

Do We Know the Operating Principles of Our Computers Better than those of Our Anatomy of Brain?

Janos Veghi¹ and Adam J. Berki^{2*}

¹Kalimanos BT, Debrecen, Hungary

²Department of Anatomy, University of Medicine, Pharmacy, Sciences and Technology of Targu Mures, Romania

Abstract

The increasing morphology of interest in understanding the behaviour of biological neural networks, and the expanding utilization of artificial neural networks in different fields and scales, both require a thorough understanding of how technological computing works. However, von Neumann in his classic "First Draft" warned that it would be unsound to use his suggested paradigm to model neural operation, furthermore that using "too fast" processors vitiates his paradigm, which was intended only to describe (the timing relations of) vacuum tubes. Thus, it is worth scrutinizing how the present technology solutions of anatomy can be used to mimic biology. Some electronic components anatomy bears a surprising resemblance to some biological structures. However, combining them using different principles can result in systems with inferior efficacy. The paper discusses how the conventional computing principles, components, and thinking about computing limit mimicking biological systems of morphology and anatomy.

Keywords: Morphology • Anatomy • computation • Biological systems

Introduction

For today, the definition of 'xeromorphic computing' started to diverge: from "very large scale integration (VLSI) with analogy components that mimicked biological neural systems" changed to "brain-inspired computing for machine intelligence" [1] or "a realistic solution whose architecture and circuit components resemble to their biological counter-parts" [2] or "implementations that are based on biologically-inspired or artificial neural networks in or using non-von Neu-mann architectures" [3]. How much resemblance to biology is required depends on the authors. Some initial resemblance indeed exists, and even some straightforward systems can demonstrate functionality in some aspects similar to that of the nervous system. "Successfully addressing these challenges of xeromorphic computing will lead to a new class of computers and systems architectures" [4] was hoped. However, as noticed by the judges of the Gordon Bell Prize, surprisingly, among the winners there have been no brain-inspired massively parallel specialized computers" [5]. Despite the vast need and investments, the concentrated and coordinated efforts, just because of mimicking the biological systems with computing adequately. On one side, "the quest to build an electronic computer based on the operational principles of biological brains has attracted attention over many years" [6]. On the other side, more and more details come to light about the brain's computational operations. As that the operating principles of the large computer systems tend to deviate not only from the operating principles of the brain but also from those of a single processor, it is worth reopening the discussion on a decade-old question "Do computer engineers have something to contribute to the understanding of brain and mind?" [6]. Maybe, and they surely have something to contribute to the understanding of computing itself. There is no doubt that the brain does computing, the key question is how? As we point

out in section II, the terminology makes it hard to discuss the terms. Section presents that in general large-scale computing systems have enormously high energy consumption and low computing efficiency. Section IV discusses the primary reasons for the issues and failures the necessity of fair imitation of biological objects' temporal behaviour in computing systems is discussed in section V. Neuromorphic computing is a particular type of workload informing the computational efficiency of computing systems, as section VI discusses it. Section VII draws parallel with classic vs. modern science and classic vs. modern computing. Section VIII provides examples, why a neuro morphic system is not a simple sum of its components.

Terminology of Neuron Morphic Computing

In his "First Draft" [7] Von Neumann mentions the need to develop the structure, giving consideration to both their design and architecture issues. Given that von Neumann discussed computing operations in parallel with neural operations, his model is biology-inspired, but because of its timing relations, it is not biology-mimicking. After carefully discussing the timing relations of vacuum tubes in section 6.3 of his draft, he made some reasonable omissions and provided an abstraction this is known as "classic computing paradigm" having a clear-cut range of validity. He warned that a technology using "too fast" processors vitiates that paradigm, and that it would be any how unsound to apply that neglect ion to describe operations in the nervous system [8]. "This is the first substantial work that clearly separated logic design from implementation. The computer defined in the 'First Draft' was never built, and its architecture and design seem now to be forgotten." [9] The comprehensive review [3] analysed the complete amount of three decades of related publications, based on frequently used keywords. Given that von Neumann constructed a brain-inspired model and defined no architecture, it is hard to interpret even the taxonomy that "Xeromorphic computing has come to refer to a variety of brain-inspired computers, devices, and models that contrast the pervasive von Neumann computer architecture." [3] If it is the architecture that his followers implemented von Neumann's brain-inspired model and abstraction meant for the "vacuum tube interpretation" only, what is its contrast, the definition of "neuro morphic computing"? Because of the "all-in-one" handling and the lack of clear definitions, the area seems to be divergent; improper methods or proper methods outside their range of validity are used. Von Neumann's model distinguished the [payload] processing time and the [non-payload, but needed] transport time. Any operation, whether neuromorphic or conventional

*Address for Correspondence: Adam J. Berki, Department of Anatomy, University of Medicine, Pharmacy, Sciences and Technology of Targu Mures, Romania, E-mail: Romaniaberki.adam@yahoo.com

Copyright: © 2022 Veghi J, et al. This is an open-access article distributed under the terms of the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

Received: 03 February, 2022, Manuscript No. Jma-21-36057; **Editor Assigned:** 05 February, 2022, PreQC No. P-36057; **QC No.** Q-36057; **Reviewed:** 17 February, 2022; **Revised:** 22 February, 2022, Manuscript No. R-36057; **Published:** 01 March, 2022, DOI: 10.37421/2684-4265.2022.6.226

or analog must comprise these two components. Von Neumann provided a perfect model valid for any kind of computing, including xeromorphic one. Notice that these components mutually block each other: an operation must not start until its operands are delivered to the input of the processing unit, and its results cannot be delivered until their processing finished. This latter point must be carefully checked, especially when using accelerators, recurrent relations of computed feedback. Importantly, the transfer time depends both on the physical distance of the respective elements in the computing chain and on the method chosen to transfer the data [10]; it can only be neglected after a meticulous analysis of the corresponding timing relations. The need for also using the time in neural spikes was early noticed. Activity is assumed to consist of neuronal ensembles – spikes clustered in space and time [11] (emphasis in original). However, if the packets are sent through a bus with congestion, the information is lost (or distorted) in most cases, especially in large systems. Here we attempt to focus on one vital aspect of biology-mimicking technical computing systems: how truly they can represent the biological time.

Issues with Large Scale Computing

Given that we can perform a “staggering amount of (1016) synaptic activity per second” [12] assuming one hundred machine instructions per synaptic activity, we arrive in the xOps region. It is the needed payload performance, rather than nominal performance (orders of magnitude between), for neural operation. However, the worst limiting factor is not large number of operations. In the bio-inspired models, up to billions of entities are organized into specific assemblies. They cooperate via communication, which increases exponentially with increasing complexity/number. The communication here means sending data and sending/receiving signals, including synchronization. The major issue is that technology cannot mimic the “private buses” of biology. To get nearer to the biological brain’s computationally and energetically efficient operation, we must mimic a complete of biological features. Only a small portion of the neurons are working simultaneously in solving the actual task; there is a massive number of very simple (‘very thin’) processors rather than a ‘fat’ processor; only a portion of the functionality and connection are pre-wired, the rest is mobile; there is an inherent redundancy, replacing a faulty neuron may be possible via systematic training. The vast computing systems can cope with their tasks with growing difficulty. The examples include demonstrative failures already known (such as the supercomputers Gyoukou and Aurora’18, or the brain simulator SpiNNaker) and many more may follow: such as Aurora’21 [13], the China mystic supercomputers and the EU planned supercomputers. Also, the present world champion (as of 2020 July) Fugaku installed [14] at some 40% of its planned capacity, and in a half year could increase only marginally. As displayed in Figure 1, the efficiency of computers assembled from parallelized sequential processors depends both on their parallelization efficiency and number of processors. It was predicted: “scaling thus put larger machines at an inherent disadvantage”, since “this decay in performance is not a fault of the architecture, but is dictated by the limited parallelism” [15]. As discussed in detail in [12-16], the efficiency of the parallelization of the large systems is essentially defined by the workload they run. The artificial intelligence is the most disruptive workload from an I/O pattern perspective. For orientation see [17] and the estimated efficiency of simulating the brain in Figure 1. The performance scaling is strongly nonlinear [16]. When targeting neuro-morphic features such as “deep learning training”, the issues start to manifest at just a couple of dozens of processors [18,19].

Limitations Due to Technical Implementation

Biology uses purely event-driven (asynchronous) computing, while modern electronics uses clock-driven (synchronous) systems. The changing technology and the ever-growing need for more computing also increased the physical size, both of the processors and the systems targeting high processor performance. Synchronizing the operation of elements of computing chain that have logical dependence, furthermore that because the finite physical size leads to “time skew” in the same signal’s arrival times, it introduces a

dispersion to the synchronization signal. The dispersion of synchronizing the computing operations vastly increases the cycle time, decreases the utilization of all computing units, and enormously increases the power consumption of computing [20], and it is one of the primary reasons for the inefficiency of the Application Specific Integrated Circuit (ASIC) circuits. Figure 2 shows how the merit “dispersion” has changed during the years due to demands against computing and the implementation technology. The definition of dispersion and its detailed discussion is given by Boahen KA [8]. The red diagram line “Proc dispersion” (i.e., the dispersion inside the processor) has grown to a value about a hundred times higher than the one at which von Neumann justified his “procedure”. The dispersion diagram line, alone vitiates applying the classic paradigm to our processors. From the other diagram lines, one can see that the technical implementation that the data must be delivered between the technology blocks “memory” and “processor”, and make the value of “system dispersion” orders of magnitude higher. This feature is mistakenly attributed to the consequence of the “von Neumann architecture” as the “von Neumann bottleneck”. A case where the classic paradigm is really “unsound” given that it neglected the transfer time. How the workload turns this “technical

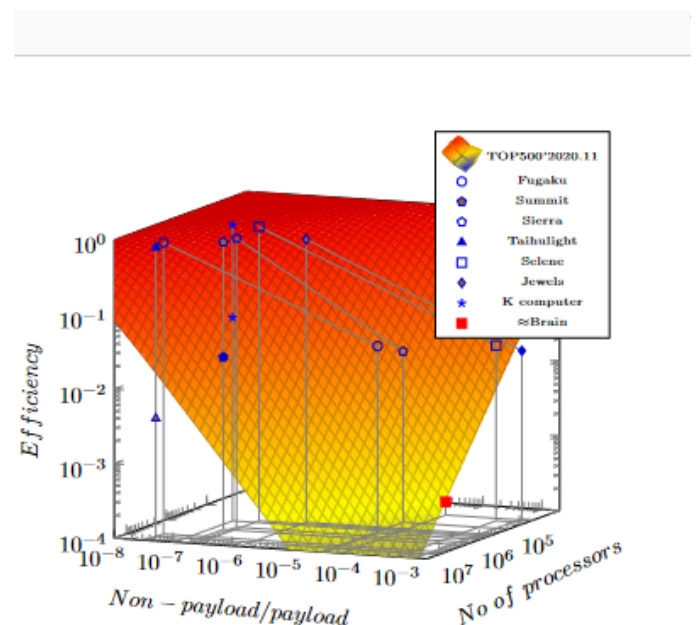


Figure 1. The 2-parameter efficiency surface (in function of parallelization efficiency measured by benchmark HPL and number of processing elements) as concluded from Amdahl's Law.

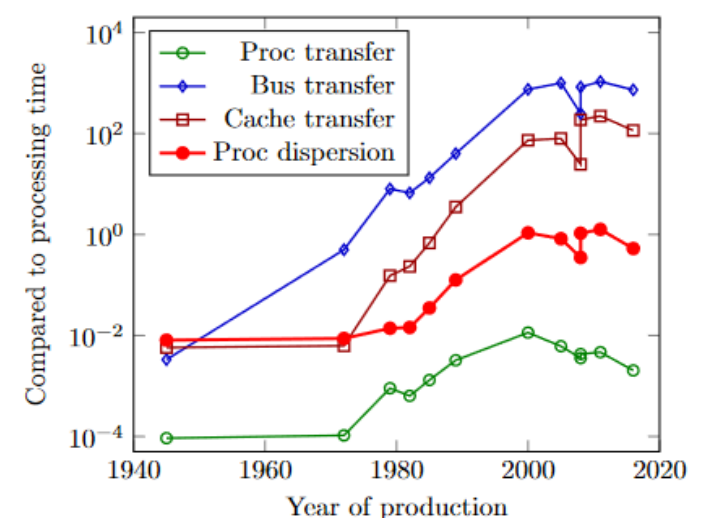


Figure 2. The history of some relative temporal characteristics of processors, in function of their year of production. Notice how cramming more transistors in a processor changed disadvantageously their temporal characteristics.

implementation" (stemming from the Single Processor Approach (SPA)) into a real bottleneck is illustrated with the case when neural-like communication shall be served. The inset in Figure 3 shows a simple xeromorphic use case: one input neuron and one output neuron are communicating through a hidden layer, comprising only two neurons. Figure 3A mostly shows the biological implementation: all neurons are directly wired to their partners, i.e., a system of "parallel buses" (the axons) exists. Notice that the operating time also comprises two non-payload times: the data input and data output, which coincide with the non-payload time of the other communication party. The diagram displays the logical and temporal dependencies of the neuronal functionality. The payload operation ("the computing") can only start after the data is delivered (by the, from this point of view, non-payload functionality: input-side communication), and the output communication can only begin when the computing is finished. Important that: i/the communication time is an integral part of the total execution time, and ii/the ability to communicate is a native functionality of the system. In such a parallel implementation, the performance of the system, measured as the resulting total time (processing + transmitting), and scales linearly with increasing both the non-payload communication speed and the payload processing speed. The present technical approaches assume a similar linearity of the performance dependence of the computing systems as "Gustafson's formulation gives an illusion that as if N [the number of the processors] can increase indefinitely" Gustafson's 'linear scaling' neglects the communication entirely (which is not the case, especially in xeromorphic computing). The interplay of the improving parallelization and the general hardware (HW) development covered for decades that the scaling was used far outside of its range of validity [16]. Figure 3B shows at technical implementation of a high-speed shared bus for communication. To the grid's right, the activity that loads the bus at the given time is shown. A double arrow illustrates the communication bandwidth, the length of which is proportional to the number of packages the bus can deliver in a given time unit. The high-speed bus is only very slightly loaded. We assume that the input neuron can send its information in a single message to the hidden layer furthermore that the processing by the neurons in the hidden layer both starts and ends at the same time. However, the neurons must compete for accessing the bus, and only one of them can send its message immediately, the other(s) must wait until the bus gets released. The neurons at a few picoseconds distance from each other must communicate through the shared bus, in the several nsec range. As the timing analysis in pointed out the resulting transfer time.

Depends linearly on the number of "neurons" in the system, and the systems spend most of their time waiting for the arbiter. The "high-speed bus" has marginal importance in this case: it is only slightly utilized. This effect is so strong, that in vast systems a "communication collapse" follows, when the "rooftline" approached. The output neuron can only receive the message when the first neuron completed it. Furthermore, the output neuron must first acquire the second message from the bus, and the processing can only begin after having both input arguments. This constraint results in sequential bus delays both during on-payload processing in the hidden layer and the payload processing in the output neuron. At this point, one can stick to synchronous computing, which increases dispersion to an intolerable level. It leads, however, to the need of mitigating the dispersion using methods such as "spikes mare dropped if the receiving process is busy over several delivery cycles" The other option is to proceed without synchronization (i.e., mixing the synchronous and asynchronous operating modes). The HW gates always have an input signal, and when started, they perform the operation they were designed for. If the correct input signal reached its input, then it works as we expect. Otherwise, the output is undefined, as its input is. Adding one more neuron to the layer introduces one more delay, which explains why "shallow networks with many neurons per layer . . . scale worse than deep networks with less neurons" [18] the system sends them at different times in the different layers (and even they may have independent buses between the layers), although the shared bus persists in limiting the communication. The neuromorphic systems that need much more communication, making their non-payload to payload ratio very wrong. The linear dependence at low nominal performance values explains why the initial successes of any new technology, material or method in the field, using the classic computing model, can be misleading: in simple cases the classic paradigm performs tolerably well thanks to that compared to biological neural networks, current neuron/dendrite models are simple, the networks small, and learning models appear to be rather basic.

The Importance of Imitating the Timely Behaviour

In both biological and electronic systems, both the distance between the entities of the network and the signal propagation speed is finite. Because of this, in the physically large-sized systems the 'idle time' of the processors

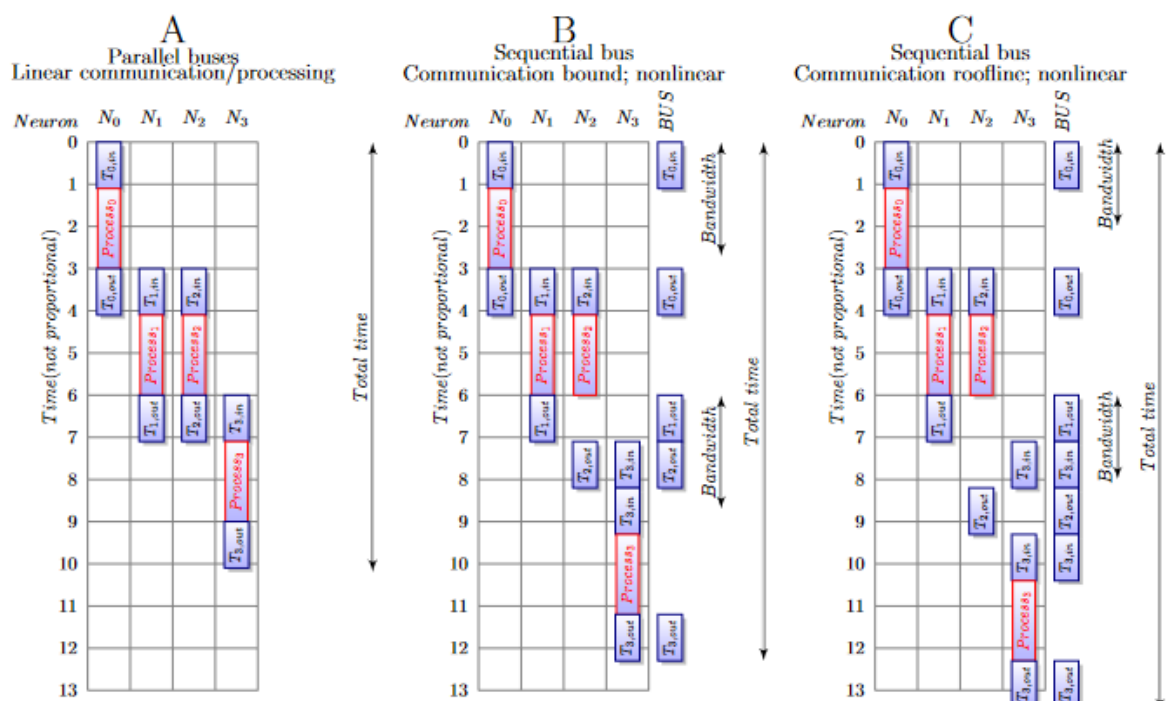


Figure 3. Implementing neuronal communication in different technical approaches. **A)** The parallel bus, **B)** and **C)** the shared serial bus, before and after reaching the communication "rooftline".

defines the final performance a parallelized sequential system can achieve. In the conventional computing systems, the 'data dependence' limits the available parallelism: we must compute the data before using it as an argument for another computation. Thanks to the 'weak scaling' this 'communication time' is neglected. In xeromorphic computing, however, as discussed in connection with Figure 3, the transfer time is a vital part of information processing. A biological brain must deploy a speed accelerator" to ensure that the control signals arrive at the target destination before the arrival of the controlled messages, despite that the former derived from a distant part of the brain. This aspect is so vital in biology that the brain deploys many cells with the associated energy investment to keep the communication speed higher for the control signal. To arrive at a less limiting architecture, additional ideas shall be borrowed from the biological architectures. Although the "private buses" cannot be reproduced technologically, introducing hierarchical buses and organized traffic, using computer network-like logical addressing and the "small world" nature of communication, more close imitation can be reached.

The approach uses the following key novel ideas:

- Implementing directly-wired connections between physically neighbouring cells
- Creating a particular hierarchical bus system
- Placing a special communication unit, the (Inter-Core Communication Block (ICCB))

Figure 4 shows (purple) between the computer cores mimicking neurons (green) creating a specialized 'fat core' neuron (B) with the extra abilities to access the local and far memories (M) and to forward messages via the gateway (G) to other similar 'fat core' neurons (similar gateways can be implemented for the inter-processor communication, and higher organizational levels). This organization enables both easy sharing of locally important state variables, keeps local traffic away from the bus and reduces wiring inside the chip. The ICCBs can closely mimic the correct parallel behaviour of biology. The resemblance between Figure 4A and in reference underlines the importance of making clear distinction between handling 'near' and 'far' signals. Furthermore, it underlines as well as the necessity of their simultaneous utilization. The conduction time in biological systems must be separately maintained in biology-mimicking computing systems. Making time-stamps and relying on the computer network delivery principles is not sufficient for maintaining correct relative timing. The timely behaviour is a vital feature of the biology-mimicking systems and cannot be replaced with the synchronization principles of computing. One possible way is to put a "time grid" on the processes simulating biology. This requirement results in the neurons continuing their calculation periodically from some concerted state. Such a synchronization method introduces a "biological clock period" that is a million-fold longer than the processor's clock period. Although this effect drastically reduces the achievable computing temporal performance the synchronization principle is so common that also the special-purpose xeromorphic chips use it as a built-in

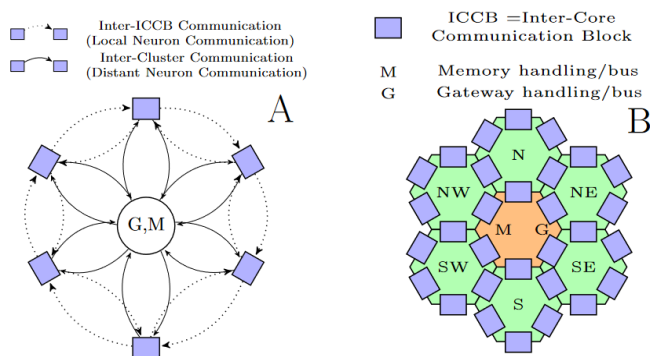


Figure 4. Subfigure-A shows how to reduce the limiting effect of the SPA, via mimicking the communication between local neurons using direct-wired inter-core communication and the communication between the farther neurons via using the inter-cluster communication bus, in the cluster head, and subfigure-B shows a specialized 'fat core' neuron.

feature. In their case, the speed of neuronal functionality is hundreds of times higher than that of the competing solutions, and the communication principles are slightly different (i.e., the non-payload/payload ratio is vastly different), the performance-limiting effect of the "quintal nature of computing time" persists when used in extensive systems.

The Role of the Workload on the Computing Efficiency

As was very early predicted [34] and decades later experimentally confirmed [15], the scaling of the parallelized computing is not linear. Even, "there comes a point when using more processors actually increases the execution time rather than reducing it" [15]. Paper by Keuper J and Preundt FJ [12] discusses first/second-order approaches to explain the issue. The first-order approach explains the experienced saturation and the second-order the predicted decrease. As Keuper J and Preundt FJ [12] discusses, the different workloads mainly due to their different communication-to-computation ratio, work with different efficiency on the same computer system. The neuromorphic operation on conventional architectures shows the same issues [19,20].

Limitations due to the classic computing paradigm

The careful analysis discovers a remarkable parallel between the proposed 'modern computing' vs. the classic computing and the modern science vs. the classic science. The parallel can help accept that what one cannot experience in every-day computing can be true when using computing under extreme conditions. Using another computing theory is a must, especially when targeting xeromorphic computing. In the frames of "classic computing", as was bitterly admitted "any studies on processes like plasticity learning and development exhibited over hours and days of biological time are outside our reach"

System is not a simple sum of its component

One must not conclude from a feature of a component to a similar feature of the system: the non-linearity discussed above are especially valid for the large-scale computing systems mimicking xeromorphic operation. We mention two prominent examples here. One can assume that if the time of the operation of a neuron can be shortened, the performance of the whole system gets proportionally better. Two distinct options are to use shorter operands (move less data and to perform lessbit manipulations) and to mimic the operation of the neuron using quick analog signal processing instead of slow digital calculation. The so-called HPL-Albenchmark used Mixed Precision rather than Double Precision operands in benchmarking their supercomputer. The name suggests as if in solving Ait asks the supercomputer can show that peak efficiency. We expect that when using half-precision (FP16) rather than double precision (FP64) operands in the calculations, four times less data are transferred and manipulated by the system. The measured power consumption data underpin our expectations. However, the computing performance is only three times higher than in the case of using 64-bit (FP64) operands. The non-linearity has its effect even in this simple case. In the benchmark, the housekeeping activity also takes time [16]. Another plausible assumption is that if we use quick analog signal processing to replace the slow digital calculation, as proposed in or using different materials/effects the system gets proportionally quicker. Adding analog components to a digital processor, however, has its price. Given that the digital processor cannot handle resources outside of its world; one must call the Operating System (OS) for help. The required context switching takes time in the order of executing 104 instructions which dramatically increases the total execution time and makes the non-payload to payload ratio much worse. Similarly, it is not reasonable to decrease processing speed if the corresponding transfer speed cannot be reduced. Although these cases seem to be very different, they share at least the common feature. They change not only one parameter they also change the non-payload to payload ratio that defines the efficiency. The analysis of their temporal behaviour (including their connection to the computing system) limits the utility of any new material/effect/technology; the detailed discussion.

Conclusion

The authors have identified some critical bottlenecks in current computational systems/neuronal networks rendering the conventional computing architectures unadoptable to large (and even medium) sized neuromorphic computing. Built with the segregated processor (SPA, wording from Amdahl) the current systems lack autonomous communication of processors and have an inefficient method of imitating biological systems; mainly their temporal behaviour

Acknowledgment

The datasets used and/or analysed during the current study are available from the corresponding author on reasonable request.

Data Availability

The datasets used and/or analysed during the current study are available from the corresponding author on reasonable request.

References

- Roy, Kaushik, Akhilesh Jaiswal, and Priyadarshini Panda. "Towards spike-based machine intelligence with neuromorphic computing." *Nature* 575 (2019): 607-617.
- Riyaz, Ahmed Mohammed, and B.K. Sujatha. "A review on methods, issues and challenges in neuromorphic engineering." In (2015) International Conference on Communications and Signal Processing (ICCSP), pp. 0899-0903.
- Végh, János, and Ádám J. Berki. "Do we know the operating principles of our computers better than those of our brain?" In (2020) International Conference on Computational Science and Computational Intelligence 668-674. IEEE (2020).
- Schuller, Ivan K., Rick Stevens, Robinson Pino, and Michael Pechan. Neuromorphic computing—from materials research to systems architecture roundtable. USDOE Office of Science (SC) (United States) (2015).
- Bell, Gordon, David H. Bailey, Jack Dongarra, Alan H. Karp, and Kevin Walsh. "A look back on 30 years of the Gordon Bell Prize." *IJHPCA* 6 (2017): 469-484.
- Furber, Steve, and Steve Temple. "Neural systems engineering." *J R Soc Interface* 13 (2007): 193-206.
- Végh, János, and Ádám J. Berki. "On the spatiotemporal behavior in biology-mimicking computing systems." arXiv preprint arXiv: 200.08841 (2020).
- Boahen, Kwabena A. "Point-to-point connectivity between neuromorphic chips using address events." *IEEE Transactions on Circuits and Systems II: Analog and Digital Signal Processing* 47 (2000): 416-434.
- Singh, Jaswinder Pal, John L. Hennessy, and Anoop Gupta. "Scaling parallel programs for multiprocessors: Methodology and examples." *Computer* 26 (1993): 42-50.
- Végh, János. "Which scaling rule applies to Artificial Neural Networks." arXiv preprint arXiv: (2005.): 08942.
- Végh, János. "How Amdahl's Law limits the performance of large artificial neural networks." *Brain Inform* 1 (2019): 1-11.
- Keuper, Janis, and Franz-Josef Preundt. "Distributed training of deep neural networks: Theoretical and practical limits of parallel scalability." In 2016 2nd Workshop on Machine Learning in HPC Environments (MLHPC), IEEE pp. 19-26.
- Waser, Rainer. Nanoelectronics and information technology: "Advanced Electronic Materials And Novel Devices". John Wiley & Sons (2012).
- Esmailzadeh, Hadi, Emily Blem, Renee St Amant, Karthikeyan Sankaralingam, and Doug Burger. "Dark silicon and the end of multicore scaling." *IEEE Micro* 32 (2012):122-134.
- Hameed Rehan, Wajahat Qadeer and Mark Horowitz, et al. "Understanding sources of inefficiency in general-purpose chips." In Proceedings of the 37th annual international symposium on Computer architecture (2010). 37-47.
- Williams, Samuel, Andrew Waterman, and David Patterson. "Roofline: an insightful visual performance model for multicore architectures." *Communications of the ACM* 52 (2009): 65-76.
- Gustafson, John L. "Reevaluating Amdahl's law." *Communications of the ACM* 31(1988): 532-533.
- Végh, János. "Introducing temporal behavior to computing science." In *Advances in Software Engineering, Education, and e-Learning*, Springer, Cham, pp. 471-491.
- Moradi, Saber, and Rajit Manohar. "The impact of on-chip communication on memory technologies for neuromorphic systems." *J Phys. D: Appl Phys* 52 (2018): 014003.
- Van Albada, Sacha J., Andrew G. Rowley and Steve B. Furber. "Performance comparison of the digital neuromorphic hardware SpiNNaker and the neural network simulation software NEST for a full-scale cortical microcircuit model." *Front Neurosci* (2018): 291.

How to cite this article: Veghi, Janos and Adam J. Berki. "Do We Know the Operating Principles of Our Computers Better than those of Our Anatomy of Brain?" *Morphol Anat* 6 (2022): 226.