

Multi-Dimensional Characterization of Temporal Data Mining on Graphics Processors

Jeremy Archuleta, Yong Cao, Tom Scogland, Wu-chun Feng

Department of Computer Science, Virginia Tech
Blacksburg, VA, USA

{jsarch, yongcao, njustn, feng}@cs.vt.edu

Abstract

Through the algorithmic design patterns of data parallelism and task parallelism, the graphics processing unit (GPU) offers the potential to vastly accelerate discovery and innovation across a multitude of disciplines. For example, the exponential growth in data volume now presents an obstacle for high-throughput data mining in fields such as neuroscience and bioinformatics. As such, we present a characterization of a MapReduce-based data-mining application on a general-purpose GPU (GPGPU). Using neuroscience as the application vehicle, the results of our multi-dimensional performance evaluation show that a “one-size-fits-all” approach maps poorly across different GPGPU cards. Rather, a high-performance implementation on the GPGPU should factor in the 1) problem size, 2) type of GPU, 3) type of algorithm, and 4) data-access method when determining the type and level of parallelism. To guide the GPGPU programmer towards optimal performance within such a broad design space, we provide eight general performance characterizations of our data-mining application.

1. Introduction

There is a growing trend in scientific computing towards the use of accelerators to reduce the time to discovery. Unlike current general-purpose multicore CPU architectures, these accelerators combine hundreds of simplified processing cores executing in parallel to achieve the high-end performance demanded by scientists. Current examples of accelerators include the Cell Broadband Engine (Cell BE), field-programmable gate arrays (FPGAs), and general-purpose graphics processing units (GPGPUs) such as the NVIDIA CUDA-enabled GPU and AMD/ATI Brook+ GPU. Furthermore, upcoming technologies like AMD Fusion and Intel Larrabee point toward a future of accelerator-based computing platforms.

Current top-of-the-line GPGPUs possess hundreds of processing cores and memory bandwidth that is 10 times higher than conventional CPUs. The significant increase in parallelism within a GPGPU and the accompanying increase in performance has been successfully exploited for scientific, database, geometric, and imaging applications, which in many cases, has resulted in an order-of-magnitude performance improvement over top-of-the-line CPUs. GPGPUs

also provide many other tangible benefits such as improved performance per dollar and performance per watt over conventional CPUs. Combining high performance, lower cost, and increasingly usable tool chains, GPGPUs are becoming more programmable and capable of solving a wider range of applications than simply three-dimensional triangle rasterization, and as such, is the target platform for our temporal data mining research.

Although temporal data mining is a relatively new area of data mining [15], this technique is becoming progressively more important and is widely used in various application fields, such as telecommunication control [7], earth science [16], financial data prediction [11], medical data analysis [6], and neuroscience [19]. Specifically, neuroscientists would like to identify how the neurons in the brain are connected and related to each other. This usually involves stimulating one area of the brain and observing which other areas of the brain “light up.” Recent technological advances in electrophysiology and imaging now allow neuroscientists to capture the timing of hundreds of neurons [14], [17]. However, the amount of data captured results in “data overload” and requires powerful computational resources. Current technology, like GMiner [18], is a step in the right direction, but being limited to a single CPU running a Java virtual machine (VM), GMiner forces output to be processed post-mortem.

What neuroscientists truly desire is real-time, interactive visualization of the effects of their probes. This capability would open up an entirely new window of opportunities whereby patients can be tested, diagnosed, and operated upon in a single, *faster* procedure – we believe that GPGPUs can provide the necessary performance. As such, this paper presents a characterization of a temporal data-mining application in a multi-dimensional environment. Specifically, we evaluate its performance across the following five dimensions: 1) parallel algorithm type, 2) data-access method, 3) problem size, 4) GPGPU generation, and 5) number of threads.

Our results show that GPGPUs can provide the requisite performance, but that a “one-size-fits-all” approach is unsuitable for temporal data mining on graphics processors. Instead, the problem size and graphics processor determine which type of algorithm, data-access pattern, and number of threads should be used to achieve the desired performance. This result both corroborates and contrasts with previous,

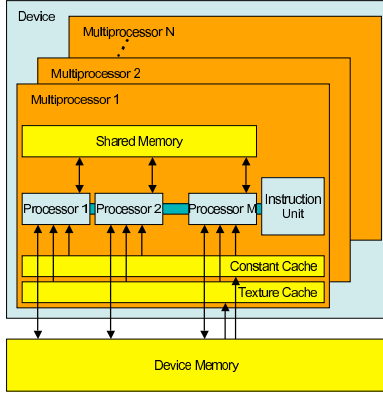


Figure 1. Overview of the NVIDIA Compute Unified Device Architecture (CUDA) [2]

similar MapReduce algorithms on graphics processors. However, while previous works only provided results in an optimal configuration, this paper presents general characterizations to help explain *how* a MapReduce-based, data-mining application should harness the parallelism in graphics processors.

To this end, we first present a background of current CUDA GPU technology and related MapReduce implementations in Section 2 followed by a detailed description of our temporal data-mining algorithm in Section 3. Section 4 lists relevant features of our testbed on which all of the tests were conducted with the results and characterizations presented in Section 5. Lastly, we offer some directions for future work and conclusions in Sections 6 and 7, respectively.

2. Background and Related Work

2.1. Graphics Processors

GPUs have been used for many years to accelerate the rendering of computer graphics. Fed by the increasing demand for improved 3-D graphics at higher frame rates and larger resolutions, GPUs diverged from standard CPUs into exotic specialized architectures. In recent years, GPUs are transitioning away from a single-purpose device into a more general-purpose architecture, capable of doing more than computing pixel values. This transition opens the doors for a range of applications to be accelerated. We describe here the architectural and programmatic features of state-of-the-art NVIDIA GPUs.

2.1.1. CUDA Architecture. State-of-the-art NVIDIA GPUs present a Compute Unified Device Architecture (CUDA) to the programmer, as shown in Figure 1. This architecture can be broadly separated into two primary features – core organization and memory hierarchy.

The execution units of a CUDA-capable device are organized into multiprocessors, where each multiprocessor contains 8 scalar processor cores. The multiprocessor architecture is called SIMT (Single Instruction, Multiple Thread) and executes in a similar manner as a SIMD (Single Instruction, Multiple Data) architecture. While optimal performance is attained when groups of 32 threads, i.e. a *warp*, follow the same execution path, individual threads can diverge along different thread paths. When divergence occurs within a warp, every instruction of every thread path is executed, with threads enabled or disabled depending on whether the thread is executing that particular thread path or not. A single instruction is completed by the entire warp in 4 cycles.

There are several forms of memory accessible by an execution core on an NVIDIA GPU. Located on-chip, each multiprocessor contains its own set of 32-bit registers along with its own shared memory region which is quickly accessible by any core on the multiprocessor. The exact number of registers available and size of the shared memory depends on the compute capability (i.e., “generation”) of the GPU. In addition to shared memory, a multiprocessor contains two read-only caches, one for textures and another for constants, to improve memory access to texture and constant memory, respectively. The device memory, containing both local and global memory, resides off-chip and furthest from the execution cores (excluding the host machine’s main memory). It may seem odd that local memory is not in fact locally on-chip, but logically it serves the same purpose in that it is an overflow space for what will not fit in registers.

2.1.2. CUDA Programming Model. To enable the developer to harness the computing power of their GPUs, NVIDIA extended the C programming language to allow developers to easily offload computationally intensive kernels to the GPU for accelerated performance. This new programming model is commonly referred to as the CUDA programming model.

When a kernel executes, N parallel threads execute simultaneously. The programmer can logically arrange the N threads into one-, two-, or three-dimensional thread blocks. The index of a thread within this block and the ID of the thread have a one-to-one mapping to simplify identification. While threads within a thread block can share data within the same address space, each thread block has its own address space. This arrangement allows for synchronization between threads but *not* between thread blocks. To increase the amount of parallelism further, M “equally-shaped” thread blocks can be executed in parallel, increasing the total amount of available parallelism to $M * N$.

As mentioned above, groups of 32 threads form a warp with multiple warps composing a thread block and multiple thread blocks forming a kernel. When a kernel executes, the thread blocks are placed on different multiprocessors, according to available execution capacity [2]. All of the

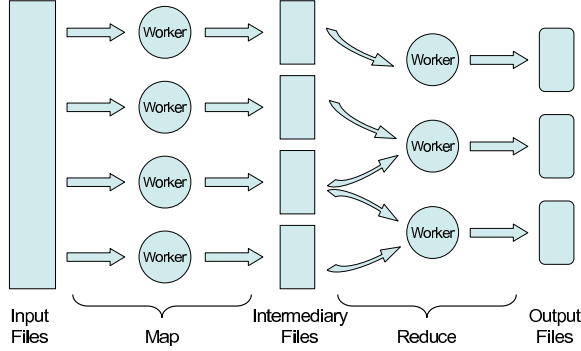


Figure 2. High-level view of the parallelism available in a MapReduce programming model

threads (and by association, warps) of one thread block will execute on the same multiprocessor. The instruction unit on each multiprocessor schedules warps with zero-cycle overhead and a single instruction of the warp executes in 4 cycles. Programmer-controlled placement and scheduling of the warps and thread blocks on the hardware is not currently available. As we will show in Section 5, this small, but rich, feature has a huge impact on the realizable performance of CUDA applications.

2.2. MapReduce

MapReduce is a programming model developed by Google to provide a convenient means for programmers to process large data sets on large parallel machines [8]. Moreover, programmers that utilize a MapReduce framework do not need prior experience using parallel systems. While there has been considerable debate over exactly how applicable this programming model is [9], [10], the ability to process large data sets in parallel is an important requirement for real-time data-mining.

The general MapReduce algorithm leverages two functional programming primitives, *map* and *reduce* in sequence. First, the *map* function is applied to a set of inputs consisting of a key/value pair to create a set of intermediate key/value pairs. Then, the *reduce* function is applied to all intermediate key/value pairs containing the same intermediate key to produce a set of outputs. Due to the functional nature of both *map* and *reduce*, each phase can be executed in parallel in order to utilize the vast resources available in large data centers. A high-level view of the parallelism available in the algorithm is shown in Figure 2.

The original implementation of MapReduce was built and optimized by Google to run on its private data centers. Providing the same functionality to the public, Hadoop is an open-source Java implementation that runs on everyday clusters and is under active development [1]. Specialized variations of the MapReduce framework also exist for mul-

ticore processors [5], the Cell processor [13], and graphics processors [4], [12], [20]. However, obtaining high performance within these frameworks is difficult (“... the cost of the Map and the Reduce function is unknown, it is difficult to find the optimal setting for the thread configuration.” [12]) and often left to the user (“... performance can be strongly affected by the number of registers ... amount of local memory ... number of threads ... algorithm ... among other factors. ... We allow the programmer to expose these choices through template parameters ...” [4]).

3. Temporal Data Mining

3.1. Overview

Association rule mining is a common data-mining technique used to discover how subsets of items relate to the presence of other subsets. Temporal data mining is a restricted variation of association rule mining, where temporal relations between items will also be considered.

A prototypical example of temporal data mining can be found in the area of market-basket analysis, where a store might want to know how often a customer buys product *B* given that product *A* was purchased earlier. In other words, the store wishes to know how often $\{peanut\ butter, bread\} \rightarrow \{jelly\}$. We note that unlike association rule mining, temporal data mining differentiates $\{bread, peanut\ butter\} \rightarrow \{jelly\}$ from $\{peanut\ butter, bread\} \rightarrow \{jelly\}$.

In this paper, we focus on one specific area of temporal data mining called frequent episode mining, where the aim is to discover frequently appearing episodes (i.e., sequences of items) in a time-ordered database [3]. We define frequent episode mining as follows.

Let $D = \{d_1, d_2, \dots, d_n\}$ be an ordered database of items where d_i is a member of the alphabet $I = \{i_1, i_2, \dots, i_m\}$. An episode, A_j , is a sequence of k items $\langle i_{j_1}, i_{j_2}, \dots, i_{j_k} \rangle$, where $\{i_{j_1}, \dots, i_{j_k}\} \in I$. There is an occurrence of A_j in database D if and only if there exists a sequence of indices $\langle r_1, r_2, \dots, r_k \rangle$ in increasing order such that $i_{j_l} = d_{r_l}$ for $l = 1, \dots, k$. The count of an episode, $count(A_j)$, is the total number of *non-overlapped* occurrences of A_j in D . The task of frequent episode mining is to find all such episodes, A_j , for which $count(A_j)$ is greater than a given threshold α .

The standard algorithm for frequent episode mining is described in Algorithm 1. As seen, this algorithm generates a set of candidate episodes for each level (i.e., length of an episode), counts the number of occurrences for each candidate episode, eliminates infrequent ones, and generates the set of candidate episodes for the next level.

The elimination step prunes infrequent episodes so that the generation step for the next level can exclude episode candidates that are guaranteed to be infrequent. Without the

Algorithm 1 *Frequent Episode Mining*(D, α).**Input:** supporting threshold α , sequential database $D = \{d_1, d_2, \dots, d_n\}$ **Output:** frequent episode set $S_A = A_1, A_2, \dots, A_m$

- 1: $k \leftarrow 1, S_A \leftarrow \emptyset$
- 2: level $k \leftarrow 1, A'_k \leftarrow \{\{i_1\}, \{i_2\}, \dots, \{i_m\}\}$, (generate candidate episode for level 1)
- 3: **while** $A'_k \neq \emptyset$ **do**
- 4: Calculate $\text{count}(A'_{k,j})$ for all episodes $A'_{k,j}$ in A'_k (**counting step**)
- 5: Eliminate all infrequent episodes, $\text{count}(A'_{k,j})/n \leq \alpha$, from A'_k (**elimination step**)
- 6: $S_A \leftarrow S_A \cup A'_k$
- 7: $k \leftarrow k + 1$
- 8: Generate candidate episode set A'_k from A'_{k-1} (**generation step**)
- 9: **end while**
- 10: **return** S_A

elimination step, the total number of episode candidates can grow exponentially as a function of the episode length k and the alphabet size N , as shown in Table 1.

Episode Length (L)	Episodes
1	N
2	$N(N-1)$
3	$N(N-1)(N-2)$
...	...
N	$\frac{N!}{(N-L)!}$

Table 1. Potential number of episodes with length L from an alphabet of size N

When the culling threshold is low, fewer episode candidates will be eliminated thereby resulting in an exponential growth of computation. To address such computational challenges, advanced algorithms are designed to distribute the computation among multiple processors [21]. With the recent advances of graphics processors to more general computing platforms with large I/O bandwidth, parallel association rule mining and GPGPUs appear to be a natural fit. While several data-mining algorithms exist on the GPU, to the best of our knowledge, this paper presents the first temporal data-mining algorithm ported to a GPU.

3.2. Core Algorithm

The core algorithm of Algorithm 1 is the counting step. To discover the presence of an episode $A_j = \langle a_1, a_2, \dots, a_L \rangle$ of level L a finite state machine can be implemented as shown in Figure 3. Each item a_i in the episode is a state with transitions to a_{i+1} , a_1 , and to itself depending on the value of the next character c . Because we are counting *all* episodes present in the database, when the *final* state is reached, a

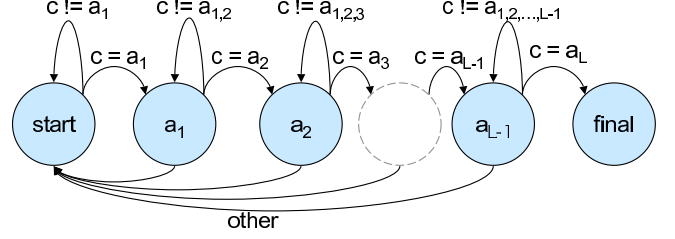


Figure 3. Finite state machine to discover the presence of an episode

counter keeping track of the total number of episodes found is incremented, and the finite state machine is reset back to *start* where the process repeats until all of the characters in the database are compared.

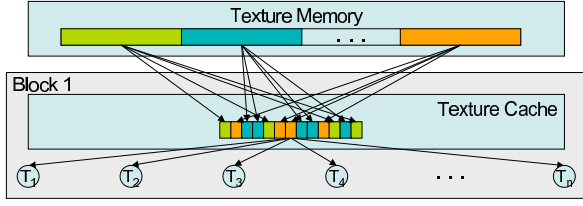
3.3. Parallelism on CUDA

Due to the wide range of parallelism available, we implemented four algorithms using the CUDA programming framework. These four algorithms can be generally classified as the cartesian product of (1) thread-level parallelism or block-level parallelism, with (2) local buffering or no local buffering, following a MapReduce programming model. The algorithms are shown graphically in Figure 4. In this paper, we do not explore block placement or scheduling as the programmer does not currently have an interface with which to alter these parameters. One can choose what code runs on which block, but one has no knowledge of what multiprocessor the block is running on or in what order it will be scheduled.

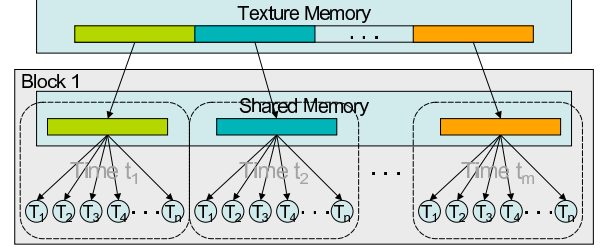
3.3.1. MapReduce. At a high level, our algorithms follow a MapReduce programming model to achieve efficient parallelism. Using the observation that counting the number of occurrences of A_k is independent from counting A_l , the *map* function returns the number of occurrences of A_j in the database portion D_i . The *reduce* function is dependent on whether thread or block parallelism is used.

3.3.2. Thread-Level Parallelism. Our first two algorithms implement strict thread-level parallelism to assign one thread to search for one episode, A_j , in the database D . Using one thread to search for one episode causes the *reduce* function to be an identity function which simply outputs the value produced by the *map* function.

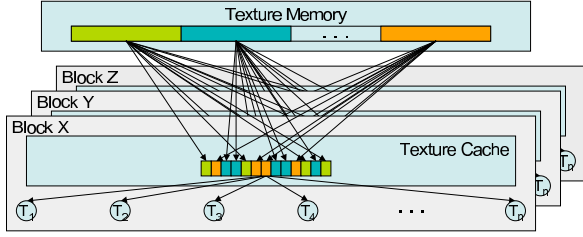
Algorithm 1: Thread-Level Without Buffering. Since each thread will scan the entire database, our first algorithm places this database in the read-only texture memory such that each thread will be able to utilize the high bandwidth available on the GPGPU. With each thread starting from the same point in the database, the spatial and temporal locality of the data-access pattern should be able to be exploited



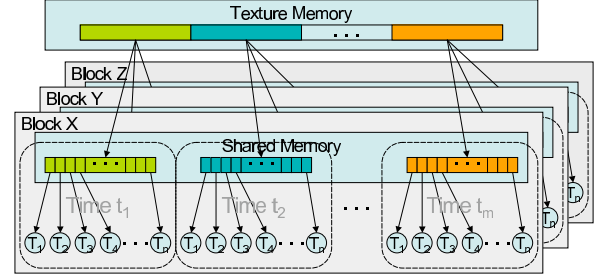
(a) Algorithm 1: each thread searches for a unique episode using data stored in texture memory



(b) Algorithm 2: each thread searches for a unique episode but uses shared memory to buffer the data first



(c) Algorithm 3: threads within a block search for the same episode with different blocks searching for different episodes using data in texture memory



(d) Algorithm 4: threads within a block search for the same episode with different blocks searching for different episodes but uses shared memory to buffer the data first

Figure 4. Available parallelism for each of the four algorithms

by the texture cache on each multiprocessor. Furthermore, threads are assigned to thread blocks in order until the maximum number of threads per thread block is reached. For example, if the maximum number of threads per thread block is 512, then threads 1-512 are assigned to the first thread block, threads 513-1024 to the second thread block, and so on, until no threads are left.

Algorithm 2: Thread-Level With Buffering. Our second algorithm also uses thread-level parallelism, but instead of using texture memory, this algorithm buffers portions of the database in shared memory in order to minimize the thread contention for the memory bus. Following this approach, a thread copies a block of data to a shared memory buffer, processes the data in the buffer, then copies the next block of data to the memory buffer, and so on, until the entire database is processed. The scheduling of threads to thread blocks follows the same approach in Algorithm 1.

3.3.3. Block-Level Parallelism. At a higher level of parallel abstraction, the CUDA programming model enables parallelism at the block level. At this level, our two block-level algorithms assign one block to search for one episode. Within a block, the threads collaborate to perform the search by having each thread search a portion of the database. With multiple threads searching for the same episode, the *reduce* function cannot be the identity function. Instead, the *reduce* function sums the counts from each thread. However, since an episode might span across threads, an intermediate step to

check for this possibility occurs between the *map* and *reduce* functions. An example of an episode spanning across threads is shown in Figure 5.

Algorithm 3: Block-Level Without Buffering. Similar to Algorithm 1, we implement this version of our block-level parallel algorithm without buffering of the database. Instead, the threads within each block access the data through texture memory. However, unlike Algorithm 1, each of the t threads within a block start at a different offset in the database while threads with the same ID in *different* blocks start at the same offset. The total number of threads available using Algorithm 3 is $t \cdot e$ where e is the number of candidate episodes to be searched for.

Algorithm 4: Block-Level With Buffering. The final algorithm we discuss in this paper uses block-level parallelism with buffering of the database to shared memory. The starting offset for each thread in Algorithm 4 is relative to the buffer size and *not* the database size as in Algorithm 3. Therefore, thread T_i will always access the exact same block of shared memory for the entire search – the data within the memory will change as the buffer is updated. The total number of available threads is identical to Algorithm 3.

4. Experimental Testbed

To analyze the performance characteristics of temporal data mining on graphics processors, we performed a series of tests on three generations of NVIDIA GPGPUs representing recent and current technology. An overview of the

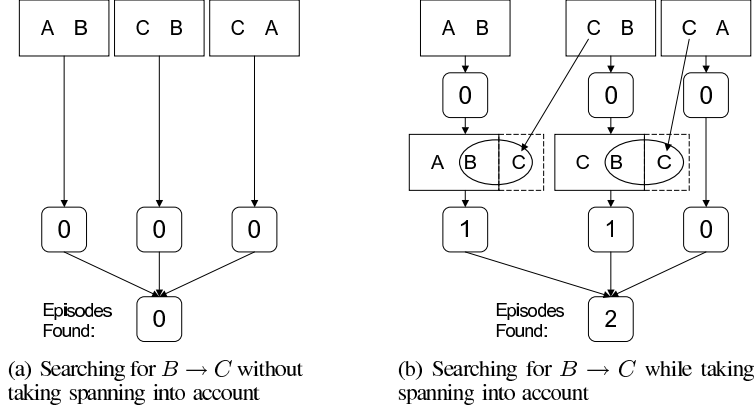


Figure 5. Example where the number of episodes found depends on whether spanning is taken into account

Graphics Card	GeForce 8800 GTS 512	GeForce 9800 GX2	GeForce GTX 280
GPU	G92	2xG92	GT280
Memory (MB)	512	2x512	1024
Memory Bandwidth (GBps)	57.6	2x64	141.7
Multiprocessors	16	16	30
Cores	128	128	240
Processor Clock (MHz)	1625	1500	1296
Compute Capability	1.1	1.1	1.3
Registers per Multiprocessor	8196	8196	16384
Threads per Block (Max)	512	512	512
Active Threads per Multiprocessor (Max)	768	768	1024
Active Blocks per Multiprocessor (Max)	8	8	8
Active Warps per Multiprocessor (Max)	24	24	32

Table 2. Feature summary of GeForce 8800 GTS 512, GeForce 9800 GX2, and GeForce GTX 280

architectural features are shown in Table 2 with a more detailed description of each graphics card along with the host machine in which the tests were performed presented below. A complete list of the specifications for the different compute capability levels of the GPUs found in these cards can be found in Appendix A of [2].

4.1. Host

The host machine consists of an E4500 Intel Core2 Duo running at 2.2 GHz with 4 GB (2x2GB) of 200-MHz DDR2 SDRAM (DDR2-800). The operating system is a 64-bit version of the Ubuntu GNU/Linux 7.04 distribution running the 2.6.20-16-generic Linux kernel as distributed through the package management system. Programming and access to the GPUs used the CUDA 2.0 toolkit and SDK with the NVIDIA driver version 177.67. Furthermore, all processes related to the graphical user interface (GUI) were disabled to limit external traffic to the GPU.

4.2. Cards

4.2.1. NVIDIA GeForce 8800 GTS 512 with G92 GPU.

To evaluate a recent generation of NVIDIA CUDA tech-

nology, we ran our tests on an GeForce 8800 GTS 512 graphics card with NVIDIA G92 GPU and 512 MB of onboard GDDR3 RAM. NVIDIA lists the G92 GPU as having compute capability 1.1, where compute capability determines the features and specifications of the hardware. Informally speaking, different compute capabilities signify different hardware generations. The GeForce 8800 GTS 512 contains 16 multiprocessors with each multiprocessor containing eight 1625-MHz execution cores, 8196 registers, and 16 KB of shared memory. The warp size is 32 threads with warps scheduled in intervals of four cycles. There can be at most 512 threads per block with 768 active threads per multiprocessor implying that two blocks of 512 threads *cannot* be active simultaneously on the same multiprocessor. Furthermore, there can be at most 8 active blocks and 24 active warps per multiprocessor. The texture cache working set is between 6 and 8 KB per multiprocessor.

Beginning with compute capability 1.1, the GPU supports atomic operations between threads on 32-bit words in shared or global memory allowing programmers to write thread-safe programs. It is worth recalling, however, that this improvement does not allow for threads in different blocks to synchronize as each block is independent of other blocks.

4.2.2. NVIDIA GeForce 9800 GX2 with G92 GPU.

We also evaluated an NVIDIA GeForce 9800 GX2, which contains *two* NVIDIA G92 GPUs and *two* units of 512MB of GDDR3 RAM. Essentially, the 9800GX2 is two 8800 GTS 512 cards merged onto a single graphics card with the execution cores running at 1500 MHz instead of 1625 MHz as in the 8800 GTS 512. Additionally, the 9800 GX2 has a modest 10% increase in memory bandwidth over the 8800 GTS 512 (64 GBps versus 57.6 GBps).

4.2.3. NVIDIA GeForce GTX 280 with GT200 GPU.

The current generation of CUDA technology has compute capability 1.3. For our tests, we used a GTX 280 graphics card with GT200 GPU. With 1024 MB of GDDR3 RAM and 30 multiprocessors, this card has the largest amount of device memory, number of processing cores (240), and memory bandwidth (141.7 GBps) of the cards tested. Furthermore, this GPU has 100% more registers per multiprocessor (16384), 33% more active warps (32), and 25% more active threads (1024) than the G92 GPUs.

5. Results

We present several performance characterizations of our algorithms running on different graphics cards at different episode levels with varying numbers of threads per block. At episode level L , an algorithm is searching for an episode A_j of length L , where $A_j = \langle a_1, a_2, \dots, a_L \rangle$. In the results presented, $L \in \{1, 2, 3\}$, a_l is a member of the set of upper-case letters in the English alphabet (i.e., $a_l \in \{A, B, \dots, Z\}$), and the database contains a total of 393,019 letters. In our experiments, level 1 contains 26 episodes, level 2 contains 650 episodes, and level 3 contains 15,600 episodes.

A single test consists of selecting a single episode level, algorithm, card, and block size with the execution time counted as the amount of time between the moment the kernel is invoked, to the moment that it returns. Although we restricted access to the GPU by disabling all non-vital graphical services to minimize the effect of errant GPU calls, each test was performed ten times with the average used to represent the test time. However, the minimum and maximum execution times are also displayed to show the range of times that the execution can take, which in some cases is quite large.

While the complete results from our tests include 12 different tests with different dimensions of criteria, we detail below some characterizations from these tests with respect to three higher-level criteria – level, algorithm, and card – and their impact on execution time. We also note that because some of the low-level architectural information of the NVIDIA GPUs is unavailable to the public, that the characterizations presented are *general* in nature; we dis-

cuss plans to uncover some of these low-level architectural features in Section 6.

5.1. Impact of Level on Execution Time

To understand the impact of the problem size on execution time, we performed a series of tests where the hardware and algorithm remained constant, but the level L varied. Because the number of episodes to search for increases exponentially as a function of L , as noted earlier in Table 1, the scalability of an algorithm with respect to problem size is important.

5.1.1. Characterization 1: Thread Parallel Algorithm has $O(C)$ Time Complexity Per Episode. Algorithm 1 and 2 are constant time algorithms per episode. Whether performing 26, 1560, or 25,320 searches, the amount of time to complete each individual search is essentially the same. Since a search for each episode is completely independent of other searches, and each search is assigned to one thread, there is no added complexity needed during the reduce phase to identify episodes that span thread boundaries. In addition, the search is based on a simple state machine the complexity of searching for a single episode in a single dataset stays constant regardless of the level. Therefore, the entire execution time can be spent performing the *map* function across the entire database. We explain in Section 5.2 whether Algorithm 1 or Algorithm 2 has an overall faster execution time for various levels.

Since these algorithms are constant time per search, we actually find that they are effectively constant time for up to a rather large number of searches when executed on the GPU as can be seen in Figures 6(a) and 6(b). By this we mean that 26, 650, or even several thousand searches complete in the same amount of time on the GPU, reasons for this will become more clear in Characterizations 4 and 7.

5.1.2. Characterization 2: Buffering Penalty in Thread Parallel Can be Amortized. Algorithm 2 uses buffering to combine the memory bandwidth of all threads in a block and reduce contention on texture memory. This does not however, come without a cost. The initial load time is high, and since only one block may be resident on a multiprocessor during this time, no computation is done during the load. As more threads are added to a block Algorithm 2 exponentially decreases in execution time as shown in Figure 6(b). This characteristic shows that Algorithm 2 is able to make use of the processing power of a greater number of threads, largely thanks to the fact that the load time is either equal to or lower than the search time for smaller numbers of threads. Since the load cost is approximately constant, and all threads can access the resulting shared memory block with minimal contention, results will be calculated faster the more threads there are in a block up to the point where scheduling overhead on the multiprocessor overwhelms the computation time.

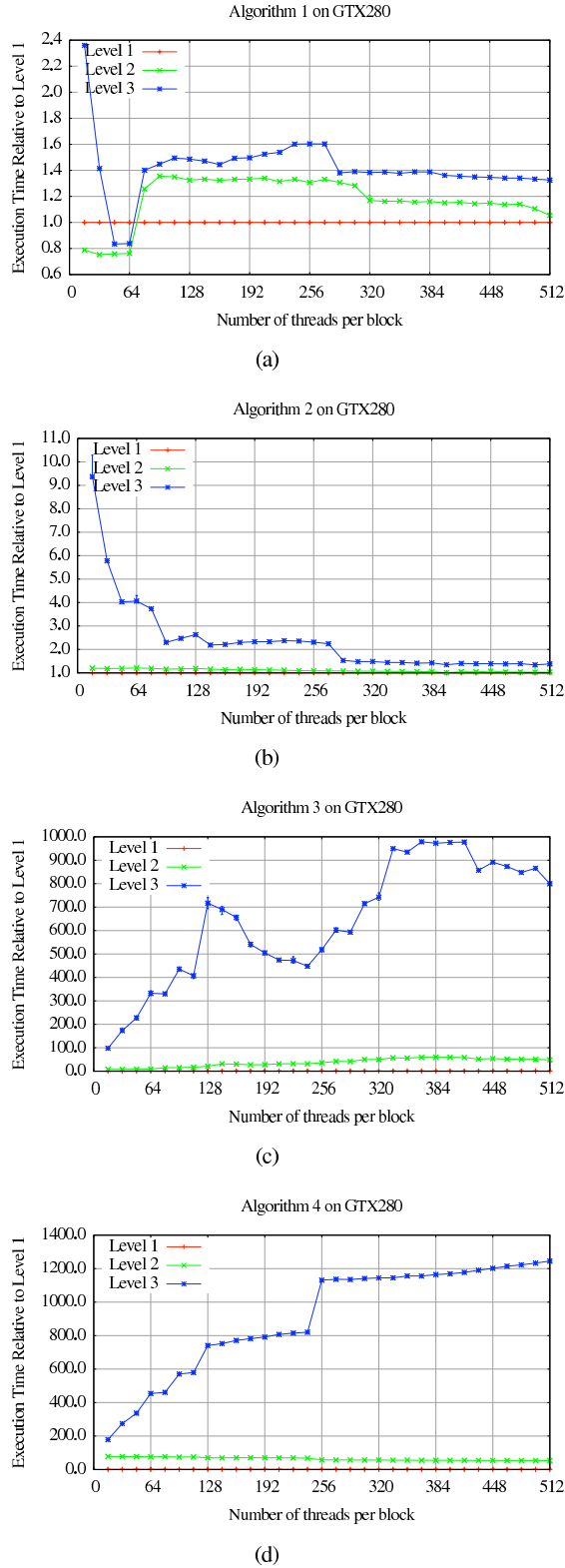


Figure 6. Impact of problem size on the GTX 280 on different algorithms

5.1.3. Characterization 3: Block Parallelism Does Not Scale with Block Size. Unlike Algorithm 2, Algorithms 3 and 4 actually lose performance per episode as the number of threads per block and level increase. Figures 6(c) and 6(d) show a general trend of larger execution times as the number of threads increases with Algorithm 4 at an almost constant slope when solving the problem size at Level 3. Furthermore, the change in execution time between Level 1 and Level 2 and between Level 2 and Level 3 is also increasing. These two trends are due to the extra complexity of finding episodes that span thread boundaries and the cost of loading more blocks than can be active on the card simultaneously.

As the number of threads increases, the number of boundaries increases. As the level (i.e., episode length) increases, the likelihood that an episode spans the boundary between threads increases. With the number of boundaries increasing and the probability that an episode will span a boundary increasing as well, the computation needed to be performed after the *map* function and before the *reduce* function increases producing longer overall execution times.

5.2. Impact of Algorithm on Execution Time

While the scalability of an algorithm with respect to problem size is important, it is often the case that a user wishes to examine a problem of specific size and only has access to one type of hardware. That is, a user wants to solve the same problem on the same card and can only vary the algorithm and number of threads to use for that algorithm. For example, a neuroscientist may want to examine a neural dataset from one experiment in minute detail before deciding to perform subsequent experiments. In this case, the user would want to use the fastest algorithm for the specific problem. Our characterizations below are in relation to the GTX 280 as it is the most recent of the cards tested.

5.2.1. Characterization 4: Thread Level Parallelism Alone is Not Sufficient for Small Problem Sizes. When evaluating small problem sizes, e.g., $L = 1$, there are not enough episodes to generate enough threads to utilize the resources of the card. It is necessary to first add parallelism at the block level and then to incorporate multiple threads within each block. Since the number of episodes is defined by the threshold and there is 1 thread per episode, having more than 26 threads active in Algorithm 1 or Algorithm 2, only increases contention for the processing cores which is why these algorithms have an uptrend as the number of threads increases (Figure 7(a)). Algorithms 3 and 4, on the other hand, are orders of magnitude faster as they first create 26 blocks and add threads to help search for same episode. As such, the execution times for these algorithms trend downwards but plateau since the 26 thread blocks are not sufficient to fully utilize all 30 multiprocessors on the GTX 280 graphics card.

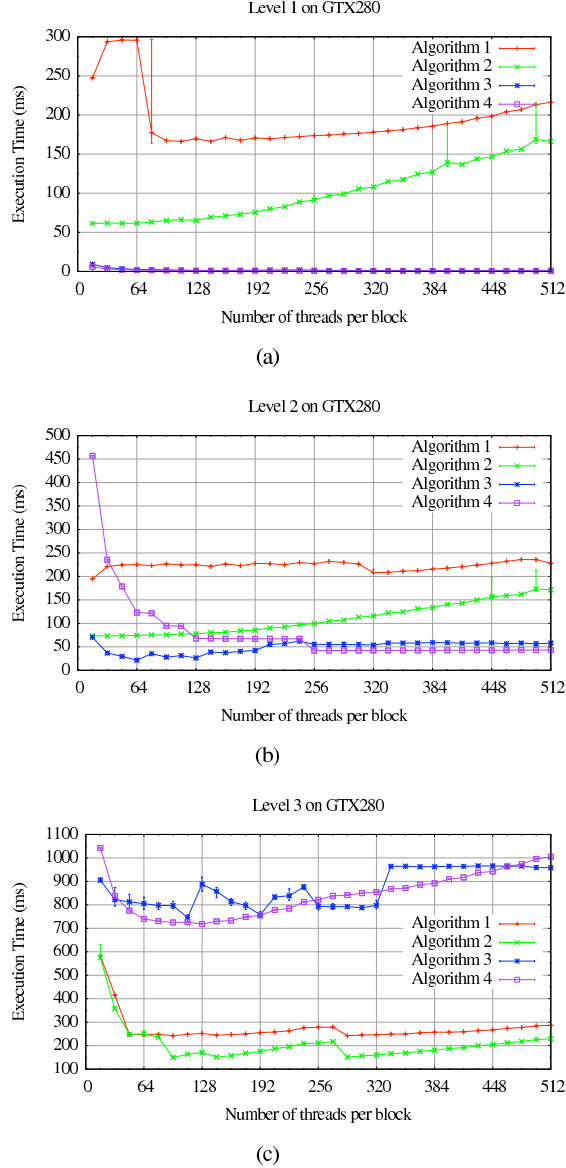


Figure 7. Impact of algorithm on the GTX 280 at different problem sizes

5.2.2. Characterization 5: Block Level Depends on Block Size for Medium Problem Sizes. As seen in Figure 7(b), a medium problem size ($L = 2$) contains enough parallelism at the thread level such that the thread-parallel Algorithm 2 outperforms the block-parallel Algorithm 4 for small numbers of threads per block and therefore larger numbers of blocks. That is, Algorithms 1 and 2 decrease the total number of thread blocks as the number of threads per block increases due to the fixed number of episodes equating to a fixed number of total threads. At Level 2, the number of blocks varies as a function of threads per block starting at $\frac{650}{16}$ and decreasing to $\frac{650}{512}$. Unlike Level 1, a value of 32

threads per block or greater results in more thread blocks contributing *positively* to the search.

For the block-level algorithms, Algorithm 4 eventually outperforms Algorithm 3 (at 64 threads per block), but it never achieves the *best* execution time which is Algorithm 3 at 64 threads. An explanation of this is hard to pinpoint exactly as the internal workings and scheduling of the NVIDIA GPGPUs are not publicly available. However, we believe that by using texture memory and heavy caching Algorithm 3 will obtain close to optimal bandwidth with fewer threads per block resulting in less contention (which will only increase as more threads are added). Additionally, the buffering to texture memory is a one-time penalty which, as we mentioned in Characterization 2, is amortized over the total number of threads and alleviated by accessing the shared memory in read-only fashion.

5.2.3. Characterization 6: Thread-Level Parallelism is Sufficient for Large Problem Sizes. The GTX 280 has 30 multiprocessors with a maximum of 1024 active threads per multiprocessor for a total of 30,720 potentially active threads available. When $L = 3$, there are 25,230 episodes to be searched. As seen in Figure 7(c), the thread-level parallel algorithms (Algorithm 1 and 2) are significantly faster than the block level algorithms (Algorithms 3 and 4). This performance difference can be attributed to the fact that with 25,230 episodes to be searched, Algorithms 1 and 2 can have more episodes being searched simultaneously than Algorithms 3 and 4 for a given number of threads per block. Algorithms 3 and 4 are limited to 240 episodes being searched due to the limitation of 8 active blocks on each of the 30 multiprocessors in the GTX 280 and each block searching for a single episode. Algorithms 1 and 2 on the other hand, can have up to 30,720 active episodes as each *thread* within a block will search for a unique episode. The actual number of active episodes for Algorithms 1 and 2 is determined by the resources that each thread consumes and the available resources on each multiprocessor.

5.3. Impact of Hardware on Execution Time

The other major decision which can affect performance is the choice of hardware on which the algorithm will be run. Some users may have a variety of hardware and wish to know which will return results the fastest, or still others may wish to determine the optimal card for their problem when considering a new purchase. The characterizations below showcase two of the major factors which come into play in determining the right card for the job.

5.3.1. Characterization 7: Thread Level Parallelism is Dependent on Processor Clock Frequency for Small and Medium Problems. Algorithms 1 and 2 are greatly dependent on the processor clock frequency for small and

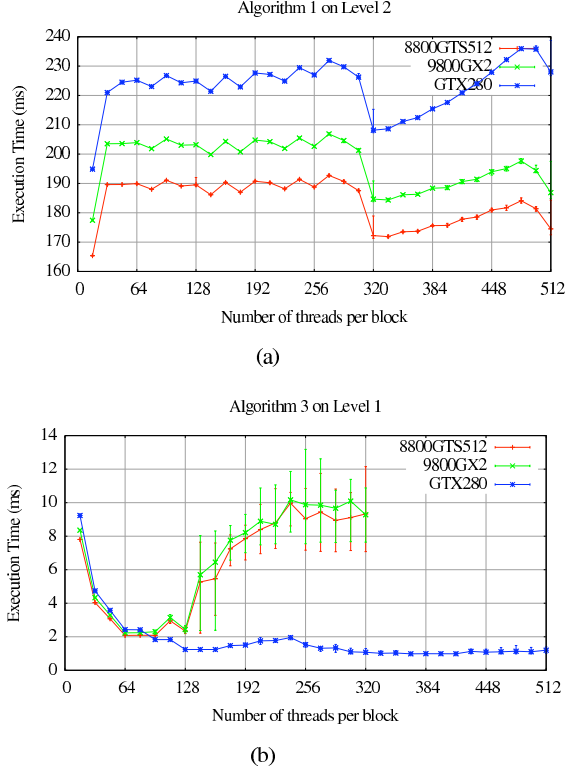


Figure 8. Impact of hardware on execution time

medium problem sizes, and scale essentially linearly by this measure as one can see in Figure 8(a). Referring back to Table 2, the frequencies of the three cards tested are 1625 MHz (8800 GTS 512), 1500 MHz (9800 GX2), and 1296 MHz (GTX 280). With this information and the results in Figure 8(a), it is clear to see that the relationships between frequency speed closely match the relationships between execution times. This trend is due to the fact that levels 1 and 2 are incapable of filling enough multiprocessors for the number of processors or contention to become a determining factor in performance. In the Appendix, one can see in that the same test at Level 3 produces a very different picture, where the 30-multiprocessor GTX 280 outperforms the 16-multiprocessor 9800 GX2 and the 8800 GTS 512 at nearly all thread counts.

5.3.2. Characterization 8: Block-Level Algorithms are Affected by Memory Bandwidth. Algorithm 3 can be greatly effected by memory contention when searching large problem sizes. This occurs due to the fact that the total number of threads accessing memory is $active_episodes * threads_per_block$. Given the massive number of blocks needed for large problem sizes, this algorithm requires a very high thread count per *multiprocessor* over a long period producing a large and constant amount of memory

contention. The two slowest performing cards, the 8800 GTS 512 and 9800 GX2, actually have fewer processors (128) contending for memory than the GTX 280 (240), but also have memory bandwidth in the range of 62 to 64 GBps as compared to the 141GBps on the GTX 280. The wide variance in individual execution times and higher overall execution times resulting from these hardware differences can be seen in Figure 8(b).

6. Future Work

Although we have successfully developed a high-performance, parallel, temporal data mining application on a GPU, we are pursuing three improvements to make real-time temporal data mining available to neuroscientists.

First, we wish to add support for arbitrarily large episodes and observe the impact on performance of both the thread level and block level algorithms. We are particularly interested in observing whether the thread level algorithms (e.g., Algorithm 1 and 2) will continue to scale.

We are also looking at the effect of feature changes on the algorithm execution time. One feature is episode expiration where $A \Rightarrow B$ iff $B.time() - A.time() < \delta$. With episode expiration, episodes will be less likely to span multiple thread boundaries as they will be constrained by the value of δ . Correspondingly, we expect the *reduce* phase in Algorithms 3 and 4 to decrease as fewer episodes will span boundaries on average.

Lastly, while being able to identify general performance characteristics, it is difficult to identify how *optimal* performance can be obtained due to the unavailability of low-level architectural information on specific hardware. To remedy this, we plan to create a series of micro-benchmarks to discover the underlying hardware and architectural characteristics such as scheduling and caching. The CUDA Occupancy Calculator made available by NVIDIA is a useful resource, but is insufficient in identifying how *optimal* performance can be obtained as it makes basic assumptions on the algorithm in use.

7. Conclusion

The ability to mine data in real-time will enable scientists to perform research in ways that have only been dreamed about. However, as the sheer volume of data grows, the algorithms used to mine this data need to keep pace or the utility of the data is lost. In the field of neuroscience the inability to analyze the data can have fatal consequences.

One approach to this problem of “data overload” is to create new parallel algorithms capable of extracting the computational performance available on GPGPUs. In this paper, we have developed and characterized the performance of a parallel temporal data mining application on NVIDIA CUDA graphics processors.

As one might expect, the best execution time for large problem sizes always occurs on the newest generation of the hardware, the NVIDIA GeForce GTX 280 graphics card. What is surprising however, is that the oldest card we tested was consistently the fastest for small problem sizes. Beyond this, our results showed that the extent to which a GPGPU MapReduce-based frequent episode mining implementation must *dynamically* adapt the type and level of parallelism with respect to episode length and algorithm in order to obtain the best performance is inconsistent. For example, when searching for episodes of length 1, an algorithm using blocks with 256 threads and buffering to shared memory achieves the best performance. However, episodes of length 2 require block sizes of 64 *without* buffering, and episodes of length 3 should use 96 threads per block with every thread searching for a unique episode.

None-the-less, due to the ever-increasing volume of data and demand for high performance in neuroscience and bioinformatics, we have provided 8 performance characterizations as a guide for future temporal data mining applications on GPGPUs.

Acknowledgments

We would like to thank Naren Ramakrishnan, Debprakash Patnaik, and Patrick Butler for their invaluable discussions concerning the theory behind temporal data mining and Sean Ponce for ensuring the correctness of the GPGPU implementation.

This research was supported in part by an NVIDIA Professor Partnership Award.

References

- [1] Hadoop: Open-Source Implementation of MapReduce, 2008.
- [2] *NVIDIA CUDA Compute Unified Device Architecture - Programming Guide, Version 2.0*, 2008.
- [3] Rakesh Agrawal, Tomasz Imielinski, and Arun N. Swami. Mining Association Rules between Sets of Items in Large Databases. In *Proceedings of the 1993 ACM SIGMOD International Conference on Management of Data*, 1993.
- [4] Bryan Catanzaro, Narayanan Sundaram, and Kurt Keutzer. A Map Reduce Framework for Programming Graphics Processors. In *Third Workshop on Software Tools for MultiCore Systems (STMCS 2008)*, 2008.
- [5] Cheng-Tao Chu, Sang Kyun Kim, Yi-An Lin, YuanYuan Yu, Gary Bradski, Andrew Y. Ng, and Kunle Olukotun. MapReduce for Machine Learning on Multicore. In *Advances in Neural Information Processing Systems (NIPS) 19*, 2006.
- [6] Rajai El Dajani, Maryvonne Miquel, Marie-Claire Forlini, and Paul Rubel. Modeling of Ventricular Repolarisation Time Series by Multi-layer Perceptrons. In *AIME '01: Proceedings of the 8th Conference on AI in Medicine in Europe*, pages 152–155, London, UK, 2001. Springer-Verlag.
- [7] Gautam Das, King ip Lin, Heikki Mannila, Gopal Renganathan, and Padhraic Smyth. Rule Discovery from Time Series. In *Proceedings of the Fourth International Conference on Knowledge Discovery and Data Mining*, pages 16–22. AAAI Press, 1998.
- [8] Jeffrey Dean and Sanjay Ghemawat. MapReduce: Simplified Data Processing on Large Clusters. In *OSDI '04*, 2004.
- [9] David J. DeWitt and Michael Stonebraker. MapReduce: A Major Step Backwards, 2008.
- [10] David J. DeWitt and Michael Stonebraker. MapReduce II, 2008.
- [11] Martin Gavrilo, Dragomir Anguelov, Piotr Indyk, and Rajeev Motwani. Mining the Stock Market (extended abstract): Which Measure is Best? In *KDD '00: Proceedings of the Sixth ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pages 487–496, New York, NY, USA, 2000. ACM.
- [12] Bingsheng He, Wenbin Fang, Qiong Luo, Naga K. Govindaraju, and Tuyong Wang. Mars: A MapReduce Framework on Graphics Processors. In *PACT08: IEEE International Conference on Parallel Architecture and Compilation Techniques 2008*, 2008.
- [13] Marc De Kruijf. MapReduce on the Cell Broadband Engine Architecture. Technical report, 2007.
- [14] Debprakash Patnaik, P. S. Sastry, and K. P. Unnikrishnan. Inferring Neuronal Network Connectivity using Time-constrained Episodes. *CoRR*, abs/0709.0218, 2007.
- [15] John F. Roddick and Myra Spiliopoulou. A Survey of Temporal Knowledge Discovery Paradigms and Methods. *IEEE Transactions on Knowledge and Data Engineering*, 14(4):750–767, 2002.
- [16] Michael Steinbach, Steven Klooster, Pang ning Tan, Christopher Potter, and Vipin Kumar. Temporal Data Mining for the Discovery and Analysis of Ocean Climate Indices. In *Proceedings of the KDD Temporal Data Mining Workshop*, pages 2–13, 2002.
- [17] K. P. Unnikrishnan, Debprakash Patnaik, and P. S. Sastry. Discovering Patterns in Multi-neuronal Spike Trains using the Frequent Episode Method. *CoRR*, abs/0709.0566, 2007.
- [18] K.P. Unnikrishnan, P.S. Sastry, and Naren Ramakrishnan. GMiner, <http://neural-code.cs.vt.edu/gminer.html>.
- [19] Pablo Valenti, Enrique Cazamajou, Marcelo Scarpettini, Ariel Aizemberg, Walter Silva, and Silvia Kochen. Automatic Detection of Interictal Spikes using Data Mining Models. *Journal of Neuroscience Methods*, 150(1):105–110, 2006.
- [20] Jackson H.C. Yeung, C.C. Tsang, K.H. Tsoi, Bill S.H. Kwan, Chris C.C. Cheung, Anthony P.C. Chan, and Philip H.W. Leong. Map-reduce as a Programming Model for Custom Computing Machines. In *IEEE Symposium on Field-Programmable Custom Computing Machines (FCCM)*, 2008.
- [21] Mohammed J. Zaki. Parallel and Distributed Association Mining: A Survey. *IEEE Concurrency*, 7(4):14–25, 1999.

Appendix

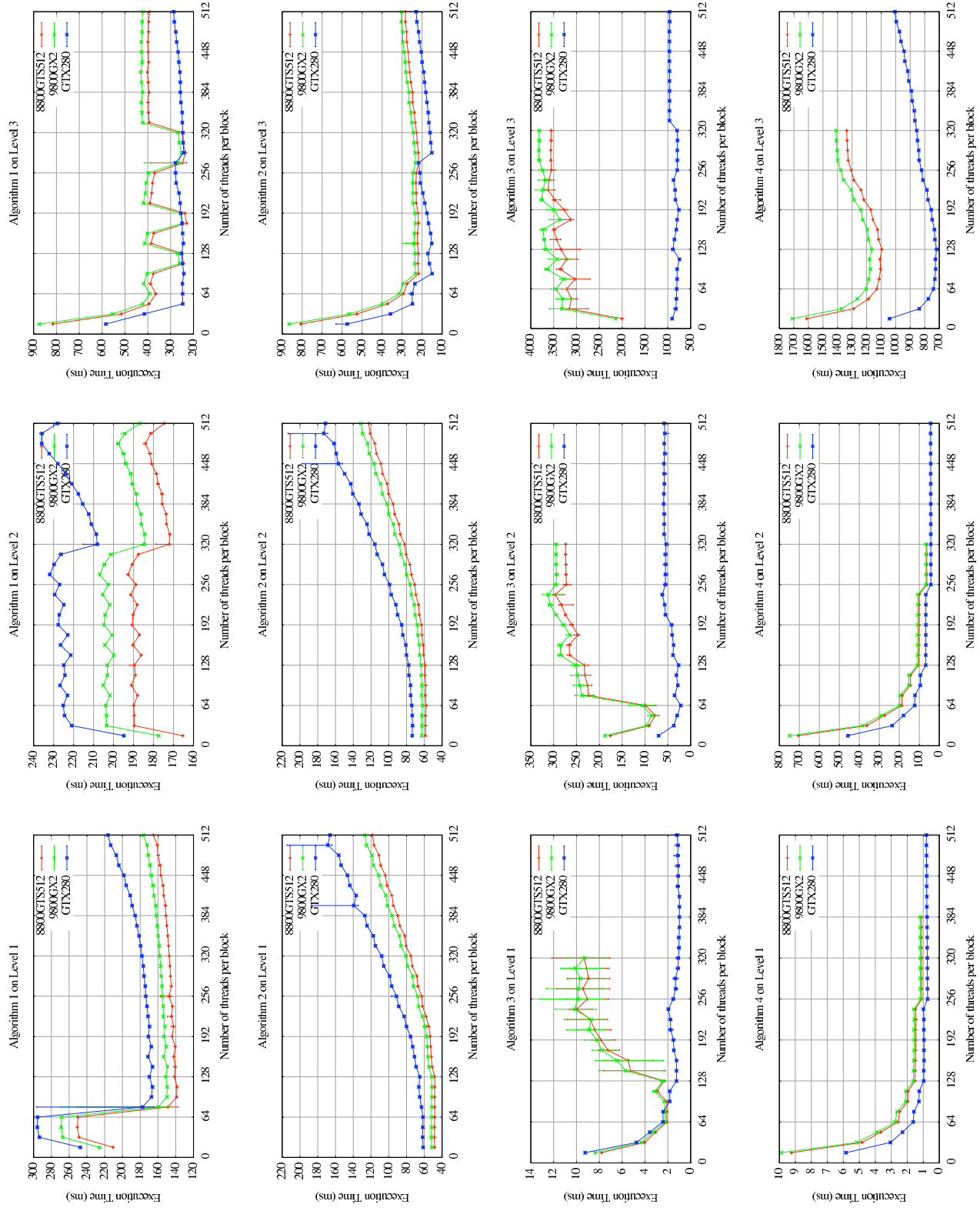


Figure 9. Complete set of tests: Row = Algorithm, Column = Level, X-axis = Number of threads per block, Curve = GPGPU hardware