

THE ECONOMICS OF ENERGY EFFICIENCY: HUMAN COGNITION VS. AI LARGE LANGUAGE MODELS

Babu GEORGE

Alcorn State University, MS - 39096, USA
bgeorge@alcorn.edu

Abstract

The human brain and Large Language Models (LLMs) exemplify two profoundly different information-processing systems, with a significant imbalance in energy consumption. This article provides a comparative analysis of the energy requirements of biological cognition and machine-based models, linking these differences to their economic consequences. The operational costs of LLMs, driven by energy-intensive computations, present financial and environmental sustainability challenges, influencing the scalability of AI adoption across industries. In contrast, the human brain's remarkable energy efficiency highlights the economic advantages of biological intelligence, which operates with minimal resource investment. We explore the economic trade-offs of AI's growing energy demands, including cost implications for firms, policy responses such as carbon regulations, and the potential labor market disruptions arising from energy-efficient automation. By integrating insights from neuroscience, engineering, and economics, we argue that sustainable AI development hinges on a fusion of brain-inspired paradigms, eco-efficient hardware solutions, and strategic policy frameworks that balance innovation with economic feasibility.

Key words: Energy economics; efficiency; Large Language Models; neuromorphic computing; sustainable AI; human cognition

JEL Classification: Q40, Q55, Q56,

I. INTRODUCTION

The metaphor “the brain is a Ferrari that sips fuel like a Prius, while LLMs are freight trains, powerful but thirsty” (Chemero, 2023) underscores a central conundrum in modern artificial intelligence (AI). As LLMs, exemplified by models such as GPT-4 and Claude 3, push toward trillions of parameters, their energy consumption becomes an increasingly significant concern for environmental sustainability and operational feasibility (IEEE, 2024).

In stark contrast, the human brain accomplishes complex cognitive tasks, from comprehending language to abstract reasoning, using roughly 20 watts, comparable to a dim lightbulb (Human Brain Project, 2023). This efficiency emerges from billions of years of evolutionary refinement. While LLMs can generate sophisticated text, translate languages, and even produce creative content, they currently do so at the cost of massive parallel processing, primarily employing graphics processing units (GPUs) or tensor processing units (TPUs) that demand exponentially higher energy inputs (arXiv, 2024).

This paper seeks to:

1. Compare the energy consumption of human cognition and LLMs through a common task, writing a 500-word essay.
2. Investigate the biological and computational underpinnings of their energy efficiency.
3. Analyze cutting-edge neuromorphic computing solutions and model optimization strategies.
4. Discuss the future of sustainable AI, focusing on the importance of embodied cognition and interdisciplinary collaboration.

II. ENERGY CONSUMPTION OF THE HUMAN BRAIN

The human brain is often described as a marvel of biological engineering, performing complex computations while consuming a surprisingly small amount of energy. Despite representing only about 2% of the body's mass, the brain utilizes 20% of the body's total energy expenditure (Herculano-Houzel, 2011). Its exceptional efficiency arises from optimized neural networks, efficient signal transmission, and effective energy metabolism, allowing the brain to maintain high levels of performance while conserving energy.

The human brain's energy efficiency is evident in its ability to allocate resources optimally. Neurons, the primary cells responsible for transmitting information, rely on minimal energy use during communication. Sparse coding, a mechanism that limits the number of active neurons during information processing, significantly reduces energy consumption (Laughlin & Sejnowski, 2003). Additionally, neural firing is

conserved by mechanisms that ensure only the most essential neurons are active at any given time, thereby limiting unnecessary energy expenditure (Herculano-Houzel, 2011). This balance between functionality and energy use highlights the brain's ability to perform sophisticated tasks without a proportional increase in energy requirements.

Neurons communicate through action potentials and synaptic transmission, both of which involve the movement of ions across cell membranes. The brain achieves efficiency in this process through the recycling of ions. Ion channels are structured to conserve energy by reducing the overall metabolic cost of maintaining resting and active states (Attwell & Laughlin, 2001). This mechanism allows the brain to handle a vast amount of data processing without overwhelming its energy reserves, further contributing to its remarkable efficiency.

The brain's energy demands are primarily met through the oxidation of glucose, a process that maximizes energy output. Glucose oxidation is highly efficient, producing more energy per molecule than other metabolic pathways (Clarke & Sokoloff, 1999). This efficient metabolism enables the brain to support activities such as reasoning, memory, and sensory processing without experiencing significant energy deficits. Moreover, the brain's reliance on glucose highlights its ability to sustain high levels of cognitive activity with minimal resource requirements.

The architecture of the human brain also plays a critical role in energy efficiency. As the brain has evolved, its neurons have been scaled to optimize energy use while maintaining high computational capacity. This optimization allows the human brain to manage a large number of neurons and synapses without proportional increases in energy consumption (Herculano-Houzel, 2016). This scaling ensures that the brain remains efficient even as it processes more complex information.

III. ENERGY CONSUMPTION OF LLMs

In recent years, large language models (LLMs) have gained widespread adoption for natural language processing (NLP) tasks. While these models have significantly advanced the state of AI, they require substantial computational resources, leading to considerable energy consumption. Understanding the energy demands of LLMs is critical for assessing their environmental impact and guiding the development of more sustainable AI technologies.

The energy consumption of LLMs primarily stems from their computational intensity during both training and inference phases. Training a large model involves processing vast amounts of data over multiple iterations, which requires extensive GPU and TPU usage. For instance, training GPT-3 required 175 billion parameters and consumed an estimated 1,287 MWh of electricity, emitting approximately 552 tons of CO₂ (Patterson et al., 2021). This scale of energy use is attributed to the sheer number of calculations needed to optimize the model across billions of parameters.

The carbon footprint of LLMs is significant due to their reliance on energy-intensive hardware. The environmental impact depends on the energy source used to power data centers. If renewable energy sources are utilized, the carbon emissions can be significantly reduced (Henderson et al., 2020). However, data centers in regions relying on fossil fuels contribute substantially to greenhouse gas emissions, raising concerns about the sustainability of widespread LLM adoption.

To address these concerns, researchers have explored ways to reduce the energy demands of LLMs. Techniques such as model distillation, parameter sharing, and pruning have been developed to make models smaller and more efficient without compromising performance (Sanh et al., 2019). Additionally, innovations in hardware design, such as the development of energy-efficient GPUs and TPUs, have also contributed to lowering energy consumption. Companies like Google and OpenAI are investing in optimizing training pipelines to further minimize environmental impact.

Despite improvements in efficiency, there is a trade-off between performance and sustainability. Larger models tend to perform better in NLP benchmarks but at the cost of higher energy consumption. Balancing these priorities requires a focus on not only creating powerful models but also addressing their environmental impact. Policies to regulate energy use in AI development and incentives for using renewable energy sources in data centers could support this goal (Bender et al., 2021).

IV. EVOLUTIONARY OPTIMIZATION OF THE HUMAN BRAIN

The human brain operates at approximately 20 watts, sustaining an estimated 80–100 billion neurons and trillions of synaptic operations (Chemero, 2023). Despite handling diverse tasks, language comprehension, motor coordination, and memory retrieval, its baseline power requirement remains remarkably low due to evolutionary adaptations for survival, metabolic efficiency, and robust learning mechanisms.

When engaging in language-based tasks (e.g., writing an essay), activation patterns are predominantly localized to language centers (Broca's and Wernicke's areas) and executive function regions, minimizing redundant energy usage (Human Brain Project, 2023). Neurons often fire sparsely, utilizing spikes only when necessary, further conserving energy.

Biological neurons operate analogously, relying on chemical gradients and ion channels, thereby avoiding the overhead inherent in digital binary operations. This analog nature permits continuous modulations of signal strength, enabling more efficient transmission and processing of information.

For a one-hour essay-writing session, the brain's total energy usage is roughly 0.02 kWh, signifying an extremely efficient computing apparatus honed by evolution to meet metabolic constraints (Human Brain Project, 2023).

V.COMPUTATIONAL INTENSITY OF THE LLMs

LLMs employ massive parallelism to generate text. Even a seemingly simple task, such as producing a 500-word essay, can demand around 2.9 Wh per inference query. Iterative refinements, common in interactive usage, compound this figure to about 0.029 kWh per user session (IEEE, 2024).

The true magnitude of LLM energy usage becomes apparent when scaled to millions of daily users. A platform serving 1 million users for essay-writing tasks each day may consume approximately 29,000 kWh, equating to the daily energy requirements of over 2,700 U.S. households (arXiv, 2024).

While inference is substantial, training costs are even higher. A single training run for GPT-3 was estimated to emit 552 tons of CO₂, comparable to the annual emissions of 123 gasoline-powered cars (IEEE, 2024). Such impacts underscore the urgency of devising more energy-efficient training and deployment methods.

VI.NEUROMORPHIC COMPUTING

Neuromorphic computing is an emerging field focused on replicating the structure and function of the human brain in computational systems. By drawing inspiration from biological neural networks, this approach aims to overcome the limitations of traditional computing, particularly in energy efficiency, adaptability, and processing power. Its promise lies in revolutionizing artificial intelligence (AI) through more flexible, efficient systems with reduced energy demands.

Pioneered by Carver Mead in the 1980s (Mead, 1990), neuromorphic computing borrows principles from neuroscience to design artificial neural systems. Unlike the conventional von Neumann architecture, which separates processing and memory and creates bottlenecks for large datasets, neuromorphic architectures integrate these functionalities in a distributed manner (Indiveri & Liu, 2015). This integration mirrors biological brains, enabling parallel and efficient data handling.

A central feature of neuromorphic systems is their use of spiking neural networks (SNNs), which emulate the asynchronous, event-driven nature of brain activity. Instead of relying on continuous activation functions like traditional artificial neural networks, SNNs communicate through discrete spikes, closely reflecting how neurons transmit signals (Roy et al., 2019). This event-driven mechanism reduces energy consumption and enhances real-time processing.

Advantages of Neuromorphic Computing

1. *Energy Efficiency*

Neuromorphic systems often consume far less energy than conventional AI hardware. Their event-driven design ensures that power is used only during active computations, thereby lowering overall energy usage. For instance, Intel's Loihi chip consumes orders of magnitude less power than traditional GPUs when running AI tasks (Davies et al., 2021).

2. *Real-Time Processing*

Thanks to their parallel architectures and low latency, neuromorphic systems excel at real-time processing of sensory data, such as visual and auditory signals. They are particularly useful in fields like robotics, autonomous vehicles, and edge computing (Schuman et al., 2017).

3. *Adaptability and Learning*

Neuromorphic systems inherently adapt to changing environments. They can perform on-device learning using local rules like spike-timing-dependent plasticity (STDP), mimicking how synapses strengthen or weaken based on activity patterns in the brain (Furber, 2016). This adaptability contrasts with static algorithms, which lack ongoing learning capabilities.

Applications of Neuromorphic Computing

1. *Artificial Intelligence*

By enabling energy-efficient, high-performance AI, neuromorphic computing is a key component of sustainable AI development. Its low-power profile makes it especially suitable for edge AI tasks requiring real-time responsiveness.

2. *Robotics*

In robotics, neuromorphic systems enhance sensory-motor integration. Event-based vision sensors, derived from neuromorphic principles, allow robots to process high-speed motion data with minimal energy consumption (Indiveri et al., 2011).

3. *Healthcare*

Neuromorphic architectures are naturally aligned with brain-machine interfaces and neural prosthetics. They can handle real-time neural signal processing, paving the way for advancements in neurorehabilitation and assistive devices (Qiao et al., 2015).

Despite its promise, neuromorphic computing faces notable challenges. The lack of standardized programming tools and frameworks complicates development and scaling, while faithfully replicating biological neurons and synapses in hardware remains a technical hurdle (Schuman et al., 2017). Progress in materials science, such as the development of memristors, and integration with emerging technologies like quantum computing will likely catalyze future breakthroughs. As the field advances, neuromorphic computing is expected to play a pivotal role in shaping next-generation AI and computing systems.

Neuromorphic Computing at the Hardware Level

Neuromorphic computing aims to imitate the brain's structure and function by adopting principles like spike-based communication and local memory storage (Human Brain Project, 2023). This biologically inspired design holds significant potential for reducing power consumption. Notable examples include:

1. *SpiNNaker2*

- Comprises 10 million cores arranged to mimic biological neural networks.
- Achieves up to a 16× reduction in energy usage compared to conventional GPUs for certain tasks (Human Brain Project, 2023).

2. *BrainScaleS*

- Employs analog circuits and spiking neural networks to replicate neuronal firing patterns.
- Minimizes data redundancy and dynamic power consumption by executing computations only when spikes occur.

3. *Algorithmic Innovations*

- Spiking neural networks and evolutionary algorithms mimic synaptic plasticity, potentially reducing energy costs by 80–95% compared to standard deep learning methods.
- These approaches also enable continuous on-chip learning, reducing the need for large-scale retraining cycles.

By aligning physical computational processes with biologically inspired models, neuromorphic chips can closely approximate the remarkable energy efficiencies observed in natural neural networks (arXiv, 2024).

VII. MISCELLANEOUS ENERGY OPTIMIZATION STRATEGIES OF LLMs

Even in the absence of fully neuromorphic hardware, several strategies can mitigate LLMs' energy-intensive profile:

Model Compression

- Techniques such as pruning, weight quantization, and knowledge distillation can decrease model size without significantly degrading performance.
- Recent developments in ultra-low-bit quantization (e.g., 1.58-bit algorithms) have demonstrated substantial power savings while retaining acceptable inference accuracy (IEEE, 2024).

Hardware-Software Co-Design

- Energy-aware inference protocols that adjust clock rates and voltage levels based on real-time computational demands can significantly cut power use.
- Customized AI accelerators (e.g., TPU ASICs) can be designed to handle specific model architectures, further reducing overhead.

Renewable Integration

- Shifting data centers to rely on solar, wind, or nuclear power can lower the carbon footprint by 40–60%.
- Intelligent scheduling of AI workloads to off-peak hours or renewable energy surpluses can further optimize resource utilization (arXiv, 2024).

VIII. THE EMBODIMENT GAP

A fundamental limitation of current large language models (LLMs) is their disembodied nature: they learn exclusively from massive textual datasets without any direct sensory or motor experiences (Chemero, 2023). In contrast, human cognition and language evolve through constant interaction with the physical and social environment, forming neural pathways optimized for context-specific understanding.

Embodiment is central to human cognition, shaping how we acquire knowledge, develop language, and navigate the world. Through sensory and motor experiences, humans ground abstract concepts in physical reality (Barsalou, 2008). For instance, understanding the term "heavy" involves not just linguistic knowledge but also the physical experience of lifting or carrying objects. This embodied learning enables humans to connect language intuitively with real-world situations.

LLMs, by comparison, lack such grounding. They learn from textual patterns alone, unable to interact with the environment or incorporate sensory-motor feedback. While LLMs can replicate human language and infer text-based patterns, they do not possess the experiential foundation that gives language its real-world meaning. As a result, LLMs often struggle with tasks requiring an understanding of spatial relationships, physical causality, or sensory experiences, frequently producing responses that lack accuracy or contextual relevance.

Critical Limitations of Disembodiment

The absence of embodied experiences in LLMs creates several key challenges:

1. Lack of Contextual Grounding

Without sensory and motor experiences, LLMs cannot fully comprehend context beyond textual patterns. For example, they may generate plausible-sounding yet incorrect descriptions of physical phenomena due to their lack of interaction with the real world (Bender & Koller, 2020). This limitation significantly hinders their application in domains like robotics or scientific reasoning, where deep contextual understanding is essential.

2. Challenges in Commonsense Reasoning

Commonsense reasoning often depends on embodied experiences, such as knowing that ice is slippery or that liquids flow downward. While LLMs can infer statistical correlations from text, they frequently fail to generalize these correlations to novel scenarios or align reasoning with real-world constraints (Marcus & Davis, 2020).

3. Ethical Implication

The lack of embodied understanding also raises ethical concerns. Without sensory or experiential grounding, LLMs lack the intuitive grasp of human emotions and values that inform decision-making. This limitation complicates efforts to ensure ethical behavior and alignment with human intentions in complex real-world contexts (Floridi & Chiriatti, 2020).

Toward Embodied AI

Addressing the limitations of disembodied LLMs involves integrating sensory and motor experiences into AI systems. Embodied AI research explores how agents equipped with sensors, cameras, and actuators can learn by interacting with their environments. Such systems develop a grounded understanding by associating sensory inputs with physical actions and outcomes. For instance, a robot could learn the concept of "fragility" by handling objects of varying durability, connecting the term with direct experiential knowledge (Chen et al., 2021).

Integrating LLMs with embodied systems offers a promising pathway to bridge the gap between language and experience. For example, pairing LLMs with robotics could enable systems to interpret textual instructions while using sensory data to ground those instructions in physical reality. Similarly, virtual environments and simulations could serve as controlled settings for AI agents to acquire embodied knowledge without the risks associated with real-world experimentation.

Scaling vs. Grounding

As LLMs grow in scale, incorporating trillions of parameters, their reliance on brute-force learning demands exponentially higher computational resources. By contrast, embodied cognition suggests that tying language comprehension to physical experience could reduce the need for massive, undifferentiated training datasets, thereby lowering energy consumption (Human Brain Project, 2023).

Hybrid Architectures

Insights from neuroscience highlight the brain's efficiency as a result of its integrative architecture, where perception, action, and cognition are deeply interconnected. Next-generation AI architectures could emulate this efficiency by combining neuromorphic hardware with embodied or "sensorized" learning frameworks. Such systems would capture real-time environmental data to refine learning processes with fewer training examples, paving the way for more efficient and contextually aware AI.

IX. THE ECONOMIC IMPLICATIONS OF ENERGY EFFICIENCY IN AI

The stark contrast in energy efficiency between the human brain and large language models (LLMs) has significant economic ramifications. Energy costs constitute a major portion of AI deployment expenses, impacting firms' operational budgets and influencing the broader AI industry's scalability. As data centers increasingly power LLMs, regions with higher energy prices may experience disproportionate financial burdens, potentially exacerbating global economic inequalities in AI accessibility. In contrast, the biological efficiency of human cognition underscores the economic advantages of natural intelligence, which requires minimal resource investment for cognitive processing.

From a macroeconomic perspective, the rapid expansion of energy-hungry AI models raises concerns about sustainable economic growth. Governments and policymakers face trade-offs between incentivizing AI-driven productivity gains and managing environmental and infrastructure costs. Carbon taxes, energy regulations, and green AI initiatives are emerging policy tools to mitigate these economic pressures. Additionally, firms developing neuromorphic computing and energy-efficient AI models may gain competitive advantages by reducing long-term operational costs, thus shaping future industry dynamics.

The labor market implications of AI's energy efficiency divide also warrant attention. While LLMs automate numerous cognitive tasks, their high energy costs could restrict widespread adoption, preserving human labor's role in knowledge-intensive industries. Conversely, advancements in energy-efficient AI might accelerate automation, necessitating policy responses such as workforce reskilling programs and AI taxation mechanisms to offset economic displacement.

X. CONCLUSION

The human brain's 20-watt paradigm starkly contrasts with the energy-intensive nature of today's LLMs. This discrepancy highlights the potential for a new wave of AI systems that prioritize efficiency alongside

accuracy and scale. Neuromorphic computing offers a promising hardware foundation, while techniques such as model compression, hardware-software co-design, and renewably powered data centers can incrementally reduce the carbon footprint of LLMs.

Addressing the embodiment gap remains pivotal. By infusing AI models with sensory and contextual grounding, we may move beyond brute-force parameter scaling, thus aligning large-scale computation with the inherent resourcefulness of biological systems. Ultimately, the path toward sustainable AI depends on interdisciplinary collaborations, spanning neuroscience, computer engineering, policy-making, and environmental science, to ensure that the “Ferrari” efficiency of the brain informs and transforms how we build the “freight trains” of modern AI.

XI. REFERENCES SECTION

1. Chemero, A. (2023). LLMs differ from human cognition because they are not embodied. *Nature Human Behaviour*, 7, 1828–1829. <https://doi.org/10.1038/s41562-023-01723-5>
2. Human Brain Project. (2023, September 4). *Learning from the brain to make AI more energy-efficient*. <https://www.humanbrainproject.eu/en/follow-hbp/news/2023/09/04/learning-brain-make-ai-more-energy-efficient/>
3. IEEE. (2024). Measuring and Improving the Energy Efficiency of Large Language Models Inference. *IEEE Xplore*. <https://ieeexplore.ieee.org/document/10549890>
4. arXiv. (2024). Towards Greener LLMs: Bringing Energy-Efficiency to the Forefront of LLM Inference. <https://arxiv.org/abs/2403.20306>
5. Attwell, D., & Laughlin, S. B. (2001). Energetic cost of neural computation. *Nature Reviews Neuroscience*, 2(11), 695–703.
6. Clarke, D. D., & Sokoloff, L. (1999). Glucose utilization and cognitive demands. *Journal of Neurochemistry*, 67(6), 2037–2043.
7. Herculano-Houzel, S. (2011). The remarkable energetic efficiency of the human brain. *PNAS*, 108(6), 2322–2327.
8. Herculano-Houzel, S. (2016). Brain scaling and energy efficiency in humans. *Frontiers in Human Neuroscience*, 10, 185.
9. Laughlin, S. B., & Sejnowski, T. J. (2003). Communication in neurons: Optimization for energy. *Nature Reviews Neuroscience*, 4(8), 683–694.
10. Bender, E. M., Gebu, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623. <https://doi.org/10.1145/3442188.3445922>
11. Henderson, P., Hu, J., Romoff, J., Brunskill, E., Jurafsky, D., & Pineau, J. (2020). Towards the systematic reporting of the energy and carbon footprints of machine learning. *Journal of Machine Learning Research*, 21(248), 1–43. <https://arxiv.org/abs/2002.05651>
12. Patterson, D., Gonzalez, J., Le, Q., Liang, C., Munguia, L., Rothchild, D., & Dean, J. (2021). Carbon emissions and large neural network training. *Proceedings of the International Conference on Learning Representations*. <https://arxiv.org/abs/2104.10350>
13. Sanh, V., Debut, L., Chaumond, J., & Wolf, T. (2019). DistilBERT, a distilled version of BERT: Smaller, faster, cheaper, and lighter. *Proceedings of the 5th Workshop on Energy Efficient Machine Learning and Cognitive Computing at NeurIPS 2019*. <https://arxiv.org/abs/1910.01108>
14. Barsalou, L. W. (2008). Grounded cognition. *Annual Review of Psychology*, 59, 617–645. <https://doi.org/10.1146/annurev.psych.59.103006.093639>
15. Bender, E. M., & Koller, A. (2020). Climbing towards NLU: On meaning, form, and understanding in the age of data. *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, 5185–5198. <https://doi.org/10.18653/v1/2020.acl-main.463>
16. Chen, T., Bao, J., & Zhu, S. (2021). Towards embodied AI: Learning to interact with environments. *Proceedings of the AAAI Conference on Artificial Intelligence*. <https://arxiv.org/abs/2106.13935>
17. Floridi, L., & Chiriatti, M. (2020). GPT-3: Its nature, scope, limits, and consequences. *Minds and Machines*, 30(4), 681–694. <https://doi.org/10.1007/s11023-020-09548-1>
18. Marcus, G., & Davis, E. (2020). GPT-3, bloviator: OpenAI’s language generator has no idea what it’s talking about. *MIT Technology Review*. <https://www.technologyreview.com/2020/08/22/1007539/gpt-3-openai-language-generator-artificial-intelligence-ai-opinion/>
19. Davies, M., Srinivasa, N., Lin, T. H., China, G., Cao, Y., Joshi, P., & Dixit, N. (2021). Loihi: A neuromorphic manycore processor with on-chip learning. *IEEE Micro*, 38(1), 82–99. <https://doi.org/10.1109/MM.2018.112130359>

19. Furber, S. (2016). Large-scale neuromorphic computing systems. *Journal of Neural Engineering*, 13(5), 051001. <https://doi.org/10.1088/1741-2560/13/5/051001>
20. Indiveri, G., & Liu, S. C. (2015). Memory and information processing in neuromorphic systems. *Proceedings of the IEEE*, 103(8), 1379–1397. <https://doi.org/10.1109/JPROC.2015.2444094>
21. Indiveri, G., Linares-Barranco, B., Hamilton, T. J., Van Schaik, A., Etienne-Cummings, R., Delbruck, T., & Liu, S. C. (2011). Neuromorphic silicon neuron circuits. *Frontiers in Neuroscience*, 5, 73. <https://doi.org/10.3389/fnins.2011.00073>
22. Mead, C. (1990). Neuromorphic electronic systems. *Proceedings of the IEEE*, 78(10), 1629–1636. <https://doi.org/10.1109/5.58356>
23. Qiao, N., Mostafa, H., Corradi, F., Osswald, M., Stefanini, F., Sumislawska, D., & Indiveri, G. (2015). A reconfigurable on-line learning spiking neuromorphic processor. *Frontiers in Neuroscience*, 9, 141. <https://doi.org/10.3389/fnins.2015.00141>
24. Roy, K., Jaiswal, A., & Panda, P. (2019). Towards spike-based machine intelligence with neuromorphic computing. *Nature*, 575(7784), 607–617. <https://doi.org/10.1038/s41586-019-1677-2>