Tests are Performed Randomly"*. J Electron Test (2013).

- [17]. Mark D. Hill, University of Wisconsin-Madison Michael R. Marty, Google. "Amdahl's Law in the Multicore Era".
- [18]. Christina Delimitrou, Christos Kozyrakis. "Amdahl's Law for Tail Latency". August 2018 | Vol. 61| No. 8| Communications of the ACM.
- [19]. Surya Narayanan Natarajan. "Modeling performance of serial and parallel sections of multi-threaded programs in many core era". pr'epar'ee `a l'unit'e de recherche INRIA – Bretagne Atlantique Institut National de Recherche en Informatique et Automatique Composante Universitaire (ISTIC).
- [20]. Philip J. et. al. "Performance analysis of single-board computer clusters". Future Generation Computer Systems. 102 (2020) 278-291. ELSEVIER.
- [21]. Gareth Halfacree. "Benchmarking the Raspberry Pi 3 B+". Mar 14, 2018.
- [22]. [22]. Roy Longbottom, UK Government. "Raspberry Pi 4B 32 Bit Benchmarks". Technical Report. June 2019.
- [23]. <https://www.pidramble.com/wiki/benchmarks/power-consumption>.
- [24]. Lucy Hattersley. "Raspberry Pi 4 vs Raspberry Pi 3B+". <https://magpi.raspberrypi.org/articles/raspberry-pi-4-vs-raspberry-pi-3b-plus>
- [25]. "Raspberry Pi 4 vs Raspberry Pi 3B+," The MagPi magazine. <https://magpi.raspberrypi.org/articles/raspberry-pi-4-vs-raspberry-pi-3b-plus>.

TÓM TẮT

PHÂN TÍCH HIỆU NĂNG CỦA SIÊU MÁY TÍNH DỰA TRÊN NỀN TẢNG RASPBERRY

Để đạt được tốc độ ở petaflops và exaflops, các siêu máy tính hiện đại hiện nay dựa trên công nghệ điện toán 2009-2018 cần phải tiêu thụ điện năng từ 6 MW đến 20 MW, hầu hết điện năng đó đều được chuyển đổi thành nhiệt, do vậy đòi hỏi chi phí cao về công nghệ làm mát và giải nhiệt. Việc xử lý giảm nhiệt vẫn là một vấn đề quan trọng đối với hầu hết các siêu máy tính hiện nay. Các siêu máy tính có hiệu suất năng lượng cao dựa trên bộ xử lý của dòng chip ARM trên các điện thoại di động thông minh là một lựa chọn mới vì nó cho phép chúng giải quyết các vấn đề về hiệu suất, năng lượng và chi phí. Với công nghệ dòng chip ARM giới thiệu gần đây các siêu máy tính có máy chủ sử dụng mô-đun cụm raspberry Pi CPU 64-bit 1.4 GHz hướng tới mục tiêu tiết kiệm năng lượng, chiếm ít không gian lưu trữ, tỏa ít nhiệt năng và hạn chế tối đa lượng CO₂ tạo ra đó là xu hướng nghiên cứu mới, đã nằm trong tầm tay. Vấn đề là cần làm gì và hiệu năng của siêu máy tính dựa trên bộ xử lý đa lõi di động như thế nào. Bài viết này đề xuất kiến trúc cụm siêu máy tính raspberry Pi và phân tích hiệu năng.

Từ khóa: Bộ xử lý Mobile ARM; Cụm siêu máy tính Raspberry Pi; Phân tích hiệu năng.

Received 7th November 2020

Revised 8th January 2021

Accepted 10th May 2021

Author affiliations:

¹ Hanoi Open University;

² Hung Yen University of Technology and Education;

³ Military Science and Technology Institute.

*Corresponding author: phamvanhai@hou.edu.vn.