

Popularity-based caching of CMS datasets

Marco Meoni ^{a,1}, Raffaele Perego ^b and Nicola Tonellotto ^b

^a *University of Pisa & INFN/CERN, Pisa, Italy*

^b *ISTI-CNR, Pisa, Italy*

Abstract The distributed monitoring infrastructure of the Compact Muon Solenoid (CMS) experiment at the European Organization for Nuclear Research (CERN) records on a Hadoop infrastructures a broad variety of computing and storage logs. They represent a valuable source of information for system tuning and capacity planning. In this paper we analyze machine learning (ML) techniques on large amount of traces to discover patterns and correlations useful to classify the popularity of experiment-related datasets. We implement a scalable pipeline of Spark components which collect the dataset access logs from heterogeneous monitoring sources and group them into weekly snapshots organized by CMS sites. Predictive models are trained on these snapshots and forecast which dataset will become popular over time. Dataset popularity predictions are then used to experiment a novel strategy of data caching, called Popularity Prediction Caching (PPC). We compare the hit rates of PPC with those produced by well known caching policies. We demonstrate how the performance improvement is as high as 20% in some sites.

Keywords. CERN CMS, Big Data, Dataset Popularity, Classification, Caching

1. Introduction

The Worldwide LHC Computing Grid (WLCG) [1] project has deployed a global computing infrastructure to store, distribute and analyze nearly 50 Petabytes of data generated by the Large Hadron Collider (LHC) at the European Organization for Nuclear Research (CERN) by the end of 2017. One of the four leading experiments at the LHC particle accelerator, the Compact Muon Solenoid (CMS) [2], has collected for nearly two decades a massive amount of monitoring information regarding the usage of its distributed computing infrastructure.

With a data volume exceeding the boundary of traditional relational database systems, CMS has recently stored into a Hadoop [3] cluster nearly 4 PB of metadata produced by its monitoring systems. Among the others, users accounting, jobs lifecycle, resources utilization, sites performance, software versions and datasets access logs. The availability of this vast but heterogeneous amount of monitoring data has fueled exploratory activities using Big Data analytics so-

¹Corresponding Author, on behalf of the CMS collaboration. E-mail: marco.meoni@cern.ch