

Machine Learning-Based Cache Miss Prediction

Edin Jelačić · Cristina Seceleanu · Ning Xiong · Peter Backeman ·
Sharifeh Yaghoobi · Tiberiu Seceleanu

the date of receipt and acceptance should be inserted later

Abstract Integrating machine learning into computer architecture simulation offers a new approach to performance analysis, moving away from traditional algorithmic methods. While existing simulators accurately replicate hardware, they often suffer from slow execution, complex documentation, and require deep CPU knowledge, limiting their usability for quick insights. This paper presents a deep learning-based approach for simulating a key CPU component, cache memory. Our model “learns” cache characteristics by observing cache miss distributions, without needing detailed manual modeling. This method accelerates simulations and adapts to different program needs, demonstrating accuracy comparable to traditional simulators. Tested on Sysbench and image processing algorithms, it shows promise for faster, scalable, and hardware-independent simulations.

Keywords Machine Learning · Cache · Simulation.

1 Introduction

As hardware costs increase, computer architecture simulation has become essential for optimizing and estimating performance, identifying bottlenecks, and providing a cost & time-effective solution without the explicit need for physical hardware. A commonly used way to perform this optimization is the analysis of large amounts of data obtained from such simulators, as illustrated by Fig. 1. Several algorithmic simulators and instrumentation tools are available for dynamic mem-

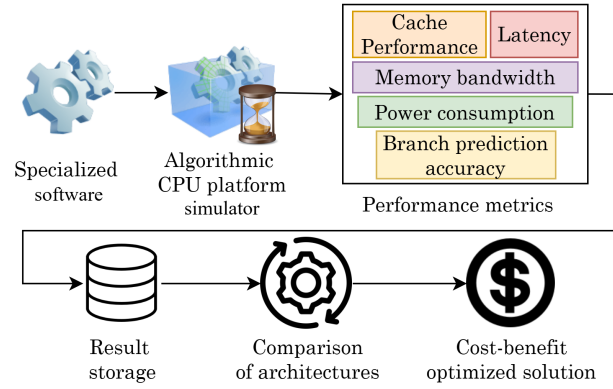


Fig. 1: Commonly utilized industrial software simulation paradigm

ory, security analysis, and performance monitoring, including DynamoRIO Cachesim [5], Valgrind [25], Intel PIN [23], and QEMU [4]. These tools serve specific purposes and use cases, but they share the drawback of requiring either algorithmic processing of each individual instruction, or instrumentation of execution on an existing CPU, which introduces a significant performance overhead in terms of computation time [31]. This overhead may lead to unreliable performance measurements and is therefore preferably avoidable.

In this work, we focus on the prediction of cache misses based on a novel statistical machine learning design. Traditional algorithmic simulators provide detailed insight into program execution, but face limitations due to high overhead and the need for detailed architectural knowledge, especially in complex systems. Two main challenges arise in instruction-accurate simulation of x86 CISC architectures: the construction of complex, error-prone models for each architecture [24], and significant performance issues, with simulators running much slower than native execution [22]. Recent advances like Ithemal [24], SimNet [22], and the latest

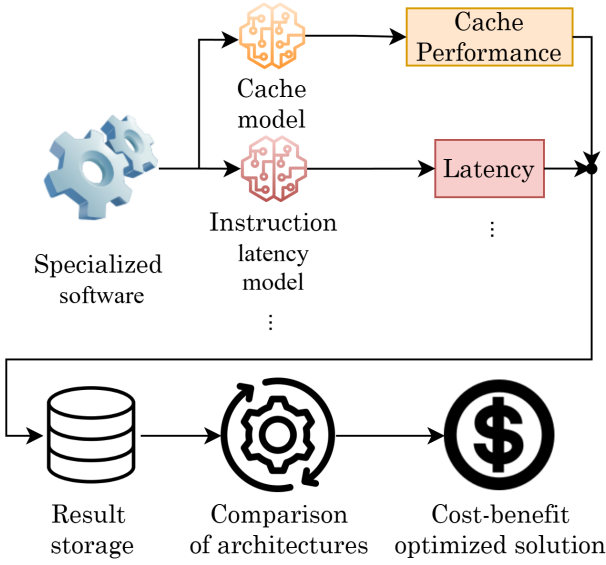


Fig. 2: Proposed industrial software simulation paradigm

TAO [28] address these challenges using deep learning to predict instruction latencies and other performance metrics in specific architectures without extensive manual system definitions.

In modern computing, optimizing the speed of memory access is critical to improving performance and energy efficiency. Cache memory plays a vital role in this by providing rapid data retrieval to the processor. Consequently, the optimization of a hardware platform concerns the organization of the cache memory utilized by the platform. Each distinct application demands a varied utilization of cache memory throughout its runtime. Dimensioning of memory resources with respect to an executing program strongly impacts cache miss occurrences that lead to potentially significant delays in processing. Therefore, system designers, particularly in high-cost industrial environments, are keen to cater to specific software needs when deciding on prospective platform investment.

However, managing cache size and structure remains a challenge due to trade-offs between cost, performance, energy consumption, and cache misses, which lead to increased memory access times and degraded system performance. While existing simulators model cache behavior, faster and more adaptable solutions are needed to better address these inefficiencies. A promising, yet underexplored approach is the use of machine learning for modeling application-specific cache behavior.

This work focuses on the design and implementation of a neural network-based model with the purpose of predicting cache performance within an existing cache

simulation framework, aiming to enhance and improve it for specific use cases, without assuming a fixed CPU architecture. Complementary to some existing work, our research here aims ultimately: (i) to replace a traditional simulator for cache prediction with a holistic machine learning model; (ii) for this model to function and generate performance metrics across various program executions using a hardware-agnostic approach; (iii) for this model to focus on the comparison of as many hardware combinations as possible while targeting a limited set of programs.

We introduce a *long-short-term memory* (LSTM) [11] deep learning approach for predictive cache miss modeling on program execution traces on multiple prospective architectures and demonstrate this on benchmark traces for testing. Our choice is justified by the fact that LSTMs are variants of recurrent neural networks, capable of adequately handling long-term sequential dependencies. To our knowledge, it is the only work that addresses the issue of simulating cache performance on various architectures in this manner. Our ML-based method offers the potential for faster predictions and greater flexibility to adapt to varying program and architecture characteristics without the need for detailed full-system architectural simulation. This ultimate goal framework, along with its utility, is illustrated

In summary, the key contributions of this paper are:

- A novel method for modeling different cache levels using a recurrent multivariate regression model;
- Pioneering the use of LSTM networks in modeling cache miss distribution, independent of the architecture;
- The ability of our method to observe the impact of core/cache configurations on cache miss distribution per individual core and core count for the corresponding first-level caches and shared unified caches.

The validation of this approach comes by testing it on widely used benchmark software and an image processing algorithm across various cache sizes and core counts, and the evaluation of the efficiency of the method is performed by comparison to an algorithmic cache simulator.

The remainder of the paper is structured as follows. Section 2 provides essential background on memory, cache, neural networks, and the DynamoRio Cachesim simulator. Section 3 details our method to predict the cache failure distribution throughout the execution of the program. We follow with Section 4 where we dive into the data, the instruction features and the process of training and deploying our model. Section 5 presents

the results of our model’s performance, and Section 6 discusses the successes and challenges of our approach, while Section 7 explores alternative approaches and state-of-the-art techniques, comparing them with our objectives. We conclude with final thoughts and suggestions for future work in Section 8.

2 Background

In contemporary microarchitectures, *cache memory* embodies a hardware implementation of small, yet high-speed memory situated in close proximity to the processor. In modern processors, it serves as a hardware-implemented memory layer positioned between individual registers and the main memory blocks. The latter are commonly identified today based on the main implementation technology, e.g. *DRAM* (Dynamic Random-Access Memory).

2.1 Memory hierarchy

Processor registers are compact memory locations, usually 32 or 64 bits, and serve as architectural components of the processor instruction set. The processor operates on these registers directly, through various instructions. Registers are used to store instructions, counters, memory addresses, operands, and the immediate results of instruction execution.

The main memory constitutes volatile read/write storage, where programs are loaded directly during runtime. It is notably large, and it is common for contemporary desktop platforms to be equipped with more than 32 gigabytes, due to the necessities of modern operating systems and increasingly complex software. However, accessing DRAM for data whenever necessary would be prohibitively slow, as DRAM lies architecturally and physically outside the processor. This would represent a bottleneck with the processor being starved waiting for incoming data, particularly given that a program’s execution often relies on reusing data.

2.2 Cache memory

Cache memory is a specialized component integrated into the processor to reduce latency by storing recently accessed data. It is structured into multiple levels, including L1, L2, and L3, each with varying capacities and speeds. L1 is the fastest but the smallest, while L3 offers greater capacity but higher access latency. The processor utilizes the cache in a hierarchical manner to minimize reliance on slower main memory, significantly

reducing program execution delays. In multicore processors, usually each core has at least one cache level (L1), with other levels usually shared between cores.

Throughout program execution, the processor employs cache memory to store frequently used data and instructions. When the processor is instructed to access a specific address by the instruction code, it initiates the process by checking the L1 cache. If the data is not found in L1, it proceeds to check L2, and so forth. If the data is not found in any of the caches, it is retrieved from the main memory and stored in one of the cache levels [16], and until data is retrieved, the execution is paused. This event is known as a *cache miss*, and various advances in modern computer design aim to minimize these occurrences and mitigate the corresponding latency penalties. Strategies include continuous increases in on-die cache memory sizes throughout history [9], cache memory restructuring, the incorporation of various optimizations [10], and the introduction of mechanisms, such as pre-fetchers, which anticipate the most likely instructions or data that the processor will need and fetch them before they are needed [20].

In addition to the above, the evolution of processors, exemplified by advancements like multithreading, can be viewed as a step towards minimizing the time-delay associated with a cache miss by implementing thread-level parallelism [17]. Predicting cache misses, implementing strategies for their avoidance or scalable mitigation techniques for compulsory cache misses is of great significance, especially in a modern computing landscape, where memory speed lags behind that of contemporary processor designs [6].

2.3 Neural Networks

Neural networks (NN) have proven to be useful as universal function approximators [15, 26]. They have been used, in various forms, for improvements or predictions in computer architectures [3]. The basic neural network is the *feed-forward* NN, called so due to its structure of a mesh of neurons with one or multiple inputs, one or multiple outputs and one or more hidden layers representing linear functions paired with nonlinear activation functions.

An NN feedforward estimates a function f , with data flowing from the input $\mathbf{x}$, processing in the hidden layer(s) defined by a parameter set Θ , the output being a vector of real-valued results, commonly denoted by the *prediction* $\mathbf{y}$. More precisely, NN represent statistical mappings $p(\mathbf{y}|\mathbf{x}; \Theta)$ in which the parameters Θ are tuned, by backpropagation, in such a way as to locally minimize some cost function that maps the weights to the outputs of the objective cost function. Feed-forward

NNs do not have feedback loops that map past quantities to internal memory when making predictions. This makes them only appropriate for calculating predictions in which the predictions have no temporal relationship with historical data.

In those problems where the dataset on which we are trying to calculate the predictions is sequential in nature, sequential NN or RNN (recurrent neural networks) are commonly used. By sequential nature, we mean that the sequence of inputs is a factor in predicting the desired distribution, as $p(\mathbf{y} | (\mathbf{x}, \mathbf{x}_{-1}, \mathbf{x}_{-2}, \dots, \mathbf{x}_{-k}); \Theta)$ where k is the number of time steps in the sequence.

The most prominent sequential architecture is the LSTM NN, composed of a memory cell and multiple non-linear gates that regulate the flow of data, represent current, and selectively memorize prior states. This minimizes the vanishing / exploding gradient problem, typically present in RNN backpropagation [11, 32]. Such an architecture is shown in Fig. 3. Due to its prominence in the field of sequential timeseries multivariate forecasting, we opted for this model as well.

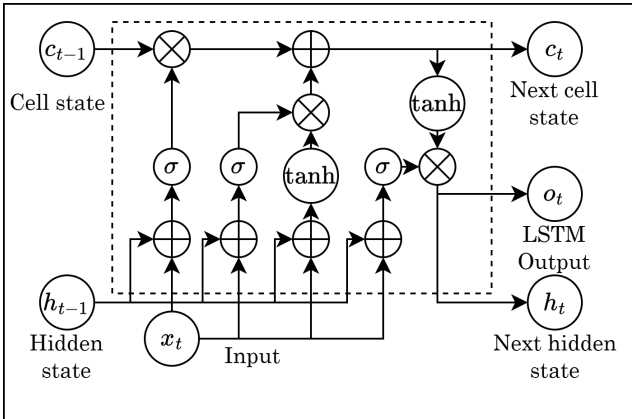


Fig. 3: LSTM Cell architecture

2.4 DynamoRIO Cachesim

DynamoRIO Cachesim is a simulator tool designed to analyze the cache performance of applications by simulating the impact of various cache configurations on memory access patterns. It functions by dynamically instrumenting running applications to collect detailed memory access data, which is then utilized to simulate cache behavior.

The simulator models L1 and L2 cache levels by default and can be extended to support more complex

configurations. It supports varying sizes, cache associativities, and block sizes to provide insights into cache performance metrics. The simulator operates on an executable by injecting instrumentation code into the application's binary at runtime, capturing each memory access that the application performs. These captures are stored in a format that Cachesim and other components in the DynamoRIO suite can use to simulate desired cache configurations and compute metrics such as cache hit rates, miss rates, performance across different cache levels, register states, and latency impacts. DynamoRIO itself is extensible and well-documented, offering the ability for users to build their own client tools for performance analysis within the DynamoRIO framework or build upon existing ones.

3 The Approach

In this section, we present our approach that relies on equating predictive modeling to the statistical estimation of a cache miss distribution across a program's execution through running aggregates of cache misses. Specifically, given a program of length K and a sequence length of N ($N < K$), to estimate such a distribution we generate $\lceil K/N \rceil$ data points, with each data point at most of value N (observe Fig. 4 for an example with a sequence length of 9 instructions).

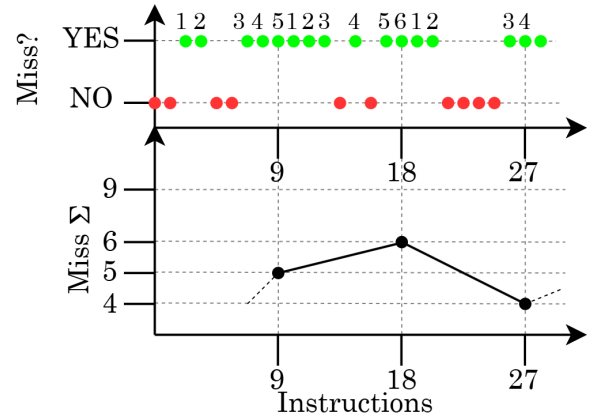


Fig. 4: Miss sub-sequence summing

Our model utilizes deep learning to take a program trace and specified cache characteristics as input, returning the cache miss distribution as output. By program traces, here we refer to x86 assembly traces of executed programs. When access to the original source code and compilation procedure is unavailable, these traces can be obtained using instrumentation tools, such

as Intel PIN [23], Valgrind [25], or other methods like debuggers and disassemblers. Hence, we turn the task of predictive modeling of cache behavior into one of multivariate sequential regression. Program traces are represented by the x86 assembly instructions (the instructions themselves, with the corresponding operands, such as constants and registers), which the execution of a program consists of, as exemplified below:

```
...
cmp    %rax $0x0000000000000022
jbe    $0x00007f9290f3e0c9
mov    %rdi -> %rsi
sub    %rax %rsi -> %rsi
cmp    %rsi $0x000000000000000f
...
```

In the above, each instruction is supplemented with additional features procured by the cache simulator that we employ for initial data gathering, that is, DynamoRIO Cachesim [5].

4 Implementation

For gathering our traces, we use Cachesim, with modifications enabling writing features of a trace to a suitable SQLite database. Namely, we implement a mechanism for writing out the results of Cachesim simulation, for each execution, to a dedicated SQLite file, in a format that could be utilized later for machine learning with reasonable overhead by the pandas [33] library for Python.

4.1 DynamoRIO Instruction Features

DynamoRIO has the ability to supplement the executed instructions with various features, of which we keep those that we deem most important for the task at hand. The features for our deep learning algorithm are shown in Table 1. The reasoning for the delta encoding is the reduction of the range of possible values by way of storing merely the relative difference between sequential addresses rather than the absolute addresses. More importantly, programs generally exhibit spatio-temporal locality [10], which the delta encoding exploits by improving overall storage efficiency and reduces the overall computational cost of the encoding as opposed to computing for 64-bit addresses.

¹ The full list may be observed as part of the `trace_entry.h` header file, which is part of the DynamoRIO repository at <https://github.com/DynamoRIO/dynamorio>.

Table 1: Table of instruction row features.

Feature Name	Values	Explanation
Disassembly String	String	The instruction itself (e.g., <code>sub %rax %rsi → %rsi</code>)
Instruction Number	Positive integers	The ordinal number of the instruction during the execution of the program.
Access Address Delta	Integers	The positive or negative offset in the memory of the address of requested data w.r.t. the previous instruction, initialized from zero and resetting at each thread switch.
PC Address Delta	Integers	The positive or negative offset of the Program Counter (PC) w.r.t. the previous instruction's PC, initialized from 0 and resetting at each thread switch.
Instruction Type	Categorical	The type of the recorded program event. Can be an Instruction Fetch (ifetch), Memory Store (write), Memory Load (read), etc. ¹
Byte Count	Positive integers	Number of bytes loaded or read.
Core	Positive integers	The current simulated core that the thread is running on.
Thread Switch	Boolean	True (1) if a thread switch occurred for this entry, false (0) if it has not.
Core Switch	Boolean	True (1) if a core switch occurred for this entry, false (0.0) if it has not.
L1D size	Positive integers	The size (in bytes) of the L1 Data cache, constant throughout a single execution.
L1I size	Positive integers	The size (in bytes) of the L1 Instruction cache, constant throughout a single execution.
LL size	Positive integers	The size (in bytes) of the unified Last Level cache, constant throughout a single execution.

4.2 Tokenization

The key aspect for applying deep learning in our approach is the tokenization process, a fundamental preprocessing step that involves converting raw text data into a format interpretable by machine learning models. This is the process of transforming the input text of the disassembled instruction string into a sequence of discrete elements, or tokens, which represent the building blocks of the text. In our case, these elements are the assembly instructions, their operands and operators such as brackets and assignments (an example is shown in Fig. 5).

For this procedure, we opt for the “Byte Level” tokenizer and the “Whitespace” preprocessor by Hugging-

Face [1]. The tools available from HuggingFace are widely used and highly regarded in the field of natural language processing [7], albeit we only use a small tool subset in our work.

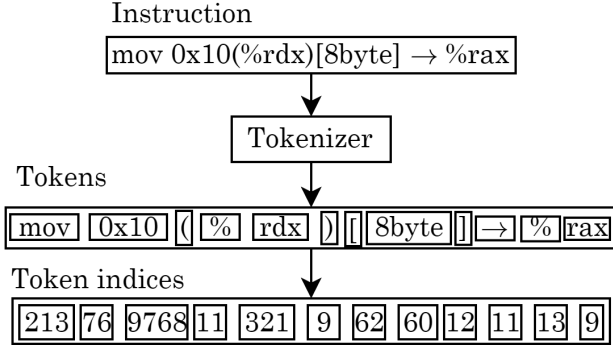


Fig. 5: x86 Instruction tokenization example (the numbers for token indices are purely exemplary)

The tokenizer is initially trained on the set of instructions present in the corresponding program traces, upon which a dictionary of tokens and corresponding token indices are generated and stored. After the tokenizer is prepared (trained, stored or loaded), each instruction is broken down into individual tokens and then the individual tokens are converted into token indices, numerical representations of tokens that a deep learning model can process. The outlined tokenization procedure is supplemented in the deep learning model with a dedicated trainable embedding layer, which contextualizes each token’s ID in the form of a specialized lookup table with preset maximum embedding length, as per Fig. 6. In our case, the maximum embedding length was 15, up until which a shorter instruction string was to be padded with [PAD] tokens, and after which longer instruction strings were truncated.

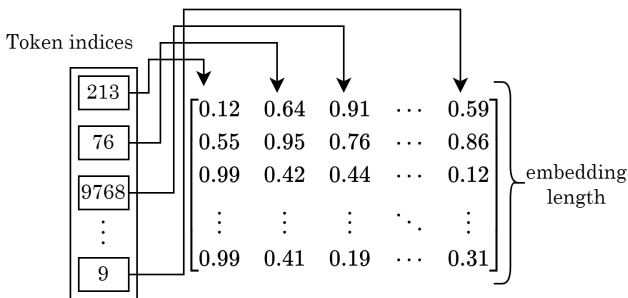


Fig. 6: Token indices lookup table, generated by the embedding layer, whereby each token is assigned a corresponding vector of a-priori fixed length (15 here)

This ensures that each instruction’s components, the operands, and the location registers are given numerical values that the model can mathematically operate on. For the desired model output, we take the cache misses for the three simulated cache levels, for which traces are executed and stored: L1 data, L1 instruction and shared LL (last-level) cache, summed up across sub-sequences. In this way, we introduce profiling of the trace without sliding the prediction window one-by-one instruction-wise, but rather sub-sequence by sub-sequence, without overlapping, which enables a significant speedup both in training and inference.

4.3 Model and Training Procedure

Our model is composed of three components: (i) a **token processing embedding layer**, (ii) an **LSTM layer** for the instructions and other features, and (iii) an **affine, fully connected dimension-reduction layer** that takes the final output of the major LSTM layer as input. The model form is described in Fig. 7 and is implemented in the PyTorch framework [29].

When a certain program is selected (typically one with parameters, such as a benchmark with variable time and stress loading limits or targets), we gather the dataset database initially by running a simulation through a custom tool within the DynamoRIO framework, with various cache parameter combinations, within a certain reasonable span of these parameters. This dataset is termed the “training/validation dataset of program X”, as it is the one in which all training will take place. In this work we opt for a training / validation ratio of 0.9.

For “testing dataset of program X”, we generate data on the same program family but with cache parameters previously unseen by the model. By *program family*, henceforth, we will refer to a set of programs with broadly similar purposes and methods, making specific cache configuration performance modeling reasonable. The models themselves are based on various hyperparameters, such as **batch size** (the number of samples that is fed into our model at each iteration of the training process), **sub-sequence length**, **LSTM hidden dimension width**, **depth**, **learning rate** and **dropout** between neighboring LSTM layers.

We also propose hyperparameter optimization for each program family, since different program types have different cache access patterns in different parts of the execution. Hyperparameter optimization is taken care of by Optuna framework [2], chosen because of the easily modifiable *Trial* class provided. The loss function that we use is MSE and we limit training to 10 epochs and the hyperparameter optimization to 5 epochs on

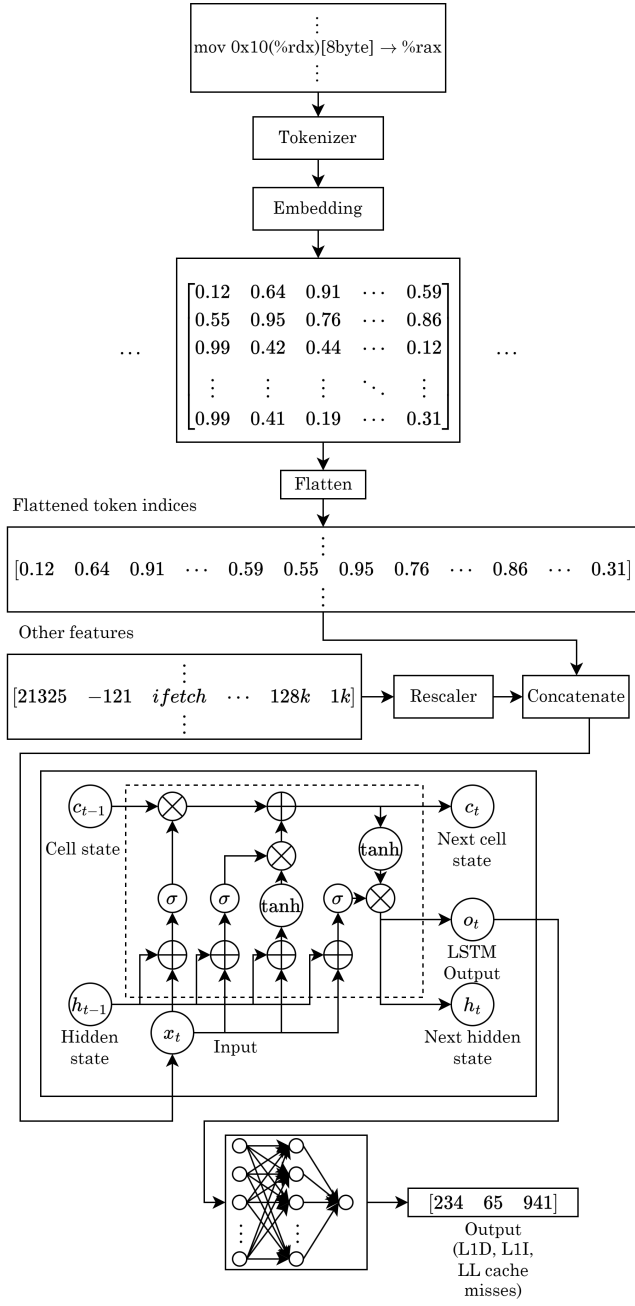


Fig. 7: Utilized LSTM model outline

a greatly reduced dataset. The reason for such limits is the relatively long training time on our hardware, reaching up to 2 hours for a single epoch on the full training dataset.

We gear the optimization towards minimizing the validation loss within the initial 5 epochs, but with an emphasis on compute time limitations. Namely, an optimization trial of 5 epochs that takes longer than a given preset amount of time is automatically considered a failure, along with its hyperparameters. This leads to a potential exclusion of otherwise satisfactory parameters

due to the constraints on execution time. Hyperparameter optimization is performed on a highly reduced data set, comprising only two databases (runs). It should be noted that several runs with improper batch size / sequence length had to be discarded during the hyperparameter optimization, as they showed to retrieve too much data for our compute machine’s RAM to handle. This was also taken care of automatically by the framework with some minor tweaks to the code.

In total, the final parameters utilized were rounded to 330 for the hidden layer dimension size, 2 LSTM layers, 0.05 for the dropout between the LSTM layers, 60 for the batch size and $N = 200$ for the sequence length.

5 Results

The experimental evaluation is carried out on an i7-13700HX, 32 GB RAM, Nvidia RTX A1000 (6 GB VRAM, 2048 CUDA cores) laptop platform. To attain a solid base for testing this approach out, a program family has to be selected such that it consists of a practical set of programs that stress specific platform components, with a broad array of algorithms. For this purpose, Sysbench [14] is chosen, along with three of its utility benchmarks: CPU, Memory, and FileIO. The CPU benchmark tests computation-intensive tasks, as efficient cache utilization, particularly L1D and L1I caches, is crucial for reducing latency. FileIO simulates disk read/write operations, with performance sensitive to cache sizes, especially the LL cache, as effective caching reduces disk access times. The Memory benchmark assesses the system’s ability to handle memory operations, with larger caches improving data retrieval speed from main memory.

Another valuable use case that we used is a highly specific image processing algorithm termed “Good Features to Track” or GFTT, exposed as a memory-intensive application [34]. Such algorithms represent a range of workloads, providing a comprehensive analysis of how cache configurations impact system performance. We introduce a wrapper utility for DynamoRIO’s cache simulator (drcachesim), to gather simulation data in the form of serverless SQLite3 databases, where one database corresponds to one program execution. The wrapper is available on GitHub (<https://github.com/ptrbman/dynamorio-missing-instructions>).

5.1 Performance Metrics

In our work, by *model performance* we refer to the ability of the model to capture the distribution of cache

misses from the original trace. We measure this ability by the values of the mean squared error and the model's R^2 , both of which are defined in further text.

In the following, let

$$\mathbf{y}_{L1D} = [y_{L1D_0}, y_{L1D_1}, \dots, y_{L1D_j}, \dots, y_{L1D_K}]^T \quad (1)$$

be the actual mini-sums of L1D cache misses across K sub-sequences, where $K = \text{program length} / \text{sub-sequence length}$ and $y_{L1D_j}, j \in [0, K]$ represents the number of cache misses in the L1 data cache in the j -th sub-sequence. Further, let

$$\hat{\mathbf{y}}_{L1D} = [\hat{y}_{L1D_0}, \hat{y}_{L1D_1}, \dots, \hat{y}_{L1D_K}]^T \quad (2)$$

be the mini-sums of L1D cache misses across K sub-sequences as predicted outputs of model inference. Similarly, we define $\mathbf{y}_{L1I}$, $\hat{\mathbf{y}}_{L1I}$, $\mathbf{y}_{LL}$, and $\hat{\mathbf{y}}_{LL}$.

The mean squared error E_{MSE} , for each of the cache types, respectively, is defined as per eq. (3).

$$E_{MSE} = \frac{1}{K} \sum_{i=1}^K (y_i - \hat{y}_i)^2 = \frac{1}{K} (\mathbf{y} - \hat{\mathbf{y}})^T (\mathbf{y} - \hat{\mathbf{y}}) \quad (3)$$

We use E_{MSE} for both training the model and as a measurement of the model's performance on the training and test datasets, since for our approach the distribution of errors throughout an execution is the most significant observation and we ideally desire to penalize the model for large outliers while retaining a smooth gradient for smaller deviations of predicted values as compared to the real values [12]. The E_{MSE} and its root square counterpart, the root mean squared error $E_{RMSE} = \sqrt{E_{MSE}}$ tell us how well the model performs with cache miss predictions throughout the execution of a program, which is one aspect of the prediction that we are interested in.

Another component of the prediction that we are interested in is the model's ability to predict total cache miss numbers, which we attain by summing up the outputs throughout the model's execution on a single program. For this measurement, we define the relative error percentage metric, based on the L1 norm, relative to the actual values, that is, E_{REP} . We let the total sum of actual cache misses be $T = \sum^K \mathbf{y}$ and the total sum of predicted cache misses be $\hat{T} = \sum^K \hat{\mathbf{y}}$, then the E_{REP} is defined as per eq. (4).

$$E_{REP} = \frac{|\sum^K \mathbf{y} - \sum^K \hat{\mathbf{y}}|}{\sum^K \mathbf{y}} \cdot 100 = \frac{|T - \hat{T}|}{T} \cdot 100 \quad (4)$$

The equation eq. (4) that we utilize here is similar to the equation for the *mean absolute error percentage*, more commonly abbreviated as MAPE. However,

while MAPE concerns the relative distribution of errors across an entire execution, E_{REP} concerns the relative discrepancy of the aggregate sums. This is of importance because the total number of cache misses is critical and a significant concern is how well the total predicted cache misses align with the total actual cache misses. For example, a E_{REP} of 0 is ideal, whereas an E_{REP} of 50 means that predictions overshoot or undershoot the real values by 50%.

The two focal values to be observed are the **mean squared error** (paired with the R^2 and E_{RMSE}), demonstrating how well the forecast of cache misses along a program execution “tracks” the actual cache misses on average, as well as the **total numbers of cache misses** at the end of executions, where each execution regards a specific cache size / core combination.

The *determination coefficient*, R^2 , is a statistical measure that indicates how well the independent variables in a regression model explain the variability of the dependent variable. It is calculated by taking the ratio of the sum of the squared differences between the predicted values and the mean of the dependent variable to the sum of the squared differences between the actual values and the mean of the dependent variable. Mathematically, $R^2 = 1 - (SS_{res}/SS_{tot})$, where SS_{res} is the sum of squared residuals (the differences between observed and predicted values), and SS_{tot} is the total sum of squares (the differences between observed values and their mean). However, it can be negative in the case that the model is highly misleading.

5.2 Execution

Training runs on several cache size, core count and program combinations:

- L1 Data cache sizes: 512 bytes, 8 kilobytes, 32 kilobytes,
- L1 Instruction cache sizes: 512 bytes, 8 kilobytes, 32 kilobytes,
- LL unified cache sizes: 1 kilobyte, 8 kilobyte, 32 kilobytes,
- Core counts: 1, 2, 4,
- Benchmarks: Sysbench CPU, Sysbench Memory, Sysbench FileIO, GFTT.

Note that we have 3 possible sizes for each L1D, L1I, LL, 3 core counts, and four benchmarks in total, giving us $3 \times 3 \times 3 \times 3 \times 4 = 324$ traces to train on.

To limit the amount of training time, due to the sheer volume of data, a limit is set for the maximum amount of instructions recorded for each execution by the DynamoRIO wrapper to $K = 25 \times 10^6$. The training datasets consist of 90% of all datasets, while 5% of the

dataset are set aside for validation during training and 5% for any residual testing.

For the full model testing, a brand new dataset is collected in a similar manner, but with different options:

- L1 Data cache sizes: 1 kilobyte, 2 kilobytes, 64 kilobytes,
- L1 Instruction cache sizes: 1 kilobyte, 2 kilobytes, 64 kilobytes,
- LL unified cache sizes: 2 kilobytes, 4 kilobyte, 16 kilobytes,
- Core counts: 1, 2, 4,
- Benchmarks: Sysbench CPU, Sysbench Memory, Sysbench FileIO, GFTT.

Of particular interest here are the 64 kB cache sizes, as they are twice the size of the maximum cache sizes from training, which potentially may produce significant issues for the model. The general executions are all of the form observed in figures 8 and 9, for the benchmark CPU L1D errors.

Fig. 8 shows an example of our predictions on one sample of executions of L1D misses of the CPU benchmark; Fig. 9 presents the first 200 sub-sequences. The trace ends at 125000 on the x-axis, due to the fact that we have $K = 25 \times 10^6$ and $N = 200$, which then corresponds to $K/N = 125 \times 10^3$.

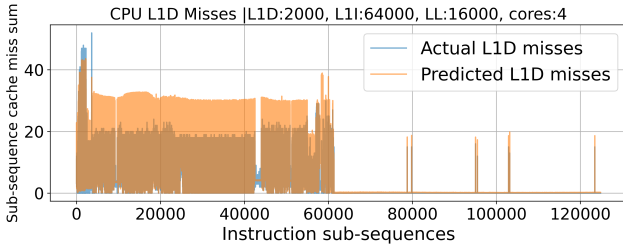


Fig. 8: Example of execution, L1D cache, Sysbench CPU Benchmark.

In this example, we can observe that the model has a tendency to exceed the number of cache misses fairly regularly. For this particular run, the E_{MSE} was 39.1 and E_{REP} was 67.4%, meaning that although the overall distribution of cache predictions is satisfactory with the E_{RMSE} around 6, the total aggregate counts leave room for improvement. The values for all E_{MSE} cache / core combinations for the tested programs, as well as E_{REP} can be observed in Figs. 10 to 13. Overall error confidence intervals with respect to core count for E_{MSE} and E_{REP} can be observed in Figs. 14 to 17. For example, for the L1I cache performance for the Memory benchmark, for the combination with L1D size 64k,

2 cores, L1I at 1k, LL at 2k, the E_{MSE} is 9.9 and the E_{REP} is 31.0, meaning that the model is 31% off the mark in terms of aggregate instructions when compared to the real value but that the overall distributed mean squared error is relatively small.

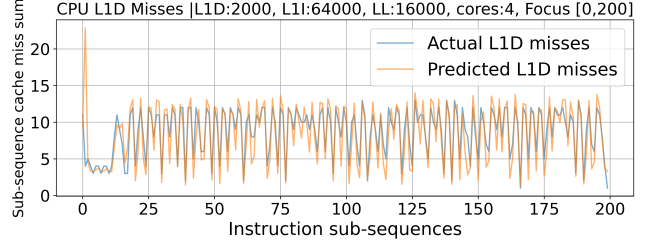


Fig. 9: Example of execution, L1D cache, Sysbench CPU Benchmark, first 200 sub-sequences

Overall, the model generally achieves satisfactory results for the distribution of cache misses, with a few notable outliers, which is to be expected given that the training procedure has been optimized to minimize MSE loss. However, when it comes to the total aggregate count of cache misses, the model consistently overestimates, particularly in the case of the GFTT program. For instance, in certain runs, especially with the 64k L1I - 16k LL cache combination, the model predicts up to 10 times more total cache misses than observed in the actual data, as shown in fig. 18. This overestimation may be due to insufficient training or the model’s difficulty fully capturing the complexity of the executing algorithm, especially for edge-case cache sizes that fall well outside the range of training data. These extreme inputs could be pushing the model to output values that significantly deviate from reality.

Regarding the analysis of output values per core count, the relatively small differences between metrics across various core counts suggest that the model performs consistently, regardless of core count. While this limits our ability to infer a strong relationship between core count and model performance, it could also indicate that the model is resilient across different core configurations. Additionally, this observation may point to potential opportunities for further enhancing the underlying DynamoRIO cache simulation to better capture the complexities of multi-core environments.

The performance of the model (in terms of E_{RMSE}) when comparing the different programs chosen may be observed at Fig. 19. The E_{RMSE} drops off significantly when raising the L1D size and L1I size, which is due to the overall fewer misses, naturally due to the cache having to retrieve data from the main memory less often. The R^2 values across different cache sizes indicate

Configuration (L1D size - Cores)	1k-1	8.0	10.5	12.0	8.0	8.9	10.7	8.9	8.9	8.9	5.9	5.6	5.6	6.3	5.6	5.8	0.1	0.1	0.1	25.3	25.4	18.9	24.5	24.4	17.0	11.7	16.4	19.6
	1k-2	8.1	10.1	11.6	8.1	8.7	10.5	8.9	9.0	9.0	5.9	5.5	5.6	6.2	5.6	5.7	0.1	0.1	0.1	25.5	26.1	19.7	24.4	24.6	18.1	11.8	16.5	20.0
	1k-4	8.4	9.5	11.1	8.3	8.5	10.2	9.1	9.4	9.3	5.8	5.4	5.4	6.2	5.6	5.7	0.2	0.2	0.2	25.7	27.0	22.8	24.1	25.0	21.7	12.0	16.7	20.3
	2k-1	31.0	25.8	24.2	32.2	26.8	25.2	38.9	30.2	31.9	5.9	5.6	5.6	6.3	5.8	6.0	0.1	0.1	0.1	25.6	22.3	19.0	25.1	23.4	17.7	11.7	15.3	16.8
	2k-2	31.1	26.4	24.6	32.3	27.3	25.9	39.0	29.9	31.6	5.9	5.5	5.6	6.3	5.7	5.9	0.1	0.1	0.1	26.0	23.4	19.2	25.2	24.2	18.3	11.8	15.3	16.8
	2k-4	31.4	27.3	25.9	32.4	28.0	27.0	39.1	29.3	30.7	5.8	5.5	5.5	6.3	5.8	5.9	0.2	0.1	0.2	26.4	25.3	21.6	25.1	25.4	21.2	12.0	15.3	16.7
	64k-1	0.9	1.0	1.0	1.0	1.0	1.0	1.1	1.1	1.1	28.6	30.3	28.9	16.2	17.5	16.4	0.1	0.1	0.1	2.5	10.6	19.2	2.3	6.6	18.1	1.3	1.3	1.3
	64k-2	1.0	1.0	1.0	1.0	1.0	1.0	1.1	1.1	1.1	25.8	27.9	26.3	14.7	15.7	14.7	0.2	0.2	0.2	2.5	10.5	19.6	2.3	6.6	18.1	1.3	1.4	1.4
	64k-4	1.0	1.0	1.0	1.0	1.0	1.0	1.1	1.1	1.1	22.5	24.6	22.9	13.2	14.2	13.3	0.2	0.2	0.2	2.5	10.5	20.7	2.3	6.7	18.3	1.3	1.4	1.4
	1k-1	13.8	4.4	1.7	15.3	8.2	5.4	21.5	11.1	13.0	26.9	24.9	25.2	30.0	28.0	27.9	96.4	56.5	72.1	55.7	20.0	4.4	54.4	23.7	6.9	60.8	31.0	49.2
	1k-2	14.1	5.6	2.8	15.4	9.0	6.2	21.3	10.8	12.6	26.5	24.5	24.7	29.2	27.0	26.9	85.9	48.4	60.7	55.0	21.3	8.1	53.1	24.5	10.3	57.8	31.1	49.9
	1k-4	14.3	7.5	4.8	15.4	10.2	7.7	21.0	9.9	11.6	25.7	23.7	23.9	27.1	24.1	24.2	81.1	46.4	55.7	52.7	23.1	17.1	49.4	25.1	18.2	51.4	30.8	50.4
	2k-1	49.8	31.6	26.4	53.8	39.0	33.7	68.2	46.8	50.8	27.2	24.8	25.1	30.8	28.4	28.5	98.5	47.6	65.4	52.9	18.7	0.4	53.1	23.6	4.2	55.8	27.6	36.2
	2k-2	50.9	33.9	28.4	54.5	40.7	35.5	68.0	46.0	49.9	26.9	24.4	24.7	30.2	27.6	27.6	87.4	39.4	54.5	52.9	20.4	4.3	52.5	24.8	7.8	52.8	27.3	36.4
	2k-4	52.3	37.9	33.0	54.9	43.5	39.3	67.4	44.0	47.4	26.3	23.9	24.1	28.5	25.2	25.2	82.0	39.2	49.5	51.7	23.1	13.6	49.7	26.1	15.9	45.8	26.3	35.7
	64k-1	4.9	3.4	2.3	2.3	1.6	0.7	31.9	23.2	24.6	145.3	152.2	145.6	96.4	105.3	98.3	9.8	0.4	1.4	26.1	53.0	145.3	22.3	32.3	131.9	27.3	21.7	20.7
	64k-2	5.4	4.3	2.9	2.4	1.6	0.5	32.3	22.8	24.4	132.5	140.5	133.4	87.4	93.6	87.6	11.5	1.2	3.0	24.9	50.8	149.0	21.1	30.6	131.7	27.8	21.3	20.5
	64k-4	6.3	5.0	3.8	3.3	1.4	0.6	33.3	22.3	24.0	116.4	123.8	116.9	79.9	84.4	79.5	6.2	3.1	1.4	22.7	47.5	159.4	18.5	29.6	133.7	28.9	20.7	20.1
	1k-16k																											
	1k-2k																											
	1k-4k																											
	2k-16k																											
	2k-2k																											
	2k-4k																											
	64k-16k																											
	64k-2k																											
	64k-4k																											
Configuration (L1I size - LL size)																												

Fig. 10: CPU Benchmark prediction measurements; top in left to right order: L1D E_{MSE} , L1I E_{MSE} , LL E_{MSE} ; bottom in left-to-right order: L1D E_{REP} , L1I E_{REP} , LL E_{REP} , green is better, red is worse

Configuration (L1D size - Cores)	1k-1	16.5	31.9	29.3	17.7	28.2	27.8	13.0	14.6	14.0	8.4	7.9	8.0	16.4	16.7	17.1	0.1	0.1	0.1	31.9	42.9	53.6	28.3	55.8	54.7	12.0	20.9	30.0	
	1k-2	18.4	34.2	31.0	18.6	29.8	28.8	13.1	15.7	15.0	8.3	7.7	7.9	16.1	16.5	16.8	0.1	0.1	0.1	31.8	45.6	54.5	28.3	57.4	56.7	12.1	21.9	30.2	
	1k-4	22.3	37.8	34.7	21.1	33.3	31.6	13.3	18.8	17.9	8.2	7.5	7.7	15.5	16.1	16.1	0.2	0.2	0.2	31.4	50.9	57.5	28.4	60.8	62.0	12.0	24.6	30.7	
	2k-1	40.4	44.2	41.6	41.6	44.5	42.4	45.6	38.8	40.0	8.4	7.8	8.0	16.1	16.7	17.4	0.1	0.1	0.1	30.7	40.3	63.4	28.5	56.5	63.4	12.0	23.3	23.6	
	2k-2	41.7	46.0	43.4	42.6	46.0	43.6	45.6	39.6	40.8	8.2	7.6	7.8	15.8	16.5	17.0	0.1	0.1	0.1	30.8	41.9	62.4	28.6	57.1	63.6	12.0	24.4	24.3	
	2k-4	44.7	48.6	46.1	44.5	48.3	46.1	45.7	41.3	42.4	8.1	7.4	7.6	15.2	16.0	16.1	0.2	0.2	0.2	30.5	45.4	60.4	28.7	58.9	65.0	12.0	26.6	25.5	
	64k-1	1.0	1.1	1.0	1.1	1.1	1.1	1.1	1.1	1.1	1.1	10.3	10.1	10.0	18.8	18.5	18.4	0.2	0.2	0.2	2.6	6.6	55.7	2.5	21.2	53.9	1.3	1.4	1.4
	64k-2	1.0	1.1	1.1	1.1	1.1	1.1	1.1	1.1	1.2	1.2	10.1	9.9	9.9	18.4	18.2	18.1	0.2	0.2	0.2	2.6	6.6	55.5	2.5	21.3	53.7	1.4	1.4	1.5
	64k-4	1.1	1.1	1.1	1.1	1.1	1.1	1.1	1.2	1.2	1.2	9.8	9.7	9.6	17.8	17.6	17.5	0.2	0.2	0.2	2.6	6.6	54.8	2.4	20.9	53.0	1.4	1.5	1.6
1k-1	3.3	15.8	13.7	3.1	10.7	10.6	17.4	5.5	7.5	20.0	20.1	21.4	16.6	21.6	24.5	88.0	77.0	94.3	86.2	0.3	5.7	74.6	4.6	1.5	67.1	18.0	51.0		
1k-2	2.3	16.7	14.3	2.7	11.3	10.6	17.4	4.8	6.7	20.4	20.0	21.2	16.2	20.9	24.2	80.9	71.6	86.4	84.9	0.9	3.6	74.5	4.0	0.6	64.2	17.4	50.2		
1k-4	0.1	17.8	15.6	1.4	12.6	11.3	17.4	2.8	4.4	20.5	19.4	20.3	24.9	18.2	21.2	73.3	64.9	79.2	80.6	2.0	0.2	73.0	1.9	4.1	57.9	15.2	47.5		
2k-1	24.2	2.7	0.7	26.9	4.6	5.8	54.4	29.8	34.0	19.8	20.0	21.2	17.3	19.0	21.8	87.8	72.2	89.8	76.5	0.6	13.2	70.8	2.2	9.8	62.7	16.5	28.9		
2k-2	23.0	2.4	1.2	26.5	4.7	6.0	54.2	27.8	31.7	20.3	20.0	21.0	16.5	19.0	21.7	81.1	65.5	82.4	75.4	0.5	11.4	70.6	2.2	7.7	59.8	14.9	27.2		
2k-4	20.1	1.4	1.4	25.1	4.6	5.4	53.7	23.0	26.6	20.4	19.5	20.2	15.1	17.4	19.2	73.0	60.1	73.7	72.2	0.2	7.2	69.2	1.2	4.6	53.1	10.7	22.8		
64k-1	1.2	0.5	0.1	0.8	3.4	3.8	21.6	12.8	14.1	31.9	31.6	31.3	60.3	59.9	59.4	29.8	25.6	26.4	26.0	19.0	314.8	23.7	61.0	305.0	17.0	9.5	8.7		
64k-2	1.7	0.8	0.2	0.3	2.9	3.3	22.8	13.7	15.1	31.4	31.0	30.7	59.4	58.9	58.5	30.5	25.6	26.4	25.4	18.7	313.9	23.2	60.9	304.2	17.9	10.2	9.4		
64k-4	2.6	1.6	1.2	0.8	1.4	1.7	25.0	15.4	16.9	30.7	30.0	29.7	57.7	57.1	56.7	29.8	23.2	24.3	24.2	17.9	310.2	22.2	59.2	300.1	19.7	11.6	10.8		
1k-16k																													
1k-2k																													
1k-4k																													
2k-16k																													
2k-2k																													
2k-4k																													
64k-16k																													
64k-2k																													
64k-4k																													
Configuration (L1I size - LL size)																													

Fig. 11: Memory Benchmark prediction measurements; top in left to right order: L1D E_{MSE} , L1I E_{MSE} , LL E_{MSE} ; bottom in left-to-right order: L1D E_{REP} , L1I E_{REP} , LL E_{REP} , green is better, red is worse

||
||
||

1k-1	30.5	38.0	36.8	31.6	39.0	38.1	30.6	25.0	27.2	20.9	17.8	18.7	20.5	17.1	18.1	0.1	0.1	0.1	18.6	27.9	21.9	16.7	30.7	23.6	13.5	9.6	10.9
1k-2	30.5	36.8	35.2	31.4	37.8	36.6	29.5	27.0	28.6	20.0	16.8	17.7	18.8	15.5	16.3	0.1	0.1	0.1	17.1	28.3	22.4	15.6	30.6	23.7	12.6	8.8	9.3
1k-4	32.4	35.0	33.7	32.6	36.0	34.7	29.9	29.6	30.6	17.8	14.7	15.5	15.1	12.5	13.1	0.1	0.1	0.1	15.2	31.5	24.9	14.3	32.0	24.9	11.7	7.3	7.0
2k-1	13.7	10.2	10.6	13.7	10.4	10.7	54.6	20.2	22.3	21.0	18.0	18.8	20.7	17.3	18.3	0.1	0.1	0.1	18.7	24.7	23.0	16.9	27.1	24.9	13.0	8.1	8.2
2k-2	13.7	10.2	10.7	13.9	10.2	10.7	58.4	20.9	24.0	20.1	17.0	17.9	19.0	15.6	16.5	0.1	0.1	0.1	17.4	25.1	23.7	15.9	27.0	25.2	12.4	7.7	7.5
2k-4	13.5	10.7	11.2	13.6	10.8	11.3	67.0	23.7	30.4	17.9	14.8	15.7	15.3	12.7	13.3	0.1	0.1	0.1	15.8	27.4	26.4	14.8	27.9	26.7	11.9	7.1	6.7
64k-1	15.9	15.1	15.2	16.0	15.1	15.3	18.4	18.0	17.9	26.2	26.8	26.7	27.5	28.0	27.9	0.1	0.1	0.1	24.2	29.7	31.6	23.4	29.1	30.1	16.2	15.2	15.1
64k-2	15.6	14.9	15.0	15.7	15.0	15.1	18.7	18.1	18.1	25.9	26.4	26.3	26.9	27.3	27.3	0.1	0.1	0.1	23.3	29.3	30.5	22.4	28.4	28.8	16.5	15.3	15.4
64k-4	15.7	14.9	15.0	15.8	14.9	15.1	19.5	18.4	18.6	25.2	25.6	25.6	25.8	26.2	26.2	0.1	0.1	0.1	22.3	28.2	28.8	21.3	26.9	27.2	17.2	15.7	15.9
1k-1	13.6	20.6	19.5	13.7	21.1	20.2	1.1	9.3	9.5	84.9	70.2	74.1	82.8	56.0	63.0	357.6	425.2	393.6	4.4	3.4	0.6	3.1	5.9	1.9	0.4	10.5	14.8
1k-2	13.2	19.2	17.7	13.2	19.9	18.6	3.6	9.8	9.1	84.3	69.2	73.4	78.8	52.0	58.9	390.0	446.4	405.9	2.8	5.3	1.4	1.5	7.8	3.9	3.6	7.1	10.5
1k-4	13.9	16.8	15.7	13.6	17.6	16.4	6.0	9.2	6.8	80.9	64.9	69.2	68.9	43.1	49.3	496.6	463.2	433.4	0.3	9.4	5.6	1.0	11.5	7.7	8.5	2.0	0.5
2k-1	7.8	2.7	3.2	7.3	2.4	2.8	28.1	15.6	15.1	85.0	70.8	74.6	84.1	57.7	64.3	517.0	421.8	418.8	4.1	1.7	2.1	2.9	4.2	4.5	1.2	6.8	7.3
2k-2	8.5	3.9	4.6	8.2	3.4	4.0	31.5	14.4	14.7	85.0	69.8	73.8	80.0	53.7	60.3	515.0	456.6	452.6	2.6	3.4	3.8	1.5	6.0	6.2	3.9	3.9	3.8
2k-4	8.8	6.0	6.7	8.6	5.2	5.9	37.4	13.3	16.1	81.7	65.9	70.3	70.3	44.5	50.8	562.2	535.3	546.8	0.3	7.2	7.6	1.0	9.5	9.7	9.1	3.2	3.7
64k-1	0.8	6.6	5.9	1.7	7.3	6.6	2.2	2.6	2.1	116.2	118.0	117.9	131.0	131.7	131.6	630.8	672.8	646.9	34.0	29.8	43.3	33.6	29.9	42.9	1.3	6.6	5.5
64k-2	1.2	7.3	6.6	2.2	7.9	7.3	4.5	0.4	0.1	116.8	118.0	118.0	130.4	130.4	130.5	726.8	699.1	677.5	33.1	28.7	42.2	32.6	28.6	41.5	0.6	4.6	3.5
64k-4	1.1	8.3	7.6	1.9	9.1	8.4	9.9	4.2	4.7	115.7	116.1	116.3	126.4	126.2	126.5	1092.8	793.5	799.8	30.9	26.5	39.9	30.0	26.4	39.2	5.0	0.8	0.2
	1k-16k	1k-2k	1k-4k	2k-16k	2k-2k	2k-4k	64k-16k	64k-2k	64k-4k	1k-16k	1k-2k	1k-4k	2k-16k	2k-2k	2k-4k	64k-16k	64k-2k	64k-4k	1k-16k	1k-2k	1k-4k	2k-16k	2k-2k	2k-4k	64k-16k	64k-2k	64k-4k

Fig. 13: GFTT prediction measurements; top in left to right order: L1D E_{MSE} , L1I E_{MSE} , LL E_{MSE} ; bottom in left-to-right order: L1D E_{REP} , L1I E_{REP} , LL E_{REP} , green is better, red is worse

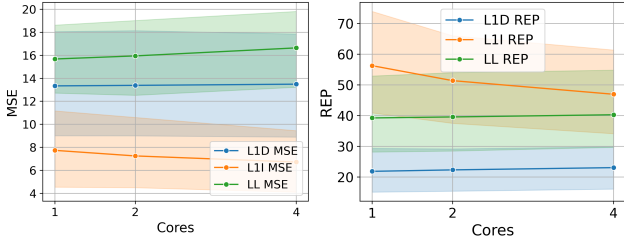


Fig. 14: Sysbench CPU Benchmark Error plot per core count

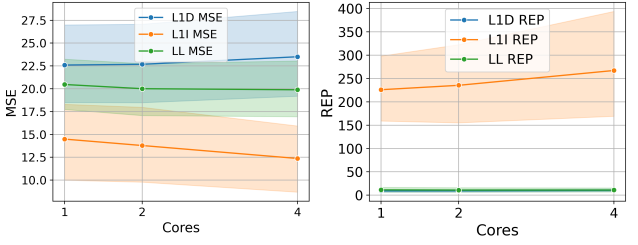


Fig. 17: GFTT Error plot per core count

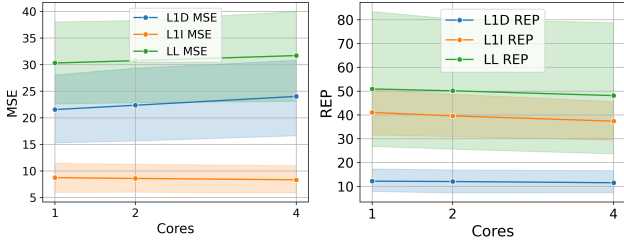


Fig. 15: Sysbench Memory Benchmark Error plot per core count

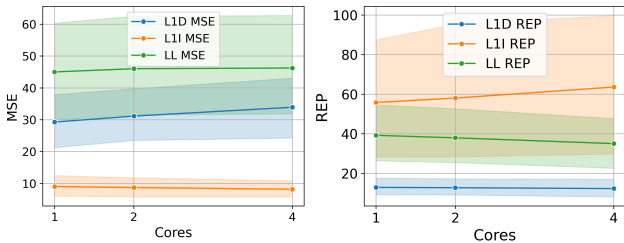


Fig. 16: Sysbench FileIO Benchmark Error plot per core count

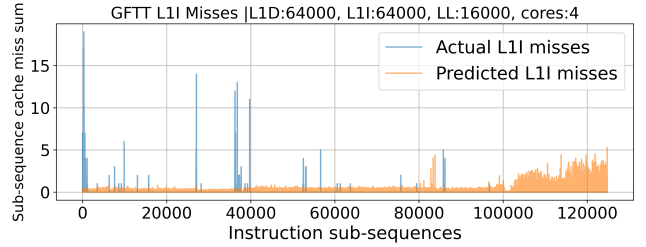


Fig. 18: GFTT, worst execution, L1I cache

CPU and FileIO show mostly positive R^2 values, indicating reasonable model performance.

The GFTT L1I shows a negative R^2 at cache sizes of 2k and 64k, suggesting that the model's performance in these specific cases may require further refinement - a topic for future work. Importantly, the CPU benchmark's cache behavior remains the most reliably well-modeled among our tests. Although FileIO and Memory benchmarks exhibit significant deviations at 2k for the L1D cache size and at 4k for the LL cache size, these findings highlight areas where the model could be improved - some improvements are suggested in section 8. This may indicate that our current model, in its present form, does not fully capture the cache dynam-

ics for these particular types of programs. Alternatively, there may be unique characteristics in these workloads sensitive to parameters beyond those explored in our study (Table 1) that the model has yet to account for.

All of the final numerical results of the executions may be observed for review purposes in the appendix, Tables 2 to 5.

6 Discussion

Our machine learning-based cache simulation model demonstrated consistent performance across the CPU, Memory, FileIO and GFTT benchmarks, with each completing 25 million instructions in approximately 45 seconds, whereas the duration was approximately 22 seconds on the i7-13700HX platform using DynamoRio. However, our model provides the premises for a critical advantage: the ability to execute in parallel multiple sequential subsets of a program trace. This capability allows for distributed processing across multiple platforms, with the potential of significantly reducing the overall simulation time compared to DynamoRIO.

Our assumption above is based on the following. Each execution trace was divided into eight parts, and the model was run and timed on each part, with the goal of analyzing the predictive power of the model on subsections of the program trace running on multiple GPUs. On average, the execution time of each section was 17.35 seconds with $\sigma = 2.14$ seconds, producing, on average, a **21% time reduction** to our model when compared to running the simulation through Cachesim.

While our model excels in scenarios with predictable data access patterns, particularly in CPU workloads, it encounters challenges with more complex I/O operations, as reflected in the negative R^2 values observed in the FileIO benchmark at smaller L1D cache sizes. This may be due to several contributing factors, including synchronous disk operations, which are known to increase memory latency and cause cache inefficiencies, leading to more frequent cache misses. Background I/O processes, such as `fsync`², can further exacerbate cache pressure, resulting in additional misses. Frequent and diverse file accesses can lead to sub-optimal data access patterns, increasing cache misses and making the modeling process more complex [19].

Despite these challenges, the flexibility of the model and the potential for parallelization represent signifi-

cant advantages. Unlike DynamoRIO, which requires running the entire program for effective simulation, our approach allows for faster, more scalable simulations. This work represents a proof-of-concept and a first step towards a more ambitious model. Future work should focus on enhancing the model’s accuracy across varied and complex workloads, particularly by improving its handling of edge cases and its generalization capabilities. These improvements could make the model a more robust tool for simulated cache performance across a wider range of applications.

7 Related Work

The predictive capabilities of LSTM-based models for cache miss occurrences on instrumentation traces have previously been investigated, as demonstrated by R. Jha et al. [13], validating the feasibility of employing LSTMs to anticipate cache miss patterns derived from program behaviors. This investigation underpins our own predictive modeling endeavors, yet our approach uniquely emphasizes the variability of cache sizes alongside the programs being analyzed. Various successful integrations of LSTM architectures into prefetchers aimed at mitigating cache misses - rather than simply modeling their distribution - have been documented, exemplified by J. Rogers [30] and Y. Zeng et al. [36]. For simulation of program performance using sequential deep learning, the Ithemal tool [24] epitomizes a sophisticated LSTM-driven technique to forecast instruction block latency; it operates, though, under the presumption of negligible cache misses within the program trace, contrasting with the focus of our method.

An alternative approach to latency prediction, predicated on parallelization rather than sequential modeling, is detailed by Li et al. [22] with the SimNet temporal convolutional network framework. Further work by Li et al. [21] showcases a higher-level simulation framework called PerfVec that is demonstrated in an instruction latency reduction objective against different combinations of cache sizes.

The work of S. Pandey et al. [27] suggests optimizations of SimNet and Ithemal to reduce the costs of moving datasets from the CPU to the GPU for overhead reduction purposes, and the most very recent work by S. Pandey et al. [28] showcases an effort in the direction of a microarchitecture simulator and dataset gathering approach that can be transferred between architectures differing by certain cache size and CPU pipeline characteristics. However, this tool focuses on the forecasting of metrics including cache misses on the granularity level of a single instruction, whereas our tool treats instruction sub-sequences in a statistical manner, which gives

² `fsync` is a system call that forces a file’s in-memory data to be written to disk, ensuring data integrity in case of a system crash. However, this operation can introduce significant performance overhead, especially in I/O-intensive workloads, due to the additional time required for disk synchronization [18].

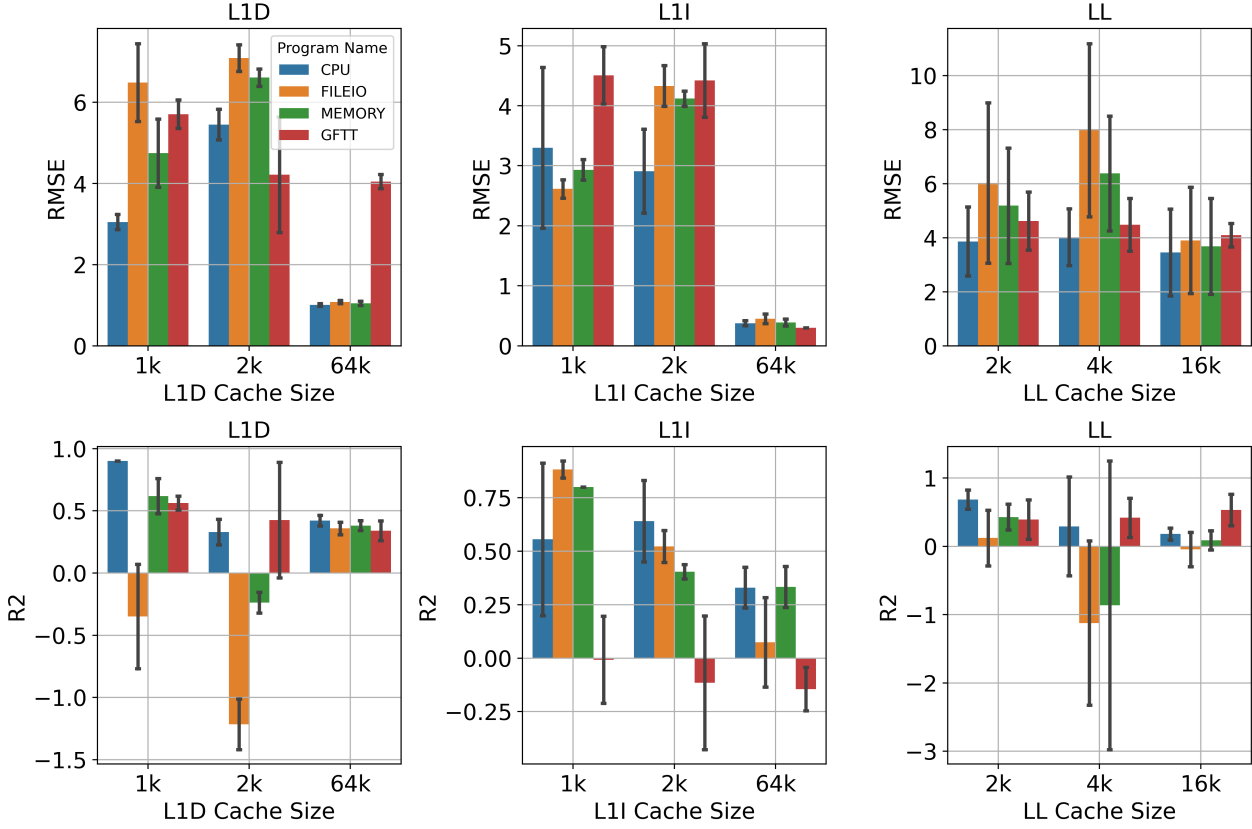


Fig. 19: E_{RMSE}/R^2 analysis of the programs used in the model testing, cache E_{RMSE}/R^2 against respective cache sizes (an R^2 of 1.0 is ideal), standard deviation for the confidence intervals

the user the choice between more granularity (fewer instructions per sub-sequence, but slower) or more speed (more instructions per sub-sequence, but less focus on individual instructions).

Additionally, our approach gives the user flexibility in terms of providing estimates for hypothetical but possible cache and multicore/single-core architectures, whereas the availability of combinations in the referenced work is constrained by the microarchitecture loaded in gem5. Moreover, our work also implements the ability of introducing higher core counts and the corresponding L1D/L1I caches, as well as the corresponding analysis that can be attained from running the model.

The robust and frequently updated software instrumentation tool, DynamoRIO, is a preferred solution among contemporary computer science researchers, as evidenced by [8,35]. We utilize DynamoRIO not solely for its prominence and ongoing advancements, but also for its open-source status, user-friendliness, and seamless integration with our tool for capturing program execution traces into its infrastructure.

8 Conclusion and Future work

This work introduces a new method for modeling cache performance tailored to specific applications using a distributable deep learning model.

We have shown that our approach not only mimics traditional simulators but also replicates results across multiple architectures. Our model handles complex Sysbench benchmarks and GFTT algorithm implementations with remarkable accuracy, even when faced with unfamiliar cache sizes and core counts. Although it tends to predict higher cache miss rates, it is excellent at pinpointing high-activity sections of program execution.

Though DynamoRIO outpaces our model on a single machine, our model has the potential to outperform it through parallelization: just three to four concurrent GPU iterations would suffice.

While relying on tools that provide hardware-specific traces, our approach itself is hardware-agnostic.

We have also provided the initial steps towards extending our approach on individual cores of multicore architectures.

Future work. We are set on fine-tuning our model on the present architecture, such as employing positional encoding for the tokens. We will also focus on minimizing cache-miss discrepancies during training by honing in on specific application benchmarks rather than broad benchmark sets. In addition, we plan to enhance the model architecture with batch normalization and regularization techniques to push performance further.

Extracting more features from Cachesim and breaking down instructions more granularly are on the agenda, as is expanding our approach using advanced simulators like gem5 and ZSim for richer data, for example extracting used registers, individual instruction latencies, branch predictor types and branch predictions and mispredictions. Richer trace data may prove to be a key component in tackling intense workloads that depend heavily on additional information about the executed instructions and the underlying access patterns, such as the FileIO benchmark where our approach faced a significant challenge, as was shown in section 6. In addition to this, we will explore features of independent cores with regard to cache modeling.

We are also exploring different recurrent architectures, such as temporal convolutional networks and gated recurrent units, to boost predictive accuracy and efficiency. Crucially, our model is non-fixed; by feeding architectures as features into the model, we can effectively compare various CPU characteristics. Our ultimate goal? Develop a flexible ML simulation framework that holistically models hardware for specialized software, allowing direct comparisons across diverse CPU configurations. The challenges are significant, but so are the opportunities, and we are committed to making substantial contributions now and in the future.

Acknowledgement. The author's work was partially supported by the Swedish Knowledge Foundation via the project PerFlex - *Performant and Flexible digital Systems through Verifiable Artificial Intelligence*, grant nr. 20220033.

References

- Summary of the tokenizers. https://huggingface.co/docs/transformers/en/tokenizer_summary. (Accessed on 03/04/2024)
- Akiba, T., Sano, S., Yanase, T., Ohta, T., Koyama, M.: Optuna: A next-generation hyperparameter optimization framework. In: Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (2019)
- Alzubaidi, L., Zhang, J., Humaidi, A.J., Al-Dujaili, A., Duan, Y., Al-Shamma, O., Santamaría, J., Fadhel, M.A., Al-Amidie, M., Farhan, L.: Review of deep learning: concepts, cnn architectures, challenges, applications, future directions. *Journal of Big Data* 2021 8:1 **8**, 1–74 (2021). DOI 10.1186/S40537-021-00444-8. URL <https://journalofbigdata.springeropen.com/articles/10.1186/s40537-021-00444-8>
- Bellard, F.: Qemu, a fast and portable dynamic translator. In: Proceedings of the Annual Conference on USENIX Annual Technical Conference, ATEC '05, p. 41. USENIX Association, USA (2005)
- Bruening, D., Zhao, Q., Amarasinghe, S.: Transparent dynamic instrumentation (2012). DOI 10.1145/2151024.2151043
- Carvalho, C.: The gap between processor and memory speeds. In: Proc. of IEEE International Conference on Control and Automation, vol. 5000, p. 15000 (2002)
- Castano, J., Martínez-Fernández, S., Franch, X., Bogner, J.: Analyzing the evolution and maintenance of ml models on hugging face. In: Proceedings of the 21st International Conference on Mining Software Repositories, MSR '24, p. 607–618. Association for Computing Machinery, New York, NY, USA (2024). DOI 10.1145/3643991.3644898
- Darashkevich, E., Rusyaev, M.R., Russia, H., Korostinskiy, M.R., Bugayenko, M.Y.: A minority of c++ objects account for the majority of allocation cpu time URL <https://github.com/EugeneDar/dynamotool>
- Etiemble, D.: 45-year CPU evolution: one law and two equations. In: Second Workshop on Pioneering Processor Paradigms. Vienne, Austria (2018). URL <https://hal.science/hal-01719766>
- Hennessy, J.L., Patterson, D.A.: Computer Architecture: A Quantitative Approach, 5 edn. Morgan Kaufmann, Amsterdam (2012)
- Hochreiter, S., Schmidhuber, J.: Long short-term memory. *Neural Computation* **9**, 1735–1780 (1997). DOI 10.1162/NECO.1997.9.8.1735
- Jadon, A., Patil, A., Jadon, S.: A comprehensive survey of regression based loss functions for time series forecasting (2022). URL <https://arxiv.org/abs/2211.02989>
- Jha, R., Karuvaly, A., Tiwari, S., Moss, J.E.B.: Cache miss rate predictability via neural networks
- Kopytov, A.: Sysbench manual. MySQL AB pp. 2–3 (2012)
- Kratsios, A.: The universal approximation property. *Annals of Mathematics and Artificial Intelligence* **89**, 435–469 (2019). DOI 10.1007/s10472-020-09723-1. URL <http://arxiv.org/abs/1910.03344><http://dx.doi.org/10.1007/s10472-020-09723-1>
- Kumar, S., Singh, P.K.: An overview of modern cache memory and performance analysis of replacement policies. Proceedings of 2nd IEEE International Conference on Engineering and Technology, ICETECH 2016 **17**, 210–214 (2016). DOI 10.1109/ICETECH.2016.7569243
- Kwak, H., Lee, B., Hurson, A., Suk-Han Yoon, Woo-Jong Hahn: Effects of multithreading on cache performance. *IEEE Transactions on Computers* **48**(2), 176–184 (1999). DOI 10.1109/12.752659. URL <http://ieeexplore.ieee.org/document/752659/>
- Lee, C.G., Byun, H., Noh, S., Kang, H., Kim, Y.: Write optimization of log-structured flash file system for parallel i/o on manycore servers. In: Proceedings of the 12th ACM International Conference on Systems and Storage, SYSTOR '19, p. 21–32. Association for Computing Machinery, New York, NY, USA (2019). DOI 10.1145/3319647.3325828
- Lee, J., Bahn, H.: File access characteristics of deep learning workloads and cache-friendly data management. In: 2023 10th International Conference on Electrical Engineering, Computer Science and Informatics (EECSI),

- pp. 328–331 (2023). DOI 10.1109/EECSI59885.2023.10295817
20. Lee, J., Kim, H., Vuduc, R.: When prefetching works, when it doesn't, and why. *ACM Trans. Archit. Code Optim.* **9**(1) (2012). DOI 10.1145/2133382.2133384
 21. Li, L., Flynn, T., Hoisie, A.: Learning generalizable program and architecture representations for performance modeling (2024)
 22. Li, L., Pandey, S., Flynn, T., Liu, H., Wheeler, N., Hoisie, A.: Simnet: Accurate and high-performance computer architecture simulation using deep learning. *Proceedings of the ACM on Measurement and Analysis of Computing Systems* **6**(2), 1–24 (2022)
 23. Luk, C.K., Cohn, R., Muth, R., Patil, H., Klauser, A., Lowney, G., Wallace, S., Reddi, V.J., Hazelwood, K.: Pin: building customized program analysis tools with dynamic instrumentation. *Acm sigplan notices* **40**(6), 190–200 (2005)
 24. Mendis, C., Renda, A., Amarasinghe, S., Carbin, M.: Ithema: Accurate, portable and fast basic block throughput estimation using deep neural networks. In: *International Conference on machine learning*, pp. 4505–4515. PMLR (2019)
 25. Nethercote, N., Seward, J.: Valgrind: a framework for heavyweight dynamic binary instrumentation. In: *Proceedings of the 28th ACM SIGPLAN Conference on Programming Language Design and Implementation, PLDI '07*, p. 89–100. Association for Computing Machinery (2007). DOI 10.1145/1250734.1250746
 26. Nishijima, T.: Universal approximation theorem for neural networks (2021). URL <https://arxiv.org/abs/2102.10993v1>
 27. Pandey, S., Li, L., Flynn, T., Hoisie, A., Liu, H.: Scalable deep learning-based microarchitecture simulation on gpus. In: *SC22: International Conference for High Performance Computing, Networking, Storage and Analysis*, pp. 1–15 (2022). DOI 10.1109/SC41404.2022.00084
 28. Pandey, S., Yazdanbakhsh, A., Liu, H.: Tao: Re-thinking dl-based microarchitecture simulation. *Proceedings of the ACM on Measurement and Analysis of Computing Systems* **8**, 1–25 (2024). DOI 10.1145/3656012
 29. Paszke, A., Gross, S., Massa, F., Lerer, A., Bradbury, J., Chanan, G., Killeen, T., Lin, Z., Gimelshein, N., Antiga, L., et al.: Pytorch: An imperative style, high-performance deep learning library. *Advances in neural information processing systems* **32** (2019)
 30. Rogers, J.: Effects of an lstm composite prefetcher (2019)
 31. Sandberg, A., Nikoleris, N., Carlson, T.E., Hagersten, E., Kaxiras, S., Black-Schaffer, D.: Full speed ahead: Detailed architectural simulation at near-native speed. In: *2015 IEEE International Symposium on Workload Characterization*, pp. 183–192. IEEE (2015)
 32. Sherstinsky, A.: Fundamentals of recurrent neural network (rnn) and long short-term memory (lstm) network. *Physica D: Nonlinear Phenomena* **404**, 132306 (2020)
 33. pandas development team, T.: pandas-dev/pandas: Pandas (2020). DOI 10.5281/zenodo.3509134
 34. Yaghoobi, S.: Leveraging machine learning for fast performance prediction for industrial systems : Data-driven cache simulator. Master's thesis, Mälardalen University, School of Innovation, Design and Engineering (2024)
 35. Yang, Y., Gao, C., Li, Z., Wang, Y., Wang, R.: Binary level concolic execution on windows with rich instrumentation based taint analysis. In: *International Symposium on Dependable Software Engineering: Theories, Tools, and Applications*, pp. 351–367. Springer (2023)
 36. Zeng, Y., Guo, X.: Long short term memory based hardware prefetcher: a case study. In: *Proceedings of the International Symposium on Memory Systems, MEMSYS '17*, p. 305–311. Association for Computing Machinery, New York, NY, USA (2017). DOI 10.1145/3132402.3132405

Appendix - PyTorch model code

```

class CombinedLSTMModel(nn.Module):
    def __init__(
        self,
        token_vocab_size, # characteristic of the tokenizer
        token_embedding_dim, # number of tokens per instruction
        access_feature_size, # number of trace features
        hidden_dim, # matrix size of the hidden dimension of the LSTM layer
        output_dim, # number of outputs of our model
        num_layers, # number of consecutive LSTM layers
        dropout, # dropout rate between the LSTM layers
    ):
        super(CombinedLSTMModel, self).__init__()
        self.token_embedding = nn.Embedding(token_vocab_size, token_embedding_dim)
        # Input size must be the sum of flattened token embedding size
        # and access feature size
        self.lstm = nn.LSTM(
            input_size=token_embedding_dim*2 + access_feature_size,
            hidden_size=hidden_dim,
            num_layers=num_layers,
            batch_first=True,
            dropout=dropout,
        )
        self.fc = nn.Linear(hidden_dim, output_dim)

    def forward(self, token_features, access_features):
        # Embed the token features
        embedded_tokens = self.token_embedding(token_features)
        # Assume access_features is already of the appropriate size and
        # requires no embedding
        # LSTM can handle variable sequence lengths if necessary, but here we
        # focus on batch size flexibility
        elongated_embeddings = embedded_tokens.view(
            token_features.size(0), token_features.size(1), -1
        )
        combined_features = torch.cat(
            (elongated_embeddings, access_features), dim=2
        ) # Concatenate along the feature dimension
        lstm_out, _ = self.lstm(combined_features)
        # Using the last time step's output
        output = self.fc(lstm_out[:, -1, :])
        return output

```

Appendix - Execution Results

Table 2: Table of results for Sysbench CPU Benchmark Model executions

L1D Size (bytes)	L1I Size (bytes)	LL Size (bytes)	Cores	L1D		L1I		LL		L1D				L1I				LL				Execution time per block							
				Total	Total	Total	Total	Total	Total	MSE	RMSE	R2	Abs. Err. Percent	MSE	RMSE	R2	Abs. Err. Percent	MSE	RMSE	R2	Abs. Err. Percent	t ₁ s	t ₂ s	t ₃ s	t ₄ s	t ₅ s	t ₆ s	t ₇ s	t ₈ s
				Actual	Predicted	Actual	Predicted	Actual	Predicted																				
1k	1k	2k	1	929.857	970.716	343.013	428.490	1,056.247	970.716	10.5	3.2	0.9	4.4	5.6	2.4	0.8	24.9	25.4	5.0	0.8	20.0	17	17	15	14	16	17	19	16
1k	1k	2k	2	929.842	981.834	343.007	426.989	1,056.225	981.834	10.1	3.2	0.9	5.6	5.5	2.4	0.8	24.5	26.1	5.1	0.8	21.3	18	16	16	18	17	16	17	18
1k	1k	4k	1	929.857	945.486	343.013	429.432	952.906	945.486	12.0	3.5	0.9	1.7	5.6	2.4	0.8	25.2	18.9	4.3	0.8	4.4	18	16	14	19	18	16	18	16
1k	1k	4k	2	929.842	955.566	343.007	427.865	952.871	955.566	11.6	3.4	0.9	2.8	5.6	2.4	0.8	24.7	19.7	4.4	0.8	8.1	17	16	16	18	18	17	17	18
1k	1k	2k	4	929.794	999.360	343.014	424.442	1,056.170	999.360	9.5	3.1	0.9	7.5	5.4	2.3	0.8	23.7	27.0	5.2	0.8	23.1	18	14	15	17	18	17	18	16
1k	1k	4k	4	929.794	974.292	343.014	424.945	952.791	974.292	11.1	3.3	0.9	4.8	5.4	2.3	0.8	23.9	22.8	4.8	0.8	17.1	17	13	13	16	16	14	18	18
1k	1k	10k	1	929.857	1,058.499	343.013	435.221	386.988	1,058.499	8.0	2.8	0.9	13.8	5.9	2.4	0.8	26.9	25.3	5.0	0.1	55.7	17	16	14	16	17	16	17	17
1k	1k	10k	2	929.842	1,061.302	343.007	433.769	386.985	1,061.302	8.1	2.8	0.9	14.1	5.9	2.4	0.8	26.5	25.5	5.0	0.1	55.0	17	14	14	16	17	16	16	18
1k	1k	10k	4	929.794	1,063.067	343.014	431.174	386.969	1,063.067	8.4	2.9	0.9	14.3	5.8	2.4	0.8	25.7	25.7	5.1	0.1	52.7	18	14	17	16	19	18	18	16
1k	2k	2k	1	929.857	1,006.300	314.724	402.809	1,027.114	1,006.300	8.9	3.0	0.9	8.2	5.6	2.4	0.8	28.0	24.4	4.9	0.8	23.7	17	16	18	18	14	17	18	16
1k	2k	2k	2	929.842	1,013.684	314.737	399.673	1,027.116	1,013.684	8.7	3.0	0.9	9.0	5.6	2.4	0.8	27.0	24.6	5.0	0.8	24.5	16	13	16	16	17	18	18	16
1k	2k	4k	1	929.857	980.388	314.724	402.579	932.916	980.388	10.7	3.3	0.9	5.4	5.8	2.4	0.8	27.9	17.0	4.1	0.8	6.9	16	17	17	15	16	14	18	17
1k	2k	4k	2	929.842	987.532	314.737	399.427	932.909	987.532	10.5	3.2	0.9	6.2	5.7	2.4	0.8	26.9	18.1	4.3	0.8	10.3	16	17	16	14	17	19	18	18
1k	2k	10k	1	929.857	1,072.163	314.724	409.211	382.912	1,072.163	8.0	2.8	0.9	15.3	6.3	2.5	0.7	30.0	24.5	4.9	0.1	54.4	16	15	17	17	16	14	18	20
1k	2k	10k	2	929.842	1,073.120	314.737	406.698	382.911	1,073.120	8.1	2.8	0.9	15.4	6.2	2.5	0.7	29.2	24.4	4.9	0.1	53.1	16	16	17	17	18	19	18	18
1k	2k	4k	4	929.794	1,024.456	314.774	390.779	1,027.090	1,024.456	8.5	2.9	0.9	10.2	5.6	2.4	0.8	24.1	25.0	5.0	0.8	25.1	16	17	13	17	17	17	15	17
1k	2k	4k	4	929.794	1,001.137	314.774	391.069	932.855	1,001.137	10.2	3.2	0.9	7.7	5.7	2.4	0.8	24.2	21.7	4.7	0.8	18.2	16	14	16	19	17	18	18	18
1k	2k	10k	4	929.794	1,072.981	314.774	400.051	382.891	1,072.981	8.3	2.9	0.9	15.4	6.2	2.5	0.7	27.1	24.1	4.9	0.1	49.4	18	15	16	16	14	18	16	18
1k	64k	2k	1	929.857	1,033.081	2.911	4.557	666.006	1,033.081	8.9	3.0	0.9	11.1	0.1	0.3	0.4	56.5	16.4	4.0	0.6	31.0	16	14	15	16	19	16	17	19
1k	64k	2k	2	929.842	1,029.859	3.167	4.700	666.269	1,029.859	9.0	3.0	0.9	10.8	0.1	0.4	0.4	48.4	16.5	4.1	0.6	31.1	16	13	14	14	16	17	18	19
1k	64k	4k	1	929.857	1,050.885	2.911	5.009	572.400	1,050.885	8.9	3.0	0.9	13.0	0.1	0.3	0.4	72.1	19.6	4.4	0.5	49.2	18	16	14	17	19	19	19	19
1k	64k	4k	2	929.842	1,047.054	3.167	5.089	572.675	1,047.054	9.0	3.0	0.9	12.6	0.1	0.4	0.4	60.7	20.0	4.5	0.5	49.9	16	15	13	17	15	17	18	19
1k	64k	10k	1	929.857	1,129.502	2.911	5.718	230.767	1,129.502	8.9	3.0	0.9	21.5	0.1	0.3	0.4	96.4	11.7	3.4	0.2	60.8	18	14	14	18	17	16	17	15
1k	64k	10k	2	929.842	1,128.150	3.167	5.889	231.036	1,128.150	8.9	3.0	0.9	21.3	0.1	0.4	0.3	85.9	11.8	3.4	0.2	57.8	17	13	14	15	18	16	18	19
1k	64k	2k	4	929.794	1,021.430	3.534	5.174	666.598	1,021.430	9.4	3.1	0.9	9.9	0.2	0.4	0.4	46.4	16.7	4.1	0.6	30.8	18	16	13	18	16	17	16	17
1k	64k	4k	4	929.794	1,037.912	3.534	5.502	573.007	1,037.912	9.3	3.0	0.9	11.6	0.2	0.4	0.4	55.7	20.3	4.5	0.5	50.4	16	15	15	17	16	18	16	18
1k	64k	10k	4	929.794	1,125.346	3.534	6.401	231.279	1,125.346	9.1	3.0	0.9	21.0	0.2	0.4	0.4	81.1	12.0	3.5	0.2	51.4	17	16	14	19	16	17	18	18
2k	1k	2k	1	653.093	859.704	343.013	427.986	989.786	859.704	25.8	5.1	0.4	31.6	5.6	2.4	0.8	24.8	22.3	4.7	0.8	18.7	18	16	15	16	17	18	18	18
2k	1k	2k	2	653.059	874.546	343.007	426.834	989.742	874.546	26.4	5.1	0.4	33.9	5.5	2.4	0.8	24.4	23.4	4.8	0.8	20.4	19	19	15	16	16	17	14	15
2k	1k	4k	1	653.093	825.456	343.013	428.950	936.508	825.456	24.2	4.9	0.5	26.4	5.6	2.4	0.8	25.1	19.0	4.4	0.8	0.4	19	13	17	16	17	18	15	18
2k	1k	4k	2	653.059	838.541	343.007	427.779	936.463	838.541	24.6	5.0	0.5	28.4	5.6	2.4	0.8	24.7	19.2	4.4	0.8	4.3	17	19	13	17	18	16	18	18
2k	1k	10k	1	653.093	978.540	343.013	436.372	393.565	978.540	31.0	5.6	0.3	49.8	5.9	2.4	0.8	27.2	25.6	5.1	0.1	52.9	16	15	13	17	15	17	16	18
2k	1k	10k	2	653.059	985.513	343.007	435.292	393.554	985.513	31.1	5.6	0.3	50.9	5.9	2.4	0.8	26.9	26.0	5.1	0.1	52.9	19	15	15	16	18	16	15	18
2k	2k	2k	1	653.093	907.625	314.724	404.139	964.751	907.625	26.8	5.2	0.4	36.0	5.8	2.4	0.8	28.4	23.4	4.8	0.8	23.6	13	11	10	11	11	12	16	17
2k	2k	2k	2	653.059	918.673	314.737	401.728	964.728	918.673	27.3	5.2	0.4	40.7	5.7	2.4	0.8	27.6	24.2	4.9	0.7	24.8	12	10	10	11	11	12	12	12
2k	2k	4k	1	653.093	872.860	314.724	404.308	921.435	872.860	25.2	5.0	0.4	33.7	6.0	2.4	0.8	28.5	17.7	4.2	0.8	4.2	19	15	13	16	17	17	15	18
2k	2k	4k	2	653.059	885.091	314.737	401.539	921.412	885.091	25.9	5.1	0.4	35.5	5.9	2.4	0.8	27.6	18.3	4.3	0.9	4.8	15	16	16	16	17	15	17	15
2k	2k	10k	1	653.093	1,004.574	314.724	411.536	389.470	1,004.574	32.2	5.7	0.3	53.8	6.3	2.5	0.7	30.8	25.1	5.0	0.1	53.1	21	18	14	15	14	17	14	17
2k	2k	10k	2	653.059	1,008.908	314.737	409.720	389.464	1,008.908	32.3	5.7	0.3	54.5	6.3	2.5	0.7	30.2	25.2	5.0	0.1	52.5	20	18	16	16	14	17	17	15
2k	64k	2k	1	653.093	958.709	2.911	4.296	640.176	958.709	30.2	5.5	0.3	46.8	0.1	0.3	0.4	47.6	15.3	3.9	0.7	27.6	17	15	13	16	17	17	18	17
2k	64k	2k	2	653.059	953.539	3.167	4.416	640.406	953.539	29.9	5.5	0.3	46.0	0.1	0.4	0.4	49.4	15.3	3.9	0.7	27.3	20							

Table 3: Table of results for Sysbench FileIO Benchmark Model executions

L1D Size (bytes)	L1I Size (bytes)	L2 Size (bytes)	Cores	L1D		L1I		L2		L1D			L1I			L2			Execution time per block										
				Total Actual	Total Predicted	Total Actual	Total Predicted	Total Actual	Total Predicted	MSE	RMSE	R2	Abs. Err. Percent	MSE	RMSE	R2	Abs. Err. Percent	MSE	RMSE	R2	Abs. Err. Percent	t ₁ s	t ₂ s	t ₃ s	t ₄ s	t ₅ s	t ₆ s	t ₇ s	t ₈ s
				1k	2k	1k	2k	1k	2k	1k	2k	1k	2k	1k	2k	1k	2k	1k	2k	1k	2k	1k	2k	1k	2k	1k	2k	1k	2k
1k	1k	2k	1	1,654,670	1,399,186	1,351,265	1,418,858	2,886,788	1,399,186	53.70	7.30	-0.60	15.40	6.30	2.50	0.90	5.00	71.00	8.40	0.20	9.10	19	16	18	17	17	18	18	19
1k	1k	2k	2	1,649,213	1,366,599	1,351,290	1,417,640	2,880,669	1,366,599	57.20	7.60	-0.80	17.10	6.30	2.50	0.90	4.90	76.00	8.70	0.20	10.00	19	16	18	17	17	17	18	19
1k	1k	4k	1	1,654,670	1,441,105	1,351,265	1,427,518	2,517,207	1,441,105	52.40	7.20	-0.60	12.90	6.40	2.50	0.90	5.60	115.50	10.70	-0.60	18.60	19	21	20	21	17	18	17	17
1k	1k	8k	2	1,649,213	1,399,176	1,351,290	1,424,298	2,509,735	1,399,176	57.30	7.60	-0.80	15.20	6.30	2.50	0.90	5.40	121.00	11.00	-0.70	17.30	20	16	18	20	17	17	17	19
1k	1k	16k	1	1,624,166	1,306,244	1,351,163	1,413,247	2,854,762	1,306,244	62.80	7.90	-1.10	19.60	6.20	2.50	0.90	4.60	86.20	9.30	0.00	11.60	19	17	17	18	17	18	18	19
1k	1k	4k	4	1,624,166	1,323,401	1,351,163	1,416,690	2,471,721	1,323,401	64.20	8.00	-1.10	18.50	6.40	2.50	0.90	4.80	124.00	11.10	-0.80	14.90	19	17	18	17	17	18	18	19
1k	1k	16k	1	1,654,670	1,833,876	1,351,265	1,431,797	465,345	1,833,876	29.70	5.40	0.10	10.80	6.60	2.60	0.90	6.00	38.20	6.20	-0.40	95.50	19	17	18	16	17	18	17	17
1k	1k	16k	2	1,649,213	1,778,815	1,351,290	1,432,972	466,087	1,778,815	29.70	5.80	0.00	7.90	6.50	2.50	0.90	6.00	37.20	6.10	-0.40	89.20	19	16	17	16	17	17	17	19
1k	1k	16k	4	1,624,166	1,648,160	1,351,163	1,430,321	467,819	1,648,160	44.20	6.70	-0.50	1.50	6.30	2.50	0.90	5.90	33.00	5.90	-0.30	75.20	19	17	17	17	16	16	17	17
1k	2k	2k	1	1,654,670	1,504,889	1,098,937	1,210,159	2,618,505	1,504,889	48.40	7.00	-0.50	9.10	23.40	4.80	0.40	10.10	87.50	9.40	-0.10	4.50	19	14	16	16	16	17	17	17
1k	2k	2k	2	1,649,213	1,477,612	1,099,776	1,218,090	2,613,402	1,477,612	51.60	7.20	-0.60	10.10	22.60	4.80	0.40	10.80	90.70	9.50	-0.20	5.20	18	15	17	17	17	16	17	17
1k	2k	4k	1	1,654,670	1,524,201	1,098,937	1,254,832	2,356,975	1,524,201	47.80	6.90	-0.50	7.90	21.40	4.60	0.50	14.20	104.20	10.20	-0.60	14.20	19	18	16	16	16	17	18	17
1k	2k	4k	2	1,649,213	1,485,032	1,099,776	1,257,581	2,347,743	1,485,032	52.60	7.30	-0.60	10.00	20.50	4.50	0.50	14.40	110.40	10.50	-0.70	13.00	19	17	16	17	18	17	17	18
1k	2k	16k	1	1,654,670	1,878,589	1,098,937	1,246,866	460,524	1,878,589	28.00	5.30	0.10	13.50	20.40	4.50	0.50	13.50	34.70	5.90	-0.40	84.40	20	14	17	17	18	17	17	19
1k	2k	16k	2	1,649,213	1,843,033	1,099,776	1,278,425	461,020	1,843,033	30.20	5.50	0.10	11.80	18.90	4.40	0.50	16.20	34.30	5.90	-0.30	81.50	19	16	18	17	17	18	18	19
1k	2k	2k	4	1,624,166	1,408,820	1,100,469	1,227,855	2,588,553	1,408,820	57.20	7.60	-0.90	13.30	20.80	4.60	0.50	11.60	95.90	9.80	-0.30	7.10	19	18	17	16	16	17	18	18
1k	2k	4k	4	1,624,166	1,412,142	1,100,469	1,252,833	2,305,122	1,412,142	59.00	7.70	-0.90	13.10	19.20	4.40	0.50	13.80	116.90	10.80	-1.00	10.70	19	15	18	17	16	18	18	18
1k	2k	16k	4	1,624,166	1,738,869	1,100,469	1,324,401	462,660	1,738,869	37.60	6.10	-0.20	7.10	16.00	4.00	0.60	20.30	32.70	5.70	-0.30	71.60	20	14	17	17	19	17	17	19
1k	64k	2k	1	1,654,670	1,794,144	3,199	9,616	1,410,778	1,794,144	31.00	5.60	1.10	8.40	0.10	0.40	0.30	200.60	22.00	4.70	-0.10	11.00	19	17	18	17	17	19	17	17
1k	64k	2k	2	1,649,213	1,749,370	3,430	10,987	1,406,335	1,749,370	33.80	5.80	0.00	6.10	0.20	0.40	0.10	220.30	24.50	5.00	0.00	8.90	19	17	18	17	17	17	18	19
1k	64k	4k	1	1,654,670	1,851,791	3,199	10,814	1,012,708	1,851,791	29.60	5.40	0.10	11.90	0.10	0.40	0.30	238.00	40.00	6.30	-0.50	54.90	19	17	18	16	17	17	17	18
1k	64k	4k	2	1,649,213	1,805,836	3,430	12,091	1,006,750	1,805,836	32.30	5.70	0.00	9.50	0.20	0.40	0.20	232.50	40.20	6.30	-0.50	52.60	19	16	17	17	17	17	18	19
1k	64k	16k	1	1,654,670	2,104,130	3,199	8,026	257,728	2,104,130	27.70	5.30	0.20	27.20	0.20	0.40	0.20	150.90	12.60	3.50	0.10	70.30	19	16	18	17	17	17	17	19
1k	64k	16k	2	1,649,213	2,094,499	3,430	8,619	258,188	2,094,499	27.70	5.30	0.10	27.00	0.20	0.40	0.20	151.30	12.70	3.60	0.10	66.30	19	16	18	18	17	17	18	19
1k	64k	2k	4	1,624,166	1,634,794	3,882	14,119	1,385,018	1,634,794	40.60	6.40	-0.30	0.70	0.40	0.60	-0.40	263.70	31.70	5.60	-0.30	3.90	19	16	18	17	17	16	17	17
1k	64k	4k	4	1,624,166	1,693,946	3,882	14,937	969,195	1,693,946	38.50	6.20	-0.30	4.30	0.30	0.60	-0.30	284.80	41.60	6.40	-0.60	50.60	19	16	18	17	16	17	18	18
1k	64k	16k	4	1,624,166	2,024,697	3,882	10,329	262,969	2,024,697	27.60	5.10	-1.30	27.70	0.30	0.60	-0.10	165.80	12.70	3.60	0.10	67.50	19	15	17	17	17	18	17	19
2k	1k	2k	1	1,322,828	1,139,818	1,351,265	1,421,099	2,661,216	1,139,818	51.50	7.20	-1.30	13.80	6.30	2.50	0.90	5.20	63.40	8.00	0.20	10.50	18	16	19	17	16	16	17	17
2k	1k	2k	2	1,310,469	1,124,203	1,351,290	1,420,521	2,647,139	1,124,203	54.20	7.40	-1.40	14.20	6.20	2.50	0.90	5.10	66.20	8.10	0.10	10.60	18	16	17	16	16	16	17	17
2k	1k	4k	1	1,322,828	1,151,576	1,351,265	1,428,865	2,414,616	1,151,576	51.50	7.20	-1.30	12.90	6.40	2.50	0.90	5.70	132.60	11.50	-0.90	25.30	18	16	18	17	16	16	16	17
2k	1k	4k	2	1,322,828	1,151,576	1,351,265	1,428,865	2,414,616	1,151,576	51.50	7.20	-1.30	12.90	6.40	2.50	0.90	5.70	132.60	11.50	-0.90	25.30	18	16	18	17	16	16	16	17
2k	1k	16k	1	1,322,828	1,570,214	1,351,265	1,434,600	468,409	1,570,214	46.10	6.80	-1.00	18.70	6.60	2.60	0.90	6.20	35.30	5.90	-0.30	78.50	20	16	18	17	17	17	18	19
2k	1k	16k	2	1,310,469	1,504,578	1,351,290	1,434,870	467,197	1,504,578	49.70	7.10	-1.20	14.80	6.50	2.50	0.90	6.20	34.70	5.90	-0.30	73.50	18	15	17	17	17	17	17	19
2k	2k	2k	1	1,322,828	1,251,750	1,098,937	1,202,260	2,416,573	1,251,750	51.20	7.20	-1.30	5.40	23.90	4.90	0.40	9.40	88.60	9.40	-0.40	7.20	18	16	17	19	18	17	18	19
2k	2k	2k	2	1,310,469	1,231,764	1,099,776	1,216,327	2,403,720	1,231,764	53.50	7.30	-1.40	6.00	22.70	4.80	0.40	10.60	88.40	9.40	-0.40	7.10	18	17	16	16	19	17	19	18
2k	2k	4k	1	1,322,828	1,251,153	1,098,937	1,241,884	2,290,630	1,251,153	51.00	7.10	-1.30	5.40	22.00	4.70	0.40	13.00	120.80	11.00	-1.00	20.70	18	15	17	16	16	17	18	17
2k	2k	4k	2	1,310,469	1,225,761	1,099,776	1,249,885	2,243,335	1,225,761	53.60	7.30	-1.40	6.50	20.90	4.60	0.50	13.60	123.30	11.10	-1.00	19.10	19	14	16	16	16	17	16	18
2k	2k	16k	1	1,322,828	1,664,623	1,098,937	1,248,036	46																					

Table 4: Table of results for Sysbench Memory Benchmark Model executions

L1D		L1H		LL		L1D										L1H										LL										Execution time per block																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																					
						Total		Total		Total		Total		Total		MSE		RMSE		R2		Abs. Err. Percent		MSE		RMSE		R2		Abs. Err. Percent		MSE		RMSE		R2		Abs. Err. Percent		t ₁ s		t ₂ s		t ₃ s		t ₄ s		t ₅ s		t ₆ s		t ₇ s		t ₈ s																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																			
L1D Size (bytes)	L1H Size (bytes)	LL Size (bytes)	Cores	Actual	Predicted	Actual	Predicted	Actual	Predicted	Actual	Predicted	M3E	RMSE	R2	Abs. Err. Percent	M3E	RMSE	R2	Abs. Err. Percent	M3E	RMSE	R2	Abs. Err. Percent	M3E	RMSE	R2	Abs. Err. Percent	M3E	RMSE	R2	Abs. Err. Percent	t ₁ s	t ₂ s	t ₃ s	t ₄ s	t ₅ s	t ₆ s	t ₇ s	t ₈ s																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																		
1k	1k	2k	1	1,375,523	1,157,707	726,633	872,551	1,879,632	1,157,707	31.9	5.6	0.5	15.8	7.9	2.8	0.8	20.1	42.9	6.5	0.6	0.3	17	14	14	16	17	16	19	18																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																												</

Table 5: Table of results for the GFTT Model executions

L1D		L1H		LL		L1D								L1H								LL								Execution time per block							
						Total	Total	MSE	RMSE	R2	Abs. Err. Percent	MSE	RMSE	R2	Abs. Err. Percent	MSE	RMSE	R2	Abs. Err. Percent	t ₁ s	t ₂ s	t ₃ s	t ₄ s	t ₅ s	t ₆ s	t ₇ s	t ₈ s										
L1D Size (bytes)	L1H Size (bytes)	LL Size (bytes)	Cores	Total Actual	Total Predicted	Total Actual	Total Predicted	Total Actual	Total Predicted	MSE	RMSE	R2	Abs. Err. Percent	MSE	RMSE	R2	Abs. Err. Percent	MSE	RMSE	R2	Abs. Err. Percent	t ₁ s	t ₂ s	t ₃ s	t ₄ s	t ₅ s	t ₆ s	t ₇ s	t ₈ s								
1k	1k	2k	1	2,023,632	1,606,586	314,146	534,562	1,819,176	1,606,586	38.0	6.2	0.5	20.6	17.8	4.2	0.1	70.2	27.9	5.3	0.5	3.4	15	14	14	15	15	15	15	15								
1k	1k	2k	2	2,023,632	1,634,488	314,146	531,671	1,819,176	1,634,488	36.8	6.1	0.5	19.2	16.8	4.1	0.2	69.2	28.3	5.3	0.5	5.3	16	14	14	14	14	15	15	15								
1k	1k	4k	1	2,023,632	1,629,368	314,146	547,014	1,691,447	1,629,368	36.8	6.1	0.5	19.5	18.7	4.3	0.1	74.1	21.9	4.7	0.5	0.6	21	21	20	21	18	19	21	21								
1k	1k	4k	2	2,023,632	1,664,880	314,146	544,691	1,691,447	1,664,880	35.2	5.9	0.5	17.7	17.7	4.2	0.1	73.4	22.4	4.7	0.5	1.4	21	23	18	19	17	20	20	20								
1k	1k	2k	4	2,023,632	1,683,170	314,146	518,085	1,819,176	1,683,170	35.0	5.9	0.5	16.8	14.7	3.8	0.3	64.9	31.5	5.6	0.4	9.4	11	13	13	13	13	14	15	15								
1k	1k	4k	4	2,023,632	1,706,171	314,146	531,413	1,691,447	1,706,171	33.7	5.8	0.6	15.7	15.5	3.9	0.2	69.2	24.9	5.0	0.5	5.6	20	21	18	19	19	19	20	20								
1k	1k	16k	1	2,023,632	1,748,195	314,146	580,855	1,317,132	1,748,195	30.5	5.5	0.6	13.6	20.9	4.6	0.0	84.9	18.6	4.3	0.6	4.4	25	25	18	19	17	20	20	20								
1k	1k	16k	2	2,023,632	1,756,234	314,146	578,944	1,317,132	1,756,234	30.5	5.5	0.6	13.2	20.0	4.5	0.0	84.3	17.1	4.1	0.7	2.8	32	20	17	19	19	20	20	20								
1k	1k	16k	4	2,023,632	1,742,038	314,146	568,377	1,317,132	1,742,038	32.4	5.7	0.6	13.9	17.8	4.2	0.1	80.9	15.2	3.9	0.7	0.3	20	24	19	19	17	19	21	21								
1k	2k	2k	1	2,023,632	1,595,808	273,047	425,870	1,779,048	1,595,808	39.0	6.2	0.5	21.1	17.1	4.1	0.0	56.0	30.7	5.5	0.4	5.9	24	15	16	16	16	17	18	18								
1k	2k	2k	2	2,023,632	1,621,258	273,047	415,058	1,779,048	1,621,258	37.8	6.2	0.5	19.9	15.5	3.9	0.1	52.0	30.6	5.5	0.4	7.8	21	22	19	19	17	19	18	20								
1k	2k	4k	1	2,023,632	1,615,371	273,047	444,989	1,662,461	1,615,371	38.1	6.2	0.5	20.2	18.1	4.3	0.0	63.0	23.6	4.9	0.5	1.9	21	16	14	15	20	18	20	20								
1k	2k	4k	2	2,023,632	1,648,167	273,047	433,853	1,662,461	1,648,167	36.6	6.0	0.5	18.6	16.3	4.0	0.1	58.9	23.7	4.9	0.5	3.9	24	18	16	17	17	16	17	16								
1k	2k	16k	1	2,023,632	1,745,551	273,047	499,265	1,307,282	1,745,551	31.6	5.6	0.6	13.7	20.5	4.5	-0.1	82.8	16.7	4.1	0.7	3.1	23	18	18	17	19	20	21	21								
1k	2k	16k	2	2,023,632	1,757,274	273,047	488,108	1,307,282	1,757,274	31.4	5.6	0.6	13.2	18.8	4.3	0.0	78.8	15.6	3.9	0.7	1.5	20	19	21	19	18	20	21	21								
1k	2k	2k	4	2,023,632	1,668,115	273,047	390,707	1,779,048	1,668,115	36.0	6.0	0.5	17.6	12.5	3.5	0.3	43.1	32.0	5.7	0.4	11.5	22	19	18	18	17	20	20	20								
1k	2k	4k	4	2,023,632	1,691,980	273,047	407,725	1,662,461	1,691,980	34.7	5.9	0.5	16.4	13.1	3.6	0.3	49.3	24.9	5.0	0.4	7.7	21	19	16	14	18	21	18	20								
1k	2k	16k	4	2,023,632	1,749,309	273,047	461,196	1,307,282	1,749,309	32.6	5.7	0.6	13.6	15.1	3.9	0.2	68.9	14.3	3.8	0.7	1.0	19	25	20	21	19	20	18	20								
1k	64k	2k	1	2,023,632	1,836,278	1,056	5,546	1,468,709	1,836,278	25.0	5.0	0.7	9.3	0.1	0.3	-0.1	425.2	9.6	3.1	0.7	10.5	20	22	20	22	20	20	20	20	20							
1k	64k	2k	2	2,023,632	1,824,607	1,056	5,771	1,468,709	1,824,607	27.0	5.2	0.6	9.8	0.1	0.3	-0.1	446.4	8.8	3.0	0.7	7.4	20	20	21	19	17	20	24	20								
1k	64k	4k	1	2,023,632	1,831,490	1,056	5,212	1,355,125	1,831,490	27.2	5.2	0.6	9.5	0.1	0.3	-0.1	393.6	10.9	3.3	0.7	14.8	20	20	18	19	18	20	21	20								
1k	64k	4k	2	2,023,632	1,838,692	1,056	5,343	1,355,125	1,838,692	28.6	5.3	0.6	9.1	0.1	0.3	-0.1	405.9	9.3	3.1	0.7	10.5	20	17	17	19	19	20	21	20								
1k	64k	16k	1	2,023,632	2,045,826	1,056	4,832	1,162,057	2,045,826	30.6	5.5	0.6	1.1	0.1	0.3	0.0	357.6	13.5	3.7	0.7	0.4	20	19	18	19	20	23	16	18								
1k	64k	16k	2	2,023,632	2,096,027	1,056	5,175	1,162,057	2,096,027	29.5	5.4	0.6	3.6	0.1	0.3	0.0	390.0	12.6	3.5	0.7	3.6	19	18	22	21	19	20	21	21								
1k	64k	2k	4	2,023,632	1,806,706	1,056	5,408	1,468,709	1,806,706	20.2	4.5	0.4	15.6	0.1	0.3	-0.1	432.7	8.1	2.8	0.8	2.1	20	18	17	18	20	19	20	20								
1k	64k	4k	4	2,023,632	1,886,448	1,056	5,633	1,355,125	1,886,448	30.6	5.5	0.6	6.8	0.1	0.3	-0.1	433.4	7.0	2.7	0.8	0.5	20	19	18	19	19	20	21	21								
1k	64k	16k	4	2,023,632	2,145,800	1,056	6,300	1,162,057	2,145,800	29.9	5.5	0.6	6.0	0.1	0.3	-0.1	496.6	11.7	3.4	0.7	8.5	21	25	18	19	20	20	20	20								
2k	1k	2k	1	1,462,803	1,501,582	314,146	536,595	1,765,780	1,501,582	10.2	3.2	0.7	2.7	18.0	4.2	0.1	70.8	24.7	5.0	0.5	1.7	20	19	18	17	16	17	15	17								
2k	1k	2k	2	1,462,803	1,519,541	314,146	533,491	1,765,780	1,519,541	10.2	3.2	0.7	3.9	17.0	4.1	0.2	69.8	25.1	5.0	0.5	3.4	19	18	17	18	25	21	22	21								
2k	1k	4k	1	1,462,803	1,510,327	314,146	548,558	1,721,743	1,510,327	10.6	3.3	0.7	3.2	18.8	4.3	0.1	74.6	23.0	4.8	0.5	2.1	20	18	17	18	20	19	20	20								
2k	1k	4k	2	1,462,803	1,530,552	314,146	546,055	1,721,743	1,530,552	10.7	3.3	0.7	4.6	17.9	4.2	0.1	73.8	23.7	4.9	0.5	3.8	20	19	21	19	19	20	27	21								
2k	1k	16k	1	1,462,803	1,576,213	314,146	582,882	1,325,555	1,576,213	13.7	3.7	0.6	7.8	21.0	4.6	0.0	85.5	18.7	4.3	0.6	4.1	20	18	18	26	18	20	20	21								
2k	1k	16k	2	1,462,803	1,587,632	314,146	581,073	1,325,555	1,587,632	13.7	3.7	0.6	8.5	20.1	4.5	0.0	85.0	17.4	4.2	0.7	2.6	20	17	19	21	20	22	20	22								
2k	2k	2k	1	1,462,803	1,497,188	273,047	430,568	1,730,394	1,497,188	13.4	3.2	0.7	2.4	17.3	4.2	0.0	57.7	27.1	5.2	0.4	4.2	22	17	20	19	18	20	19	20								
2k	2k	2k	2	1,462,803	1,511,886	273,047	419,732	1,730,394	1,511,886	10.2	3.2	0.7	3.4	15.6	4.0	0.1	53.7	27.0	5.2	0.4	6.0	22	18	19	18	18	20	21	22								
2k	2k	4k	1	1,462,803	1,503,892	273,047	448,695	1,693,860	1,503,892	10.7	3.3	0.7	2.8	18.3	4.3	0.0	64.3	24.9	5.0	0.5	4.5	19	17	19	18	19	20	23	20								
2k	2k	4k	2	1,462,803	1,521,894	273,047	437,777	1,693,860	1,521,894	10.7	3.3	0.7	4.0	16.5	4.1	0.1	60.3	25.2	5.0	0.5	6.2	2	21	19	19	19	20	22	19								
2k	2k	16k	1	1,462,803	1,569,702	273,047	502,588	1,315,805	1,569,702	13.7	3.7	0.6	7.3	20.7	4.5	-0.2	84.1	16.9	4.1	0.7	2.9	23	17	18	19	19	21	22	19								
2k	2k	16k	2	1,462,803	1,582,855	273,047	491,525	1,315,805	1,582,855	13.7	3.7	0.6	8.2	19.5	4.5	-0.2	80.0	15.9	4.0	0.7	2.5	23	18	18	19	19	20	21	21								
2k	64k	2k	1	1,462,803	1,600,298	1,056	5,510	1,449,749	1,600,298	20.2																											