

UNDERSTANDING ENERGY CONSUMPTION TRENDS IN HIGH PERFORMANCE COMPUTING NODES

Jonathan Muraña¹, Juan J. Durillo², and Sergio Nesmachnow¹

¹High Performance Computing group, Universidad de la República, Montevideo, URUGUAY

²Leibniz Supercomputing group, Leibniz-Rechenzentrum, Munich, GERMANY

ABSTRACT

This article presents a study of energy consumption behavior in high performance computing nodes in relation to the usage of computing resources. Linear models are constructed for identifying common patterns or differences in energy consumption across different architectures. The study is significant as it provides insights into the energy consumption of computing nodes, helping to build simple, yet useful and transparent models. Moreover, the methodology provides building blocks for building new models with high predictive quality and broad applicability across different architectures. The results reveal similarities in energy consumption across different architectures when compared in terms of CPU cycles and cache misses. Additionally, employing linear models based on CPU cycles and cache misses allows for both the explanation of energy consumption behavior and the achievement a reasonable predictive quality. Overall, a partition-based linear model outperforms a global linear and a mean partition-based model by up to 5.7%.

1 INTRODUCTION

Nowadays, the energy consumption of medium and large-scale computing facilities is a matter of increasing concern for governments and society at large, due to its impact on the environment and climate change. Therefore, efficient electricity consumption is a key aspect when planning new computing infrastructure or for maintaining its daily operations (Katal et al. 2023). The effective management and planning of energy consumption of high performance computing nodes, hardware components, and cooling systems is crucial for computing infrastructures aimed at ensuring operational efficiency (Zhang et al. 2021).

Within energy management and planning, energy consumption models to predict operational behaviour are crucial. This modeling enables intelligent utilization, such as leveraging off-peak hours in the electrical grid or minimizing idle or underutilized periods of infrastructure, which may result in inefficient energy consumption. Building realistic energy models for computing nodes is a costly endeavor. Complex models are constructed by collecting data on specific hardware and software, which are then used to build, mainly based on computational intelligence, highly accurate but less explainable models. Furthermore, models need to be periodically updated or discarded because software and hardware changes of the computing infrastructure make them less effective in their predictive capacity.

This article addresses the study and construction of energy models from the perspective of explainability, where simple models are generated for several high performance computing nodes to examine trends and common patterns. Furthermore, the convenience in terms of gains and losses regarding the predictive quality of the models is progressively analyzed, as the design becomes more complex. The trade-off between the complexity and predictive quality of the models is analyzed.

Results reveal similarities in energy consumption patterns across various architectures. Additionally, using linear models based on CPU cycles and cache misses allows for the explanation of energy consumption behaviors and achieves a reasonable predictive quality. The partition-based linear model outperformed both the global linear and mean partition-based models by up to 5.7%.

The article is organized as follows. Section 2 introduces the energy modeling in general and the characteristics of the one developed in this work. Section 3 reviews related work. The methodology is explained in Section 4. The experimental results are presented and discussed in Section 5. Finally, Section 6 presents the conclusions and formulates the main lines for future work.

2 ENERGY CONSUMPTION IN HIGH PERFORMANCE COMPUTING NODES

The energy models of high performance computing nodes studied in this article are classified as software-centric linear models (Ismail and Materwala 2020). Software-centric models are built by collecting data from the performance counters of computing nodes, which measure CPU cycles, memory access, disk usage, among other statistics. These values are recorded during the execution of synthetic tasks called benchmarks, which are specifically designed to capture a wide range of data domain samples. In a subsequent stage, the data is used to train models following a technique that can range from linear regression models to sophisticated neural networks. In particular, linear models aim to express the energy consumption value through a linear combination of the measured performance counters values. Equation 1 illustrates the general structure of the presented linear models, where f_i denotes a specific measured performance counters associated with a linear coefficient γ_i , representing its contribution to the overall value of the energy consumed by the task (E). The constant ε is the independent term. The values of γ_i and ε are the parameters to be adjusted by applying the process known as model training or model fitting.

$$E = \varepsilon + \sum_{i=1} \gamma_i \times f_i \quad (1)$$

The training process requires a training dataset. The training dataset used for developing the proposed energy consumption models comprises data collected from the execution of particular applications across various high performance computing nodes. The set of executed applications extends the Apex-MAP benchmark (Strohmaier and Shan 2005). The applications were designed to explore different combination of memory accesses and computational demands, thereby encompassing several typical memory access patterns.

The synthetic benchmark data collected during the DEEP project was used for constructing the models in our study. The Performance Application Programming Interface (Terpstra et al. 2010) was used for collecting performance counter values, which are directly linked to energy consumption (Hähnel et al. 2012). The recorded data comprises values for ten performance counters: the number of CPU instructions, the count of L1, L2 and L3 cache misses, the quantity of double precision float operations, the ratio of vectorization per cycle of double precision floating point operations (employing 128-bit, 256-bit, and 512-bit vectorization per CPU cycle), the total number of CPU cycles and the execution time. In this work, we refer to performance counters as features.

The data collected during the executions is referred to as a trace. A trace is generated by recording the values of each of the performance counters studied during the executions of the Apex-MAP benchmark, taking into account various parameter combinations, on a single computing node. The result is one trace per computing node with 6000 distinct executions, meaning 6000 points per trace.

Four high performance computing nodes are considered in the experiments:

- CN node: two Intel Xeon Golden 6146 processors, Skylake microarchitecture, 12 cores (24 threads) at 3.2GHz and 192 GB RAM.
- ESB node: one Intel Xeon Silver 4215 processor, Cascade Lake microarchitecture, 8 cores at 2.50 GHz and 48 GB RAM.
- DAM node: two Intel Xeon Platinum 8260M processors, Cascade Lake microarchitecture, 24 cores at 2.4 GHz and 384 GB RAM.
- SKN node: two Intel Xeon Phi CPU 7210F processors, Knights Landing microarchitecture, 64 cores at 1.30GHz and 92 GB RAM

3 RELATED WORK

In the last years, several articles have studied the prediction of energy consumption of computational infrastructures. The survey by Ismail and Materwala (2020) reviewed and compared recent articles on software-based modeling of high performance computing nodes. All reviewed articles proposed models using performance counters as the independent variables and energy consumption as the dependent variable. Next, four works closely related to the proposal of this article are examined.

Zhou et al. (2017) presented models based on performance counters to predict the energy consumption of high performance computing nodes. Among the collected features were CPU time, memory usage, page faults, disk reads, and bandwidth usage. The authors used benchmarks of three characteristics: CPU-intensive, disk-intensive, and transactional web-intensive. The contribution of each feature was analyzed using Principal Component Analysis (PCA), and features related to network usage were discarded. Four regression models were built: linear regression, exponential regression, power regression, and polynomial regression. The four proposed regressions outperformed two methods from the literature, mainly in disk-intensive tasks, where a 2% improvement on the predictive quality was achieved.

Liang et al. (2020) presented energy consumption models based on feature selection and deep learning. Three types of benchmarks were employed for training and testing the models focusing on CPU, disk, and transaction web utilization. For feature selection, authors utilized a method from information theory known as information entropy, which measures the dispersion of variable distributions. Features were prioritized for each benchmark type, and their importance were averaged to create a unified priority list. In determining the quantity of features, neural networks went through several training and evaluation phases, with an importance threshold fixed to analyze model error. Twelve features meeting the importance threshold were chosen, achieving a proper trade-off between predictive quality and computational time. The first three selected features were CPU utilization, bytes consumed per CPU, and context switches rate. The proposed model obtained a mean reduction in relative error of over 1.1%, compared to five models in the literature.

Park et al. (2021) introduced a tree-based piecewise linear model to predict the optimal frames per second considering the total energy consumption on CPU and GPU, tailored for use in a Dynamic Voltage And Frequency Scaling governor integrated into mobile gaming. The model required a design for balancing predictive quality and interpretability, ensuring quick utilization by the governor. The authors compared various machine learning techniques, ultimately selecting a tree-based model that uses linear functions at the leaves. The metrics applied for the evaluation of the studied techniques included Mean Absolute Error and a metric for measuring model interpretability based on mathematical formulas considering model size and computational complexity for inference. The proposed model achieved improvements in energy efficiency of approximately 10% compared to state-of-the-art governors based on simple linear regression.

Our previous article (Muraña et al. 2021) presented three energy models based on features and machine learning techniques to build generalizable models across different architectures, i.e., using models trained on one computing node to predict the consumption of another computing node. The experimental validation was performed over part of the same experimental raw data used in the research reported in the current article, allowing for a direct comparisons between the results. Feature selection was conducted using PCA, and six features were selected to be used in the models. The convolutional neural network outperformed the predictive quality of linear regression by 2.58% and the dense neural network by 1.29% when evaluated on the same high performance computing nodes they were trained on. Regarding the evaluation of model generalization for use on high performance computing nodes different from those they were trained on, the results were promising but not conclusive. The model using a convolutional neural network was better than linear regression, but the interpretability of the results was not clear. The research reported in this article was motivated by the search for alternatives to achieve the generalization of energy models.

In summary, energy consumption modeling of high performance computing nodes has been extensively studied in the literature and accurate models have been built based on sophisticated techniques. However, to the best of our knowledge, no studies examined energy consumption from the perspective of characterization and explanation, comparing their behavior across different high performance computing nodes.

4 METODOLOGY

This section describes the methodology followed to address the problem, including the description of the metrics applied to evaluate the predictive quality of the models, the feature selection procedure, and the use of heatmaps and numbering. Finally, the design of the implemented models is detailed.

4.1 Metrics for Predictive Quality Evaluation

Several metrics have been proposed for evaluating the quality of predictive models. The Root Mean Squared Error (RMSE) is a widely used metric for evaluating regression models. Equation 3 shows the RSME formulation, where y_i is the real energy consumption and $\hat{y}_i$ is the estimation of the evaluated model.

Normalizing the RMSE by considering the maximum and minimum values is a common practice for the metric to be not influenced by varying magnitudes. In the reported analysis, it is a very useful technique when comparing the RMSE values from two high performance computing nodes with markedly different maximum energy consumption values. Both the features and energy values are scaled to the [0,1] interval before training and testing models. Consequently, the RMSE metric is normalized using the minimum and maximum energy values of each computing node, facilitating the evaluation and comparison across different high performance computing nodes. For the sake of clarity, when the RMSE metric is normalized with the maximum and minimum energy values of the computing node, considering all data within the domain, it is referred to as $RMSE_G$ (global RMSE). When the RMSE metric is normalized with the maximum and minimum energy values of the computing node, but it is computed considering only data in a certain subdomain, it is referred to as $RMSE_Z$ (RMSE by zone).

When comparing RMSE values calculated for certain subdomain where the maximum energy values differ from that of the computing node, another normalized RMSE metric, $RMSE_Z^*$ (normalized RMSE by zone), is used. Equation 3 defines this metric, where the real energy values y_i correspond to points which are in a certain subdomain and $\max(y_i)$ represents the maximum value in the same subdomain.

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2} \times 100 \quad (2)$$

$$RMSE_Z^* = \frac{\sqrt{\sum_{i=1}^n (y_i - \hat{y}_i)^2 / n}}{\max(y_i)} \times 100 \quad (3)$$

4.2 Feature Selection Method

The proposed approach focuses on identifying features that provide the best explanation of energy consumption while also understanding the reasons behind the predictive capabilities of the studied models. In this line, a methodology is proposed to study the relationship between features and energy consumption in all traces and evaluate which features are advantageous to prioritize when constructing a simplified model with only two independent variables, which allows for a three-dimensional visual analysis.

The feature selection is based on linear modeling and is described as follows. First, linear models for energy consumption estimation are built considering all possible feature pairs in the four input traces. Second, the models are ranked based on the achieved mean values of $RMSE_G$. Finally, the features that obtain the best position in the ranking are selected.

4.3 Heatmaps and Zone Numbering

The reported values are graphically presented using heatmaps, to analyze the relationship between features and energy consumption. Heatmaps make it easy to find patterns and trends, enable quick comparisons, and facilitate the identification of correlation, among other advantages.

In the designed heatmaps, the domain is partitioned and the values of the studied functions are averaged within each partition. Figure 1(a) and 1(b) show how a heatmap is constructed from tabular function data.

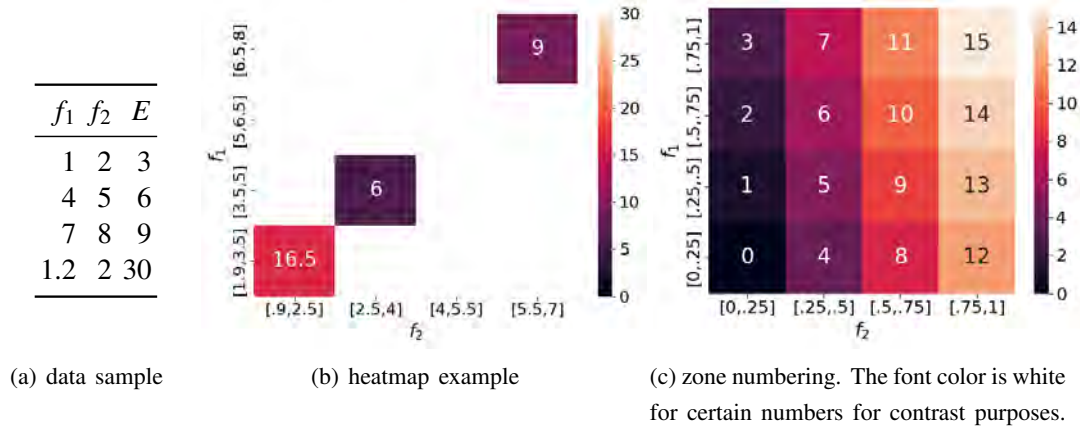


Figure 1: Heatmap example and zone numbering.

The example in Figure 1 shows two features f_1 and f_2 corresponding to independent variables, along with the dependent variable E (e.g., energy consumption). There are four points. The second and the third belong to two different partitions, whereas the other two points belong to the same partition. In this study, the dependent variable is the energy consumption, f_1 corresponds to the feature that best contribute to improve the predictive quality of the model and f_2 is the second best. The partitions for each axis are defined by normalizing the values of each feature to the interval $[0,1]$, which is then divided into four subintervals: $[0,0.25]$, $(0.25,0.5]$, $(0.5,0.75]$, and $(0.75,1]$. Figure 1(c) shows the axis partition considered in this study.

The partitions of the domain are referred to as *zones*. In the proposed study, 16 zones are considered. The number of zones provide an appropriate range of variation for the analysis of energy consumption. In turn, using 16 zones provides a reasonable number of partitions for the analysis, simplifying the modeling process. Figure 1(c) shows the numbers assigned to each zone.

4.4 Design of Energy Models considering Zones

The proposed approach aims to study simple energy models for a better understanding of energy consumption characteristics in high performance computing nodes. To achieve this goal, a series of models are proposed to identify common patterns across different computers. These models are intended to design strategies that enable the construction of precise and generalizable models. The comparative analysis of the built models also allows for evaluating, among other aspects, the possible advantages of using less complex models.

The built models are characterized by the combination of three aspects: type of model (based on mean or based on linear combinations), training data (corresponding to a zone or to the entire domain), and testing data (corresponding to a zone or to the entire domain).

Three models were built:

1. *Mean model by zones (MZ)*: In this model, the energy values are calculated by averaging the points within a specific zone, as applying a decision tree with fixed partitions. It is a straightforward model that provides a reference for energy values across the domain, used as a baseline for the comparison of the results obtained using more sophisticated models. The model is defined by Equation 4, where $\bar{E}(x)$ is the predicted energy value of a point x that corresponds to a zone Z_i . $\mathbb{1}_{x \in Z_i}$ is an indicator function that equals 1 when x belongs to zone Z_i , and 0 otherwise. $\bar{E}(x)$ is calculated as the mean of the energy values $E(x_j)$ considering all points within zone Z_i .

$$\bar{E}(x) = \mathbb{1}_{x \in Z_i} \frac{\sum_{x_j \in Z_i} E(x_j)}{|Z_i|}, \quad i = 0, 1, 2, \dots, 15 \quad (4)$$

2. *Global linear model (GL)*: In this model, all points in the domain are fitted to the same linear model, regardless of their zone. This model is similar to the one presented in our previous work (Muraña et al. 2021), but using fewer features as independent variables. The model is defined by Equation 4, where $\bar{E}(x)$ represents the predicted energy value of point x . $\bar{E}(x)$ is calculated as a linear combination of $f_1(x)$ and $f_2(x)$, the values of features f_1 and f_2 of point x . The parameters β_0 , β_1 , and β_2 are determined to minimize the RMSE_G metric.

$$\bar{E}(x) = \beta_0 + \beta_1 f_1(x) + \beta_2 f_2(x) \quad (5)$$

3. *Linear model by zones (LZ)*: In this model, a linear regression is built for each zone, considering only points from that zone during training. For estimating the energy value, the model corresponding to the zone of the point being evaluated is used. This model allows adjusting the linear model to the corresponding zone, taking advantage of the assumption that nearby points, specifically those from the same zone, have similar energy values and behave similarly in response to variations in computing resource consumption. Additionally, the use of piecewise functions helps modeling non-linear (i.e., logarithmic or exponential) behaviors. The model is defined by Equation 6, where $\bar{E}$ is the predicted energy value of a point x . The parameters γ_0 , γ_1 , and γ_2 for each zone Z_i , are determined to minimize the RMSE_Z metric corresponding to zone Z_i .

$$\bar{E}(x) = \mathbb{1}_{x \in Z_i} (\gamma_0 + \gamma_1 f_1(x) + \gamma_2 f_2(x)), \quad i = 0, 1, 2, \dots, 15 \quad (6)$$

Models use 70% of the data for training and the remaining 30% for testing. Data selection for both training and testing is performed by zone, to include points from each zone in both training and testing sets.

5 EXPERIMENTAL ANALYSIS

This section presents the results of the energy consumption study. First, the results of the feature selection process are reported. Then, considering the selected features, the results of a preliminary analysis of the raw data are presented. This analysis helps to understand its characteristics and guide further analysis and modeling. Finally, the predictive quality of the models is assessed, and the results are compared.

5.1 Feature Selection Results

Regarding the RMSE_G metric, the first five positions of the ranking were as follows:

1. features L2 and cycles, with an RMSE_G of 0.127
2. features L3 and cycles, with an RMSE_G of 0.130
3. features cycles and instructions, with an RMSE_G of 0.151
4. features L3 and L2, with an RMSE_G of 0.160
5. features flopsDP (scalar) and cycles, with an RMSE_G of 0.168.

RMSE_G values obtained using features (L2 and cycles) and (L3 and cycles) were significantly better than for other combinations. Cycles was identified as a relevant feature. The improvement when using L2 over L3 was not significant, so a linear combination between L2 and L3 was used, defined as $L = \alpha_{L2} L2 + \alpha_{L3} L3$. The weights were computed applying Ordinary Least Squares to fit the energy consumption for each computing node, resulting in: for CN, $\alpha_{L2} = 0.389$ and $\alpha_{L3} = 0.218$; for ESB, $\alpha_{L2} = 0.367$ and $\alpha_{L3} = 0.433$; for DAM, $\alpha_{L2} = 0.431$ and $\alpha_{L3} = 0.323$; and for SKN, $\alpha_{L2} = 0.353$ and $\alpha_{L3} = 0.122$.

The related literature has extensively highlighted the significance of CPU and memory usage as the primary resources influencing energy consumption in high-performance computing nodes (Arshad et al. 2022). The obtained results suggest that cycles emerge as the most indicative feature of CPU usage. In turn, the results indicate that cache misses, specifically in level 2 and level 3 caches, serve as reliable indicators of memory usage.

5.2 Preliminary Analysis of the Traces: Number of Points by Zone and Mean Energy per Zone

The heatmaps in Figure 2 present the number of points for each zone in the four studied traces. Analyzing the distribution of points in each zone is relevant because benchmarks, although designed with the main idea of covering diverse resource usage combinations, may exhibit varying levels of data samples in different zones. The number of points in each zone influences the statistical validity of the models.

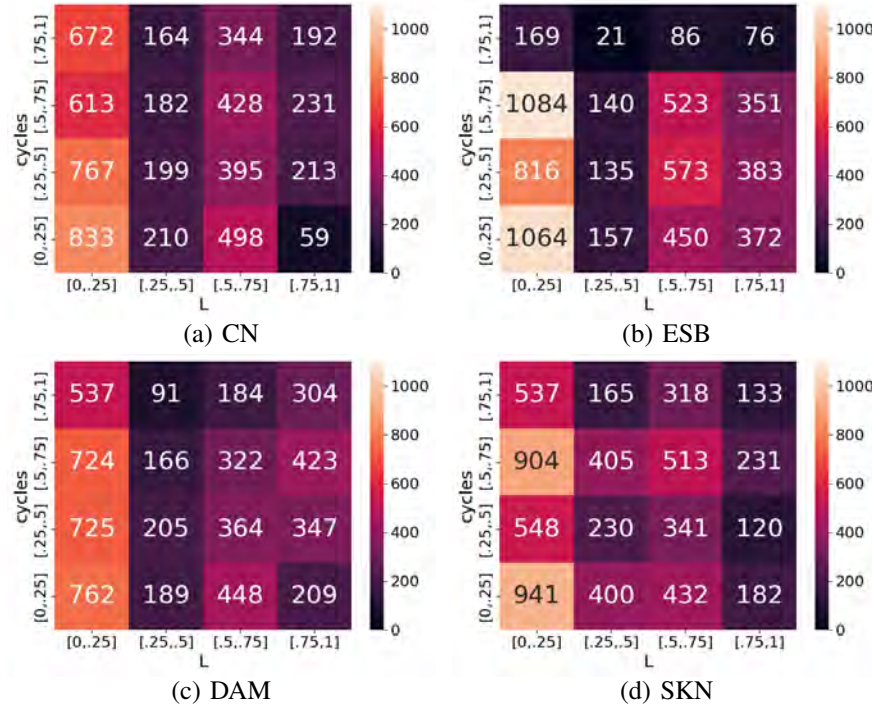


Figure 2: Number of points in each zone.

Across all traces, zones 0, 1, and 2 exhibited the highest number of points compared to the other zones, ranging from 548 to 1084 points. These three zones, characterized by low memory usage (indicated by low values of the "L" feature), contain the majority of data points. Zone 7 consistently had a limited number of points across all traces, with a minimum of 21 points in the ESB trace and a maximum of 165 points in the SKN trace. Zone 7 is associated with low memory usage but intense CPU usage. Specifically, for the ESB trace, zones 7, 11, and 13 exhibited a sparse distribution of data points.

The heatmaps in Figure 3 presents the mean energy per zone for the four studied traces. The reported values are normalized considering the maximum energy value of each trace. An initial observation reveals a consistent pattern in the mean energy consumption per zone across the traces. This consistency suggests a strong correlation between energy consumption and the three studied features (cycles, L2, and L3), seemingly unaffected by the specific characteristics of the computing nodes. A second observation that strengthens the reliability of the measurements in the experiments is the consistent trend where energy consumption either increases or remains unchanged when moving from left to right and from bottom to top. This trend corresponds to the ascending pattern of the features, indicating that higher CPU or memory usage consistently leads to higher energy consumption across all traces.

Finally, the comparative analysis showed a reduced difference of energy values between zones with simultaneous high CPU usage and low memory usage (zone 3), and low CPU usage and high memory usage (zone 12), across all traces. The largest difference observed was 14% in the DAM trace. The slight differences suggest that at extreme values, intensive CPU and memory usage have similar impacts on energy consumption. Regarding the energy consumption variation, on the one hand, in the direction of increasing

CPU usage, energy consumption proportionally increases from one zone to the next (e.g., in the CN trace, energy consumption rose from 0.087 in zone 0 to 0.21 in zone 1, 0.36 in zone 2, and 0.6 in zone 3). On the other hand, in the direction of increasing memory usage, there is a more significant jump from zone 0 to zone 4 (e.g., in CN, energy consumption increased from 0.087 in zone 0 to 0.35 in zone 4), followed by a more gradual increase. The logarithmic behavior of energy consumption with increasing memory usage aligns with the findings of our previous study (Muraña et al. 2019).

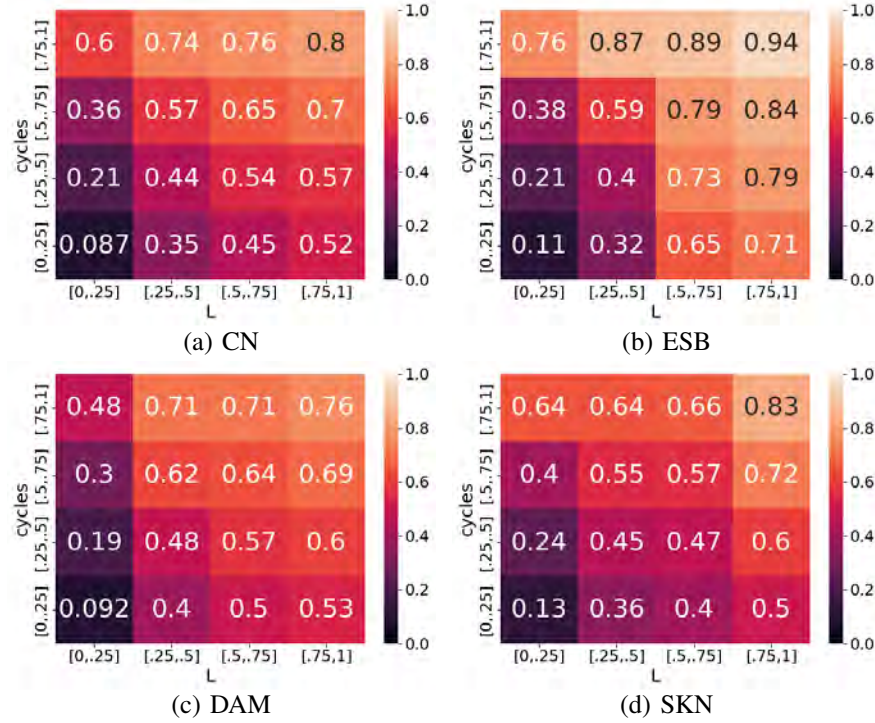


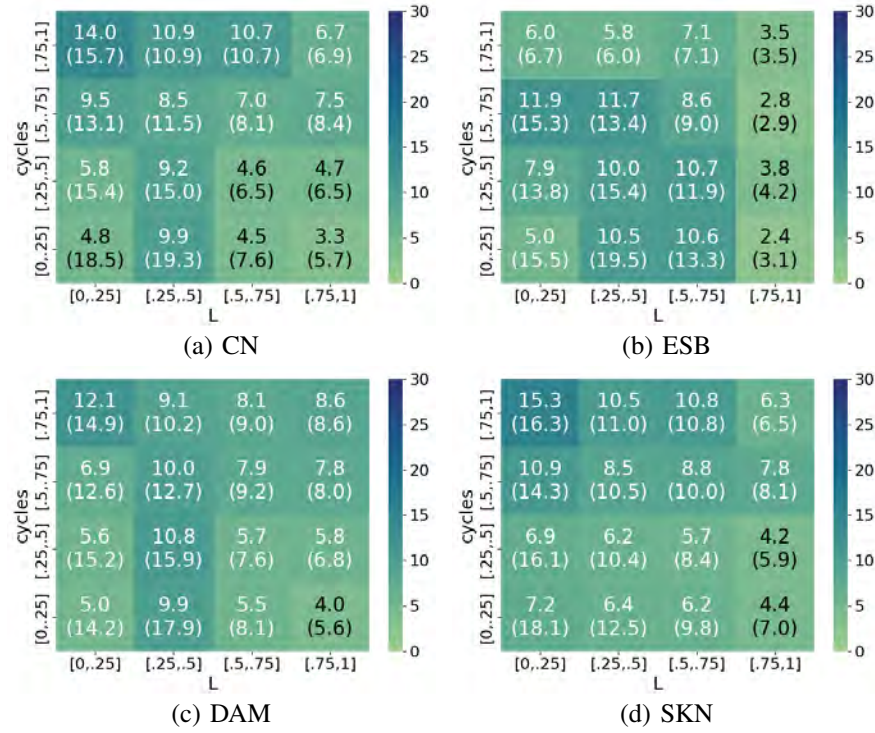
Figure 3: Mean energy consumption by zone.

5.3 Predictive Quality in Energy Models

This section reports the computed metrics for the studied models, discusses the trends, and presents a comparative analysis.

5.3.1 Mean Model by Zones

The heatmaps in Figure 4 present the $RMSE_Z$ and $RMSE_Z^*$ results for the MZ model in all traces. In general, the MZ model obtained similar $RMSE_Z$ values across all traces. However, differences are detected in the ESB trace, in which MZ obtained a better predictive quality than in the other traces. When comparing errors between zones using the $RMSE_Z^*$ metric, in general, zones with more intensive memory usage, i.e., the fourth column of the heatmaps, achieved the most accurate estimates across all traces. This observation is supported by the fact that the $RMSE_Z^*$ values in the fourth column range from a minimum of 2.4 to a maximum of 8.6 across all traces, which is markedly lower than the errors of other zones. In only five cases, considering all traces, a non-fourth column zone did not achieve superior $RMSE_Z^*$ values compared to the fourth column zone within the same trace. This overall trend is observed specifically in zones 8, 9, and 10 of CN, as well as zones 8 and 9 of the DAM trace. However, in all these cases, the improvements obtained by these non-fourth column zones did not exceed 2.

Figure 4: RMSE_Z and (RMSE_Z^{*}) of MZ model evaluated in each zone.

5.3.2 Global Linear Model

Model GL was evaluated globally and by zone. Global evaluation means that the RMSE is calculated over the entire set of points of the testing dataset, regardless of the zone they belong to. This evaluation considers the metric RMSE_G, as explained in Section 4.1. The evaluation by zone involves calculating 16 values of the RMSE_Z metric, one for each zone using the testing points belonging to the respective zone. This procedure allows evaluating the predictive quality of the model in each zone. Results are reported and commented on the next paragraphs.

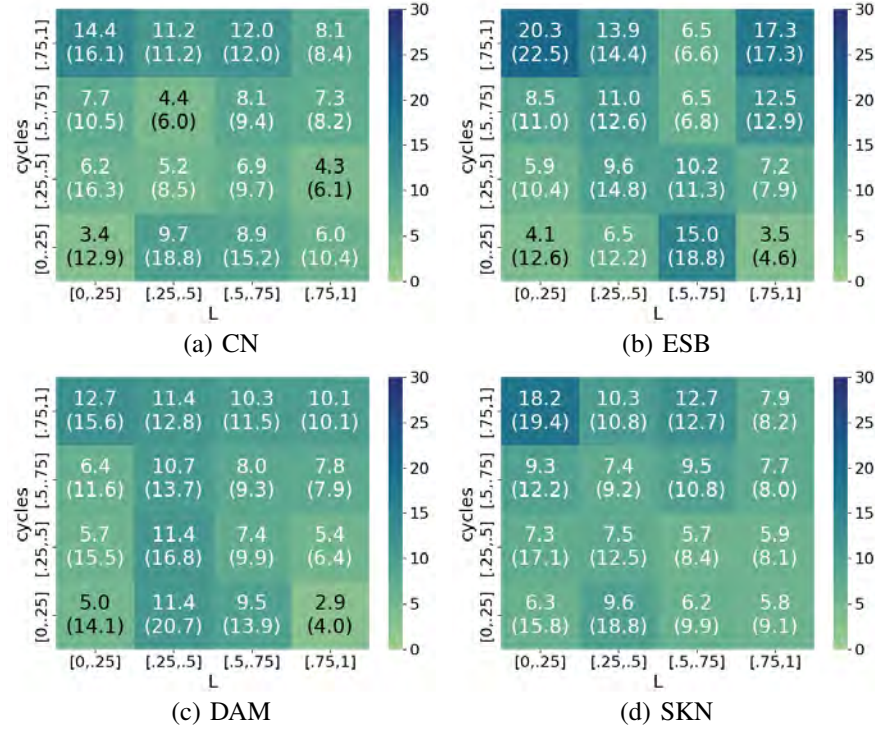
Global Evaluation. Table 1 reports the values of the RMSE_G metric obtained for the GL model, evaluated globally across all traces and compare them with those of the with the linear LRG model, which considers all features for modeling. The LRG model was presented in our previous work (Muraña et al. 2021). No results are reported for LRG over the SKN trace, since it was not studied in the previous work.

Table 1: Results of the RMSE_G metric for the Global Linear model.

	CN	ESB	DAM	SKN
GL	8.5	9.0	8.3	9.4
LRG	5.8	7.6	7.2	-
Difference	2.7	1.4	1.1	-

Results in Table 1 show that using fewer features in the linear model reduced the predictive quality approximately 3 in the CN and ESB traces, and a 2 loss in the DAM trace when simplifying the model.

Evaluation by Zone. The heatmaps in Figure 5 shows RMSE_Z and RMSE_Z^{*} of the GL model evaluated in each zone in all traces.

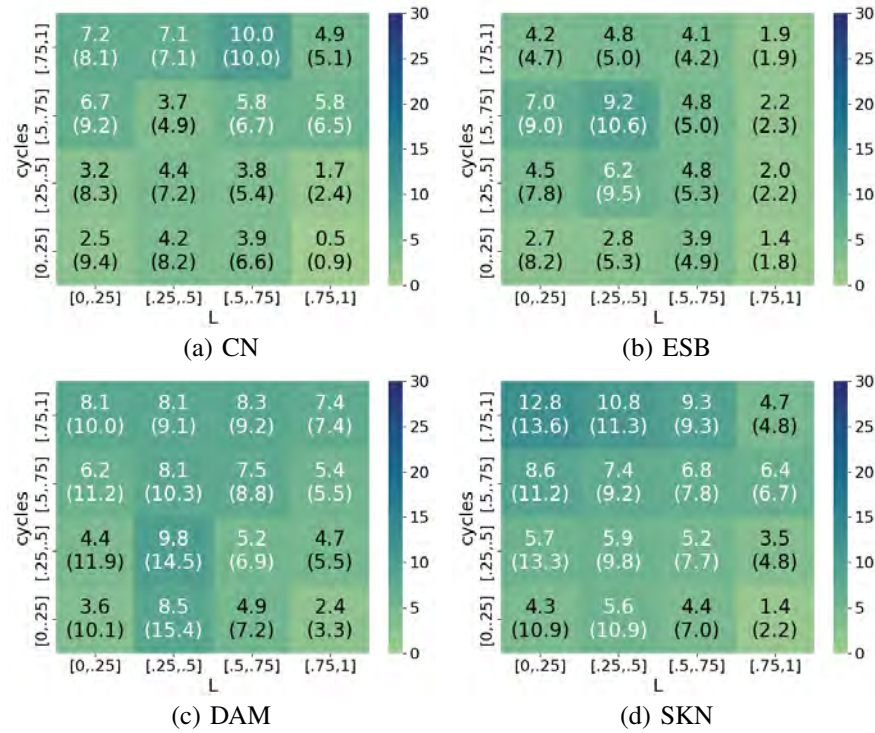
Figure 5: RMSE_Z and (RMSE_Z^{*}) of GL model evaluated in each zone.

Analyzing the zones where GL had the lower predictive quality is relevant to identify potential improvements. The worst RMSE_Z^{*} values were obtained in zones 4 (18.8) and 3 (16.1) for CN, in zones 3 (22.5) and 8 (18.8) for ESB, in zones 4 (20.7) and 5 (16.8) for DAM, and in zones 3 (19.4) and 4 (18.8) for SKN. Results on those worst-performing zones suggest that the GL model had difficulties in adjusting in zones 3 and 4. The best-performing zones of the GL model were found in zones 6 (6.0%) and 13 (6.1) for CN, zones 12 (4.6) and 11 (6.6) for ESB, zones 12 (4.0) and 13 (6.4) for DAM, and zones 14 (8.0) and 13 (8.1) for SKN. Overall, GL performed best in zones 12 and 13.

Results showed that MZ also faced difficulties in zones 3 and 4. When considering the zones where one model outperformed the other by more than 1 of the RMSE_Z metric, results indicate that: In CN, GL outperformed the MZ model in 4 zones and MZ outperformed the GL model in 6 zones. In ESB, GL outperformed the MZ model in 4 zones and MZ outperformed the GL model in 7 zones. In DAM, GL outperformed the MZ model in 1 zones and MZ outperformed the GL model in 7 zone. In SKN, GL outperformed the MZ model in 2 zones and MZ outperformed the GL model in 7 zones. The zone-based comparisons suggest that the MZ model achieved better results than the GL model.

5.3.3 Linear Model by Zones

The heatmaps in Figure 6 shows the values of the RMSE_Z and RMSE_Z^{*} metrics obtained for the MZ model in all traces. Considering the RMSE_Z metric, the LZ model outperformed the MZ model in all zones of all traces, except in zone 11 of DAM and zone 7 of SKN. In these two zones, the differences in RMSE_Z were slightly: 0.2 in zone 11 of DAM and 0.3 in zone 7 of SKN. The mean difference in RMSE_Z, considering only the zones across all traces where LZ performs better than MZ, was 2.28 and the maximum difference was 7.6, in zone 4 ESB. When compared to the GL model, the LZ model is better in all zones of all traces except in zones 6 and 7 of SKN. In zone 6, the RMSE_Z values are identical, whereas in zone 7, there is a difference of 0.4. The mean difference in RMSE_Z, considering only the zones across all traces where LZ outperformed GL, is 3.37, with the maximum difference being 16.1 (zone 3 ESB).

Figure 6: $RMSE_Z$ and $(RMSE_Z^*)$ of LZ model.

5.4 Summary, Model Comparison and Discussion

Table 2 summarizes the comparison between the studied models, broken down by trace. The metrics reported in Table 2 are $RMSE_G$, mean value of $RMSE_Z$ (denoted as $RM\bar{SE}_Z$) considering all zones, and the three best and worst zones for each model in each trace, according to the metric $RMSE_Z^*$. Regarding $RMSE_G$ and $RM\bar{SE}_Z$, MZ slightly outperformed GL in both metrics and the most significant difference occurred in the ESB trace with a 2.5 margin in $RM\bar{SE}_Z$. In turn, LZ outperformed MZ, with the most significant difference also observed in the ESB trace with a 5.7 margin in $RM\bar{SE}_Z$.

Table 2: $RMSE_G$ and $RM\bar{SE}_Z$ metrics, and best/worst zones of each model in all traces.

	CN			DAM		
	MZ	GL	LZ	MZ	GL	LZ
$RMSE_G$	8.2	8.5	5.3	7.6	8.3	6.1
$RM\bar{SE}_Z$	7.6	7.7	4.7	7.7	8.5	6.4
Best zones	12,9,13	6,13,14	12,13,6	12,13,9	12,13,14	12,13,14
Worst zones	4,0,3	4,1,3	11,0,2	4,5,1	4,5,3	4,5,1
	ESB			SKN		
	MZ	GL	LZ	MZ	GL	LZ
$RMSE_G$	8.5	9.0	4.6	8.9	9.4	7.2
$RM\bar{SE}_Z$	7.4	9.9	4.2	7.9	8.6	6.4
Best zones	14,12,15	12,11,10	12,15,13	13,15,12	14,13,15	12,13,15
Worst zones	4,0,5	3,8,15	6,5,2	0,3,1	3,4,1	3,1,7

Zones 3 (high CPU and low memory use) and 4 (low CPU and low to moderate memory use) were the most challenging. Conversely, zones 13 and 12 had a consistent predictive quality. These zones are characterized by high memory utilization and low to moderate CPU usage. The trend of best and worst-performing zones is observed to a greater or lesser extent across all models and traces.

6 CONCLUSIONS AND FUTURE WORK

This article presented a study of energy consumption patterns in four high performance computing nodes. The most significant features in terms of their contribution to energy consumption were CPU cycles and L2/L3 cache misses. Linear and mean models were built based on domain partitions (zones), and compared using heatmaps, which allowed for visual pattern detection and explanation of the obtained results.

The different high performance computing nodes studied exhibited similarity in energy behavior across zones. The most challenging zones for energy estimation shared similar characteristics across nodes (high CPU usage and low memory consumption, and low CPU usage and low to moderate memory utilization). The partition-based mean model MZ slightly improved over the global linear model GL. The partition-based linear model LZ achieved a maximum improvement of 5.7% over MZ.

Future lines of work include developing new models for zones where linear models are less accurate, to improve the predictive quality, and studying how the estimation quality of the models scales from one computing node to another, incorporating the findings of the reported research.

REFERENCES

- Arshad, U., M. Aleem, G. Srivastava, and J. C.-W. Lin. 2022. "Utilizing Power Consumption and SLA Violations using Dynamic VM Consolidation in Cloud Data Centers". *Renewable and Sustainable Energy Reviews* 167:112782.
- Hähnel, M., B. Döbel, M. Völp, and H. Härtig. 2012. "Measuring Energy Consumption for Short Code Paths Using RAPL". *ACM SIGMETRICS Performance Evaluation Review* 40(3):13–17.
- Ismail, L. and H. Materwala. 2020. "Computing Server Power Modeling in a Data Center: Survey, Taxonomy, and Performance Evaluation". *ACM Computing Surveys (CSUR)* 53(3):1–34.
- Katal, A., S. Dahiya, and T. Choudhury. 2023. "Energy Efficiency in Cloud Computing Data Centers: a Survey on Software Technologies". *Cluster Computing* 26(3):1845–1875.
- Liang, Y., Z. Hu, and K. Li. 2020. "Power Consumption Model Based on Feature Selection and Deep Learning in Cloud Computing Scenarios". *IET Communications* 14(10):1610–1618.
- Muraña, J., C. Navarrete, and S. Nesmachnow. 2021. "Machine Learning for Generic Energy Models of High Performance Computing Resources". In *High Performance Computing*, 314–330. Cham: Springer International Publishing.
- Muraña, J., S. Nesmachnow, F. Armenta, and A. Tchernykh. 2019. "Characterization, Modeling and Scheduling of Power Consumption of Scientific Computing Applications in Multicores". *Cluster Computing* 22(3):839–859.
- Park, J., N. Dutt, and S. Lim. 2021. "An Interpretable Machine Learning Model Enhanced Integrated CPU-GPU DVFS Governor". *ACM Transactions on Embedded Computing Systems* 20(6):1–28.
- Strohmaier, E. and H. Shan. 2005. "Apex-Map: A Global Data Access Benchmark to Analyze HPC Systems and Parallel Programming Paradigms". In *2005 ACM/IEEE Conference on Supercomputing*, 49–49.
- Terpstra, D., H. Jagode, H. You, and J. Dongarra. 2010. "Collecting Performance Data with PAPI-C". In *Tools for High Performance Computing*, 157–173.
- Zhang, Q., Z. Meng, X. Hong, Y. Zhan, J. Liu, J. Dong *et al.* 2021. "A Survey on Data Center Cooling Systems: Technology, Power Consumption Modeling and Control Strategy Optimization". *Journal of Systems Architecture* 119:102253.
- Zhou, Z., J. Abawajy, F. Li, Z. Hu, M. Chowdhury, A. Alelaiwi *et al.* 2017. "Fine-Grained Energy Consumption Model of Servers Based on Task Characteristics in Cloud Data Center". *IEEE Access* 6:27080–27090.

AUTHOR BIOGRAPHIES

JONATHAN MURAÑA is Ph.D. candidate in the High-Performance Computing group at Universidad de la República, Montevideo, Uruguay. His research interests include computational intelligence and green computing. His email is jmurana@fing.edu.uy.

JUAN J. DURILLO is an Scientific Employee at Leibniz Supercomputing Centre, Munich, Germany. His research interests include Machine Learning, GPU Computing, and Cloud Computing. His email address is juan.durillobarriouveau@lrz.de.

SERGIO NESMACHNOW is an Full Professor and Researcher at Universidad de la República, Montevideo, Uruguay. His research interests include Evolutionary Computing and Metaheuristics. His email address is sergion@fing.edu.uy.