

Energy-Efficient Prompt Placement for Green Artificial Intelligence in Large-Scale NLP Systems

Jiake Lei

Department of Computer Science, University of New Hampshire, Durham, NH, USA.
jiakelei89@unh.edu

Abstract

The rapidly expanding computational demands of large-scale natural language processing systems have elevated energy consumption into a central challenge for sustainable artificial intelligence. While parameter-efficient fine-tuning methods such as prompt tuning and prefix tuning have substantially reduced the costs associated with adapting foundation models, the placement and activation of these prompts within inference pipelines remain under-explored from a system-level energy perspective. This paper presents a holistic examination of energy-efficient prompt placement as a lever for green AI in large-scale deployment environments. By framing prompt activation as a dynamic resource allocation problem, we analyze the structural trade-offs between inference latency, model quality, and carbon footprint across cloud, edge, and hybrid infrastructures. We explore architectural design patterns that support selective prompt insertion, caching strategies for reusable prompt states, and carbon-aware scheduling of prompt-dependent computations. The discussion extends to governance challenges surrounding fairness, transparency, and accountability when energy-saving policies differentially impact user populations or linguistic groups. Through cross-domain comparisons with adaptive computation mechanisms such as mixture-of-experts routing and early-exit architectures, we identify shared principles that can guide the development of energy-proportional prompt management frameworks. The paper concludes with forward-looking perspectives on how prompt placement strategies can be integrated into broader sustainability roadmaps for industrial-scale NLP systems, emphasizing the need for standardized energy reporting, open benchmarking databases, and multi-stakeholder policy coordination.

Keywords

green artificial intelligence, prompt placement, energy-efficient NLP, large language models, selective prompt tuning, sustainable computing, system-level optimization.

1. Introduction

The past half-decade has witnessed an unprecedented expansion in the scale and pervasiveness of natural language processing systems powered by large pre-trained language models. Deployments that routinely serve billions of inference requests per day have transformed sectors ranging from information retrieval and conversational agents to code generation and scientific literature mining. The computational appetite of these systems, however, has raised serious concerns about their environmental sustainability. Energy consumption, carbon emissions, and the life-cycle impacts of hardware accelerators have become central topics in both academic discourse and industrial governance. Early investigations into the carbon footprint of deep learning for NLP revealed that training a single large transformer model can emit as much carbon dioxide as several transatlantic flights, and that the operational phase of continuous serving can far exceed the environmental

impact of one-time training when deployments are scaled globally [1]. These findings catalyzed the green AI movement, which calls for a deliberate shift in research priorities from accuracy maximization at any computational cost to efficiency, inclusivity, and environmental awareness [2].

In parallel with the sustainability agenda, the emergence of prompt-based adaptation techniques has fundamentally altered the economics of customizing foundation models. Instead of fine-tuning all parameters, practitioners can now achieve competitive performance on downstream tasks by learning a small number of continuous prompt vectors or soft tokens that are prepended or interleaved with input sequences. This paradigm of parameter-efficient fine-tuning has been instrumental in lowering the barrier to entry for model customization and has substantially reduced the storage and memory footprint of task-specific artifacts. Yet the deployment of prompt-based systems at scale introduces a new class of system-level energy trade-offs. Each prompt activation consumes computation proportional to the length of the prompt, the number of prompt tokens inserted, and the architectural choices governing how prompts interact with the model’s hidden states. When thousands of prompt variants are served simultaneously across diverse tasks and user contexts, the aggregate energy expenditure of prompt computation becomes a first-order concern for infrastructure operators.

This paper argues that prompt placement—the set of decisions about where, when, and how many prompts to insert into the inference stream—constitutes a powerful and currently underutilized control knob for energy-efficient NLP system design. By drawing on insights from adaptive computation, workload scheduling, and carbon-aware computing, we develop a system-level perspective that reframes prompt placement as a dynamic resource allocation problem embedded in a complex socio-technical infrastructure. The remainder of the paper is organized as follows. Section 2 surveys the energy landscape of contemporary NLP deployments, highlighting the disproportionate role of inference workloads and the limitations of hardware-only efficiency gains. Section 3 conceptualizes prompt placement as a resource allocation mechanism, connecting selective prompt insertion strategies to broader efficiency objectives and positioning them within the landscape of parameter-efficient methods. Section 4 explores architectural and infrastructural design patterns that enable energy-proportional prompt management, including caching hierarchies, edge-cloud orchestration, and carbon-aware scheduling. Section 5 addresses governance, fairness, and policy implications, arguing that energy optimization must be accompanied by rigorous accountability frameworks to prevent the uneven distribution of service quality degradation. Section 6 concludes with a synthesis of findings and an outlook on future research directions.

2. The Energy Footprint of Contemporary NLP Deployments

Understanding the role of prompt placement in energy efficiency first requires a clear characterization of where energy is spent in modern NLP serving pipelines. The total energy footprint of a deployed NLP system encompasses hardware manufacturing, data center infrastructure overhead, model training, and online inference. While early discourse focused heavily on the one-time cost of training ever-larger models, evidence increasingly suggests that inference dominates the cumulative energy expenditure over a system’s operational lifetime. A single high-traffic language model service may process orders of magnitude more tokens during serving than were used in pre-training, and this gap widens as models are deployed for years rather than days. Moreover, the advent of model-as-a-service offerings and the integration of language models into ubiquitous productivity tools means that inference

workloads are no longer sporadic batches but continuous, globally distributed streams of requests with strict latency requirements [3].

Inference energy consumption is determined by a complex interplay of model architecture, hardware characteristics, request batching strategies, and workload patterns. Transformer-based models, which remain the backbone of large-scale NLP, exhibit computational costs that scale quadratically with sequence length under standard self-attention, although recent innovations in efficient attention mechanisms have alleviated some of this burden. The deployment of extremely deep and wide models further amplifies energy costs, leading system designers to explore a variety of compression and acceleration techniques including knowledge distillation, quantization, pruning, and architecture search [3]. These approaches, while effective, typically operate at the level of the model itself and do not exploit the unique structural properties of prompt-based adaptation. Consequently, even a heavily compressed model may still expend significant energy on prompt-related computation when prompt tokens are treated as ordinary input tokens and processed through the full model depth.

A particularly consequential aspect of NLP serving energy is its geographical and temporal heterogeneity. Data centers in different regions draw electricity from grids with vastly different carbon intensities, and these intensities fluctuate on hourly and seasonal timescales. Recent work on carbon-aware computing has demonstrated that shifting flexible computation in time or space can dramatically reduce the carbon footprint of AI workloads without compromising service quality [14]. However, the rigidity of conventional prompt configurations—where every request activates a fixed set of prompt tokens regardless of context—prevents such systems from participating in carbon-aware orchestration at a fine granularity. A more dynamic and selective approach to prompt placement could enable the system to modulate its energy consumption in response to real-time signals about grid carbon intensity, thereby aligning NLP serving with broader sustainability objectives.

Furthermore, the energy profile of NLP deployments cannot be divorced from the economic incentives that drive over-provisioning and redundancy. To meet stringent service-level agreements on latency and availability, operators routinely deploy multiple replicas of models across geographically distributed clusters. Each replica must maintain in memory not only the base model parameters but also the prompt embeddings for all supported tasks. When the number of tasks is large, the memory footprint of prompts alone can become a bottleneck, forcing the system to rely on higher-precision formats or more frequent memory accesses that increase energy per operation. Thus, any strategy that reduces the number of simultaneously active prompts or their average length can yield compounding energy savings across the entire serving infrastructure.

3. Prompt Placement as a System Resource Allocation Mechanism

The conceptual shift from viewing prompts as static task identifiers to treating them as dynamically allocated computational resources opens up a rich design space for energy-efficient system operation. In traditional prompt tuning, a fixed number of soft prompt tokens are learned for each downstream task and prepended to the input sequence at every inference call [5]. Similarly, prefix tuning attaches learned continuous vectors to the keys and values of each transformer layer, creating a more expressive adaptation signal but also increasing the per-token computational cost since these prefix activations must be processed alongside the input at every layer [6]. Adapters insert small bottleneck modules between transformer layers and have been shown to be highly parameter-efficient, though they modify the forward pass structure and introduce additional latency [4]. More recently, the concept of selective prompt

insertion has been formalized, wherein the model learns not only the content of prompts but also a policy governing which tokens or layers should receive prompt augmentation [7]. This approach allows the system to skip prompt insertion altogether for inputs that are already handled well by the base model, thereby creating a direct link between prompt usage and energy expenditure.

Framing prompt placement as a resource allocation problem invites comparisons with other domains where dynamic resource management has yielded substantial efficiency gains. In mixture-of-experts architectures, a gating network routes each input token to a subset of expert modules, ensuring that only a fraction of the total model capacity is activated per token and thus achieving sublinear scaling of computation with model size [10]. Similarly, early-exit mechanisms allow simple inputs to terminate their forward pass at intermediate layers, avoiding the full computational cost of deep networks. Prompt placement can be seen as a complementary technique operating at the level of adaptation rather than core model capacity: instead of deciding which experts or layers to activate, the system decides which adaptation signals to inject and at which positions. When combined with inference-time routing, a prompt placement policy can effectively modulate the energy profile of a model without altering its underlying parameters, offering a lightweight and deployment-friendly path to green AI.

The energy implications of selective prompt placement are multifaceted. At the most immediate level, omitting prompt tokens for a fraction of requests reduces the total number of floating-point operations performed during inference. This reduction translates directly into lower energy consumption on GPUs, TPUs, or custom AI accelerators. However, the net system-level energy impact depends on additional factors. The decision logic that determines whether to insert prompts must itself be executed, and if this logic is complex—involving, for instance, a learned gating network or an uncertainty estimation module—it can offset some of the savings. System designers must therefore balance the granularity of the prompt placement policy against the overhead of evaluating it. Policies that make coarse-grained per-request decisions incur minimal overhead but miss opportunities for finer, token-level optimizations. Conversely, token-level policies can achieve near-optimal energy-proportionality but require careful engineering to avoid becoming a bottleneck. This tension mirrors classic trade-offs in computer architecture between the sophistication of branch prediction or cache replacement policies and their implementation cost.

Another dimension of the resource allocation framing concerns the coupling between prompt computation and the rest of the inference pipeline. In standard transformer implementations, prompt tokens are embedded and then processed through all layers identically to input tokens. If a prompt placement policy decides that certain layers should not receive prompt signals, the computational savings are realized only if the underlying serving system can efficiently skip the corresponding matrix multiplications or attention operations. This requirement places demands on the software stack and hardware that go beyond the model architecture itself. Recent advances in compiler-level optimization for dynamic neural networks, including shape inference, operator fusion, and just-in-time compilation, are beginning to make such fine-grained control feasible in production environments [18]. As these infrastructure capabilities mature, the energy benefits of selective prompt placement will become increasingly realizable.

4. Architectural and Infrastructure Design for Green Prompt Management

Realizing the vision of energy-efficient prompt placement requires more than novel learning algorithms; it calls for a comprehensive rethinking of the system architecture that supports

prompt-based inference at scale. One foundational element is a prompt caching hierarchy that mirrors the multi-level memory systems common in modern processors. At the highest level, a hot set of frequently used prompts can be retained in GPU high-bandwidth memory to minimize access latency and the energy cost of repeated embedding lookups. Less frequently accessed prompts can reside in CPU main memory or even in remote storage, with prefetching mechanisms that anticipate future requests based on temporal usage patterns and task locality. Such a hierarchy allows the system to maintain a large catalog of task-specific prompts while keeping the active working set small, thereby reducing the memory footprint and associated static power dissipation of accelerators.

The orchestration between edge devices and cloud infrastructure introduces additional opportunities and challenges for green prompt placement. Edge-deployed language models are increasingly capable of performing on-device inference for latency-sensitive applications, but they are constrained by battery capacity, thermal envelopes, and limited memory. In a hybrid architecture, simple inputs can be handled entirely on-device with a minimal set of prompts, while complex or uncertain queries are offloaded to a cloud-based model that can afford a richer prompt configuration. This offloading decision itself constitutes a form of prompt placement: the system must decide not only which prompts to insert but where the computation should occur. The energy calculus now includes the transmission cost of prompts and intermediate activations over wireless networks, which can be substantial for large prompt vectors. Recent work on split computing and early-exit models provides a useful blueprint for designing such hybrid pipelines, and integrating prompt placement policies into these frameworks remains an open research problem.

Carbon-aware scheduling adds a temporal dimension to prompt management. In a data center equipped with real-time carbon intensity signals, the system could dynamically adjust the aggressiveness of prompt selection in response to grid conditions. During periods of high carbon intensity, a more conservative prompt policy might be applied, skipping prompts for borderline cases that the base model can handle with slightly lower accuracy. Conversely, during periods of low carbon intensity or when renewable energy is abundant, the system can afford to be more liberal with prompt insertion to maximize quality. Such a policy represents a form of energy-quality trade-off that can be formalized as a multi-objective optimization problem, with service-level objectives encoded as constraints. Implementing carbon-aware prompt scheduling requires integration with data center orchestration platforms that already support workload shifting and temporal flexibility for batch jobs [14]. Extending these platforms to handle the fine-grained, latency-critical nature of inference workloads is a significant infrastructure challenge but one that aligns with broader trends toward carbon-intelligent computing.

Beyond scheduling, the hardware substrate on which prompts are executed plays a decisive role in energy efficiency. While general-purpose GPUs remain dominant, specialized accelerators designed for transformer inference are increasingly prevalent and often incorporate features that could be exploited by prompt placement policies. For instance, some accelerators offer native support for variable-length sequences, sparse attention patterns, or conditional computation pathways. A prompt placement system that is aware of these hardware capabilities can align its insertion decisions with the accelerator's efficiency sweet spots, avoiding operations that trigger expensive fallback mechanisms. This co-design approach echoes the philosophy of hardware-software co-optimization that has driven progress in mobile processors and embedded systems, and it suggests that close collaboration

between NLP researchers and hardware architects will be essential for the next generation of green AI infrastructure.

5. Governance, Fairness, and Policy Implications

The introduction of energy-saving mechanisms into large-scale NLP systems inevitably raises questions about who bears the cost of efficiency and who reaps its benefits. When a prompt placement policy selectively reduces the amount of adaptation applied to certain inputs, the resulting degradation in task performance is unlikely to be distributed uniformly across user populations, languages, or use cases. For example, inputs from low-resource languages may depend more heavily on task-specific prompts to compensate for the base model's weaker zero-shot capabilities. If the energy-saving policy inadvertently skips prompts more frequently for such languages due to their increased uncertainty or lower aggregate usage patterns, the system would systematically worsen service quality for already marginalized linguistic communities. This scenario illustrates the tension between global energy metrics and local fairness criteria that permeates much of the sustainable AI discourse [15].

Addressing such disparities requires that energy-efficiency policies be designed with fairness as a first-class constraint, not an afterthought. System operators must monitor disaggregated performance metrics across demographic and linguistic segments and establish feedback loops that can detect and correct regressions introduced by prompt placement adjustments. Techniques from the field of algorithmic fairness, including equality of opportunity and minimax fairness frameworks, can be adapted to the context of energy-AI trade-offs. For instance, a fairness-aware prompt placement policy might enforce that the energy savings achieved on high-resource language requests do not come at the expense of unacceptable accuracy drops on low-resource languages, even if this means a less aggressive energy reduction target overall. Such policies would represent a concrete instantiation of the principle that sustainability in AI must be socially sustainable as well as environmentally sustainable.

Transparency and accountability mechanisms are critical for building trust in energy-optimized NLP systems. Users and regulators are increasingly demanding visibility into the environmental impact of digital services, as evidenced by the proliferation of carbon labeling initiatives and corporate sustainability reports. However, the complexity of prompt placement systems—where the insertion policy may itself be learned and opaque—poses challenges for meaningful transparency. Drawing inspiration from model cards and datasheets for datasets, we advocate for the development of energy impact statements that accompany each task-specific prompt configuration and its associated placement policy [15]. Such statements would disclose the expected energy consumption per request under various operating conditions, the carbon intensity assumptions used, and the equity analyses conducted. Over time, standardized benchmarks for prompt-related energy consumption could emerge, analogous to the MLPerf benchmark suite for training and inference performance, but augmented with sustainability and fairness dimensions.

On the policy front, regulatory frameworks for AI are gradually incorporating environmental considerations alongside traditional concerns about privacy, safety, and non-discrimination. The European Union's AI Act, for example, includes provisions for high-risk AI systems to demonstrate appropriate levels of robustness and accuracy, and future amendments may require reporting on energy consumption and greenhouse gas emissions. Compliance with such regulations will necessitate precise instrumentation of prompt-related computation, as well as the governance processes to audit and verify the claims made by system providers. Industry consortia and standardization bodies, including the Green Software Foundation and

the Partnership on AI, are well positioned to develop the shared taxonomies and measurement protocols needed to turn energy-efficient prompt placement from a research concept into an operational reality. Without such multi-stakeholder coordination, there is a risk that energy efficiency will be pursued in a fragmented and unverifiable manner, undermining both environmental benefits and public trust.

6. Conclusion

This paper has argued that prompt placement represents a strategically important design dimension for achieving green artificial intelligence in large-scale NLP systems. By treating prompt activation as a dynamic resource allocation decision rather than a static configuration, system designers can unlock significant energy savings without sacrificing the adaptability and performance gains that make prompt-based methods attractive. We have examined the structural trade-offs inherent in selective prompt insertion, drawing parallels with established techniques such as mixture-of-experts routing and early exiting, and have situated prompt placement within the broader context of inference energy dominance and carbon-aware computing. The architectural and infrastructural enablers of energy-proportional prompt management—including caching hierarchies, edge-cloud orchestration, and hardware-software co-design—have been analyzed alongside the governance and fairness imperatives that must accompany any efficiency-driven degradation of service. Looking ahead, the integration of prompt placement into standardized energy benchmarks and regulatory frameworks will be essential for translating laboratory prototypes into trustworthy, sustainable infrastructure. The path forward challenges the NLP community to broaden its evaluative criteria beyond accuracy and latency, embracing energy consumption and carbon accountability as co-equal design objectives. In doing so, the field can align its remarkable technical capabilities with the urgent societal imperative of environmental stewardship.

References

1. Strubell, E., Ganesh, A., & McCallum, A. (2019). Energy and policy considerations for deep learning in NLP. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics* (pp. 3645–3650). Association for Computational Linguistics.
2. Schwartz, R., Dodge, J., Smith, N. A., & Etzioni, O. (2020). Green AI. *Communications of the ACM*, 63(12), 54–63.
3. Tay, Y., Dehghani, M., Bahri, D., & Metzler, D. (2020). Efficient transformers: A survey. *ACM Computing Surveys*, 53(6), Article 120.
4. Houlisby, N., Giurgiu, A., Jastrzebski, S., Morrone, B., de Laroussilhe, Q., Gesmundo, A., Attariyan, M., & Gelly, S. (2019). Parameter-efficient transfer learning for NLP. In *Proceedings of the 36th International Conference on Machine Learning* (pp. 2790–2799). PMLR.
5. Lester, B., Al-Rfou, R., & Constant, N. (2021). The power of scale for parameter-efficient prompt tuning. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing* (pp. 3045–3059). Association for Computational Linguistics.
6. Li, X. L., & Liang, P. (2021). Prefix-tuning: Optimizing continuous prompts for generation. In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing* (pp. 4582–4597). Association for Computational Linguistics.

7. Zhu, W., & Tan, M. (2023, December). SPT: Learning to selectively insert prompts for better prompt tuning. In Proceedings of the 2023 conference on empirical methods in natural language processing (pp. 11862-11878).
8. Zaken, E. B., Ravfogel, S., & Goldberg, Y. (2022). BitFit: Simple parameter-efficient fine-tuning for transformer-based masked language-models. In Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (pp. 1–9). Association for Computational Linguistics.
9. Hu, E. J., Shen, Y., Wallis, P., Allen-Zhu, Z., Li, Y., Wang, S., Wang, L., & Chen, W. (2022). LoRA: Low-rank adaptation of large language models. In Proceedings of the 10th International Conference on Learning Representations. OpenReview.
10. Shazeer, N., Mirhoseini, A., Maziarz, K., Davis, A., Le, Q. V., Hinton, G. E., & Dean, J. (2017). Outrageously large neural networks: The sparsely-gated mixture-of-experts layer. In Proceedings of the 5th International Conference on Learning Representations. OpenReview.
11. Sanh, V., Debut, L., Chaumond, J., & Wolf, T. (2019). DistilBERT, a distilled version of BERT: Smaller, faster, cheaper and lighter. arXiv preprint arXiv:1910.01108.
12. Jacob, B., Kligys, S., Chen, B., Zhu, M., Tang, M., Howard, A. G., Adam, H., & Kalenichenko, D. (2018). Quantization and training of neural networks for efficient integer-arithmetic-only inference. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (pp. 2704–2713). IEEE.
13. Jiao, X., Yin, Y., Shang, L., Jiang, X., Chen, X., Li, L., Wang, F., & Liu, Q. (2020). TinyBERT: Distilling BERT for natural language understanding. In Findings of the Association for Computational Linguistics: EMNLP 2020 (pp. 4163–4174). Association for Computational Linguistics.
14. Acun, B., Lee, K., Kazhamiaka, F., Maeng, K., Gupta, U., Chakkaravarthy, M., Brooks, D., & Wu, C.-J. (2023). Carbon Explorer: A holistic framework for carbon-efficient datacenter operations. In Proceedings of the 2nd Workshop on Sustainable Computer Systems (HotCarbon '23). ACM.
15. Mitchell, M., Wu, S., Zaldivar, A., Barnes, P., Vasserman, L., Hutchinson, B., Spitzer, E., Raji, I. D., & Gebru, T. (2019). Model cards for model reporting. In Proceedings of the Conference on Fairness, Accountability, and Transparency (FAT '19) (pp. 220–229). ACM.
16. Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? In Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT '21) (pp. 610–623). ACM.
17. Jin, D., Jin, Z., Zhou, J. T., & Szolovits, P. (2020). Is BERT really robust? A strong baseline for natural language attack on text classification and entailment. In Proceedings of the AAAI Conference on Artificial Intelligence (Vol. 34, pp. 8018–8025). AAAI Press.
18. Fowers, J., Ovtcharov, K., Papamichael, M. K., Massengill, T., Liu, M., Lo, D., Alkalay, S., Haselman, M., Adams, L., Ghandi, M., Heil, S., Cox, P., Shah, A., Bhardwaj, R., Caulfield, A. M., Chung, E. S., & Burger, D. (2018). Serving DNNs in real time at

datacenter scale with Project Brainwave. In Proceedings of the 45th Annual International Symposium on Computer Architecture (ISCA '18) (pp. 1–14). IEEE.

19. Jiang, Z., Lin, H., Zhong, Y., Huang, Q., Chen, Y., Zhang, Z., Peng, Y., Li, X., Xie, C., Nong, S., Jia, Y., He, S., Chen, H., Bai, Z., Hou, Q., Yan, S., Zhou, D., Sheng, Y., Jiang, Z., ... & Zheng, Y. (2024). MegaScale: Scaling large language model training to more than 10,000 GPUs. In Proceedings of the 21st USENIX Symposium on Networked Systems Design and Implementation (NSDI '24) (pp. 745–760). USENIX.
20. Rajbhandari, S., Rasley, J., Ruwase, O., & He, Y. (2020). ZeRO: Memory optimizations toward training trillion parameter models. In Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis (SC '20) (pp. 1–16). IEEE.