

Online Tracking Algorithms on GPUs for the PANDA Experiment at FAIR

This content has been downloaded from IOPscience. Please scroll down to see the full text.

2015 J. Phys.: Conf. Ser. 664 082006

(<http://iopscience.iop.org/1742-6596/664/8/082006>)

View [the table of contents for this issue](#), or go to the [journal homepage](#) for more

Download details:

IP Address: 134.94.122.17

This content was downloaded on 29/01/2016 at 12:52

Please note that [terms and conditions apply](#).

Online Tracking Algorithms on GPUs for the $\bar{\text{P}}\text{ANDA}$ Experiment at FAIR

L Bianchi¹, A Herten¹, J Ritman¹, T Stockmanns¹ for the $\bar{\text{P}}\text{ANDA}$ Collaboration

A. Adinetz², J Kraus³, D Pleiter²

¹ IKP, Forschungszentrum Jülich, Germany

² JSC, Forschungszentrum Jülich, Germany

³ NVIDIA GmbH

E-mail: l.bianchi@fz-juelich.de | a.herten@fz-juelich.de

Abstract. $\bar{\text{P}}\text{ANDA}$ is a future hadron and nuclear physics experiment at the FAIR facility in construction in Darmstadt, Germany. In contrast to the majority of current experiments, $\bar{\text{P}}\text{ANDA}$'s strategy for data acquisition is based on event reconstruction from free-streaming data, performed in real time entirely by software algorithms using global detector information. This paper reports the status of the development of algorithms for the reconstruction of charged particle tracks, optimized online data processing applications, using General-Purpose Graphic Processing Units (GPU). Two algorithms for trackfinding, the Triplet Finder and the Circle Hough, are described, and details of their GPU implementations are highlighted. Average track reconstruction times of less than 100 ns are obtained running the Triplet Finder on state-of-the-art GPU cards. In addition, a proof-of-concept system for the dispatch of data to tracking algorithms using Message Queues is presented.

1. Introduction

1.1. The $\bar{\text{P}}\text{ANDA}$ Experiment

The $\bar{\text{P}}\text{ANDA}$ (Anti-**P**roton **A**nnihilation at **D**armstadt) experiment is one of the main experiments at FAIR (**F**acility for **A**ntiproton and **I**on **R**esearch in **E**urope), currently in construction in Darmstadt, Germany. $\bar{\text{P}}\text{ANDA}$ will study collisions between a cooled antiproton beam ($\Delta p/p \sim 10^{-5}$) and a fixed proton target, in the momentum range 1–15 GeV/c. The unique $p\bar{p}$ initial-state configuration allows $\bar{\text{P}}\text{ANDA}$ to study a wide range of research topics in the area of Quantum Chromodynamics (QCD) including hadron spectroscopy, hypernuclei, and nuclear structure [1].

$\bar{\text{P}}\text{ANDA}$ shares with many other experiments the challenge of studying rare physics processes in a collision environment dominated by background events. The high luminosity needed to accumulate sufficient statistics means that an enormous amount of uninteresting background events, whose cross-section is many order of magnitude greater than the signal, is produced as well. The total data rate coming from the detector greatly exceeds the capabilities of offline storage, so some form of online signal/background rejection needs to be implemented.

The conventional strategy to reduce the incoming data is to perform a fast, low-level trigger selection, implemented in hardware, based on information from a subset of the detector. If the



outcome of the decision is positive, data from the whole detector are collected and stored offline for further analysis.

In the case of $\bar{\text{PANDA}}$, this approach is not feasible. Due to the nature of hadronic interactions in the energy range of interest, background and signal events are very similar, and information from the full reconstructed event is needed to perform the selection. $\bar{\text{PANDA}}$'s strategy for data acquisition is instead based on an online event filtering scheme. Data from the whole detector is read continuously. An initial stage of event building is performed by FPGA-based compute nodes. Then, software algorithms perform full reconstruction of each event and background/signal identification based on multiple software trigger lines. Events passing the trigger selection are saved for offline storage. The incoming data rate from the detectors is approximately 200 GB/s. To match the available offline storage capabilities of 3 PB/year, an online background rejection factor of about 1000 is required.

1.2. The $\bar{\text{PANDA}}$ Central Tracking Sub-detectors

The $\bar{\text{PANDA}}$ detector complex can be divided into two main sections, the Target Spectrometer (TS) and the Forward Spectrometer (FS). In the TS, the focus of this paper, a 2 T solenoidal magnetic field bends the trajectory of charged particles into three-dimensional helices, with their axis parallel to the beam direction. The two-dimensional projection of tracks on the transverse plane are thus circles. The tracking sub-detector closest to the beam is the Micro Vertex Detector (MVD) [2]. The MVD is a solid-state detector composed of a combination of pixel detectors (10.3×10^6 channels) and double-sided silicon strips (2×10^5 channels), capable of a vertex resolution of $< 100 \mu\text{m}$. Surrounding the MVD is the Straw Tube Tracker (STT) [3]. The STT is made up of 4636 drift tubes, or *straws*, filled with a 9:1 Ar/CO₂ gas mixture, and it has a spatial resolution of $\sigma_{xy} \sim 150 \mu\text{m}$, $\sigma_z \sim 2\text{--}3 \text{ mm}$. Three Gas Electron Multiplier (GEM) [1] disks are located in the forward direction, with about 35000 readout channels and a spatial resolution of $< 100 \mu\text{m}$.

1.3. Online Tracking on GPUs

Online tracking is an essential step in the $\bar{\text{PANDA}}$ online event reconstruction chain. Online tracking algorithms take discrete hit data from the tracking detectors, and process them to produce continuous particle trajectories, or tracks. The information from each track is used as input for online event building and event selection phases.

A Graphic Processing Unit (GPU) is a type of hardware architecture, designed for data-parallel computation. Originally created for accelerating graphic-intensive application in personal computers, GPU-based processors are being increasingly used for general computation, thanks to the introduction of high-level programming paradigms such as CUDA or OpenCL and the release of dedicated hardware. This paradigm is known as General-Purpose GPU (GPGPU). For computational tasks that can be formulated in a data-parallel fashion, using GPUs in alternative to CPUs can result in a significant increase in performance, in terms of both absolute and relative performance, i.e. considering the computational power per units of cost of hardware and power consumption. However, not all applications are well-suited for porting to the GPU: *massive parallelism* is typically needed to exploit the full potential of the GPU and reach substantial speedups w.r.t. the CPU version. In addition, the different hardware architecture results in programming strategies different from their CPU counterparts.

Solutions for event processing based on GPUs have been studied for many experiments in high-energy and nuclear physics [4, 5]. Among them, ALICE [6] at LHC successfully implemented GPU-based reconstruction algorithms in its High-Level Trigger stage of the DAQ chain. The goal of $\bar{\text{PANDA}}$ is, however, to include GPU-based online track reconstruction at an earlier stage of the online reconstruction chain, without a previous stage of hardware trigger selection.

2. Triplet Finder Algorithm

2.1. Algorithm Concept

The Triplet Finder algorithm is a track finding algorithm, developed specifically for the PANDA STT detector. A detailed description of the algorithm be found in [7, 8]. The core idea of the Triplet Finder (Figure 1) revolves around the fast reconstruction of a circle based on a restricted subset of hit combinations. Three sets of pivot straws are arranged in contiguous configuration called pivot layers, as illustrated in Figure 2. As soon as a hit is detected in a pivot straw, the neighboring straws are checked for hits. A virtual hit, called *triplet*, is then created as the center of mass of the three hit points. Once two triplets have been calculated, a circular track is computed from the two triplets and the interaction point in (0, 0). The last step of the algorithm consists of associating remaining hits lying on the trajectory of the circle with the track.

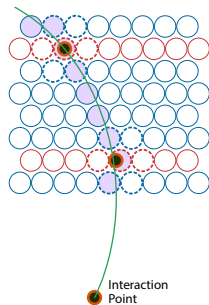


Figure 1: Sketch of the principle of operation of the Triplet Finder algorithm.

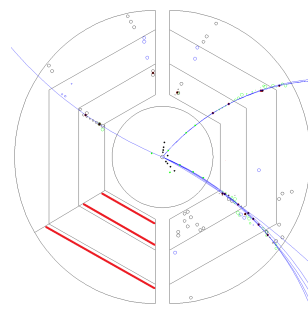


Figure 2: Schematic section of the STT in the (x, y) plane. Pivot layers are shown for one of the six symmetrical straw sectors. Triplets (orange-black dots) and reconstructed tracks (blue) are also shown.

2.2. GPU Implementation

The GPU version of the Triplet Finder algorithm is implemented using the CUDA programming language [9]. In this section, techniques and improvements of more general applicability are highlighted. A detailed description of the GPU implementation is available in [10].

Bunching The main challenge in the GPU implementation of the TF algorithm comes from its computational complexity of $\mathcal{O}(N^2)$, where N is the number of hits processed at the same time. This drawback is especially relevant for the GPU version, where $\mathcal{O}(10^4)$ hits should be processed in parallel to ensure full utilization of the GPU. One strategy to circumvent the complexity is to subdivide hits in sets, called *bunches*, that can be processed independently. This restricts the $\mathcal{O}(N^2)$ complexity to the number of hits in each bunch, while the full occupancy of the GPU can still be ensured by running multiple bunches in parallel. The subdivision is implemented as follows. First, the simulation is divided in segments of duration T_C , the core time. An additional time T_D , corresponding to the maximum straw tube drift time, is added at the end of each bunch (Figure 3). Hits are then assigned to bunches on the basis of their timestamp t_0 . If N_C and N_D are the average number of hits in the core time and the drift time, respectively, the resulting complexity is then $\mathcal{O}(N \frac{(N_C + N_D)^2}{N_C})$, which is linear in the total number of hits N . Duplicate tracks can occur as a result of partially overlapping bunches. This can be addressed by a later track coalescing stage. Figure 4 shows the effects of the adoption of bunching. It is the most valuable improvement in terms of performance for the GPU version of the Triplet Finder algorithm.

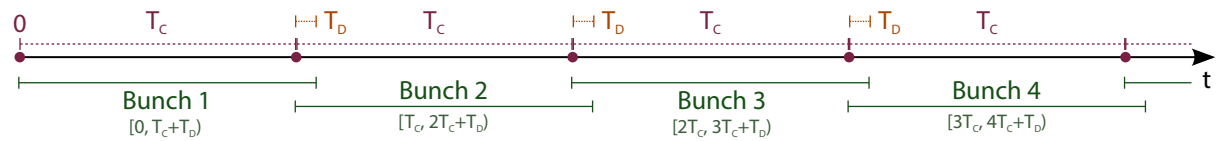


Figure 3: Scheme of bunching.

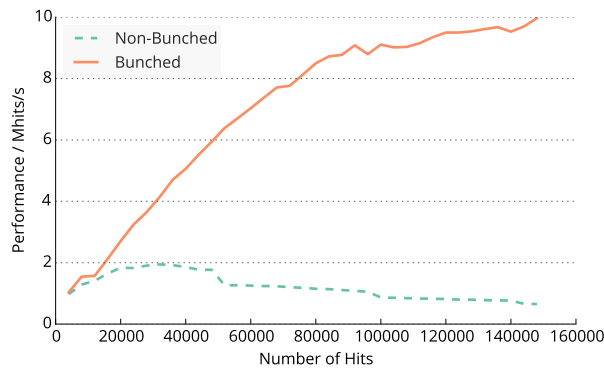


Figure 4: Comparison of the effect of bunching in terms of performance (number of hits processed per unit time) as a function of the total number of hits.

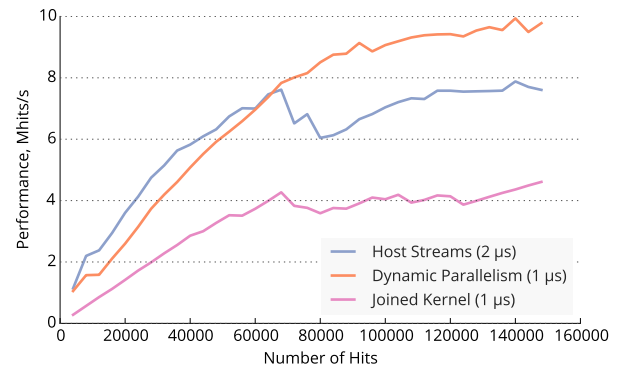


Figure 5: Performance as a function of the total number of processed hits for different kernel dispatch strategies.

Strategies for bunch dispatch on the GPU Three different strategies for dispatching multiple bunches on the GPU have been tested. The *dynamic parallelism* approach makes use of nested kernels. The host launches one master kernel, with one thread for each bunch. Each *master* kernel instance then launches separate individual *child* kernels for each phase of the algorithm, with one thread per hit. In the *host streams* approach, individual kernels for each kernel are similarly used, but they are launched by the host, with one CUDA stream per bunch. In the *joined kernel* approach, all phases of the algorithm are performed within a single kernel. The host invokes the joined kernel with one thread block for each bunch. A performance comparison of the bunch dispatch strategies is shown in Figure 5. For a high number of processed hits, the best performance results from the dynamic parallelism approach.

Effect of Data Packing Ensuring good memory access patterns is particularly important for GPU applications, and more efficient memory access patterns almost always result in improved performance. In the case of Triplet finder, a performance improvement of up to 20% can be reached when using Structures of Arrays (SoA) instead of the more intuitive Arrays of Structures (AoS).

Comparison of Server-grade and Consumer-grade cards GPU-based devices currently available on the market can be divided in two broad categories: consumer-grade cards, designed for accelerating graphics on personal computers, and server-grade cards, built specifically for use in High-Performance Computing. Although not their primary purpose, recent consumer-grade cards are able to run CUDA applications. Their inferior performance in absolute terms w.r.t. server-grade cards is accompanied by a much lower retail price. As an indicative comparison, the same version of Triplet Finder algorithm is run for one server-grade card, NVIDIA Tesla K20X, with 2688 CUDA cores, a peak single-precision performance of 3950 GFLOPS and a price of about 3200 EUR, and a consumer-grade card, NVIDIA GeForce GTX 750 Ti, with 640 CUDA

cores, a peak single-precision performance of 1306 GFLOPS and a price of about 140 USD. In absolute terms, the peak performance for the Triplet Finder algorithm is about two times higher on the K20X than on the GTX 750 Ti. In terms of Mhits/s/USD, the price-normalized performance of the GTX 750 Ti w.r.t. the K20X is about 10 times higher. However, further studies are needed to quantify how the drawbacks of using consumer-grade cards, e.g. the lack of support of ECC memory correction features, can affect reliability and data integrity in an online data processing scenario.

2.3. Summary and Outlook

Taking into account all optimizations and pre-processing steps, the Triplet Finder algorithm processes 8.3 Mhits/s on a Tesla K20X card. On a Tesla K40X, a more modern card overclocked with the GPUBoost technology, the performance is 10.3 Mhits/s. Although the GPU version of the algorithm can be considered deeply optimized, a number of possible future improvements are identified, including the porting to the GPU of the pre-processing phases currently being performed on the CPU, and the introduction of hit bookkeeping to reduce the computational complexity of the hit association phase.

3. Circle Hough Transform

The Circle Hough algorithm is a novel algorithm for trackfinding, developed at $\bar{\text{P}}\text{ANDA}$ during the last year. Compared to the Triplet Finder, it is not exclusive to the $\bar{\text{P}}\text{ANDA}$ STT. In the current version, it can be applied to hits originated in all central tracking detectors: MVD, STT, and GEM.

3.1. Algorithm Concept

The Circle Hough algorithm is based on the principle of the Hough transform, a technique for feature extraction used e.g. for detecting straight edges in a picture. In this method, an input dataset is checked against a model, described by a number of parameters. The goal is to find the set of parameters corresponding to the model that fits best with the input data. A family of prototype models is generated for each point of the input dataset; the resulting parameters are collected and the resulting parameter space, the Hough space, is analyzed. Points in the Hough space with the highest density, corresponding to the parameters occurring with the highest frequency in the input dataset, coincide with the model that fits best to the data.

In the case of track finding, the input dataset consists of hit points, and the feature to be extracted are tracks that fit best to the hit points. In the Circle Hough algorithm, tracks are represented as circles in the xy plane, passing through both the interaction point in (0,0) (IP) and the hit point in $(x_{\text{hit}}, y_{\text{hit}})$ (Figure 6). A circle satisfying these requirements is referred to as *Hough circle*. The condition of passing through the hit point is valid for point-like hit points, such as hits originated in the MVD or the GEMs. For extended hit points, such as hits in the STT, Hough circles must be tangent to the circle centered in $(x_{\text{hit}}, y_{\text{hit}}, r_{\text{iso}})$. r_{iso} is the isochrone radius, calculated from the drift time of the hit. For each hit point, all possible Hough circles are computed. This is equivalent to defining all possible primary tracks compatible with the hit point. The process is repeated for all hits, sampling the Hough space. Since Hough circles are uniquely described by two parameters, e.g. the coordinates of the circle center (x_C, y_C) , the Hough space will be two-dimensional. The Hough space is then scanned for peaks; the (x_C, y_C) coordinates of the peaks correspond to the parameters of the tracks in the real space.

3.2. GPU Implementation

Although the Circle Hough algorithm was originally tested on the CPU, its native degrees of parallelism, in both the hits and the Hough circle calculation for each hit, make the development

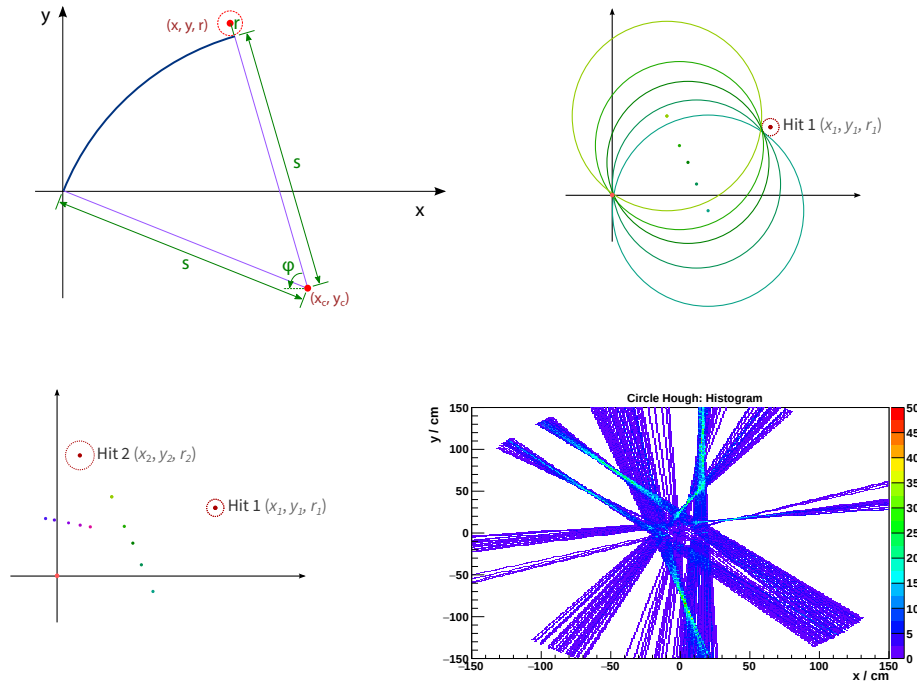


Figure 6: Schematic diagrams for the Circle Hough algorithm. From left to right, top to bottom: (1) Definition of angle φ for a Hough circle around the STT hit point in (x, y, r) . (2) Several Hough circles are created. (3) The operation is repeated for other hit points. Some Hough circle centers are pictured. (4) Example of representation of the discretized 2D Hough space.

of a GPU version a natural choice. The GPU version is implemented in CUDA C/C++, and tested on different NVIDIA cards. The computational phases of the algorithm can be outlined as follows. First, the sampling values, used as the input of the Hough circle generation, are defined. In the simplest approach, this is a constant set of values for the angle φ between the IP and the hit point in the range $[0^\circ, 360^\circ)$, with a fixed sampling granularity, $\Delta\varphi$, resulting in N_φ angles in total.

Then, the Hough circles are calculated. Hit data are copied on the device; then one kernel is invoked, with each thread computing one set of Hough circle center coordinates for each hit and each value of φ . Block size is determined to ensure the maximum occupancy of the GPU, with results reached for block sizes of 128–256.

The set of Hough circle center coordinates is copied back to the host, where the subsequent phases of the algorithm are performed. The Hough space is discretized in a 2D histograms, which is then filled with the Hough circle center coordinates. A peak-finding algorithm identifies peaks in the histogram, and their coordinates are then saved to extract track parameters.

3.3. Summary and Outlook

Analysis of the performance of the CUDA code show that the step of copying back the Hough circle center coordinates to the host is the current limiting factor. Even using pinned host memory, it accounts for 80% and 90% of the total runtime on an NVIDIA GeForce GTX 750 Ti and an NVIDIA Tesla K40X, respectively. The kernel-only performance of the Hough circle computation of the algorithm is 30 Mhits/s on a Tesla K40X. Development in the near future is thus targeted towards implementing on the GPU the phases of histogram filling and the subsequent peak-finding. The possibility to use ad-hoc sampling values for each hit instead than the constant set of angle values, and the adoption of more advanced schemes for the kernel

dispatch are also being explored.

4. Use of Message Queues in Combination with GPUs

The development of a suitable communication infrastructure (CI) is one of the central challenges in the design of a complex system for data acquisition and processing. The CI must provide the optimal balance between performance, stability, and use of resource. Additionally, in a heterogeneous computing scenario, where many types of hardware devices (CPUs, Many Integrated Cores (MIC) co-processors, GPUs, FPGAs) are used for data processing, the CI can provide the necessary compatibility layer among them. Furthermore, the CI should be flexible enough to accommodate changes both in the architecture of the system, and the possible future evolutions of low-level data transport protocols and interfaces.

One possible approach is to use Message Queues (MQ) as the basis for the exchange of data and control. In an MQ environment, both data and control packages (messages) are exchanged between senders and receivers through a control structure (queue), making for an inherently flexible, easily scalable architecture. The actual data processing can be distributed among different independent tasks, running in parallel on different machines, networks and processor types, exchanging data through the MQ.

FairMQ [11] is the implementation of MQs within the FairRoot framework. The main idea of FairMQ is to provide MQ functionality by means of a unified, high-level abstract interface. The same classes are used to create and access messages regardless of the type of communication (intra- or inter-process), the protocol, or the channel (PCIe, network, ...), providing a high degree of versatility by design. The choice of the underlying low-level library is also open, leaving the possibility to change it or to adopt optimized versions. At the moment, ZeroMQ [12] and nanomsg [13] are supported.

4.1. Circle Hough Test System

To explore the possibilities of using FairMQ in association with GPU-based tracking, a standalone, proof-of-concept system integrating FairMQ and an implementation of the Circle Hough algorithm described in Section 3 has been developed. The functionality of the system is subdivided into independent modules (Fig. 7), structured as follows. An instance of **FairMQDevice**, the *Sampler*, processes the input data and initiates the FairMQ stream. First, the **ReadFile** module to read hit data from an input file is called. The data are serialized into arrays, and copied to **FairMQMessages**. The message transmission is managed by an instance of **FairMQTransportFactory**. A second instance of **FairMQDevice**, the *Processor*, receives the FairMQ stream and performs the Circle Hough computations. First, the hit data is extracted from the **FairMQMessages** and converted to the appropriate data structures. The Circle Hough transform calculations are performed either on the CPU or on the GPU by the **CircleHoughCPU** and **CircleHoughGPU** modules, respectively. Lastly, the Hough circle center coordinates are collected by the **CreateHisto** module.

Each instance of **FairMQDevice** runs independently from the others, and automatically initiates the data processing as soon as a FairMQ communication channel is established. Once the system is properly set up, alterations to the architectures of the system, e.g. by adding processing nodes or changing the communication channel, can be done at runtime, with no alteration of the source code.

Transmission Rate Test The transmission rate in FairMQ depends on many factors, such as the transport library, the communication protocols, and the communication channel. A bare-bones system to study the transmission rate without any overhead from computation consists of two **FairMQDevices**, a *Sender* and a *Receiver*, running on the same machine. A data packet is generated once either on the CPU or on the GPU, and messages are sent continuously over

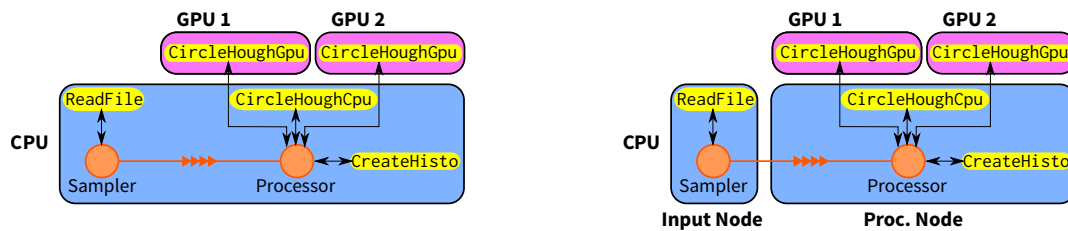


Figure 7: Scheme of the system. The orange circles represent instances of `FairMQDevice`, with the orange arrows representing the transmission of FairMQ messages within the same machine (left), and between two machines over a network (right).

a socket using the TCP/IP protocol. The size of the message is varied and the time-averaged transmission rate is measured from the FairMQ logging functionality. The results of the tests show that a transmission rate of over 5 GB/s is reached for a message size of about 100 kB. While these results should not be considered as a formal benchmark of FairMQ, they are indicative of the fact that the maximum achievable is limited by the low-level data transport infrastructure rather than by FairMQ itself.

4.2. Summary and Outlook

The test system described in this section constitutes a first look into the inclusion of FairMQ with GPU-based tracking algorithms. Future work is aimed at its integration with the PandaRoot framework, eventually evolving into a prototype of a full-featured implementation of a online reconstruction chain.

References

- [1] Lutz M *et al.* (PANDA) 2009 (*Preprint* 0903.3905)
- [2] Erni W *et al.* (PANDA) 2012 (*Preprint* 1207.6581)
- [3] Erni W *et al.* (PANDA) 2013 *Eur.Phys.J.* **A49** 25 (*Preprint* 1205.5441)
- [4] Emelianov D and Howard J (ATLAS) 2012 *J.Phys.Conf.Ser.* **396** 012018
- [5] Halyo V, LeGresley P, Lujan P, Karpusenko V and Vladimirov A 2014 *JINST* **9** P04005 (*Preprint* 1310.7556)
- [6] Gorbunov S *et al.* (ALICE) 2011 *IEEE Trans.Nucl.Sci.* **58** 1845–1851
- [7] Mertens M *et al.* (PANDA) 2014 *J.Phys.Conf.Ser.* **503** 012036
- [8] Mertens M C, Brinkmann K T, Ritman J and Wintz P (PANDA) 2014 *Hyperfine Interactions* **229** 153–158 ISSN 0304-3843 URL <http://dx.doi.org/10.1007/s10751-014-1062-3>
- [9] NVIDIA CUDA C Programming Guide URL <http://docs.nvidia.com/cuda/cuda-c-programming-guide/>
- [10] Adinetz A, Herten A, Kraus J, Mertens M C, Pleiter D, Stockmanns T and Wintz P 2014 *Procedia Computer Science* **29** 113 – 123 ISSN 1877-0509 2014 International Conference on Computational Science URL <http://www.sciencedirect.com/science/article/pii/S1877050914001884>
- [11] Al-Turany M, Klein D, Manafov A, Rybalchenko A and Uhlig F 2014 *J.Phys.Conf.Ser.* **513** 022001
- [12] 2015 ZeroMQ URL <http://zeromq.org>
- [13] 2015 nanosmg URL <https://github.com/nanosmg/nanosmg>