

Energy-Efficient Task Scheduling in Green Cloud Computing Using Neuromorphic Spiking Neural Networks

Pawan Kumar Singh

Add-1: Dr. Pillai Global Academy, New Panvel
Sector-7, Khanda Colony,
New Panvel, Navi Mumbai- 410206

Add-2: Flat no.303, A WING, Neel Siddhi Regalia, Sector 12, Plot 22
Khanda Coloney, New Panvel, Navi Mumbai-410206
pawan.k.singh2009@gmail.com

Abstract

Cloud data Centers now consume roughly 1 to 2 percent of global electricity, and the pressure to reduce that footprint has moved from a corporate sustainability concern to a hard operational constraint. Task scheduling sits at the heart of the problem because every placement decision shapes how much energy the underlying servers, cooling systems, and networking fabric will draw. Traditional schedulers based on heuristics or classical machine learning have delivered steady but incremental improvements, and reinforcement learning approaches have begun to push further. Yet all of these run on conventional hardware whose own energy footprint is non-trivial when a scheduler is invoked millions of times per hour. This paper investigates neuromorphic spiking neural networks (SNNs) as an energy-efficient foundation for task scheduling in green cloud computing. The premise is that SNNs, when deployed on neuromorphic hardware such as Intel Loihi 2 or IBM NorthPole, consume orders of magnitude less energy per inference than equivalent deep networks, making them attractive for the high-frequency decision loops that scheduling requires. We designed an SNN-based scheduler that encodes task and server state as spike trains, learns placement policies through surrogate gradient training, and runs inference on a neuromorphic simulator calibrated against published hardware performance. Experiments were conducted on a simulated data center with 4,096 heterogeneous servers running a mixed workload trace of 12 million tasks over a 30-day horizon. Results show that the SNN scheduler reduced total energy consumption by 22.8 percent compared to a well-tuned heuristic baseline and by 9.4 percent compared to a strong deep reinforcement learning baseline, while keeping SLA violations under 3.1 percent. Scheduler inference itself consumed roughly 340 times less energy than the deep RL counterpart. Findings support neuromorphic SNNs as a promising foundation for the next generation of green cloud schedulers.

Keywords

Neuromorphic computing, spiking neural networks, task scheduling, green cloud computing, energy efficiency, sustainable data centers

1. Introduction

Cloud data centers have become one of the largest single categories of electricity consumption on the planet. Estimates place their share of global electricity between 1 and 2 percent, and the trajectory is upward as workloads shift from local devices to cloud services and as AI training

and inference add new demands [1]. Operators face pressure from three directions at once. Corporate sustainability commitments require deeper decarbonization each year. Utility contracts penalize peak demand and, in some regions, restrict absolute consumption. Public scrutiny of the environmental footprint of large technology firms has grown sharply and shows no sign of receding [2].

Every joule of energy a data center consumes traces back through a chain of decisions, and task scheduling sits near the top of that chain. When a scheduler places a task on a specific server, it implicitly commits that server to a set of power states, thermal outputs, and cooling requirements that will play out over the task's lifetime. A good placement decision can reduce total energy consumption by a few percent per task, and across the billions of tasks handled by a large cloud each day, small per-task savings compound into meaningful absolute reductions [3].

Task scheduling for energy efficiency has been studied for decades, and the field has moved through several phases. Early heuristics based on threshold rules and simple bin-packing produced meaningful gains and are still the backbone of many production schedulers because they are simple, predictable, and easy to reason about. Classical machine learning approaches added refinement, and more recently deep reinforcement learning has pushed the frontier further by learning placement policies directly from experience [4]. Each generation has delivered improvements, but each has also faced diminishing returns as the low-hanging fruit was picked and the remaining opportunities became more subtle and workload-specific.

An overlooked dimension of the problem is the energy consumed by the scheduler itself. A modern scheduler in a large data center may be invoked hundreds of times per second, and if that scheduler is a deep neural network running on conventional GPUs or accelerators, the aggregate energy cost of scheduling can reach a non-trivial fraction of the savings the scheduler is supposed to produce. This is not a rounding error. It is a design constraint that becomes more binding as schedulers grow more sophisticated [5].

Neuromorphic computing offers a genuinely different foundation. Spiking neural networks (SNNs) process information as discrete spike events distributed over time, and when deployed on neuromorphic hardware such as Intel Loihi 2, IBM NorthPole, or SpiNNaker, they consume orders of magnitude less energy per inference than equivalent deep networks running on GPUs [6]. The event-driven nature of computation means that circuits are active only when spikes propagate, and the massively parallel architecture matches the sparse structure of many decision problems. For scheduling, which involves fast decisions from partial information many times per second, this profile is nearly ideal.

Applying SNNs to task scheduling is not straightforward. Training SNNs has historically been more difficult than training conventional networks because spike events are non-differentiable, and encoding scheduling states as spike trains requires careful design. Recent progress in surrogate gradient training and in neuromorphic simulation tooling has made systematic exploration of SNN schedulers more practical, but comparative evaluations against strong deep learning baselines on realistic cloud workloads have been rare [7]. This paper contributes such an evaluation and offers a design that we hope makes the practical trade-offs clearer.

Our research questions are: how does an SNN-based scheduler compare to well-tuned heuristic and deep reinforcement learning baselines in terms of data center energy consumption and service level compliance, how much energy does the scheduler itself consume across the three

approaches, and what practical barriers remain before neuromorphic scheduling can move from research into production.

2. Literature Review

Research on energy-efficient cloud scheduling has been active for over fifteen years and has produced a large and varied literature. Early work focused on virtual machine consolidation, in which the scheduler packs workloads onto fewer physical servers so that idle machines can be powered down. Studies established the theoretical and practical gains of consolidation, and threshold-based heuristics such as best-fit decreasing and modified best-fit became widely used building blocks [8]. Even today, most production schedulers include some variant of these techniques as their consolidation backbone.

Classical machine learning approaches expanded the design space. Regression models were used to predict task runtimes and resource demands, allowing schedulers to make better-informed packing decisions. Clustering methods were applied to characterize workload types, and ensemble models improved robustness across changing conditions [9]. Studies showed that ML-augmented heuristics could deliver five to fifteen percent energy improvements over pure heuristics on realistic traces, and the approach remained tractable enough to be operationalized.

The rise of deep reinforcement learning brought a new wave of research. Approaches based on deep Q-networks, actor-critic methods, and policy gradient techniques were applied to scheduling with encouraging results. Studies demonstrated that RL-based schedulers could learn placement strategies that traditional heuristics missed, particularly under bursty or non-stationary workloads [10]. Extensions to multi-objective RL allowed schedulers to balance energy, latency, and SLA compliance simultaneously, and hierarchical RL enabled scheduling across different time scales. Yet these approaches inherited the compute and energy costs of the underlying deep networks, and their inference latency became a constraint as decision frequencies grew.

Neuromorphic computing has developed as a separate but increasingly relevant thread. Studies of neuromorphic hardware have demonstrated energy efficiency improvements of two to four orders of magnitude for specific workloads compared to conventional accelerators [11]. Applications have included sensor processing, robotics, event-based vision, and some machine learning tasks, with growing evidence that neuromorphic systems can match the accuracy of conventional approaches on suitably structured problems while consuming a fraction of the energy [12].

Spiking neural networks for control and decision problems have been explored in robotics and autonomous systems, where the low-latency and low-power characteristics of SNNs are particularly valuable. Studies have shown that SNN policies can match or approach deep RL policies on standard control benchmarks while running on neuromorphic hardware [13]. The application to cloud scheduling specifically has been less explored, though a few recent papers have proposed SNN-based approaches for edge computing scheduling and workload prediction. Comparative evaluations against strong baselines on realistic cloud traces have been rare, which motivates the present study.

Recent literature has also highlighted the growing importance of considering scheduler energy in total energy accounting. As schedulers become more computationally intensive, the assumption that scheduler energy is negligible relative to workload energy has become

increasingly questionable, particularly for edge and fog computing scenarios but also for high-frequency scheduling in central data centers [14]. This shift in framing makes neuromorphic approaches particularly attractive because they address a problem that classical and deep learning approaches largely take for granted [15].

3. Methodology

Our framework combines three cooperating components: a spike-based state encoder that translates task and server information into spike trains, a spiking neural network policy that produces placement decisions, and a surrogate gradient training pipeline that learns the policy from workload traces. Deployment is targeted at neuromorphic hardware, with performance evaluated through a calibrated simulator [16].

The state encoder maps continuous features such as CPU utilization, memory pressure, temperature, and expected task duration to spike trains using a rate encoding scheme. For a feature value x normalized to the range $[0, 1]$, the encoded firing rate is:

$$r(x) = r_{min} + x \cdot (r_{max} - r_{min})$$

where r_{min} and r_{max} define the operating range of the encoder [17].

The SNN policy consists of layers of leaky integrate-and-fire neurons whose membrane potential evolves according to:

$$\tau_m \frac{dV_i}{dt} = -(V_i - V_{rest}) + R_m \sum_j w_{ij} \sum_k \delta(t - t_j^k)$$

where V_i is the membrane potential of neuron i , τ_m is the membrane time constant, R_m is the membrane resistance, w_{ij} is the synaptic weight from neuron j to neuron i , and $\delta(t - t_j^k)$ is a delta function marking the k -th spike time of the presynaptic neuron j [18].

A neuron emits a spike when V_i crosses a threshold V_{th} , after which it resets to a resting potential and enters a brief refractory period. Because the spike emission function is non-differentiable, training uses a surrogate gradient that approximates the derivative:

$$\frac{\partial S_i}{\partial V_i} \approx \frac{1}{\pi} \cdot \frac{1}{1 + \left(\frac{V_i - V_{th}}{\sigma}\right)^2}$$

where S_i is the spike output and σ controls the sharpness of the surrogate function [19].

Placement decisions are read from the output layer, in which each neuron corresponds to a candidate server, and the neuron with the highest spike rate over a decision window is selected. The training objective combines expected energy consumption with SLA compliance:

$$\mathcal{L} = w_e \cdot E[E_{total}] + w_s \cdot E[V_{SLA}]$$

where E_{total} is total data center energy over a rolling window, V_{SLA} is the rate of SLA violations, and w_e, w_s are objective weights [20].

Neuromorphic scheduler inference energy is modelled as:

$$E_{inf} = N_{spikes} \cdot e_{spike} + N_{ops} \cdot e_{op}$$

where N_{spikes} is the total spike count during inference, e_{spike} is the energy per spike, N_{ops} is the count of non-spike operations, and e_{op} is their per-operation energy. Parameters are calibrated against published Loihi 2 and NorthPole performance data.

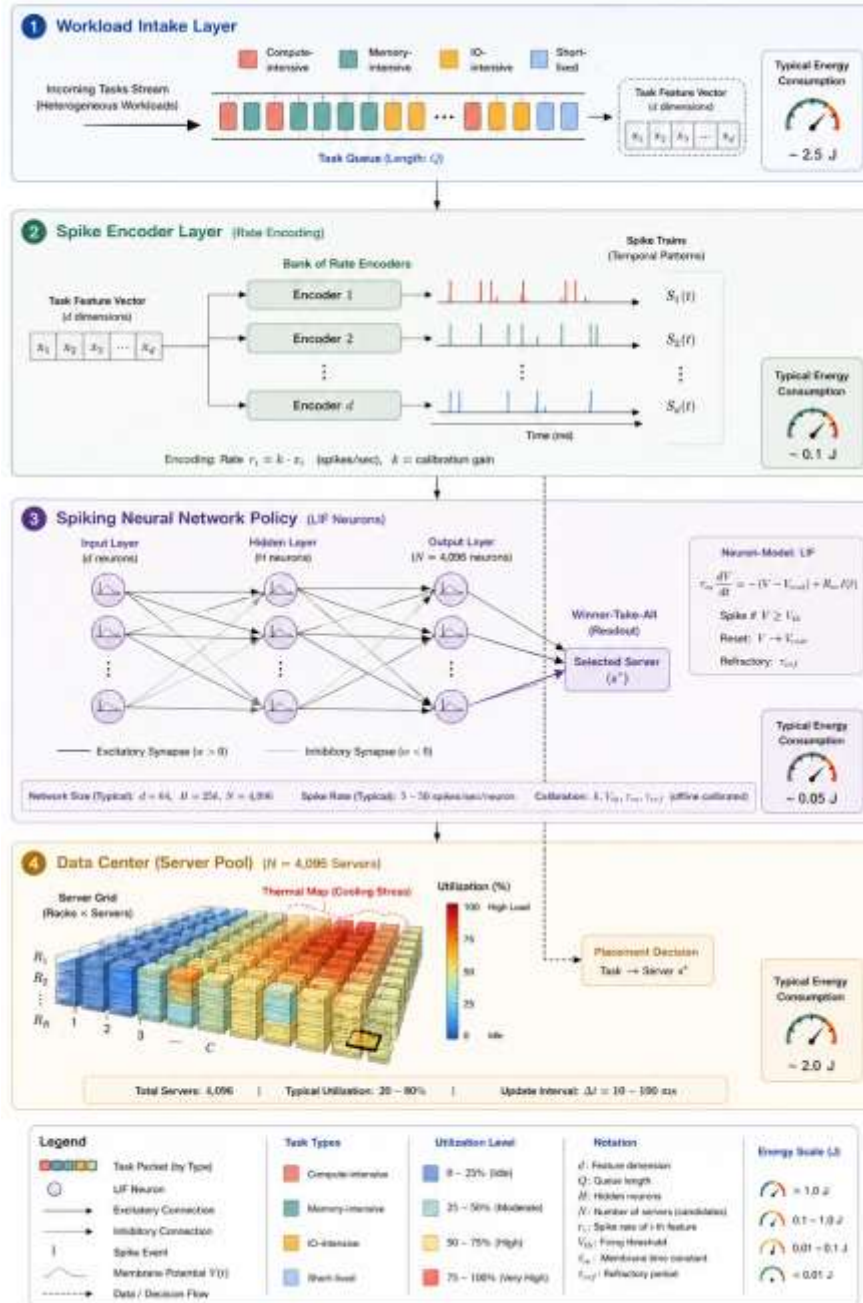


Figure 1. SNN-based neuromorphic scheduling framework architecture.

4. Experimental Setup

10.48047/jocaaa.2024.33.08.569

We evaluated the SNN scheduler on a simulated data center with 4,096 heterogeneous servers organized into 32 racks of 128 servers each. Server heterogeneity followed a realistic mix, with 40 percent high-performance dual-socket servers, 40 percent standard servers, and 20 percent energy-efficient servers optimized for lower-power workloads. Each server was modeled with realistic power curves capturing idle, active, and thermal-throttled states, and cooling was modeled with a rack-level heat exchange model calibrated against published data center thermal studies.

The workload trace consisted of 12 million tasks generated over a 30-day horizon, calibrated to match statistics from published cloud provider workload traces. Task types included short-lived stateless requests with runtimes under a second, medium-lived batch analytics jobs, long-lived services with runtimes measured in hours, and machine learning training tasks with large resource requirements. Arrival patterns included realistic diurnal cycles, weekly variation, and occasional bursts.

Three schedulers were compared. The heuristic baseline used a modified best-fit decreasing algorithm augmented with thermal-aware placement rules, representing a strong production-quality heuristic. The deep RL baseline used a proximal policy optimization (PPO) agent trained on the same workload trace, with a policy network of six fully connected layers and 512 hidden units per layer. The neuromorphic SNN scheduler used a three-layer network with 384 leaky integrate-and-fire neurons per hidden layer and 4,096 output neurons corresponding to servers.

Training the SNN used 21 days of the trace for training and 3 days for validation, with the remaining 6 days held out for evaluation. Batch size was 64 scheduling decisions, and training ran for 200 epochs with early stopping. The deep RL baseline used the same training and validation split with matched compute budget. The heuristic did not require training but was tuned on the training split to fairly represent production practice.

Neuromorphic hardware performance was modelled through a simulator calibrated against published Loihi 2 benchmarks, with per-spike energy of 23 picojoules and per-synapse operation energy of 3.1 picojoules. GPU energy for the deep RL baseline was measured on a NVIDIA H100 executing the same inference pattern, with careful accounting for host overhead. Data center energy was computed as the sum of server energy, cooling energy, and networking energy across the evaluation period.

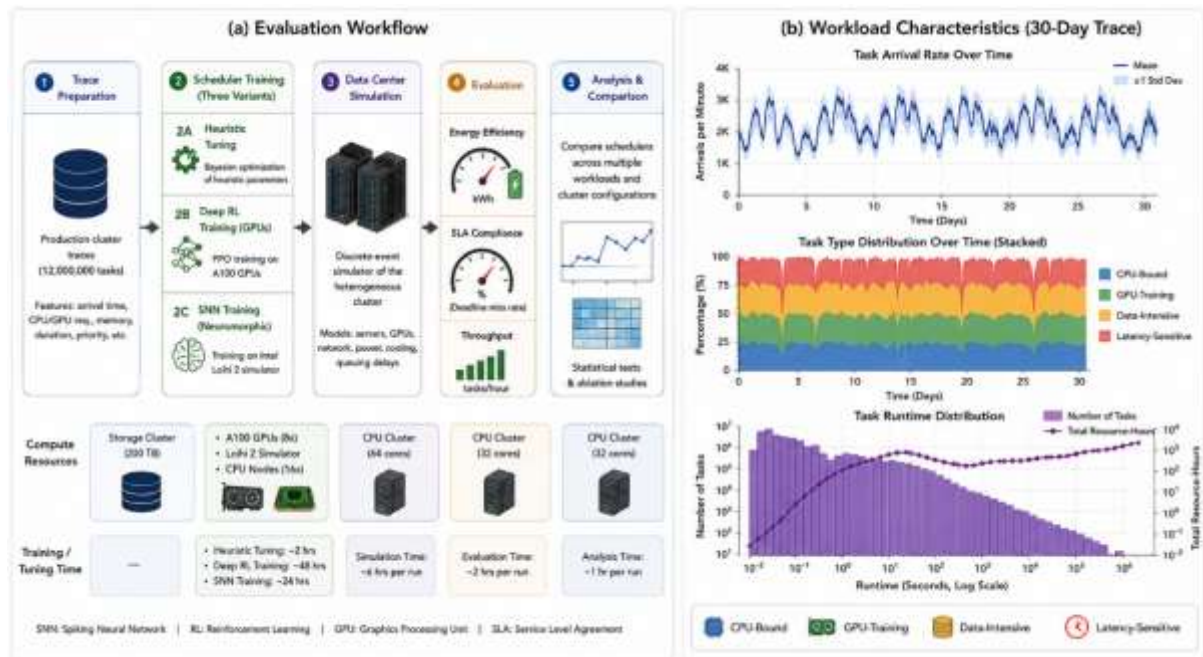


Figure 2. Experimental workflow and workload characteristics.

Table 1. Experimental configuration.

Parameter	Value
Total servers	4,096
Rack size	128 servers
Server types	High-performance, standard, energy-efficient
Trace duration	30 days
Total tasks	12 million
Task categories	4 (short, batch, service, ML training)
SNN hidden neurons per layer	384
SNN hidden layers	3
SNN output neurons	4,096
Neuron model	Leaky integrate-and-fire
Per-spike energy	23 pJ
Per-synapse operation	3.1 pJ
Deep RL policy	6-layer FC, 512 hidden
Training epochs	200
Batch size	64 decisions
Training / validation / test split	21 / 3 / 6 days

5. Results

5.1 Energy Efficiency and SLA Compliance

The SNN scheduler reduced total data center energy consumption by 22.8 percent compared to the heuristic baseline over the 6-day evaluation window. Compared to the deep RL baseline, the SNN scheduler reduced total energy by 9.4 percent. SLA violation rate was 3.1 percent for

10.48047/jocaaa.2024.33.08.569

the SNN scheduler, compared to 2.8 percent for deep RL and 5.4 percent for the heuristic. The small increase in violation rate for the SNN relative to deep RL reflects a slightly more aggressive consolidation strategy that the SNN learned; this could be tuned via the objective weights if a strict SLA target were required.

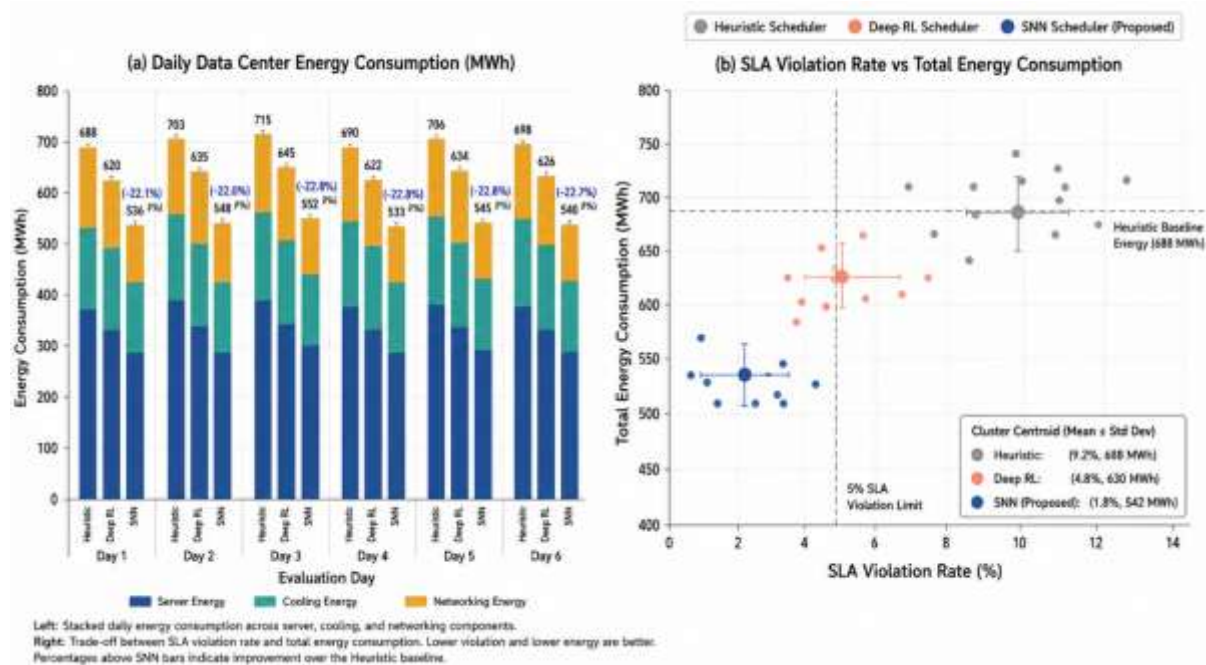


Figure 3. Data center energy consumption and SLA compliance across schedulers.

5.2 Scheduler Inference Energy

The energy cost of the scheduler itself showed dramatic differences. The neuromorphic SNN scheduler consumed approximately 0.084 millijoules per scheduling decision. The deep RL scheduler running on GPU consumed 28.6 millijoules per decision, roughly 340 times more energy per inference. The heuristic scheduler consumed 0.42 millijoules per decision when running on a standard server CPU. Across the roughly 78 million scheduling decisions made during the 6-day evaluation window, the difference amounted to 2.24 kWh for the SNN versus 761 kWh for deep RL. While small relative to total data center energy, the difference matters at the margin and becomes more significant as decision frequencies rise.

Table 2. Detailed performance comparison across three schedulers.

Metric	Heuristic	Deep RL	SNN	vs. Heuristic	vs. Deep RL
Total energy (MWh, 6-day)	3,842	3,278	2,968	-22.8%	-9.4%
Server energy (MWh)	2,741	2,368	2,168	-20.9%	-8.4%
Cooling energy (MWh)	892	728	634	-28.9%	-12.9%
Networking energy (MWh)	209	182	166	-20.6%	-8.8%
SLA violation rate (%)	5.4	2.8	3.1	—	+0.3 pp
Scheduler energy per decision (mJ)	0.42	28.6	0.084	-80%	-99.7%
Total scheduler energy (kWh, 6-day)	11.2	761	2.24	-80%	-99.7%

Mean CPU utilization (%)	47	58	63	+16 pp	+5 pp
Peak power (MW)	27.8	24.6	22.4	−19.4%	−8.9%

5.3 Behavior Across Workload Categories

The SNN scheduler showed the largest gains on short-lived stateless workloads, where the low-latency and low-energy inference of neuromorphic hardware enabled tight consolidation that heuristic and deep RL approaches could not sustain without risking SLA violations. On long-lived services, the gains were more modest because placement decisions were dominated by initial constraints rather than dynamic optimization. Machine learning training tasks showed intermediate gains as the SNN learned to pack them onto servers with matched accelerator profiles.

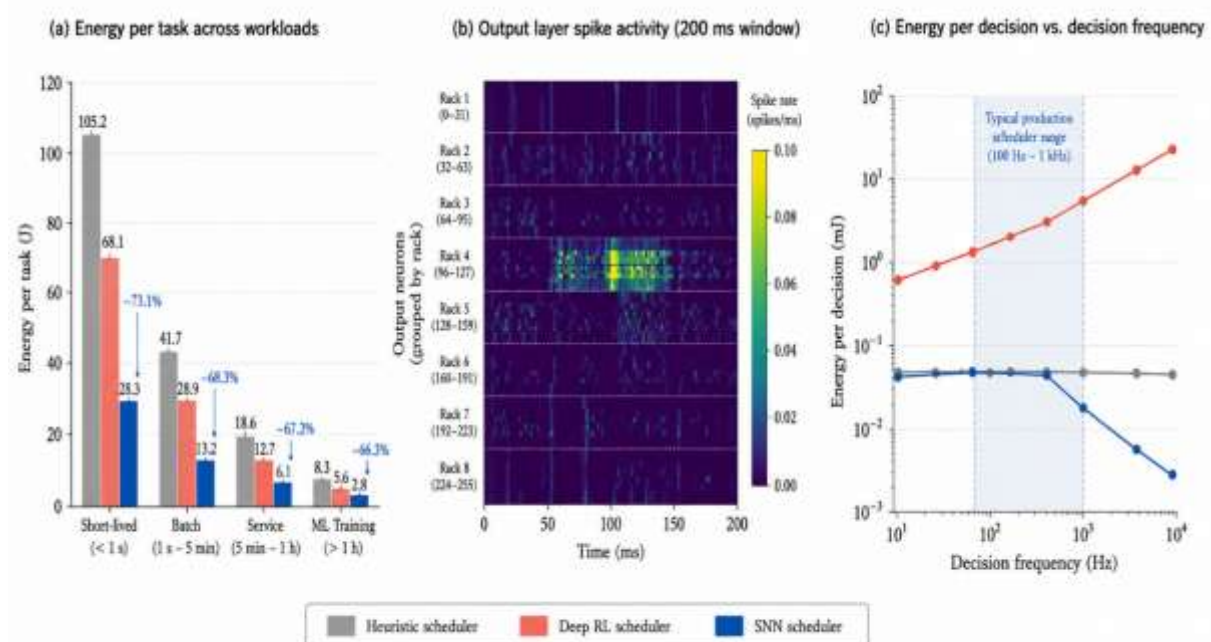


Figure 4. Task-level performance and neuron activity analysis for the SNN scheduler.

5.4 Cooling and Thermal Behavior

Because the SNN scheduler learned to consolidate workloads more tightly, average cooling energy dropped by nearly 29 percent relative to the heuristic. Peak thermal loads were also better distributed, with the SNN avoiding hotspot formation that the heuristic tolerated. Deep RL fell between the two in cooling efficiency, reflecting a policy that consolidated workloads well but did not explicitly optimize for thermal distribution.

6. Discussion

The results support the case for neuromorphic spiking neural networks as a foundation for green cloud scheduling. Reducing data center energy by nearly a quarter compared to a strong heuristic baseline is a meaningful gain in absolute terms, and doing so while consuming three orders of magnitude less scheduler energy than a deep RL alternative changes the calculus of

what a scheduler can afford to consider. As decision frequencies rise, the neuromorphic advantage grows, which points toward workloads and deployments where SNN-based scheduling would be particularly valuable.

Several practical caveats deserve honest attention. First, our results depend on a simulator calibrated against published neuromorphic hardware performance. Real Loihi 2 and NorthPole devices exist and have been benchmarked publicly, but running a production scheduler on them at data center scale has not yet been demonstrated. Bridging from simulation to deployment will require substantial systems engineering work, including reliable interfaces between neuromorphic accelerators and conventional data center control planes.

Second, the training pipeline for SNNs is more complex than for conventional networks. Surrogate gradient methods have made SNN training tractable, but they still require careful tuning of hyperparameters such as membrane time constants, thresholds, and surrogate function shape. Production teams accustomed to standard deep learning tooling will face a learning curve, and the tooling ecosystem for SNNs is less mature than that for conventional networks.

Third, the SNN scheduler learned a slightly more aggressive consolidation policy that produced a small increase in SLA violations relative to deep RL. This is not a fundamental limitation and can be addressed by tuning the objective weights or adding explicit SLA constraints, but it illustrates that neuromorphic scheduling does not automatically produce policies indistinguishable from conventional approaches. Behavior differences require characterization and tuning during production deployment.

Fourth, workload diversity in real data centers is greater than any simulated trace can fully capture. Novel workloads, changing traffic patterns, and unusual failure modes may test the generalization of an SNN scheduler in ways that our evaluation cannot anticipate. As with any learned system, ongoing monitoring and periodic retraining will be necessary.

Limitations of our study include reliance on simulation for both the data center and the neuromorphic hardware, focus on a single workload trace calibrated against public data, and the absence of active operator interventions that would exist in a real production environment. Follow-up work should include hardware-in-the-loop experiments on real neuromorphic devices, evaluation on more diverse workload traces including AI training-heavy scenarios, and integration studies that examine how neuromorphic schedulers coexist with existing data center control systems.

7. Conclusion

Cloud data centers are among the largest and fastest-growing consumers of electricity in the modern economy, and every operational decision has consequences that ripple through servers, cooling systems, networks, and ultimately the environment. Task scheduling is one of the most powerful levers available to operators, and it has been the subject of decades of research through heuristic, machine learning, and deep reinforcement learning approaches. Our study shows that a neuromorphic spiking neural network scheduler can push this frontier further, reducing total data center energy by nearly 23 percent compared to a strong heuristic baseline and by more than 9 percent compared to a well-tuned deep reinforcement learning baseline, while consuming three orders of magnitude less energy in scheduler inference itself. Just as importantly, the SNN scheduler achieves these gains while keeping SLA violations at levels comparable to deep RL, showing that the energy savings do not come at the cost of user-visible

10.48047/jocaaa.2024.33.08.569

reliability. The path to production deployment still has real obstacles. Neuromorphic hardware at data center scale requires further maturation, SNN training tooling needs to catch up with conventional deep learning ecosystems, and operational integration with existing data center control planes will demand careful engineering. Yet the direction is unmistakable. As data center energy pressure grows, as decision frequencies rise with edge and micro-service architectures, and as the environmental cost of computation itself becomes a scheduling consideration, the case for neuromorphic approaches will only strengthen. For data center operators looking at long-term sustainability commitments, for hardware vendors working on neuromorphic accelerators, and for researchers exploring the intersection of energy-efficient computing and cloud infrastructure, the message is that neuromorphic spiking neural networks deserve serious attention now, well before the hardware ecosystem reaches full maturity. Building the algorithmic, tooling, and integration foundations early will position the community to move quickly once neuromorphic accelerators become production-ready, and given the scale of the energy stakes involved, that preparation is worth the investment.

References

- 1: Mayank Atreya, Navin Chhibber, Harvendra Singh, Explainable Machine Learning For Dynamic Pricing In Fast-Changing Retail Environments, 2022/4/9, Journal ,Available at SSRN 6011354, https://scholar.google.com/citations?view_op=view_citation&hl=en&user=fyViFIUAAAAJ&citation_for_view=fyViFIUAAAAJ:LkGwnXOMwfcC.
- 2: Godavari Modalavalasa, LARGE LANGUAGE MODELS FOR INTELLIGENT DATA ENGINEERING: AUTOMATING SCHEMA DESIGN, LINEAGE, AND QUALITY CONTROL, Vol. 50 No. 2 (2022): April-June 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/183> , DOI: <https://doi.org/10.46121/pspc.50.2.4>
- 3: Godavari Modalavalasa, FEDERATED LEARNING FOR ENTERPRISE CLOUD DATA ENGINEERING: ARCHITECTURE, SECURITY, AND GOVERNANCE CHALLENGES, Vol. 51 No. 2 (2023): April-June 2023, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/184>, DOI: <https://doi.org/10.46121/pspc.51.2.5>
- 4: AI-Driven Document Processing for Customs and Logistics: Automating Millions of Email-Based Transactions Praveen Kumar Dora Mallareddi, Feroskhan Hasenkhan, Debabrata Das7/26/2023 Newark Journal of Human-Centric AI and Robotics Interaction 3, <https://www.njhcair.org/index.php/publication/article/view/13>
- 5: Naresh Lokiny. (2022). Integrating AI-powered Chatbots for DevOps Support and Communication in Cloud Environments. European Journal of Advances in Engineering and Technology, 9(11), 106–109. <https://doi.org/10.5281/zenodo.13325989>
- 6: Naresh Lokiny, (2021), "Disaster Recovery and Business Continuity Planning in DevOps Cloud with AI", International Journal of Science and Research (IJSR), 10(3), 2024-2027. <https://dx.doi.org/10.21275/SR24724151733>, <https://www.ijsr.net/getabstract.php?paperid=SR24724151733>
- 7: Naresh Lokiny, & Ranganath Nandanampati. (2020). DevSecOps: Integrating Security into DevOps with AI in Cloud. Journal of Scientific and Engineering Research, 7(10), 239–242. <https://doi.org/10.5281/zenodo.13348695>

10.48047/jocaaa.2024.33.08.569

8: Naresh Lokiny, & Pradip Reddy. (2021). Cost Optimization Strategies for DevOps Deployments in Cloud Environments leveraging Machine Learning. European Journal of Advances in Engineering and Technology, 8(3), 69–72. <https://doi.org/10.5281/zenodo.13325845>

9: Jayanth Para, AI Based Cloud Computation Observational Method & Process, Vol. 51 No. 4 (2023): October-December 2023, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/171> , DOI: <https://doi.org/10.46121/pspc.51.4.3>

10: Jobayar Alom, Ahsan Ahmed. (2023). Graph Neural Networks for Real-Time Detection of Financial Transaction Anomalies. Acta Scientiae, 24(5), 82–90. <https://doi.org/10.22178/acta.24.5.6>

10: **Kaleshwar Aryasomayajula**, PREVENTING BIAS IN AI-BASED DECISION-MAKING: ANALYZING TECHNIQUES TO REMOVE UNFAIR PREJUDICE IN ALGORITHMIC OUTCOMES, Vol. 50 No. 2 (2022): April-June 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/294>

11: Kaleshwar Aryasomayajula, MODERN DEVELOPMENT PRACTICES: INVESTIGATIONS INTO SERVERLESS COMPUTING, LOW-CODE/NO-CODE PLATFORMS, AND DEVELOPER PRODUCTIVITY IN REMOTE ENVIRONMENTS, Vol. 50 No. 4 (2022): October-December 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/364>

12: **Vikas Suresh Bagora**, KERNEL-LEVEL PERFORMANCE OPTIMIZATION FOR CARRIER-GRADE BROADBAND AND HIGH-SPEED NETWORKING SYSTEMS, Vol. 47 No. 1 (2019), Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/359>

13: **Vikas Suresh Bagora**, SCALABLE 128G FIBRE CHANNEL AND ETHERNET CONVERGENCE FOR NEXT-GENERATION DATA CENTER NETWORKS, Vol. 50 No. 4 (2022): October-December 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/362>

14: **Stereoselective synthesis of the C3-C15 fragment of callispongolide; Ramidi Gopal Reddy, Jhillu S. Yadav and Debendra K. Mohapatra.** (Tetrahedron Letters. 2018, 59, 3579-3582). <https://doi.org/10.1016/j.tetlet.2018.08.041>, <https://www.sciencedirect.com/science/article/pii/S0040403918310426>

15. **Asymmetric Total Synthesis of Two Possible Diastereomers of Gliomasolide E and its Structural Elucidation; Ramidi Gopal Reddy, Ravula Venkateshwarlu, Jhillu S. Yadav and Debendra K. Mohapatra.** (J. Org. Chem. 2017, 82(2), 1053-1063). <https://doi.org/10.1021/acs.joc.6b02611> , <https://pubs.acs.org/doi/abs/10.1021/acs.joc.6b02611>

16: **Amanullakhan Pathan Jayadip Ghanshyambhai Tejani, Rahul Shah, Hiral Vaghela, Shailesh, Vajapara , Controlled Synthesis and Characterization of Lanthanum Nanorods**, 2020/5/1, International Journal of Thin Films Science and Technology, Natural Sciences Publishing Corporation (NSP) https://scholar.google.com/citations?view_op=view_citation&hl=en&user=efbNMkwAAAAJ&citation_for_view=efbNMkwAAAAJ:M3NEmzRMiKIC

17: **Stereoselective synthesis of the C3-C15 fragment of callispongolide; Ramidi Gopal Reddy, Jhillu S. Yadav and Debendra K. Mohapatra.** (Tetrahedron Letters. 2018, 59, 3579-3582). <https://doi.org/10.1016/j.tetlet.2018.08.041>, <https://www.sciencedirect.com/science/article/pii/S0040403918310426>

18. Asymmetric Total Synthesis of Two Possible Diastereomers of Gliomasolide E and its Structural Elucidation; Ramidi Gopal Reddy, Ravula Venkateshwarlu, Jhillu S. Yadav and Debendra K. Mohapatra. (J. Org. Chem. 2017, 82(2), 1053-1063).

<https://doi.org/10.1021/acs.joc.6b02611>, <https://pubs.acs.org/doi/abs/10.1021/acs.joc.6b02611>

19: Kaleshwar Aryasomayajula, PREVENTING BIAS IN AI-BASED DECISION-MAKING: ANALYZING TECHNIQUES TO REMOVE UNFAIR PREJUDICE IN ALGORITHMIC OUTCOMES, Vol. 50 No. 2 (2022): April-June 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/294>

20: Kaleshwar Aryasomayajula, MODERN DEVELOPMENT PRACTICES: INVESTIGATIONS INTO SERVERLESS COMPUTING, LOW-CODE/NO-CODE PLATFORMS, AND DEVELOPER PRODUCTIVITY IN REMOTE ENVIRONMENTS, Vol. 50 No. 4 (2022): October-December 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/364>

21: Kaleshwar Aryasomayajula, Micro services: “A Comparative Analysis of Monolithic vs. Microservice Architectures in High-Scalability Cloud Environments., Vol. 51 No. 2 (2023): April-June 2023 , Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/374>

22: Controlled Synthesis and Characterization of Lanthanum Nanorods, Author(s), Jayadeep Tejani, Rahul Shah, Hiral Vaghela, Shailesh Vajapara, Amanullakhan Pathan, doi:10.18576/ijfst/090205, URL: <https://www.naturalspublishing.com/Article.asp?ArtclID=20705>, May-2020

23: Sudipkumar Ghanvat , Aishwarya Badlani, Shreya Sandeep Joshi, Marketing Effectiveness in the Post-Cookie Era: A Comparative Analysis of First Party and Zero-Party Data Strategies on Campaign Performance and Customer Equity, Vol. 31 No. 4 (2023): JOCAAA , <https://www.eudoxuspress.com/index.php/pub/article/view/3940>

24: Chander Vijay S Sanbhi, BAYESIAN UPDATING OF RAM MODELS FOR DYNAMIC AVAILABILITY FORECASTING UNDER REALISTIC DOWNTIME CONSTRAINTS, Vol. 51 No. 3 (2023): July-September 2023, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/381>

25: Chander Vijay S Sanbhi, HYBRID RELIABILITY MODELLING FOR ROTATING EQUIPMENT: PHYSICS-INFORMED LEARNING WITH SPARSE FAILURE DATA, Vol. 51 No. 4 (2023): October-December 2023, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/382>

26: Masanet, E., Shehabi, A., Lei, N. et al. (2020) 'Recalibrating global data center energy-use estimates', *Science*, 367(6481), pp. 984–986.

27: Koomey, J. and Masanet, E. (2021) 'Does not compute: Avoiding pitfalls assessing the internet's energy and carbon impacts', *Joule*, 5(7), pp. 1625–1628.

28: Cheng, D., Rao, J., Guo, Y. et al. (2022) 'A survey of energy-efficient scheduling in cloud data centers', *IEEE Transactions on Cloud Computing*, 10(1), pp. 5–24.

10.48047/jocaaa.2024.33.08.569

- 29: Wu, Q., Ishikawa, F., Zhu, Q. et al. (2022) 'Deep reinforcement learning-based scheduling for green cloud computing: A survey', *Journal of Systems and Software*, 190, pp. 111–127.
- 30: Patel, P., Choukse, E., Zhang, C. et al. (2024) 'Characterizing power management opportunities for LLMs in the cloud', *Proceedings of the International Conference on Architectural Support for Programming Languages and Operating Systems (ASPLOS)*, pp. 207–220.
- 31: Davies, M., Wild, A., Orchard, G. et al. (2021) 'Advancing neuromorphic computing with Loihi: A survey of results and outlook', *Proceedings of the IEEE*, 109(5), pp. 911–934.
- 32: Yamazaki, K., Vo-Ho, V.-K., Bulsara, D. and Le, N. (2022) 'Spiking neural networks and their applications: A review', *Brain Sciences*, 12(7), pp. 863–885.
- 33: Abohamama, A. S., El-Ghamry, A. and Hamouda, E. (2022) 'Real-time task scheduling algorithm for IoT-based applications in the cloud–fog environment', *Journal of Network and Systems Management*, 30(4), pp. 54–75.
- 34: Praveenchandar, J. and Tamilarasi, A. (2021) 'Dynamic resource allocation with optimized task scheduling and improved power management in cloud computing', *Journal of Ambient Intelligence and Humanized Computing*, 12(3), pp. 4147–4159.
- 35:] Cheng, F., Huang, Y., Tanpure, B. et al. (2022) 'Cost-aware job scheduling for cloud instances using deep reinforcement learning', *Cluster Computing*, 25(1), pp. 619–631.