

Lineage first computing: towards a frugal userspace for Linux

Michael Dales, Patrick Ferris, Anil Madhavapeddy
University of Cambridge, UK

Data-science is a vital tool in tackling the ongoing climate-crisis, but it is one that has a significant cost to it also, in terms of energy used during execution and in terms of significant hardware investment required to run it. These are then amplified considerably in the exploratory nature of scientific research. We argue this is because the affordances of the operating system, particularly filesystems, make little effort to support reuse of computed artifacts directly or promote reuse through trusted lineage data.

However, it turns out that many of the userspace interfaces exposed by the Linux kernel entirely support a radically more efficient – a more *frugal* – model of computation than is currently the default. In this talk, we explore a new OS architecture which defaults to deterministic, reusable computation with the careful recording of side effects. This in turn allows the OS to guide complex computations towards previously acquired intermediate results, but still allowing for recomputation when required.

We use these interfaces to build a new Linux userspace where we put the workflow graph—containing relationships between tools, provenance and labelling—as the core of a system that drives how processes, data, and users interact. Our prototype shell, dubbed Shark, is a data-science first operating environment that is designed to both ensure efficient use of computation and storage resources, and to make it easy for non-experts to create pipeline descriptions from end results post-hoc. It does this by making the lineage graph of a data-pipeline the key concept that ties everything else together: by tracking how the data-pipeline evolves from experimental practice, and tracking what data has already been built and what hasn't we can prevent re-execution both at development and after publication.

1 WASTE IN DATA-SCIENCE

Whilst it is not intrinsic to the domain, we observe that in practice data-science is often wasteful of resources, driven in a large part by *uncertainty* in the process [5]: uncertainty in data (e.g., which datasets and versions have been used in the past), uncertainty in code (e.g., who ran which version of a stage), and uncertainty in the methodology (e.g., which existing published results can be verified and trusted in the future). The typical solution to resolve this uncertainty is expensive re-computation, which causes excess scope 2 and 3 emissions [6].

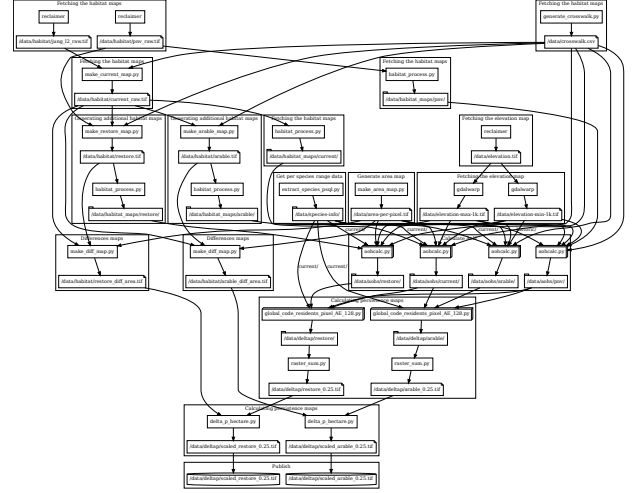


Figure 1: The LIFE biodiversity metric pipeline, showing the processing stages and intermediary datasets required. This pipeline is run for multiple scenarios and results in around 92 petabytes of image data.

We have seen this first hand being embedded in two large ecology projects at the University of Cambridge: the LIFE biodiversity metric [4] and the PACT forest assessment methodology [3]. Both pipelines are large multi-stage workflows like that shown in Figure 1, where the state of the computing is larger than a single person's working set and is distributed across multiple people with different specialities. Combined with the need to process petabytes of data and stages that can take weeks to run, this is a recipe for significant waste.

The current common set of OSs offer no way to query the lineage graph of results, and so it is easier to re-run a pipeline than investigate what went before, leading to wastage.

2 AUTOMATED DATA LINEAGE

A data lineage graph [1], that is a DAG of all inputs and processing stages in a pipeline, serves two purposes for us: it provides a way to attest what happened to humans, and it provides a way to a system run time to both execute the pipeline and to ensure parts already run are not run again.

Whilst data-scientists could manually build lineage graphs, these domain experts are typically focused on the primary

- [2] What is eBPF. <https://ebpf.io/what-is-ebpf/>.
- [3] BALMFORD, A., COOMES, D., DALES, M., FERRIS, P., HARTUP, J., JAFFER, S., KESHAV, S., LAM, M., MADHAVAPEDDY, A., MESSAGE, R., RAU, E.-P., SWINFIELD, T., WHEELER, C., AND WILLIAMS, A. Pact tropical moist forest accreditation methodology v2.1. Tech. rep., Cambridge Open Engage, 2024.
- [4] EYRES, A., BALL, T., DALES, M., SWINFIELD, T., ARNELL, A., BAISERO, D., DURÁN, A. P., GREEN, J., GREEN, R. E., MADHAVAPEDDY, A., AND BALMFORD, A. LIFE: A metric for quantitatively mapping the impact of land-cover change on global extinctions. Tech. rep., Cambridge Open Engage, 2024.
- [5] FERRIS, P., DALES, M., SWINFIELD, T., JAFFER, S., KESHAV, S., AND MADHAVAPEDDY, A. Uncertainty at scale: how CS hinders climate research. *Undone Computer Science* (2024).
- [6] GHG PROTOCOL. Corporate standard. Tech. rep., GHG Protocol, 2015.
- [7] KNUTH, D. E. Literate Programming. *The Computer Journal* 27, 2 (01 1984), 97–111.
- [8] SHAW, M. Myths and mythconceptions: what does it mean to be a programming language, anyhow? *Proceedings of the ACM on Programming Languages* 4, 234 (2020), 1–44.