



Serving DNNs like Clockwork: Performance Predictability from the Bottom Up

Arpan Gujarati*

Max Planck Institute for Software Systems

Safya Alzayat

Max Planck Institute for Software Systems

Antoine Kaufmann

Max Planck Institute for Software Systems

Reza Karimi*

Emory University

Wei Hao

Max Planck Institute for Software Systems

Ymir Vigfusson

Emory University

Jonathan Mace

Max Planck Institute for Software Systems

Abstract

Machine learning inference is becoming a core building block for interactive web applications. As a result, the underlying model serving systems on which these applications depend must consistently meet low latency targets. Existing model serving architectures use well-known reactive techniques to alleviate common-case sources of latency, but cannot effectively curtail tail latency caused by unpredictable execution times. Yet the underlying execution times are not fundamentally unpredictable—on the contrary we observe that inference using Deep Neural Network (DNN) models has deterministic performance.

Here, starting with the predictable execution times of individual DNN inferences, we adopt a principled design methodology to successively build a fully distributed model serving system that achieves predictable end-to-end performance. We evaluate our implementation, Clockwork, using production trace workloads, and show that Clockwork can support thousands of models while simultaneously meeting 100 ms latency targets for 99.9999% of requests. We further demonstrate that Clockwork exploits predictable execution times to achieve tight request-level service-level objectives (SLOs) as well as a high degree of request-level performance isolation.

1 Introduction

With the proliferation of machine learning (ML), model inferences are now not only commonplace but increasingly on the critical path of web requests [29, 71]. Inference requests are handled by underlying model serving services [16, 26, 51, 58] responsible for supporting scores of different pre-trained ML models (including personalized models and experimental A/B tests), ideally at low latency, high throughput, and low cost. These are demanding goals to meet at scale—Facebook alone serves over 200 *trillion* inference requests each day [48]. Furthermore, at least 100 companies are creating hardware chips for accelerated ML inference [48], which underscores the high stakes in this industry.

Yet significant software bottlenecks continue to hamper the efficient utilization of hardware accelerators, such as GPUs, for high-performance model serving. Consider an inference request passing through a model serving system. The request has an inherent deadline after which the answer ceases to be useful to the end-user, and so the system should seek to bound the latency of the request, or even provide service level objectives (SLOs) for consistently achieving low tail latency. The canonical approach for building such a low-latency system is to reduce potential wait times for resources through over-provisioning, since a larger pool of available resources makes it more likely to find a resource on which a pending request can be immediately scheduled. Increased resource provisioning, however, comes at the expense of efficiency and utilization.

Existing systems fundamentally assume that the constituent system components have *unpredictable* latency performance [16, 58]. Moreover, the best-effort techniques employed to tolerate such variability, such as fair queuing, further cascade the unpredictability to other system components and propagate tail latency to higher layers. While some performance volatility of a model serving system is due to external factors, such as a bursty or skewed workload, much variability in execution times stems from design decisions internal to the service, ranging from caching decisions over conditional branching behavior to concurrency from other processes, the OS, and the hypervisor. The challenge, then, is to tame the internal unpredictability.

In this paper, we present the design and implementation of Clockwork, a distributed system for serving models with predictable performance. With an explicit focus on the ubiquitous deep neural network (DNNs) architectures we first show that DNN inference is fundamentally a deterministic sequence of mathematical operations that has a predictable execution time on a GPU. To leverage this observation in designing a responsive model serving system, our approach is to preserve predictability wherever possible by *consolidating choice*: eschewing reactive and best-effort mechanisms and centralizing all resource consumption and scheduling decisions. Clockwork will only execute an inference request if it is confident that the request can meet its latency SLO. To

* Equal contribution

support such proactive scheduling, Clockwork is composed of *workers* that each handle one or more GPUs, and a centralized *controller* that schedules requests. Each Clockwork worker, responsible for the exclusive model loading and inference execution on the GPUs, achieves predictable performance. If a worker cannot execute a particular schedule, because of external factors, the request is immediately aborted and the worker resumes execution of the next request at the specified time. The Clockwork controller manages the resources of each worker and maintains a minimal advance schedule for the worker’s operations, including model placement and replication.

We have implemented Clockwork in C++ and evaluated it using a wide range of DNN models on production workload traces. In comparison to Clipper [16] and INFaaS [58], two prior model serving systems, Clockwork more effectively meets latency goals while providing comparable or better goodput. Clockwork more effectively shares resources between different models, and scales to thousands of models per worker. For realistic workloads comprising unpredictable, bursty, and cold-start clients, Clockwork consistently meets low-latency response times of under 100ms.

The main contributions of this paper are as follows:

- We demonstrate that predictability is a fundamental trait of DNN inference that can be exploited to build a predictable model serving system.
- We propose a system design approach, *consolidating choice*, to preserve predictable responsiveness in a larger system comprised of components with predictable performance.
- We present the design and implementation of Clockwork, a distributed model serving system that mitigates tail latency of DNN inference from the bottom up.
- We report from an experimental evaluation on Clockwork to show that the system supports thousands of models concurrently per GPU and substantially mitigates tail latency, even while supporting tight latency SLOs. Clockwork achieves close to ideal goodput even under overload, with unpredictable and bursty workloads, and with many contending users.

2 Background and Motivation

The state of machine learning. The meteoric rise of applications driven by machine learning (ML), ranging from computer vision [28, 78] to ad-targeting [3, 17] to virtual assistants [13, 64], has prompted significant interest into making both ML training and inference faster. These efforts have targeted the underlying ML models, hardware accelerators, and software infrastructure. Chief among the ML modeling approaches are *deep neural networks* (DNNs), which are composed of multiple layers of artificial neurons tuned through non-linear convolution and pooling operations [25].

A plethora of specialized hardware are being developed and deployed for ML training and inference [48], such as ASIC and FPGA chips, Google’s TPUs [41], and Facebook’s Big

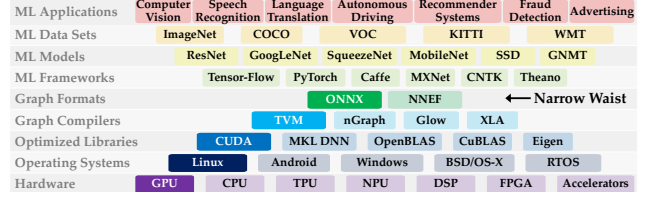


Fig. 1: Model serving targets the narrow waist of the ML software stack (adapted from Reddi *et al.* [48]). Clockwork targets the shaded blocks on the left.

Basin [29] chips. The dominant machine learning hardware in data centers, however, is the GPU, representing a third of the global market in 2020 [5], and will be our focus here.

Interposed between the emerging DNN applications and hardware accelerators, an ecosystem of ML software frameworks is flourishing. Fig. 1 displays several prominent projects in today’s ML software stack. Layered protocol stacks in complex systems and competitive environments tend to evolve into hourglass-shaped architectures [4]. We are witnessing the ONNX and NNEF graph exchange formats for DNNs [49, 52] emerging as the “narrow waist” of the ML stack, acting as an interface between high-level ML model development and low-level software and hardware concerns.

Model serving. Operators increasingly deploy machine learning on the critical path of nascent interactive applications [71]. This has elevated machine learning inference to separate, managed *model serving* services [16, 26, 58]. From the vantage point of an operator, the model serving users (customers or internal applications) upload their pre-trained DNN ahead of time (the natural format for which is ONNX/NNEF). Their applications can then submit *inference requests* to an API. The model serving back-end manages the users’ models and the hardware accelerator resources, and provides timely responses to inference requests. Upon receiving an inference request, it loads the appropriate model into hardware if not already loaded, runs the DNN on the input, and returns the resulting output to the user. Model serving has similar concerns to other datacenter services [2]: it multiplexes workloads of different users concurrently and load balances requests across multiple workers and GPU hardware accelerators.

Low-latency inference. Model serving users require a timely response to their queries. Most cloud and data center services have *service-level objectives* (SLOs) that codify the performance that clients can expect from the service [40]. The most common type is a *latency SLO*, which specifies the service’s acceptable request latencies, typically on the order of milliseconds [14, 32, 41]. For example, a latency SLO might specify a 10ms average response time, or a 40ms 99th percentile response time, or both. If a service fails to meet its SLOs – for example, by being too slow for too many requests – the service provider may risk a penalty.

Model serving further operates under hard cost constraints. Specialized ML hardware is necessary to achieve interactive

latencies [41], but it is comparatively expensive to procure and operate, and must thus be used efficiently [60, 65]. Existing model serving systems achieve efficient inferences for specific heavily used models by dedicating them entire GPUs and using copious batching [41]. However, many use cases cannot justify dedicated hardware resources: applications with insufficient request volume; specialization (*e.g.* location-specific search or language-to-language translation); and experimentation (*e.g.* retrained models and A/B testing) [63]. Efficiently serving models with low request rates requires a large number of models to share accelerators; no existing model serving system supports this.

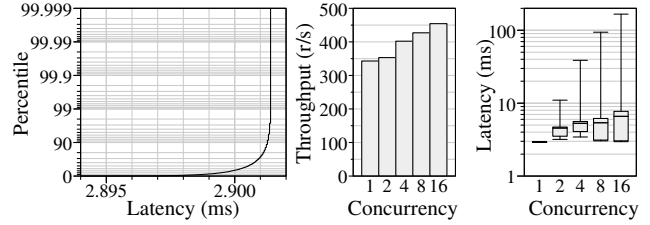
While it is already difficult for model serving operators to meet latency SLOs under these constraints, the bigger challenge lies in minimizing *tail latency*, the insidious bane of interactive performance. Numerous sources of latency variability in complex individual [46] and distributed [18, 56] systems have been identified and studied, including out-of-order scheduling, interference from concurrency, power saving modes, and network queuing delays.

The crux of tail latency lies in performance variability of both the constituent system/network components and the encompassing architecture. To tame it, the system designer can either seek to (quoting Dean and Barroso [18]) “*create a predictably responsive whole out of less-predictable parts*”, or to expend significant effort to systematically unshroud and mitigate the performance variability of these underlying components. To meet tight tail-latency SLOs under resource constraints, the latter approach is necessary.

Observation: DNN inference is predictable. We observe that DNN executions exhibit negligible latency variability, a result both intuitive in concept — DNN inferences involve no conditional branches — and demonstrable in practice. Although we describe our observations in the context of GPU execution, they extend to other accelerators such as TPUs, and also to CPU execution where appropriate.

Conceptually, a DNN inference is a fully deterministic execution. Each DNN inference request carries a fixed-size *input* tensor argument; in practical terms this is a statically-sized array of bytes. A worker receives this input over the network into main memory. To execute on a GPU, the input is copied from main memory to GPU memory over the PCIe interconnect. The DNN is then executed on the GPU. Abstractly, a DNN is a pre-defined sequence of tensor multiplications and activation functions. Concretely, the DNN code applies these operations to the input tensor one-at-a-time to transform the input into an output. DNN code lacks conditional branching; input choices such as batching size and RNN sequence length are specified ahead of time as parameters. The output is also a statically-sized array of bytes, and it is copied from GPU memory back to main memory over the PCIe interconnect.

We compiled ResNet50v2 [78] with TVM 0.7 [15] and executed 11 million inferences in isolation on a state-of-the-art



(a) CDF of 1-thread latency (b) Inference throughput and latency. Whiskers show min and max.

Fig. 2: Inference is predictable in isolation (left). Running inferences concurrently gains up to 25% throughput (middle), at a cost of substantially increased latency variability (right) due to interleaved GPU and OS executions.

NVIDIA Tesla v100 GPU using random inputs and batch size 1. We measured the latencies of each inference and show the median and high-percentile latencies in Fig. 2a. The 99.99th percentile latency was within 0.03% of the median latency.

If DNN execution times can be measured and then accurately predicted for future inferences on that model, the next question is whether a distributed model serving system can preserve the predictable responsiveness of the core inference execution.

3 Predictable Performance

To build a responsive system through principled design, we further study the factors that can cause or amplify performance variability. Importantly, components at any level of the modern system stack can contribute to variable request latency, whether at the application layer, in the operating system, or even in the hardware [46]. Network effects and workload fluctuations add two more sources of unpredictability to distributed systems.

The whole is more than the sum of its parts. The overall system performance variability is primarily governed by how the system is assembled from its constituent components. We can handle variable latency of a software component in several ways. First, we can ignore the problem and allow the volatility to propagate to later requests or percolate to other components of the system. Even performance-conscious code that is optimized to improve throughput or average latency does not fix tail latency [19]. An example of this contagiousness of unpredictability, known as the “straggler” problem in data analytics frameworks [7, 56], is when a worker executes a request that takes unusually long and the other requests that were enqueued on the worker in the meantime then incur the extra delay from the unexpected wait-time. Ignoring the variability can further compound the problem across the system, such as when the request handler itself has variable latency [69].

Second, we can mitigate the volatility by ensuring all requests match the worst-case latency, thus exchanging lower resource utilization for predictability—often a steep price when worst-case latency is significantly higher than the median.

Third, we can minimize variability by expending more resources, again in trade for lower utilization. Some networked systems, for instance, are designed to submit the same job to mul-

multiple workers in parallel and then to cancel unneeded jobs upon successfully receiving a result from the fastest worker [18].

Fourth, upon detecting an unusual delay, we can notify a feedback mechanism to adjust the environment to lower the impact on future requests. Such “best-effort” methods are typically reactive and aimed at longer-term effects, such as by temporarily adding more resources (auto-scaling [23]), throttling requests, or balancing load.

Consolidating choice. We take a fundamentally different approach: *designing a predictable system from the bottom up*. Our strategy is to restrict the choices available to lower system layers as much as possible—a philosophy based on our observation that when executing an essentially predictable task, performance variability only arose when a lower layer in the system was given choices regarding how to execute its task. Examples from all layers of the systems stack abound, including:

- **Hardware level:** when a GPU is passed multiple CUDA kernels to execute in parallel, the GPU has the choice of how to allocate resources, including execution units and memory bandwidth, between kernels. The GPU makes these choices based on its internal state and undocumented, proprietary policies.
- **OS level:** when we create multiple threads that the operating system can execute on the same core, the OS has the choice of what threads to execute when, based on internal scheduling policies and state.
- **Application level:** when the worker processes of a distributed application each manage their own cache independently, the workers have the choice of what to cache and for how long, leading to unpredictable hit rates and latency variability [38]; similarly, when worker processes implement their own thread pools and queuing policies, they have the choice of which requests to execute first, leading to unpredictable queuing times.

Fig. 2b illustrates this: a standard design for building a worker would use thread pools serving inference requests in parallel to saturate the GPU. While concurrent threads indeed increase inference throughput by up to 25%, the factors above cause tail latency to increase by 100x.

Our approach is to *consolidate choices* in the upper layers: once a layer implements choices for lower layers based on internal state, it forces the lower layer to follow a narrow path of possible executions, causing the performance of the resulting layer to be nearly deterministic. The upper layer can then sufficiently predict the performance of the lower layers and reason with foresight about resource utilization and the anticipated execution times for all requests. The price of this strategy, however, is a tighter coupling of components and a less modular architecture.

Imperfect predictability. Notably, we can consolidate choice without requiring *perfect* predictability. Real systems will retain some unpredictable components, such as managing CPU caches or workload shifts, even after consolidating choices in its upper layers. Instead, the chief goal of concentrat-

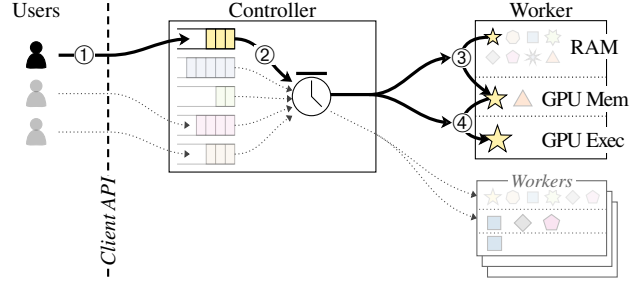


Fig. 3: Clockwork comprises multiple Workers and a centralized Controller. Models (★) reside on Workers; inference requests are queued and scheduled centrally on Clockwork’s Controller. See §4.1 for a detailed description.

ing these choices is to make predictable executions the common case. This frees us from implementing best-effort mechanisms to tolerate the occasional, rare instance of unpredictability; instead unpredictability can be directly treated as an error.

4 Design

By recursively restricting choice from lower layers, we converge on a design where the most performance-critical execution choices are made in the topmost layer. In the context of a model serving service, this process converges to an architecture, which we call Clockwork, with a centralized controller and workers with predictable performance.

4.1 Overview

Architecture. Fig. 3 illustrates Clockwork’s architecture. Users submit inference requests (1) which are queued centrally on Clockwork’s controller. Each worker has a set of DNN models (★) loaded into RAM and maintains exclusive control over one or more GPUs. The centralized scheduler has a global view of system state, including all workers, and decides when to execute each request (2). To execute a request, the scheduler explicitly decides when to load models into GPU memory (3) and when to execute requests on the GPU (4). At any time, the scheduler makes accurate, high-quality caching, scheduling, and load balancing decisions. The controller can perform these actions proactively because execution on workers is highly predictable. The controller transmits continual scheduling information to the workers that, by design, will execute schedules exactly as directed.

Illustrative example. To elucidate the Clockwork architectural components with more detail, including the choices that were consigned to the controller, consider the key steps for serving the inference requests illustrated in Fig. 4.

① Upon receiving an inference request r_1 for model ★, the controller is aware that a target worker has yet to copy the model weights from RAM into GPU memory. It estimates the time required to load the model weights (LOAD), plus the time to subsequently execute the inference (INFER), and concludes that the request will complete within its specified SLO. The controller instructs the worker to copy the model weights to

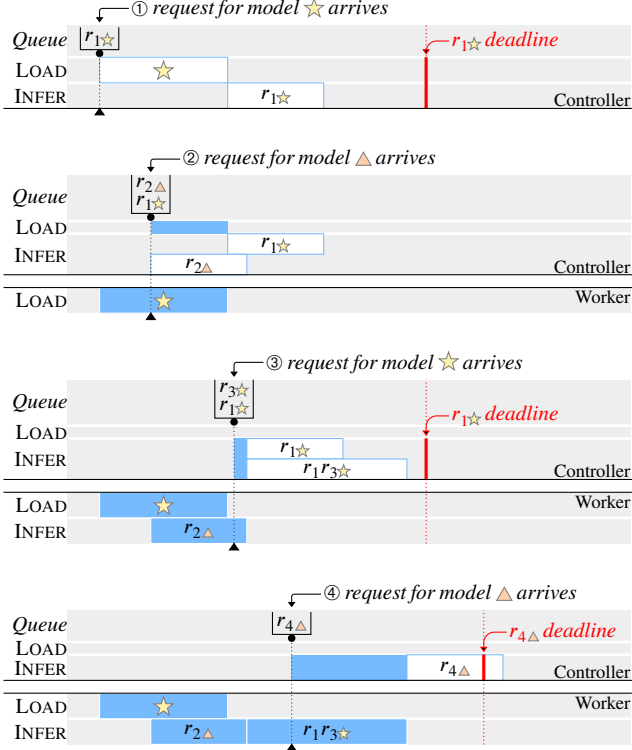


Fig. 4: Timeline of four illustrative inference requests.

GPU memory via a LOAD action. Since the controller is aware of all timings, it does not yet need to submit the subsequent INFER action until the LOAD has completed.

② While $\star$ is loading, a request r_2 for model $\triangle$ arrives. The controller is aware that, unlike $\star$, $\triangle$ is already loaded into GPU memory. The controller can choose to either INFER r_2 immediately, or wait for $\star$ to complete loading then INFER r_1 . Since the worker would be otherwise idle, the controller instructs the worker to execute the inference for r_2 immediately via an INFER action.

③ Clockwork workers only execute one INFER action and one LOAD action at a time, so the controller can wait until r_2 has nearly completed before submitting an INFER action for r_1 . In the meantime, another request r_3 for model $\star$ arrives. This gives the controller a choice between INFER for r_1 by itself, or to batch r_1 and r_3 . Batched execution is more efficient, but takes longer. In this case a batched INFER action will still complete before r_1 's deadline, so the controller instructs the worker to batch the inferences for r_1 and r_3 .

④ While r_1 and r_3 execute, a request r_4 for $\triangle$ arrives with a tight SLO. The controller is aware that r_4 will miss its deadline, even if it executes immediately after the worker becomes free. The controller does not proceed to schedule an INFER action, and cancels the request before performing any fruitless work.

Each step of the above execution is fast, *e.g.* for ResNet50, LOAD and INFER take approximately 8 ms and 3 ms respectively. Table 1 outlines representative measurements for 8 of the 61 models used for Clockwork experiments.

4.2 Consolidating Choice

Our design consolidates choice in three main ways. First, changes in the worker's state, for instance evicting a DNN from GPU memory, can influence the performance for future requests in a way that makes performance estimation complex. We therefore require that no worker operation should have implicit performance side-effects on any future operation. Second, we must ensure that a predictable component either delegates scheduling decisions that may impact performance to the centralized controller, or otherwise makes schedules deterministic. Third, when a predictable component is unable to execute a schedule as instructed, it is treated as an error to enable workers to get back on schedule. Workers do not attempt best-effort remediation, so as to avoid a cascade of mispredictions.

We enforce these three properties in Clockwork through an action command abstraction between the controller and workers that, in lieu of traditional RPC calls, either communicates a change in a worker's state or a task for a worker to execute. Each action the controller issues to a worker, such as LOAD and INFER, has predicted execution time and a designated execution window. These are derived using the known state of the worker, previously submitted actions, and known transitions in controller-maintained worker state.

4.3 Challenges for Predictable Inference

To consolidate choice we must first identify where performance-critical choices arise in system components. We have established that DNN inference itself on a GPU has deterministic performance; we next study the challenges in extending this result to a full-fledged inference system.

Managed memory and caches can be unpredictable (C1). RAM and GPU memory on a worker constitute state that impacts the performance of future requests. Additionally, some memory allocators exhibit variable timing for allocation and deallocation requests due to internal trade-offs between memory fragmentation and amortized performance. Memory that is used as a cache specifically introduces performance variability between cache hits and misses, with an internal cache replacement policy influencing performance of future items. To maintain predictability, we must instead consolidate choice by managing cache admission and eviction for each worker at the central controller. Fortunately, caching of DNN weights is coarse-grained and per-model.

Hardware interactions can be unpredictable (C2). Many system resources are implicitly administered by hardware schedulers that operate at very fine time-scales and produce different schedules under even minute shifts in the arrival times of other requests. The volatility of timing coupled with proprietary and un-documented scheduling policies make it onerous to accurately predict completion times for concurrent requests. The remedy for non-determinism is to strip away the ability for schedulers to reorder requests by forcing only a single request to be executed at a time, at the cost of spending

Model Family	Model	IO Size (kB)		Weights		GPU Execution Latency (ms)				
		Input	Output	Size (MB)	Transfer (ms)	B1	B2	B4	B8	B16
DenseNet [36]	densenet169	602	4	56.5	4.50	5.18	6.29	8.57	12.82	21.85
Inception v3 [68]	inceptionv3	1073	4	95.3	7.77	4.46	6.85	10.99	16.45	26.17
Mobile Pose [72]	mobile_pose_mobilenetv3	590	209	19.0	1.55	1.29	1.92	3.13	5.71	11.62
ResNet [30]	resnet18	602	4	46.7	3.81	1.27	1.86	2.73	4.06	7.02
	resnet50	602	4	102.3	8.33	2.61	3.78	5.61	9.13	15.67
	resnet152	602	4	240.9	19.58	7.71	11.14	16.21	26.48	44.60

Table 1: Measurements of a representative subset of the 61 models used for Clockwork experiments. Pre-trained models were sourced from the ONNX Model Zoo [53] and the GluonCV Model Zoo [28], and optimized for NVIDIA Tesla v100 GPUs using TVM v0.7 [15].

greater effort on keeping the resource fully utilized. Mercifully, one-at-a-time execution of DNN inferences on GPUs has closely comparable throughput to concurrent execution (Fig. 2b) and many classes of DNNs (*e.g.* convolutional neural networks) can saturate GPUs with small batch sizes.

External factors can trigger performance variance (C3).

Even after systematically removing the key internal sources of unpredictability by consolidating choice, there will always remain external sources outside of the controller’s purview. These include performance interference through shared network bottlenecks, thermal throttling of CPUs and GPUs, and others. The only option is to minimize their effects by building sufficient tolerance into the system.

4.4 Predictable DNN Worker

At a high-level, Clockwork workers maintain DNNs in memory and execute inference requests on one or more GPUs. The workers interface with the controller to receive actions.

Memory management. Model weights must be present in GPU memory to execute an inference. However, GPU memory capacity is small ($\leq 32\text{GB}$) relative to host memory ($\leq 4\text{TB}$), and host-to-GPU memory transfers ($\approx 8.3\text{ms}$ for ResNet50) typically take longer than running the DNN inference on the GPU ($\approx 2.9\text{ms}$). Consequently, Clockwork treats GPU memory as a cache, letting commonly or recently used models avoid expensive loads. To overcome C1, workers explicitly expose LOAD and UNLOAD actions to the controller for copying models to and removing models from worker’s GPU memory with deterministic latency. These actions also update the state that the controller tracks for the worker.

Inference execution. The controller only sends an INFER action when a model is present in GPU memory or a LOAD action will momentarily complete. The worker internally divides INFER actions into three steps. First, INPUT transfers the input vector from host to GPU memory. Next, EXEC performs the actual heavy-weight DNN GPU calculations, which dominate the total inference time. Finally, OUTPUT transfers the resulting output vector from the GPU back to host memory. These steps may coincide: the previous request’s outputs can be copied at the same time as the current request’s input is being transferred. However, multiple concurrent EXEC calls cause the GPU hardware scheduler to behave unpredictably (C2). Fortunately,

a DNN inference call by itself can efficiently utilize the GPU while also restricting the hardware scheduler to a single, predictable option (Fig. 2b). Clockwork workers therefore run a single EXEC at a time, a design choice that reduces performance variability by two orders of magnitude while only minimally decreasing inference throughput (Fig. 2b).

Interface with the controller. Clockwork workers receive LOAD, UNLOAD, and INFER actions from the controller with detailed timing expectations attached:

type	INFER, LOAD, or UNLOAD
earliest	the time when this action may begin executing
latest	when this action will be rejected

Rather than executing actions in a work-conserving, best-effort manner, workers strictly follow the schedule of actions imposed by the controller. The controller communicates two timestamps with every action, *earliest* and *latest*, to designate a time interval during which the worker may begin executing the action. Actions that cannot start within the prescribed window are cancelled and never executed. This allows workers to quickly get back on schedule after an individual action is delayed unexpectedly (C3) by skipping one or more actions, minimizing the impact of the delay on other actions. Workers communicate the result of each action back to the controller, including whether the command was successful and the measured execution time.

4.5 Central Controller

All decision-making in Clockwork occurs in the central controller. The controller receives inference requests from users and decides worker actions while striving to meet SLOs.

Modeling worker performance. The controller maintains a per-worker, per-model performance profile comprising processing time measurements of recent requests; profiles are updated continuously to tolerate shifts due to external factors (C3). The controller also tracks the outstanding actions and memory state at every worker. Since actions have inherently deterministic latency by design, the controller can deduce the earliest time that a worker could begin executing a new action (queuing time).

Action scheduler. The Clockwork controller proactively manages action schedules for workers. It utilizes a global view

of system requests, up-to-date worker performance profiles, and accurate predictions for when outstanding actions will complete. The controller attempts to pack worker schedules tightly by making narrow, realistic estimates for the `earliest` and `latest` time interval. The interval width balances a trade-off between Clockwork SLO fulfillment and system goodput. On one hand, making the interval too narrow increases the risk of an action not being executed by a worker because it could not be completed in time (C3), potentially triggering an SLO violation. On the other hand, underestimating the window length can create periods of inactivity and decrease worker utilization, thus affecting Clockwork goodput.

The scheduler lazily decides which worker should execute the inference. The controller only submits a minimal amount of work to keep workers utilized; it is in no hurry to commit because it can accurately predict action timings. Delaying choices on the controller improves schedules by providing more options, permitting the Clockwork controller to re-order and *batch* inference requests to the same model, significantly improving resource efficiency and throughput.

In our design, any worker can process any request since they all store every model in host memory; however, workers have different sets of models loaded into their GPU memory. A worker that executes only cold inferences must transfer weights for each model from host memory to the GPU and may saturate the available PCIe bandwidth, whereas a worker that executes only hot inferences may be bottlenecked by the GPU. The Clockwork scheduler balances load by mixing and matching hot and cold inferences among all workers.

5 Implementation

Clockwork’s implementation, comprising 26KLOC of C++, contains various decisions that enable Clockwork to consolidate choice on its controller.

5.1 Models

Predictable model execution. Prior model serving systems such as Clipper [16] and INFaaS [58] act as orchestration layers atop existing model execution frameworks such as TensorFlow [1] and TensorRT [50]. This decoupling makes it difficult to consolidate choice, since the model execution frameworks encapsulate scheduling and memory management decisions that we wish to make with Clockwork. Instead, Clockwork implements its own model runtime, reusing key components of the TVM optimizing compiler [15]. Clockwork’s model runtime enables fine-grained control over each stage of a model’s execution. For models provided to Clockwork (e.g. in ONNX form), we compile a binary representation using TVM and postprocess the model to produce the following:

- **Weights:** A model’s weights are a binary blob (10s to 100s of MB (cf. Table 1)).
- **Kernels:** The CUDA kernels that execute a model (10s to 100s of kB). These are not provided by the user; they are derived from the abstract model definition, and kernels

from different users can safely execute within the same process. Clockwork uses the kernels compiled by TVM. Clockwork compiles kernels for multiple configurable batch sizes; by default 1,2,4,8, and 16. Kernels for different batch sizes can use the same weights without modification.

- **Memory metadata:** At runtime, models do not directly allocate memory; instead, Clockwork will pre-allocate and manage all GPU memory and pass pointers as arguments to function calls. The memory requirements for a model are static, and Clockwork precalculates the required workspace memory and offsets required for each kernel.
- **Profiling data:** Clockwork runs a brief profiling step to produce a seed estimate for model execution times.

Model loading. Models are stored in an efficient serialized form on disk. Clockwork workers pre-load models from disk into main memory on worker startup. For the worker machines used in our evaluation, 768GB RAM can support thousands of models (cf. §6.5). Once a model is in main memory, Clockwork extracts and links the CUDA modules needed for its execution. To improve predictability, Clockwork disables JIT compilation and the caching of CUDA kernels.

5.2 DNN Workers

Each machine runs one worker process that receives and executes actions from Clockwork’s controller. We do not run Clockwork in a container or VM to avoid the performance interference such sharing can impose.

Managing model weights in memory. Clockwork pre-allocates all GPU memory and divides it into three categories:

- **Workspace:** Models require a variable amount of GPU memory for intermediate results. This memory is transient and only needed during execution; once an output has been produced, it is no longer needed. Clockwork only executes models one-at-a-time, so it allocates 512MB workspace memory.
- **IOCache:** Although Clockwork only executes models one-at-a-time, Clockwork asynchronously copies inputs to the GPU prior to execution, and outputs to host memory after execution. Clockwork allocates 512MB device memory for temporary storage of inputs and outputs before and after execution.
- **PageCache:** The remaining device memory is used for storing model weights, divided into 16MB *pages*. Multiple tensors can occupy the same 16MB page and the mapping of tensors to pages is determined statically at model-compile time. At runtime, page pointers are passed as kernel arguments and tensors are read from pre-defined offsets.

Clockwork’s PageCache has several advantages. First, avoiding repeated memory allocation calls leads to more predictable executions, since memory allocation can be an unpredictable source of overheads (C1). Second, paging *simplifies choice*: external memory fragmentation issues are

eliminated, and the controller need only track the number of total free pages to completely capture the worker’s memory state. Paging slightly increases memory utilization; however, model memory requirements are static and known ahead of time, and can be bucketed on to pages to reduce internal fragmentation. Paging does not affect the latency of memory transfers.

Actions. To orchestrate workers, the controller uses the previously described *action* abstraction. Actions contain a unique `id` and an action-dependent payload (*e.g.* INFER inputs). Each worker runs a dedicated *executor* for each action type and each worker-GPU. An executor runs a thread that dequeues actions chronologically by *earliest* timestamp, and waits until *earliest* is reached before proceeding with an action. Executors reject actions whose *latest* timestamp has passed. To reduce interference between threads and other processes, each executor is pinned to a dedicated core and runs at real-time priority. Both INFER and LOAD execute asynchronous work in their own CUDA streams. Each executor is bottlenecked by a different resource (*e.g.* GPU execution and PCIe transfers) and can run concurrently with negligible interference.

Results. A network thread maintains a persistent connection with the controller for receiving actions and sending results. A result comprises the following:

status	success or an error code
timing	start and end times, and on-device execution duration for any asynchronous work

LOAD actions acquire pages from the PageCache, then copy weights to those pages. If no pages are available then LOAD aborts. The controller explicitly frees pages with UNLOAD; this only updates in-memory metadata and always succeeds.

INFER actions comprise INPUT, EXEC and OUTPUT, each of which have dedicated executors. INPUT executes immediately on receipt of INFER; it acquires IO memory from the IOCache then copies inputs. EXEC inherits the INFER action’s *earliest* and *latest* timestamps; it checks weights and inputs are present then executes kernels on the GPU, using Workspace for intermediate calculations. OUTPUT immediately copies outputs back to main memory then releases the IO memory. To simplify controller decision making, INPUT and OUTPUT are not exposed as actions since they are orders of magnitude faster than EXEC and LOAD (10s of microseconds) for our workloads. Clockwork’s memory management allows for back-to-back INFER actions for the same model.

5.3 Central Controller

On startup, Clockwork’s controller establishes persistent connections to all workers and exchanges metadata about the size of each worker’s PageCache, the models present on each worker, and their initial pre-profiled execution times. The core duty of the controller is to satisfy requests received from clients by submitting actions to workers. This decision making is encapsulated in the *Scheduler* interface:

onRequest	client request received, specifying a model ID, SLO, and providing inference inputs
onResult	a result is received from a worker

A scheduler implements this interface, and can invoke `sendAction` to send an action to a worker, and `sendResponse` to respond to a client. A separate layer of the controller implements common tasks such as networking, forwarding inputs to workers, setting timestamps, and handling timeouts. This design concentrates all choice in a single place, and enables different scheduler implementations to be easily dropped in.

Managing worker state. The controller maintains an accurate representation of workers’ execution state, which is threefold: *memory state*, in which the scheduler tracks what models are present in the worker PageCaches and when LOAD will be required; *action profiles*, which are measurements of past 10 actions duration, stratified by model, worker, and batch size, to predict the duration of future action; and *pending actions*, which tracks submitted actions and estimates when each executor will next be available. Taken together, these enable the scheduler to accurately predict when candidate actions will complete, and avoid submitting work that cannot complete before the request’s deadline. Worker state is not a significant scalability bottleneck; action profiles require only 40 bytes for each model, worker and batch size combination.

Scheduling INFER. Upon arrival, requests are enqueued into per-model request queues. For each INFER executor, a new action must be scheduled whenever the executor has less than 5 ms of outstanding work. To schedule an INFER action, a model and batch size must be selected. The batch size can differ action-to-action, though the scheduler prioritizes larger batch sizes for efficiency.

At any point in time, a model will have zero or more queued requests. However, not every request is suitable for every batch size. Higher batch sizes take longer to execute, so a request close to its deadline might only be satisfiable using a small batch size. To handle this, each model has a request queue per batch size (we term this a *batch queue*). New requests are enqueued into *every* batch queue. Requests are dropped from batch queues when they cease to be satisfiable; *e.g.* a request in the batch size of 16 queue will be dropped sooner than it is dropped from the batch size of 8 queue.

To decide which model and batch size to schedule, we use *strategies*. A strategy specifies a *model*, a *latest* timestamp, and a *batch size*. Each INFER executor has a separate strategy queue, ordered by *latest*, containing only strategies for models it has loaded. The scheduler dequeues strategies until it finds one that is *valid*: *latest* has not elapsed, and the batch queue for the specified batch size has sufficient requests. If a strategy is valid, the scheduler will also speculatively increase the batch size as long as extra requests are available.

When a valid strategy is found, an INFER action is created and requests are dequeued to fill the batch. Old strategies for this model are removed from the strategy queue, and new

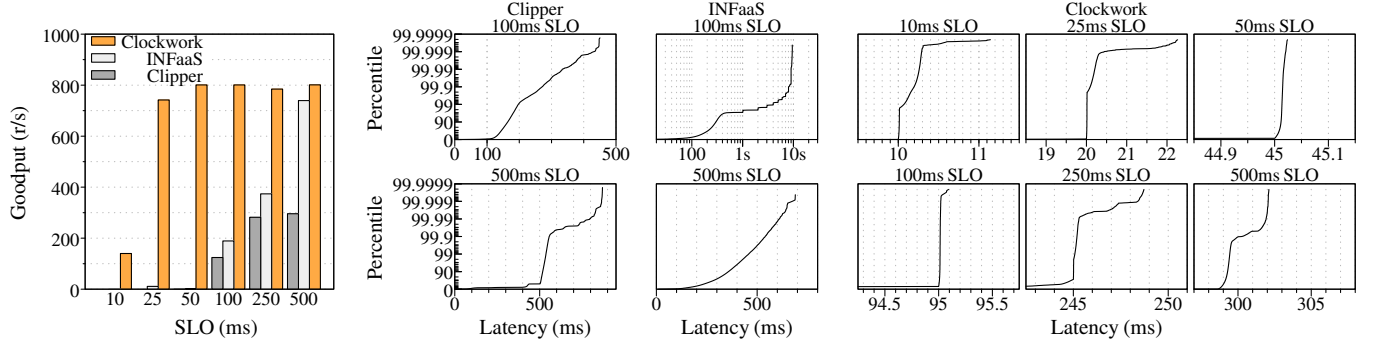


Fig. 5: Goodput and latency measurements for Clipper, INFaaS, and Clockwork. We deploy 15 instances of ResNet50 on 1 worker; each model submits 16 concurrent requests in a closed loop. (Left) Request goodput. Goodput only counts requests that succeed within the SLO. (Right) Request latency CDFs across all requests (including those rejected due to missed deadlines). Latency CDFs are scaled to highlight tail latency.

strategies are then created and enqueued. A strategy is created per batch queue; *latest* is calculated by subtracting the batch execution time from the deadline of the request at the head of the queue. Empty batch queues are skipped.

Scheduling LOAD. Each LOAD executor also schedules up to 5 ms of outstanding work. For a LOAD executor, the scheduler selects a model by estimating each model’s SLO violations given the model’s current state and outstanding requests. To do this efficiently, the scheduler maintains and incrementally updates *load* and *demand* statistics for models and GPUs:

- d_m the total demand for each model m
- $a_{m,g}$ the demand allocation of model m on GPU g .
- $\ell_g = \sum_m a_{m,g}$ the total load on each GPU g

A model’s total demand d_m is the total estimated execution time of m ’s outstanding requests; we update d_m when requests for that model arrive and complete. The demand allocations $a_{m,g}$ for m on GPU g are also updated when requests arrive and complete; they are calculated such that $\sum_g a_{m,g} = d_m$. Demand allocations are 0 for GPUs where the model is not loaded. On GPUs where the model is loaded, demand allocations are inversely proportional to the GPU’s load, since overloaded GPUs will be able to execute proportionally less of the total demand. Each GPU’s total load ℓ_g is the sum of its allocations across all models. With these estimates, each model’s load priority is defined as

$$p_{m,g} = d_m - \sum_g a_{m,g} \cdot \frac{\text{capacity}_g}{\ell_g}.$$

A model’s load priority estimates its unfulfilled work. For example, a model that is not loaded on any GPUs has priority equal to its outstanding work; a model loaded on a GPU that sits mostly idle has negative priority since the GPU can serve more work than the model demands.

Clockwork does not attempt to converge to a perfect demand allocation each time the system’s state changes. Rather, Clockwork incrementally updates each model’s demand allocation and load priority (i) when new requests arrive for that model; (ii) when an INFER is initiated for that model; (iii) when LOAD and UNLOAD affect a model; and (iv) when a request crosses the point where it can benefit from LOAD before its deadline.

The scheduler selects LOAD actions by choosing the highest priority model that is not already loaded. Notably, models with negative priority need not be loaded since their demands are already met. Clockwork uses a least-recently-used (LRU) eviction policy when selecting models to UNLOAD.

6 Evaluation

We next assess Clockwork’s ability to reliably serve DNNs under a variety of workload conditions. We begin our experimental evaluation with simple workloads in controlled settings, before expanding to heterogeneous models and diverse workloads. Our evaluation shows that Clockwork’s assumptions about predictability hold, and result in a system that can effectively meet SLOs and drastically reduce tail latency.

Experimental setup. We deploy Clockwork in a private cluster of 12 Dell PowerEdge R740 Servers. Each server has 32 cores, 768 GB RAM, and 2xNVIDIA Tesla v100 GPUS with 32 GB memory. The servers are connected by 2x10 Gbps Ethernet on a shared network. In all experiments, we run the controller, clients, and workers on separate machines.

6.1 How Does Clockwork Compare?

We begin with a comparison to two prior model serving systems, Clipper [16] and INFaaS [58]. For Clipper and Clockwork, we provision a single cluster machine to use 1 GPU to serve 15 separate copies of ResNet50. ResNet50 is the *de facto* model used for comparison previously by these systems; we chose 15 models as this reached the memory limit of Clipper¹. To evaluate INFaaS, we deployed an m5.24xlarge and a p3.2xlarge EC2 instance as the master and the worker, respectively. These are not identical experiment conditions; however, INFaaS is tightly integrated with EC2, and could not be deployed on our cluster infrastructure. We include these results for qualitative comparison.

Offered load. For each model, we run 16 closed-loop clients². The serving systems may batch requests for the same model instance, but requests to different instances cannot be

¹INFaaS memory limits were reached at 64 models

²Open-loop clients yielded similar results

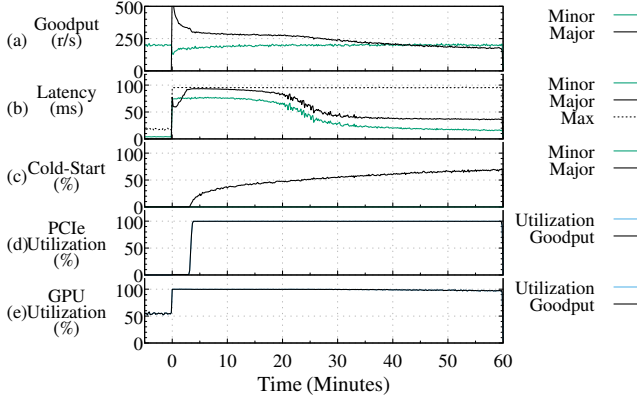


Fig. 6: Clockwork can serve thousands of models from a single worker. From $t = 0$, the Major workload adds an additional model per second, to a total of 3,600 models at $t = 60$ (cf. §6.2.).

batched. We run multiple experiments, varying the target SLO from 10 ms to 500 ms.

Goodput. Fig. 5 plots the *goodput* achieved by each system as the target SLO varies from 10 ms to 500 ms. Goodput is the number of successful requests that completed within the target SLO; it excludes timed out requests and requests that responded after the SLO.

With a high SLO of 500 ms, Clockwork and INFaaS meet their SLOs and have comparable goodput of approximately 800 r/s. Clipper’s goodput is substantially lower, as Clipper only treats SLOs as an average latency target, not a strict threshold, and converges to this target over time without bounding latency variability. As SLOs tighten, goodput and tail latency deteriorate for both Clipper and INFaaS, and their goodput collapses below a 100 ms SLO. Like Clipper, INFaaS uses the SLO as a coarse-grained goal for reactive policies. Consequently, only Clockwork can continue serving SLOs below 100 ms.

Fig. 5 also plots latency CDFs for Clipper, INFaaS, and Clockwork. We scale the CDFs to emphasize tail latency. The figure illustrates how both Clipper and INFaaS allow latency higher than their SLOs. However, of note, with a 500 ms SLO, INFaaS successfully finds a configuration that can serve this SLO, and meets its SLO for 99% of its requests. By comparison, Clockwork’s tail latency remains very close to the SLO in all cases. For the 500 ms SLO, Clockwork’s latency remains at ≈ 300 ms because it schedules each model’s entire batch of 16 requests at a time, round-robin across models. With 15 models and a 20 ms batch-16 execution duration, Clockwork does not exceed the optimal 300 ms latency.

6.2 Can Clockwork Serve Thousands?

The previous experiment represented an idealized scenario, with only a small number of models, each with a steady sustained workload. We now examine the serving limits of a single worker. We deploy 3,601 copies of ResNet50 to a worker, and set a 100 ms SLO. We submit two workloads: a Major workload and a Minor workload. The Major workload

comprises 3,600 model instances; we vary the number of instances that are active at any point in time, and evenly distribute a workload of 1,000 r/s across all active models. The Minor workload is a single model instance that maintains a fixed 200 r/s request rate throughout the experiment.

Figure Fig. 6 (a) plots the goodput achieved by the major and minor workloads. From $t = -5$ to $t = 0$ (we denote t in minutes) only the Minor workload is present, achieving its full 200 r/s. At $t = 0$, we activate one model instance of the Major workload; the addition of 1000 r/s fully saturates the GPU (e). After that, we activate an additional model of the Major workload every 1 second. As more model instances become active, the Major workload’s goodput drops since each additional model forgoes batching opportunities. At $t = 60$ all 3,600 models are active, each submitting approximately 0.28 r/s.

By $t = 3.5$, 201 models have been activated, reaching the capacity of GPU device memory. To continue serving requests, Clockwork begins swapping models on and off GPU; Fig. 6 (d) shows PCIe utilization rapidly rises to 100%. As more models activate, an increasing number of requests in the Major workload find that their model is not loaded; Fig. 6 (c) plots the rise in *cold-starts*, reaching 70% by the end of the experiment. The minor workload, with its sustained request rate of 200 r/s, does not experience any cold starts because its demand dwarfs every other model after the first 5 seconds. As the number of cold-starts increases, the demand on GPU execution decreases, enabling the Minor workload’s goodput to gradually grow back to 200 r/s. At approximately $t = 20$, the bottleneck for the Major workload shifts to PCIe utilization, enabling the Minor workload’s latency to drop back to an average of 20 ms (b).

This experiment illustrates how bottlenecks in Clockwork can shift as workload demand changes. Clockwork can deal with shifting bottlenecks even while serving a large number of models. As illustrated in Fig. 6 (b), the maximum request latency across the experiment did not exceed the 100 ms SLO.

6.3 How Low Can Clockwork Go?

Clockwork’s predictability and centralized decision-making enables it to satisfy low-latency SLOs. In this experiment, we use six Clockwork workers and evaluate the lower limit on SLOs that Clockwork can achieve by measuring the proportion of successful requests while varying the SLO. We repeat the experiment for six different workloads, varying the number of ResNet50 instances ($N = 12$ or 48) and cumulative request rate ($R = 600$ r/s, 1200 r/s, or 2400 r/s). For each experiment run, we begin with an SLO of 2.9 ms ($1 \times$ the execution latency of batch-1 ResNet50 inference). Every 30 seconds, we extend the SLO by 50%; by the end of the experiment the SLO reaches 74ms. We run a separate open-loop client for each model with a Poisson inter-arrival time distribution, and as before, all models are independent (requests cannot be batched across models).

Workload satisfaction. Fig. 7 plots the *workload satisfaction* for each experiment run. Workload satisfaction is the ratio of goodput to offered load. A workload satisfaction of

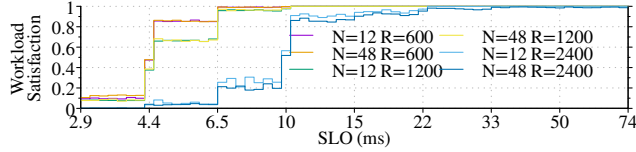


Fig. 7: Workload satisfaction rates as we vary N , the number of clients, and R , the request rate.

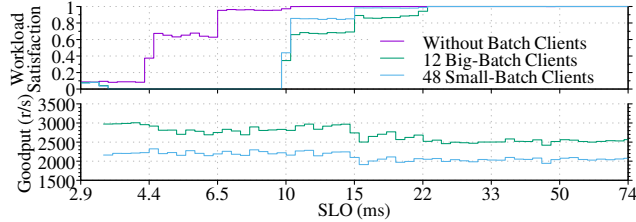


Fig. 8: Workload satisfaction rates for latency-sensitive clients (top) and workload goodput for batch clients (bottom).

1 means all requests received a successful response within their SLO. For a load of $R=600$ r/s and 1200 r/s, irrespective of the number of models, Clockwork successfully satisfied tight SLOs 10 and 22 ms. Even at $R=2400$ r/s, Clockwork comfortably managed an SLO of 74ms.

6.4 Can Clockwork Isolate Performance?

Clockwork can satisfy tight SLOs for latency-sensitive clients in isolation; we next consider when the system is shared with other users serving batch requests without latency SLOs. As before, we use six Clockwork workers, and all clients use instances of ResNet50. We provision six *latency-sensitive* clients, each submitting a 200 r/s open-loop workload. We also provision several *batch clients*, which submit sustained closed-loop workloads and do not have latency SLOs. *Big-batch* clients have a concurrency of 16, while *small-batch* clients have a concurrency of 4. Varying the concurrency affects the maximum batch size Clockwork can achieve for batch client requests. We considered three scenarios: (a) baseline without batch clients; (b) 12 big-batch clients; and (c) 48 small-batch clients

Fig. 8 illustrates the workload satisfaction rates for latency-sensitive clients and the total goodput achieved for the batch clients. Clockwork successfully prioritizes latency-sensitive requests over batch requests. Through SLO-aware scheduling, it ensures that the workload satisfaction rates are unaffected by the presence of other pending, less time-critical requests. At the same time, Clockwork does not throttle batch requests entirely, but schedules them during idle times or expected idle times. However, when the SLOs are too tight (<15 ms), many latency-sensitive requests are rejected in advance, allowing pending batch requests to pass through.

6.5 Are Realistic Workloads Predictable?

We now ask whether executions remain predictable under realistic workloads that comprise many concurrent users and models. We also investigate whether Clockwork effectively

Model Family	Count	Model Variants
DenseNet [36]	4	121, 161, 198, 201
DLA [75]	1	34
GoogLeNet [67]	1	
Inception [68]	1	v3
Xception [68]	1	
MobilePose [33]	4	SPRN18, MNv3, RN18, RN50
ResNeSt [78]	4	14, 26, 40, 101
ResNet [30]	22	18, 18b, 34, 34b, 50, 50b, 50c, 50d, 50s, 50-1.8x, 101, 101b, 101c, 101d, 101s, 101-1.9x, 101-2.2x, 152, 152b, 152c, 152d, 152s
ResNet-v2 [31]	5	18, 34, 50, 101, 152
ResNeXt [73]	3	50-32, 101-32, 101-64
SENet [35]	2	50-32, 101-32
TSN [70]	7	iv1, iv3, r18, r34, r50, r101, r152
Wide ResNet [76]	3	16-10, 28-10, 40-8
Winograd [45]	3	RN18, RN50, RN101

Table 2: List of models used in experiments.

exploits this predictability.

To answer these questions, we deploy Clockwork on 12 workers and replay a workload trace of Microsoft Azure Functions (MAF) [61]. The trace records approximately 46,000 function workloads, counting the number of invocations of each function, every minute, for two weeks. It interleaves a wide range of workloads, including heavy sustained workloads, low utilization cold workloads, bursty workloads that fluctuate over time, and workloads with periodic spikes [61]. We believe this to be a representative workload for evaluation since serverless platforms enable a wide range of applications and supporting ML inference on serverless is an active area of research [10, 39].

In this experiment, we replay six hours of the MAF trace in real-time. We use 61 different models (Table 2) taken from the ONNX Model Zoo [53] and the GlueCV Model Zoo [28]. We duplicate each model 66 times, resulting in a total of 4,026 instances and reaching the main-memory capacity of our worker machines. We replay ten or eleven function workloads for each model instance. We configure Clockwork with a 100 ms SLO.

Clockwork with realistic workloads. The time series in Fig. 9 (a) shows the offered load and goodput achieved across all models. For the 6 hour experiment, both the offered load and goodput averaged 9,638 r/s – out of a total of 208 million requests, only 58 failed due to action timing mispredictions, and no requests timed out. All GPUs were fully utilized throughout the experiment, yet no request exceeded the 100ms SLO.

Fig. 9 (b) plots the median, 99th percentile, and maximum request latency over the course of the experiment. Latency spikes occur every 5, 15, and 60 minutes, due to the presence of numerous periodic workloads within the trace [61]. Workload spikes do not cause SLO violations because of latency headroom; Fig. 9 (c) shows the average batch size for the experiment, and with each workload spike, Clockwork can schedule larger batches, with higher latency. To evaluate the cold-start behavior of this workload, we categorize a request as a cold-start if its model is not already loaded into GPU’s memory before arrival. For each 1-minute interval, Fig. 9 (d) counts the number of

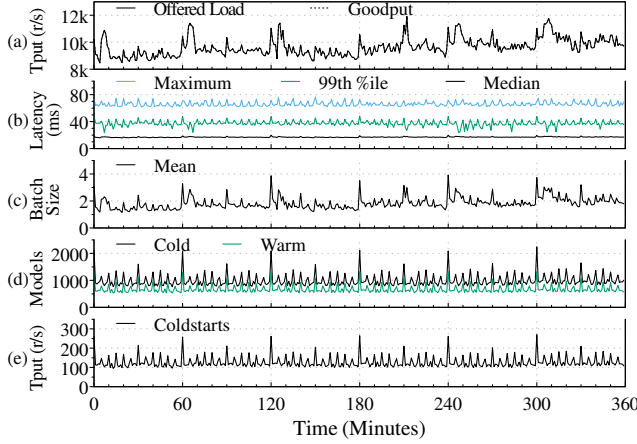


Fig. 9: Microsoft Azure Functions (MAF) over Clockwork; see §6.5 for a description.

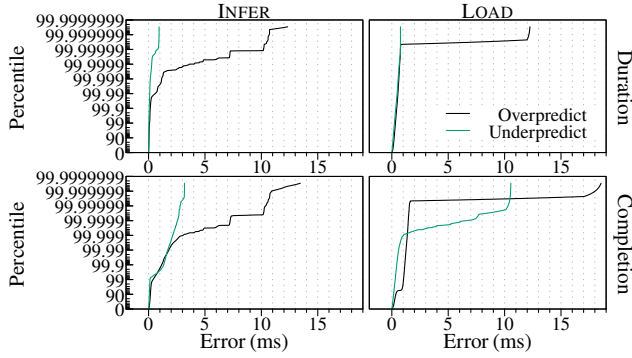


Fig. 10: Clockwork prediction and completion errors for MAF trace.

unique models that have at least one cold-start, and at least one warm-start. On average, 987 unique models perform cold-starts each minute; or approximately 25% of all models. However, while many models perform cold-starts, they only represent a small fraction of all requests. Fig. 9 (e) plots the throughput of cold-start requests, averaging 126 r/s, or 1.3% of all requests. These results show that Clockwork can sustain significant load for varied, realistic workloads comprising thousands of models.

Predictable executions. Clockwork’s scheduler relies on accurate predictions of action latency, so to assess Clockwork’s underlying assumptions of predictability, we next evaluate the accuracy of Clockwork’s predictions. We measure the latency of INFER and LOAD actions on Clockwork’s workers and compare it to the time estimated by Clockwork’s controller to derive a *prediction error*. Prediction errors comprise two types: *overprediction*, when the real execution latency is faster than predicted; and *underprediction*, when the real execution latency is slower than predicted. Consistent overpredictions can lead to idle resources, while consistent underpredictions can cause SLO violations. Fig. 10 (top) plots the prediction errors for INFER and LOAD actions. For INFER actions, the 99th percentile of overpredictions and underpredictions is 144 μ s and 55 μ s, respectively. Thereafter, the tail latency grows

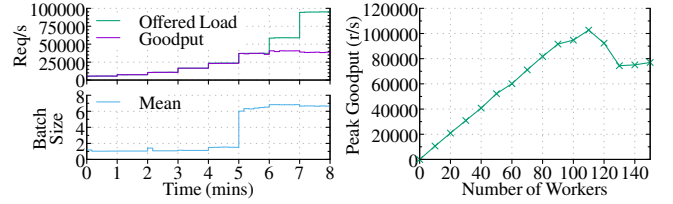


Fig. 11: (Left) With 40 emulated workers, goodput is approximately equal to offered load; peak goodput is achieved at appx. 40,000 r/s, when all workers are fully utilized. (Right) Peak goodput achieved with different numbers of emulated workers.

to exceed 10 ms in a few extremely rare cases. Clockwork consistently overpredicts more than it underpredicts, as it uses a rolling 99th percentile measurement to make its predictions. For LOAD actions, the 99th percentile of overpredictions and underpredictions is 431 μ s and 348 μ s, respectively.

Fig. 10 (bottom) plots the *completion time error*. Clockwork must accurately predict when a given action will complete, taking into account any previously submitted actions (*i.e.* queuing time). Individual prediction errors can compound, leading to increased completion time error. For INFER actions, the error compounds 4 $\times$, with a 99th percentile completion error of $\approx$ 1 ms. In extreme cases, Clockwork’s completion error also grows to more than 10 ms. However, the completion error does not substantially exceed the action duration error, implying that for Clockwork, erroneous predictions of outliers are statistically independent.

6.6 Can Clockwork Scale?

Centralized scheduling presents a potential scalability bottleneck, though prior work has demonstrated that centralized schedulers can reach impressive scale [24, 57]. Our final experiment examines the scalability of Clockwork’s controller.

To venture beyond the capacity of our testbed, we leverage a specially-developed *emulated worker* that implements Clockwork’s action interface. The emulated worker behaves identically to a bona fide Clockwork worker, except the LOAD and INFER actions perform no meaningful work; instead, they wait for a period of time according to the pre-profiled model measurements before returning a response. The emulated worker is indistinguishable from a real worker from the vantage point of Clockwork’s controller. To bypass the limited network capacity of our testbed, we modified our clients to send zero-length inputs (network is not a fundamental limitation; see §7 for discussion).

We measure the peak goodput achieved as we vary N , the number of emulated workers. We run multiple experiments, each with a different value of N , from 10 to 150 in increments of 10. We use the same models as described in §6.5, and a similar workload. Instead of replaying the trace at a fixed rate, we scale the trace and gradually offer more load in 60-second intervals. Fig. 11 (Left) illustrates one experiment run with $N=40$. Goodput follows the offered load almost perfectly up to about 40,000 r/s, at which point all workers are fully utilized and the goodput saturates.

Fig. 11 (Right) reports the peak goodput achieved with different numbers of workers. We report the median values across three experiment repetitions. The figure shows a linear increase in the peak goodput as the number of workers increases. Below $N=110$, goodput is limited by workers reaching full utilization. At $N=110$, we reach a maximum goodput of 103,387 r/s. At this point worker utilization stops being the limiting factor; instead, the bottleneck shifts to Clockwork’s controller. Beyond $N=110$ peak goodput declines.

6.7 Summary

In comparison with prior model serving systems, Clockwork achieves superior goodput, serves considerably more models concurrently, and violates substantially fewer SLOs. Owing to a lack of performance variability, Clockwork can achieve much tighter latency SLOs without sacrificing tail latency. Clockwork’s underlying assumptions about predictable executions bear out in reality: by consolidating choice, a predictable system that substantially curtails tail latency can be built.

Clockwork extends to a diverse range of workload conditions not supported by prior systems, including supporting thousands of models concurrently per GPU. Slow cold starts can run alongside high-throughput workloads without interference. Under all workload conditions, including cold starts and even under overload, Clockwork meets most SLOs without degrading service, and maintains close to maximal possible goodput. Finally, Clockwork isolates users of different models, enabling low-latency workloads to share the same system with background batch workloads.

7 Discussion

Why consolidate choice? Philosophically, the encapsulation, abstraction, and loose coupling of components are essential design practices while the building blocks and use cases of large systems are still in flux. Over time, the true use cases for the system settle and the entire system may in turn be replaced by a simpler, refined system that avoids the over-engineering and generality of its constituent parts—components that transpired to either be unnecessary in practice or to impede the common use case of the system. The squashing of layers through such specialization, effectively transforming systems into abstract units, can counteract the infamous bloat of modern software stacks. We designed Clockwork to be such an abstract unit for model serving systems.

Machine learning. Clockwork focuses on DNN inference, and excludes data preprocessing and postprocessing steps that are user-defined and CPU-bound. Safely and predictably executing these in Clockwork is a current research topic.

Individual DNN inferences are the atomic unit of work for Clockwork. Increasingly, modern ML applications are composed of pipelines or cascades of DNNs [34, 43, 62]. For these applications, performance predictability is strongly desired. We believe there are opportunities to leverage Clockwork’s

properties and perform more sophisticated pipeline scheduling that provides end-to-end guarantees. Similarly, performance predictability can influence system designs in other areas, such as large language embedding models [11] that may require dedicated or distributed accelerators. Expanding Clockwork into other ML paradigms, such as deep reinforcement learning and DNN training, raises philosophical questions about the nature and limits of predictability.

Inference accelerators. The Clockwork approach generalizes readily beyond GPUs to other inference-specific hardware accelerators [48], whose performance is arguably even more predictable. TPUs [41], for instance, are explicitly built around the idea of delegating control to software, while also eschewing general purpose processing engines with flexible control logic and generic memory hierarchies in favor of high-level operations and explicit memory hierarchies.

On the other extreme, inferences can also be executed in software on the CPU. While many models are heavily parallel in nature and execute orders of magnitude slower on CPUs, there are other models where execution on CPU is acceptable. One such example are recurrent neural networks (RNNs) which are fundamentally more sequential and often cannot effectively leverage the available parallelism on GPUs or other accelerators.

Limitations of predictability. Consolidating choice is only possible when you have control of, or guarantees about, the system’s major bottleneck resources. For example, Clockwork assumes workers have exclusive control over their machine, and dedicated GPUs. Clockwork does not assume exclusive control over the network, but does assume that the network has mostly-predictable latency between the controller and workers. In a shared setting, preserving predictability becomes more challenging – though not impossible – and this is an active area of research due to a general need to co-locate latency critical datacenter services [42, 47].

Network. Clockwork does not explicitly consider the network in its scheduling decisions; the occasional network latency spikes of dozens of ms during our experiments had negligible impact on our results. Our prototype routes all inputs and outputs through the central controller which will become a bottleneck at scale. We were able to reach the limits of our testbed network with 12 workers and a sustained, single-model workload; to test beyond this we disabled inputs as described in §6.6. This limitation is not fundamental; Clockwork’s controller only requires request *metadata* to schedule requests, and we are working to remove this limitation with a tier of load balancers.

Security. Security is important for all multi-user systems, since there are no container or hypervisor boundaries separating the workloads of different users. Clockwork does not explicitly address security; however, Clockwork does not execute arbitrary user code. Users must submit models in an abstract format that we then compile to binary code under the covers. Clockwork’s threat model resembles shared storage or database

systems, where system correctness is the chief concern; we have not verified any safety properties of Clockwork.

Fault tolerance. While Clockwork is a distributed system, we do not address the challenges of tolerating failures when serving models at large scale. This will require implementing a fault-tolerant centralized scheduler; however, we note that Clockwork’s predictable worker design will make pernicious phenomena like grey failure [27, 37] far easier to detect.

Other benefits of predictability. Concentrating choice makes it easier to implement other guarantees, such as SLOs related to burstiness or per-request cost. The Azure trace in our evaluation, for instance, contained regular, periodic spikes; exploiting advanced knowledge is an appealing future avenue for Clockwork. A further benefit of predictable system components is *performance clarity* [55]: performance bottlenecks and upcoming tasks in Clockwork are easy to reason about. Clockwork’s controller also provides a central point for *explanation*, since the controller has complete visibility of the expected and actual request behavior.

8 Related Work

Model serving. We directly compared Clockwork to Clipper [16] and INFaaS [58] in §6.1; here we provide additional comments. Both Clipper and INFaaS are designed as wrappers around existing model execution frameworks: Clipper, in order to provide a unifying abstraction; INFaaS, in order to exploit heterogeneous execution strategies. Being agnostic to the underlying execution engine sacrifices predictability and control over model execution. Both systems treat latency SLOs as long-term, reactive targets; by contrast, Clockwork is explicitly designed to consolidate choice, and exploit predictability by making proactive decisions. Clipper and INFaaS propose several orthogonal concepts that are compatible with Clockwork. Clipper’s model selection layer could be superimposed on Clockwork. INFaaS’s model variant concept could be integrated into Clockwork; we found similar predictability properties held for DNNs executing on dedicated CPU cores.

Several other projects investigate model serving in virtualized cloud environments and on serverless platforms, where predictability is in the hands of the cloud provider [10, 44, 77]. Like INFaaS, these model throughput, latency, and accuracy together for optimal model selection, but, unlike Clockwork, they do not use the backend predictability and latency SLOs for making proactive scheduling decisions. In industry, TFS² [51] is a proprietary model hosting service at Google, about which public information is not available. Amazon SageMaker [59] and Google AI Platform [26] are public cloud DNN serving systems with a similar interface to Clockwork: upload your model, then make inference requests. Both use containers under the covers as an isolation mechanism, and users suffer the associated cold-start latency. Beyond these details, further design information is not publicly known.

Real-time systems. Performance predictability, especially temporal safety, is also an important concern for safety-critical real-time systems. However in general, real-time systems are designed for *periodic* or *sporadic* workloads [8] with known minimum inter-arrival times and worst-case execution times, or for scenarios where the set of all inference requests is known in advance [66]. *Soft-real-time* systems [12] consider weaker notions of timeliness similar to the latency SLOs considered in this paper, but mainly target periodic or sporadic workloads. Clockwork, in contrast, makes no *a priori* assumptions about its workloads. Prior real-time systems work has also proposed mechanisms to tame the unpredictability inside GPUs [6, 9, 20, 22, 54]. Elliott and Anderson [21], for example, proposed interrupt handling mechanisms to circumvent the proprietary GPU drivers that ignore scheduling priorities, while Yang *et al.* [74] suggested avoiding synchronization anomalies through more careful use of CUDA synchronization primitives. These mechanisms are designed to facilitate an *a priori schedulability analysis*—mathematically bounding the blocking delays due to contention. Such bounds are orthogonal to Clockwork, which does not require strict worst-case guarantees.

9 Conclusion

As DNN inferences become increasingly central to interactive applications, the requirements for fast response tighten, the volume of requests expands, and the number of models grows. Our model serving system, Clockwork, meets these challenges. Clockwork efficiently fulfills aggressive tail-latency SLOs while supporting thousands of DNN models with different workload characteristics concurrently on each GPU, and scaling out to additional worker machines for increased capacity. The system also successfully isolates models from performance interference caused by other models served on the same system. Our results derive from our design methodology of recursively ensuring all internal architecture components have predictable performance by concentrating all choices in the centralized controller. Notably, our approach required us to either circumvent canonical best-effort mechanisms or orchestrate them to become predictable, and illustrates how consolidating choice can be applied in practice to achieve predictable performance.

Acknowledgements

We thank our shepherd Junfeng Yang and the anonymous reviewers for their insightful feedback that helped improve our work. Our work was partially supported by NSF CAREER Grant #1553579.

References

- [1] Martín Abadi, Paul Barham, Jianmin Chen, Zhifeng Chen, Andy Davis, Jeffrey Dean, Matthieu Devin, Sanjay Ghemawat, Geoffrey Irving, Michael Isard, Kudlur Manjunath, Josh Levenberg, Rajat Monga, Sherry Moore, Derek G. Murray, Benoit Steiner, Paul Tucker, Vijay Vasudevan, Pete Warden, Martin Wicke, Yuan

- Yu, and Xiaoqiang Zheng. TensorFlow: A System for Large-Scale Machine Learning. In *Proceedings of the 12th USENIX Symposium on Operating Systems Design and Implementation (OSDI)*, 2016.
- [2] Atul Adya, Daniel Myers, Jon Howell, Jeremy Elson, Colin Meek, Vishesh Khemani, Stefan Fulger, Pan Gu, Lakshminath Bhuvanagiri, Jason Hunter, et al. Slicer: Auto-Sharding for Datacenter Applications. In *Proceedings of the 12th USENIX Symposium on Operating Systems Design and Implementation (OSDI)*, 2016.
- [3] Deepak Agarwal, Bo Long, Jonathan Traupman, Doris Xin, and Liang Zhang. LASER: A Scalable Response Prediction Platform for Online Advertising. In *Proceedings of the 7th ACM International Conference on Web Search and Data Mining (WSDM)*, 2014.
- [4] Saamer Akhshabi and Constantine Dovrolis. The Evolution of Layered Protocol Stacks leads to an Hourglass-Shaped Architecture. In *Proceedings of the 2011 Conference of the ACM Special Interest Group on Data Communication (SIGCOMM)*, 2011.
- [5] Allied Market Research. Global machine learning chip market to garner \$37.85 Billion by 2025, at 40.8% CAGR. <https://www.globenewswire.com/news-release/2020/02/18/1986370/0/en/Global-Machine-Learning-Chip-Market-to-Garner-37-85-Billion-by-2025-at-40-8-CAGR.html>, February 2020.
- [6] Tanya Amert, Nathan Otterness, Ming Yang, James H Anderson, and F Donelson Smith. Gpu scheduling on the NVIDIA TX2: Hidden details revealed. In *Proceedings of the 38th IEEE Real-Time Systems Symposium (RTSS)*, 2017.
- [7] Ganesh Ananthanarayanan, Ali Ghodsi, Scott Shenker, and Ion Stoica. Effective Straggler Mitigation: Attack of the Clones. In *Proceedings of the 10th USENIX Symposium on Networked Systems Design and Implementation (NSDI)*, 2013.
- [8] Theodore P Baker and Sanjoy K Baruah. Schedulability analysis of multiprocessor sporadic task systems. *Handbook of Real-Time and Embedded Systems*, pages 3–31, 2007.
- [9] Joshua Bakita, Nathan Otterness, James H Anderson, and F Donelson Smith. Scaling Up: The Validation of Empirically Derived Scheduling Rules on NVIDIA GPUs. In *14th Workshop on Operating Systems Platforms for Embedded Real-Time Applications (OSPert)*, 2018.
- [10] Anirban Bhattacharjee, Ajay Dev Chhokra, Zhuangwei Kang, Hongyang Sun, Aniruddha Gokhale, and Gabor Karsai. Barista: Efficient and Scalable Serverless Serving System for Deep Learning Prediction Services. In *Proceedings of the 7th IEEE International Conference on Cloud Engineering (IC2E)*, 2019.
- [11] Tom B Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. Language Models are Few-Shot Learners. *arXiv preprint arXiv:2005.14165*, 2020.
- [12] Giorgio Buttazzo, Giuseppe Lipari, Luca Abeni, and Marco Caccamo. *Soft Real-Time Systems: Predictability vs. Efficiency: Predictability Vs. Efficiency*. Springer Science & Business Media, 2005.
- [13] Giovanni Campagna, Rakesh Ramesh, Silei Xu, Michael Fischer, and Monica S Lam. Almond: The Architecture of an Open, Crowdsourced, Privacy-Preserving, Programmable Virtual Assistant. In *Proceedings of the 26th International World Wide Web Conference (WWW)*, 2017.
- [14] Wai Chee Yau. How Zendesk Serves TensorFlow Models in Production. <https://medium.com/zendesk-engineering/how-zendesk-serves-tensorflow-models-in-production-751ee22f0f4b>, February 2017.
- [15] Tianqi Chen, Thierry Moreau, Ziheng Jiang, Lianmin Zheng, Eddie Yan, Haichen Shen, Meghan Cowan, Leyuan Wang, Yuwei Hu, Luis Ceze, et al. TVM: An Automated End-to-End Optimizing Compiler for Deep Learning. In *Proceedings of the 13th USENIX Symposium on Operating Systems Design and Implementation (OSDI)*, 2018.
- [16] Daniel Crankshaw, Xin Wang, Guilio Zhou, Michael J Franklin, Joseph E Gonzalez, and Ion Stoica. Clipper: A Low-Latency Online Prediction Serving System. In *Proceedings of the 14th USENIX Symposium on Networked Systems Design and Implementation (NSDI)*, 2017.
- [17] Brian Dalessandro, Daizhuo Chen, Troy Raeder, Claudia Perlich, Melinda Han Williams, and Foster Provost. Scalable Hands-Free Transfer Learning for Online Advertising. In *Proceedings of the 20th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD)*, 2014.
- [18] Jeffrey Dean and Luiz André Barroso. The Tail at Scale. *Communications of the ACM*, 56(2):74–80, 2013.
- [19] Christina Delimitrou and Christos Kozyrakis. Amdahl’s Law for Tail Latency. *Communications of the ACM*, 61(8):65–72, 2018.

- [20] Glenn A Elliott and James H Anderson. Real-world Constraints of GPUs in Real-Time Systems. In *Proceedings of the 17th IEEE International Conference on Embedded and Real-Time Computing Systems and Applications (RTCSA)*, 2011.
- [21] Glenn A Elliott and James H Anderson. Robust real-time multiprocessor interrupt handling motivated by GPUs. In *Proceedings of the 24th Euromicro Conference on Real-Time Systems (ECRTS)*, 2012.
- [22] Glenn A Elliott and James H Anderson. An Optimal k-Exclusion Real-Time Locking Protocol Motivated by Multi-GPU Systems. *Real-Time Systems*, 49(2):140–170, 2013.
- [23] Anshul Gandhi, Mor Harchol-Balter, Ram Raghunathan, and Michael A Kozuch. Autoscale: Dynamic, robust capacity management for multi-tier data centers. *ACM Transactions on Computer Systems*, 30(4):1–26, 2012.
- [24] Ionel Gog, Malte Schwarzkopf, Adam Gleave, Robert NM Watson, and Steven Hand. Firmament: Fast, Centralized Cluster Scheduling at Scale. In *Proceedings of the 12th USENIX Symposium on Operating Systems Design and Implementation (OSDI)*, 2016.
- [25] Ian Goodfellow, Yoshua Bengio, and Aaron Courville. *Deep Learning*. MIT Press, 2016. <http://www.deeplearningbook.org>.
- [26] Google AI Platform. Retrieved May 2020 from <https://cloud.google.com/ai-platform/>, 2020.
- [27] Haryadi S Gunawi, Riza O Suminto, Russell Sears, Casey Golliher, Swaminathan Sundararaman, Xing Lin, Tim Emami, Weiguang Sheng, Nematollah Bidokhti, Caitie McCaffrey, et al. Fail-slow at scale: Evidence of hardware performance faults in large production systems. *ACM Transactions on Storage*, 14(3):1–26, 2018.
- [28] Jian Guo, He He, Tong He, Leonard Lausen, Mu Li, Haibin Lin, Xingjian Shi, Chenguang Wang, Junyuan Xie, Sheng Zha, et al. GluonCV and GluonNLP: Deep Learning in Computer Vision and Natural Language Processing. *Journal of Machine Learning Research*, 21(23):1–7, 2020.
- [29] Kim Hazelwood, Sarah Bird, David Brooks, Soumith Chintala, Utku Diril, Dmytro Dzhulgakov, Mohamed Fawzy, Bill Jia, Yangqing Jia, Aditya Kalro, et al. Applied Machine Learning at Facebook: A Datacenter Infrastructure Perspective. In *Proceedings of the 24th IEEE International Symposium on High-Performance Computer Architecture (HPCA)*, 2018.
- [30] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep Residual Learning for Image Recognition. In *Proceedings of the IEEE 2016 Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016.
- [31] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Identity Mappings in Deep Residual Networks. In *Proceedings of the 14th European Conference on Computer Vision (ECCV)*, 2016.
- [32] Jeremy Hermann and Mike Del Balso. Meet Michelangelo: Uber’s Machine Learning Platform. <https://eng.uber.com/michelangelo/>, September 2017.
- [33] Andrew Howard, Mark Sandler, Grace Chu, Liang-Chieh Chen, Bo Chen, Mingxing Tan, Weijun Wang, Yukun Zhu, Ruoming Pang, Vijay Vasudevan, et al. Searching for MobileNetV3. In *Proceedings of the IEEE 2019 Conference on Computer Vision (ICCV)*, 2019.
- [34] Kevin Hsieh, Ganesh Ananthanarayanan, Peter Bodik, Shivaram Venkataraman, Paramvir Bahl, Matthai Philipose, Phillip B Gibbons, and Onur Mutlu. Focus: Querying Large Video Dataset with Low Latency and Low Cost. In *Proceedings of the 13th USENIX Symposium on Operating Systems Design and Implementation (OSDI)*, 2018.
- [35] Jie Hu, Li Shen, and Gang Sun. Squeeze-and-Excitation Networks. In *Proceedings of the IEEE 2018 Conference on Computer Vision and Pattern Recognition (CVPR)*, 2018.
- [36] Gao Huang, Zhuang Liu, Laurens Van Der Maaten, and Kilian Q Weinberger. Densely Connected Convolutional Networks. In *Proceedings of the IEEE 2017 Conference on Computer Vision and Pattern Recognition (CVPR)*, 2017.
- [37] Peng Huang, Chuanxiong Guo, Lidong Zhou, Jacob R Lorch, Yingnong Dang, Murali Chintalapati, and Randolph Yao. Gray Failure: The Achilles’ Heel of Cloud-Scale Systems. In *Proceedings of the 16th Workshop on Hot Topics in Operating Systems (HotOS)*, 2017.
- [38] Qi Huang, Ken Birman, Robbert Van Renesse, Wyatt Lloyd, Sanjeev Kumar, and Harry C Li. An Analysis of Facebook Photo Caching. In *Proceedings of the 24th ACM Symposium on Operating Systems Principles (SOSP)*, 2013.
- [39] Vatche Ishakian, Vinod Muthusamy, and Aleksander Slominski. Serving Deep Learning Models in a Serverless Platform. In *Proceedings of the 6th IEEE International Conference on Cloud Engineering (IC2E)*, 2018.
- [40] Chris Jones, John Wilkes, Niall Murphy, and Cody Smith. *Site Reliability Engineering: How Google Runs Production Systems*. O’Reilly Media, 2016. <https://landing.google.com/sre/sre-book/chapters/service-level-objectives/>.

- [41] Norman P Jouppi, Cliff Young, Nishant Patil, David Patterson, Gaurav Agrawal, Raminder Bajwa, Sarah Bates, Suresh Bhatia, Nan Boden, Al Borchers, et al. In-Datcenter Performance Analysis of a Tensor Processing Unit. In *Proceedings of the 44th ACM/IEEE International Symposium on Computer Architecture (ISCA)*, 2017.
- [42] Kostis Kaffes, Dragos Sbirlea, Yiyan Lin, David Lo, and Christos Kozyrakis. Leveraging Application Classes to Save Power in Highly-Utilized Data Centers. In *Proceedings of the 11th ACM Symposium on Cloud Computing (SoCC)*, 2020.
- [43] Daniel Kang, John Emmons, Firas Abuzaid, Peter Bailis, and Matei Zaharia. NoScope: Optimizing Neural Network Queries over Video at Scale. *Proceedings of the VLDB Endowment*, 10(11), 2017.
- [44] Ram Srivatsa Kannan, Lavanya Subramanian, Ashwin Raju, Jeongseob Ahn, Jason Mars, and Lingjia Tang. GrandSLAM: Guaranteeing SLAs for Jobs in Microservices Execution Frameworks. In *Proceedings of the 14th European Conference on Computer Systems (EuroSys)*, 2019.
- [45] Andrew Lavin and Scott Gray. Fast Algorithms for Convolutional Neural Networks. In *Proceedings of the IEEE 2016 Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016.
- [46] Jialin Li, Naveen Kr Sharma, Dan RK Ports, and Steven D Gribble. Tales of the Tail: Hardware, OS, and Application-Level Sources of Tail Latency. In *Proceedings of the 5th ACM Symposium on Cloud Computing (SoCC)*, 2014.
- [47] David Lo, Liquan Cheng, Rama Govindaraju, Parthasarathy Ranganathan, and Christos Kozyrakis. Heracles: Improving Resource Efficiency at Scale. In *Proceedings of the 42nd Annual International Symposium on Computer Architecture (ISCA)*, 2015.
- [48] Peter Mattson, Vijay Janapa Reddi, Christine Cheng, Cody Coleman, Greg Diamos, David Kanter, Paulius Micikevicius, David Patterson, Guenther Schmuelling, Hanlin Tang, et al. MLPerf: An industry standard benchmark suite for machine learning performance. *IEEE Micro*, 40(2):8–16, 2020.
- [49] Neural Network Exchange Format (NNEF). Retrieved May 2020 from <https://www.khronos.org/nnef/>, 2020.
- [50] NVIDIA TensorRT. Retrieved May 2020 from <https://developer.nvidia.com/tensorrt>, 2020.
- [51] Christopher Olston, Noah Fiedel, Kiril Gorovoy, Jeremiah Harmsen, Li Lao, Fangwei Li, Vinu Rajashekhar, Sukriti Ramesh, and Jordan Soyke. TensorFlow-Serving: Flexible, High-Performance ML Serving. *Workshop on ML Systems at NeurIPS 2017*, 2017.
- [52] Open Neural Network Exchange Format: The new open ecosystem for interchangeable AI models. Retrieved May 2020 from <https://onnx.ai/>, 2020.
- [53] The ONNX Model Zoo. Retrieved May 2020 from <https://github.com/onnx/models>, 2020.
- [54] Nathan Otterness, Ming Yang, Tanya Amert, James Anderson, and F Donelson Smith. Inferring the scheduling policies of an embedded cuda gpu. In *13th Workshop on Operating Systems Platforms for Embedded Real-Time Applications (OSPRT)*, 2017.
- [55] Kay Ousterhout, Christopher Canel, Sylvia Ratnasamy, and Scott Shenker. Monotasks: Architecting for Performance Clarity in Data Analytics Frameworks. In *Proceedings of the 26th ACM Symposium on Operating Systems Principles (SOSP)*, 2017.
- [56] Kay Ousterhout, Ryan Rasti, Sylvia Ratnasamy, Scott Shenker, and Byung-Gon Chun. Making Sense of Performance in Data Analytics Frameworks. In *Proceedings of the 12th USENIX Symposium on Networked Systems Design and Implementation (NSDI)*, 2015.
- [57] Jonathan Perry, Amy Ousterhout, Hari Balakrishnan, Devavrat Shah, and Hans Fugal. FastPass: A Centralized “Zero-Queue” Datacenter Network. In *Proceedings of the 2014 Conference of the ACM Special Interest Group on Data Communication (SIGCOMM)*, 2014.
- [58] Francisco Romero, Qian Li, Neeraja J Yadwadkar, and Christos Kozyrakis. INFaaS: A Model-less Inference Serving System. *arXiv preprint arXiv:1905.13348*, 2019.
- [59] Deploying a Model on Amazon SageMaker Hosting Services. Retrieved May 2020 from <https://docs.aws.amazon.com/sagemaker/latest/dg/how-it-works-hosting.html>, 2020.
- [60] Roy Schwartz, Jesse Dodge, Noah A Smith, and Oren Etzioni. Green AI. *arXiv preprint arXiv:1907.10597*, 2019.
- [61] Mohammad Shahradd, Rodrigo Fonseca, Íñigo Goiri, Gohar Chaudhry, Paul Batum, Jason Cooke, Eduardo Laureano, Colby Tresness, Mark Russinovich, and Riccardo Bianchini. Serverless in the Wild: Characterizing and Optimizing the Serverless Workload at a Large Cloud Provider. In *Proceedings of the 2020 USENIX Annual Technical Conference (ATC '20)*, 2020.
- [62] Haichen Shen, Lequn Chen, Yuchen Jin, Liangyu Zhao, Bingyu Kong, Matthai Philipose, Arvind Krishnamurthy, and Ravi Sundaram. Nexus: a GPU Cluster Engine for Accelerating DNN-based Video Analysis. In

Proceedings of the 27th ACM Symposium on Operating Systems Principles (SOSP), 2019.

- [63] Julien Simon. Amazon Elastic Inference – GPU-Powered Deep Learning Inference Acceleration. <https://aws.amazon.com/blogs/aws/amazon-elastic-inference-gpu-powered-deep-learning-inference-acceleration/>, November 2018.
- [64] Kacper Sokol and Peter A Flach. Glass-Box: Explaining AI Decisions With Counterfactual Statements Through Conversation With a Voice-enabled Virtual Assistant. In *Proceedings of the 27th International Joint Conference on Artificial Intelligence (IJCAI)*, 2018.
- [65] Emma Strubell, Ananya Ganesh, and Andrew McCallum. Energy and Policy Considerations for Deep Learning in NLP. *arXiv preprint arXiv:1906.02243*, 2019.
- [66] Jinghao Sun, Jing Li, Zhishan Guo, An Zou, Xuan Zhang, Kunal Agrawal, and Sanjoy Baruah. Real-Time Scheduling upon a Host-Centric Acceleration Architecture with Data Offloading. In *Proceedings of the 26th IEEE Real-Time and Embedded Technology and Applications Symposium (RTAS)*, 2020.
- [67] Christian Szegedy, Wei Liu, Yangqing Jia, Pierre Sermanet, Scott Reed, Dragomir Anguelov, Dumitru Erhan, Vincent Vanhoucke, and Andrew Rabinovich. Going Deeper with Convolutions. In *Proceedings of the IEEE 2015 Conference on Computer Vision and Pattern Recognition (CVPR)*, 2015.
- [68] Christian Szegedy, Vincent Vanhoucke, Sergey Ioffe, Jon Shlens, and Zbigniew Wojna. Rethinking the Inception Architecture for Computer Vision. In *Proceedings of the IEEE 2016 Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016.
- [69] Ymir Vigfusson, Hussam Abu-Libdeh, Mahesh Balakrishnan, Ken Birman, Robert Burgess, Gregory Chockler, Haoyuan Li, and Yoav Tock. Dr. Multicast: Rx for Data Center Communication Scalability. In *Proceedings of the 5th European Conference on Computer systems (EuroSys)*, 2010.
- [70] Limin Wang, Yuanjun Xiong, Zhe Wang, Yu Qiao, Dahua Lin, Xiaoou Tang, and Luc Van Gool. Temporal Segment Networks: Towards Good Practices for Deep Action Recognition. In *Proceedings of the 14th European Conference on Computer Vision (ECCV)*, 2016.
- [71] Carole-Jean Wu, David Brooks, Kevin Chen, Douglas Chen, Sy Choudhury, Marat Dukhan, Kim Hazelwood, Eldad Isaac, Yangqing Jia, Bill Jia, et al. Machine Learning at Facebook: Understanding Inference at the Edge. In *Proceedings of the 2019 IEEE International Symposium on High Performance Computer Architecture (HPCA)*, 2019.
- [72] Bin Xiao, Haiping Wu, and Yichen Wei. Simple Baselines for Human Pose Estimation and Tracking. In *Proceedings of the 16th European Conference on Computer Vision (ECCV)*, 2018.
- [73] Saining Xie, Ross Girshick, Piotr Dollár, Zhuowen Tu, and Kaiming He. Aggregated Residual Transformations for Deep Neural Networks. In *Proceedings of the IEEE 2017 Conference on Computer Vision and Pattern Recognition (CVPR)*, 2017.
- [74] Ming Yang, Nathan Otterness, Tanya Amert, Joshua Bakita, James H Anderson, and F Donelson Smith. Avoiding Pitfalls when using NVIDIA GPUs for Real-Time Tasks in Autonomous Systems. In *Proceedings of the 30th Euromicro Conference on Real-Time Systems (ECRTS)*, 2018.
- [75] Fisher Yu, Dequan Wang, Evan Shelhamer, and Trevor Darrell. Deep Layer Aggregation. In *Proceedings of the IEEE 2018 Conference on Computer Vision and Pattern Recognition (CVPR)*, 2018.
- [76] Sergey Zagoruyko and Nikos Komodakis. Wide Residual Networks. *arXiv preprint arXiv:1605.07146*, 2016.
- [77] Chengliang Zhang, Minchen Yu, Wei Wang, and Feng Yan. Mark: Exploiting Cloud Services for Cost-Effective, SLO-aware Machine Learning Inference Serving. In *Proceedings of the 2019 USENIX Annual Technical Conference (ATC)*, 2019.
- [78] Hang Zhang, Chongruo Wu, Zhongyue Zhang, Yi Zhu, Zhi Zhang, Haibin Lin, Yue Sun, Tong He, Jonas Mueller, R Manmatha, et al. ResNeSt: Split-Attention Networks. *arXiv preprint arXiv:2004.08955*, 2020.

A Artifact Appendix

A.1 Abstract

The artifact consists of Clockwork’s prototype source code, instructions for building from source, and directions for preparing the environment. The instructions for launching a Docker instance that has all dependencies pre-installed is provided as well. The artifact also contains scripts, descriptions, and instructions to run the experiments automatically or manually for reproducing the graphs and results presented in the paper.

A.2 Artifact check-list

- **Program:** dnn-model-serving, multi-tenant
- **Compilation:** cmake, g++
- **Binary:** worker, controller, client
- **Model:** distributed, multi-tenant
- **Data set:** azure-functions-trace-2019, poisson-distribution
- **Run-time environment:** Linux, CUDA, network
- **Hardware:** NVIDIA, Tesla-V100
- **Execution:** automated, manual
- **Metrics:** throughput, latency, SLO-violation, tail-latency
- **Output:** telemetry-measurements, table, graph
- **Experiments:** throughput-latency, scalability, predictability, SLO, tail-latency
- **Required disk space:**
Clockwork: 210MB
Total including compiled models and dataset: 12GB
- **Expected experiment run time:** About 17 hours in total
- **Public link:**
<https://gitlab.mpi-sws.org/cld/ml/clockwork>
- **Code licenses:**
Clockwork: Apache License 2.0
TVM: Apache License 2.0
CUDA Common Library: Apache License 2.0
Catch2: Boost Software License 1.0
- **Data licenses:**
Azure Functions Trace 2019: CC-BY Attribution

A.3 Description

A.3.1 How to access

The artifact is publicly available at
<https://gitlab.mpi-sws.org/cld/ml/clockwork>

A.3.2 Hardware dependencies

To reproduce the exact experiment results, worker machines must have 768GB RAM or higher, 16 CPU cores or more, at least one 32GB Tesla v100 GPU and 10Gbps network. The large-scale experiment with Azure Functions (Fig. 9) requires 12 worker machines. Most other experiments require fewer worker machines; details on the number of machines for each experiment and environment customization guide is provided in each experiment’s documentation.

A.3.3 Software dependencies

- **Clockwork:**
Ubuntu 18.04 or later, CUDA v9.0+, libtbb-dev, libasio-dev, libconfig++-dev, libboost-all-dev, g++-8, make, cmake, automake, autoconf, libtool, curl, unzip, clang, llvm, and protobuf.
- A Dockerfile is provided to facilitate the build process.
- **Data analysis and plotting scripts:**
Python 3.x and the numpy, pandas, matplotlib, and seaborn libraries.

A.3.4 Data sets

- Publicly released Azure Functions 2019 trace [61]
<https://gitlab.mpi-sws.org/cld/trace-datasets/azure-functions>

A.3.5 Models

The DNN models pre-compiled for NVIDIA Volta V100 GPUs are accessible at
<https://gitlab.mpi-sws.org/cld/ml/clockwork-modelzoo-volta>

A.4 Installation

- **Installation pre-requisites:**
<https://gitlab.mpi-sws.org/cld/ml/clockwork/-/blob/master/docs/prerequisites.md>
- **Building Clockwork:**
<https://gitlab.mpi-sws.org/cld/ml/clockwork/-/blob/master/docs/building.md>
- **Setting-up the environment:**
<https://gitlab.mpi-sws.org/cld/ml/clockwork/-/blob/master/docs/environment.md>
- **Clockwork configuration:**
<https://gitlab.mpi-sws.org/cld/ml/clockwork/-/blob/master/docs/configuration.md>

A.5 Experiment workflow

Experiments can be run using the scripts provided in the repository. We have also provided instructions to run the experiments manually. To get started with Clockwork, we recommend getting the system running manually, in order to understand the pieces involved, and to ensure the system has been configured appropriately for your machines. Afterwards, you might choose to run the experiments using the provided scripts or manually. The experiments repository is available at
<https://gitlab.mpi-sws.org/cld/ml/clockwork-results>

Experiment	Related figure	Execution time (hr)	Documentation and scripts
How Does Clockwork Compare?	Fig. 5	3	https://gitlab.mpi-sws.org/cld/ml/clockwork-results/-/tree/master/sec61_fig5
Can Clockwork Serve Thousands?	Fig. 6	1.5	https://gitlab.mpi-sws.org/cld/ml/clockwork-results/-/tree/master/sec62_fig6
How Low Can Clockwork Go?	Fig. 7	1	https://gitlab.mpi-sws.org/cld/ml/clockwork-results/-/tree/master/sec63_fig7
Can Clockwork Isolate Performance?	Fig. 8	1	https://gitlab.mpi-sws.org/cld/ml/clockwork-results/-/tree/master/sec64_fig8
Are Realistic Workloads Predictable?	Fig. 10	8	https://gitlab.mpi-sws.org/cld/ml/clockwork-results/-/tree/master/sec65_fig9_fig10
Can Clockwork Scale?	Fig. 11	2	https://gitlab.mpi-sws.org/cld/ml/clockwork-results/-/tree/master/sec66_fig11

Table 3: The experiments reproducing the presented results in this paper, their related figures, execution time, and links to the extensive documentation and scripts for each experiment.

A.6 Evaluation and expected results

The experiments repository is structured based on §6. We have provided the experiment titles, their related figures on the paper, execution time of each experiment, and the links to directories containing the respective descriptions, scripts and instructions in [Table 3](#).

A.7 Experiment customization

The directions for running each experiment manually is provided in each experiment’s documentation. Instructions for customizing the experiment environment is provided at <https://gitlab.mpi-sws.org/cld/ml/clockwork/-/blob/master/docs/customizing.md>

A.8 AE Methodology

Submission, reviewing and badging methodology:

<https://www.usenix.org/conference/osdi20/call-for-artifacts>