

Finding a Needle in a Field of Haystacks: Metadata Search for Distributed Research Repositories

Anna Blue Keleher
University of Maryland
bluek@cs.umd.edu

Kyle Chard, Ian Foster
(Advisors)
University of Chicago

Ioan Raicu, Alex Orhean
(Advisors)
Illinois Institute of Technology

ABSTRACT

Fast, scalable, and distributed search services are commonly available for single nodes, but lead to high infrastructure costs when scaled across tens of thousands of filesystems and repositories, as is the case with Globus. Endpoint-specific indexes may instead be stored on their respective nodes, but while this distributes storage costs between users, it also creates significant query overhead. Our solution provides a compromise by introducing two levels of indexes: a single centralized “second-level index” (SLI) that aggregates and summarizes terms from each endpoint; and many endpoint-level indexes that are referenced by the SLI and used only when needed. We show, via experiments on Globus-accessible filesystems, that the SLI reduces the amount of space needed on central servers by over 96% while also reducing the set of endpoints that need to execute user queries.

ACM Reference format:

Anna Blue Keleher, Kyle Chard, Ian Foster (Advisors), and Ioan Raicu, Alex Orhean (Advisors). 2017. Finding a Needle in a Field of Haystacks: Metadata Search for Distributed Research Repositories. In *Proceedings of SC '17, Denver, CO, USA, November 12-17, 2017*, 3 pages.
<https://doi.org/>

1 INTRODUCTION

Free-text search across a large-scale filesystem is an old and often-addressed problem [5, 7, 9]. However, when that system is distributed and multiplied by several thousand, traditional methods are no longer feasible [10]. This is the landscape inherent to Globus [2], a service that facilitates high performance and reliable data access, transfer, and synchronization across a network of more than 10,000 storage systems and repositories (called “endpoints”). The semi-structured nature and scale of these endpoints makes manual search difficult for users, especially as they wish to search across many endpoints.

To avoid having to build a new access-specific inverted index for each query, we consider indexes built across the entire Globus network, irrespective of access privileges (which may be subsequently accounted for when filtering results). We also limit the scope of search to file metadata (i.e., a “find” query) for two reasons: first, it makes indexing the entire network (centrally or otherwise)

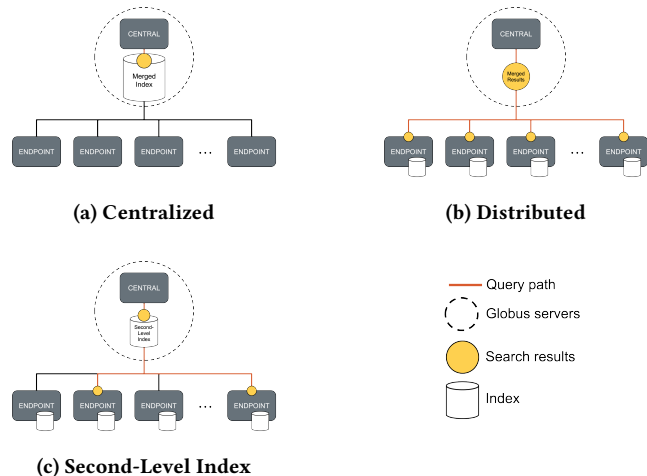


Figure 1: The three indexing architectures

feasible by reducing data sizes; and second, it is appropriate for many scientific file types which contain little free text [8].

Traditional methods suggest two indexing approaches: a single complete centrally-stored index (the *Centralized Model*) [3, 6]; and distinct indexes stored at each endpoint (the *Distributed Model*) [1, 12]. Search engines like Solr [11] and Elasticsearch [4] can accomplish the former by maintaining complete indexes at a centralized location; however, scaling to the Globus network requires significant investment in infrastructure. The alternative is building indexes at each endpoint and implementing support for distributed querying. While this approach reduces central storage requirements, it adds significant query overhead as many indexes must be queried.

Our solution provides a compromise by introducing two levels of indexes composed of many endpoint-specific indexes and a single, approximate, central *Second-Level Index* (SLI). The SLI is much smaller than the full centralized index because it reduces the number of terms in the index (through aggressive stemming and range consolidation) and the number of objects referenced for each term (endpoints rather than documents). Queries are then only redirected to the relevant subset of endpoints. The three indexing models are illustrated in Figure 1.

2 METHODOLOGY

Irrespective of which indexing model is used, each endpoint must be crawled to obtain metadata. We use the Xapian search engine [6] for indexing and extend its tokenizer to parse terms from metadata-specific string formats and numerical fields.

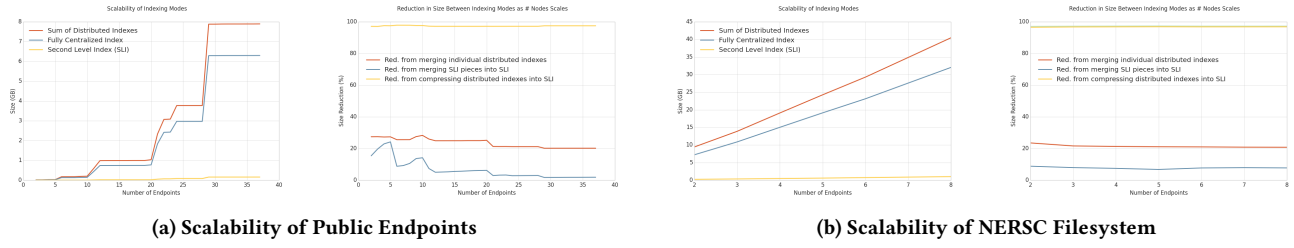


Figure 2: Scalability of Size of Indexing Modes and Reduction Between Modes

In the SLI model, we construct a standard endpoint-level index and then collapse fields and records into the SLI. For example, string fields (e.g., name, pathname) are split by separators (e.g., underscore, camel-case) and then aggregated into common tokens. Numerical fields (e.g., size, modification date) are collapsed into ranges. Queries at the SLI are stemmed in the same way as its index terms so that they are fully represented by the index, eliminating false negatives. Full queries are then passed to identified endpoint indexes.

3 RESULTS

We explore the performance of the three indexing models using two different sets of real Globus data: the 37 largest public Globus endpoints (as of June 2017), which are generally small and diverse; and a snapshot of a NERSC filesystem, representative of large, private Globus endpoints (comprising 164.4 TB of scientific data). For the NERSC dataset, we simulate a Globus network structure by grouping documents by path proximity, and binning them into eight roughly equal endpoints on separate virtual machine instances.

Table 1 shows the index size for each indexing model and dataset. The central SLI index is more than 96% smaller than that required in the Centralized Model. Because the SLI model also stores traditional indexes on each endpoint, its aggregate space requirements are slightly higher, but the burden on any single node is greatly reduced.

Table 1: Index Size Comparison Between Models

Public Endpoints			
Model	Size of Index at Location (MB)		
	Central	Endpoints	Per-Endpoint Avg.
Centralized	10,370	0	0
SLI	240	13,001	351
Distributed	0	13,001	351
Size Reduction of Central Index			97.69%
NERSC Filesystem			
Model	Size of Index at Location (MB)		
	Central	Endpoints	Per-Endpoint Avg.
Centralized	30,553	0	0
SLI	970	38,579	4,822
Distributed	0	38,579	4,822
Size Reduction of Central Index			96.83%

In order to evaluate the SLI model’s scalability, we track the sizes of the centralized index and SLI as endpoints are added. As shown in Figure 2, both the aggregate (sum) and merged indexes for the SLI model grow at a much lower rate than the traditional centralized index. The reduction percentage for the SLI model does not decay as data grows.

We also note that the incidence of terms existing on all endpoints is low, even between the artificially similar NERSC “endpoints.” We conclude that the SLI model improves on the Distributed Model for query performance whenever the SLI accurately narrows the set of relevant endpoints, which Figure 3 indicates is likely to be common.

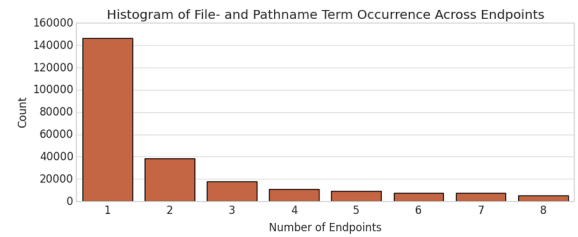


Figure 3: Frequency of filename or pathname terms occurring on multiple endpoints

4 CONCLUSION

We present a two-level index model as a compromise between the two traditional large-scale indexing architectures. We evaluate the efficacy of the approach on data across a collection of Globus endpoints. We find that the SLI algorithm stores far less data at the central node than a purely centralized approach while also greatly reducing search overhead relative to a purely distributed approach.

REFERENCES

- [1] R. Ahmed and R. Boutaba. 2011. A survey of distributed search techniques in large scale distributed systems. *IEEE Communications Surveys & Tutorials* 13, 2 (2011), 150–167.
- [2] K. Chard, S. Tuecke, and I. Foster. 2014. Efficient and Secure Transfer, Synchronization, and Sharing of Big Data. *IEEE Cloud Computing* 1, 3 (Sept 2014), 46–55. <https://doi.org/10.1109/MCC.2014.52>
- [3] Miguel Costa, Daniel Gomes, Francisco Couto, and Mário Silva. 2013. A survey of web archive search architectures. In *Proceedings of the 22nd International Conference on World Wide Web*. ACM, 1045–1050.
- [4] Clinton Gormley and Zachary Tong. 2015. *Elasticsearch: the definitive guide*. O’Reilly.
- [5] Ed Greengrass. 2000. *Information retrieval: A survey*. University of Maryland.
- [6] Vesna Hassler. 2005. Open source libraries for information retrieval. *IEEE Software* 22, 5 (2005), 78–82.